

CS224v

Agentic AI

Lecture 1: Introduction

Monica Lam

Core User-Facing Capabilities of LLMs



Conversational Q&A

Direct chat with multi-turn context retention across a session



RAG & Retrieval

Grounding answers in external docs and live web search, with citations



Function Calling

The model decides when to call an API, database, or external tool



Agentic Workflows

Multi-step planning, tool use, and self-correction with less supervision



Computer Use

Operating a GUI directly — clicking, typing, navigating on a user's behalf



Code Execution

Running and iterating on code in a live sandbox, not just generating text



Extended Reasoning

Step-by-step chain-of-thought for math, logic, and complex problem solving



Multimodal I/O

Understanding and generating images, audio, and video alongside text

TODAY'S AI AGENTS ARE NOT DEPLOYED IN HIGH-STAKES SERVICES
(MEDICINE, SCIENCE, FINANCE, LAW)

FUNDAMENTAL TECHNIQUES TO IMPROVE:
ACCURACY, ACCOUNTABILITY, LONG HORIZON, LONG CONTEXT

IMPACT
AFFORDABLE HIGH-STAKES SERVICES
KNOWLEDGE DISCOVERY ACCELERATION

LLM Hallucinations

Date	Incident	Why it matters	Citation
May 2023	Mata v. Avianca	ChatGPT-assisted legal filing included fake cases and citations ; court imposed sanctions.	NYT, 5/27/2023
Dec 2023	Michael Cohen filing	AI-generated fabricated legal citations appeared in a high-profile legal matter.	Reuters, 12/29/2023
Dec 2023	Chief Justice Roberts report	Chief Justice Roberts urged caution when using AI in legal contexts.	U.S. Supreme Court, 12/31/2023
Jan 2024	Stanford legal hallucination study	Found 69%–88% hallucination rates on legal tasks across GPT-3.5, PaLM 2, and Llama 2 .	Dan Ho, Stanford arXiv, 1/11/2024
Feb 2024	Moffatt v. Air Canada.	Company held responsible after chatbot gave incorrect fare/refund guidance .	The Guardian 2/16/2024
Feb 2024	Colorado sanctions case	Lawyers sanctioned for AI-generated filings with fake court citations .	Reuters, 2/22/2024

HALLUCINATION IS STILL RAMPANT WITH RAG
TODAY

SCALED CONSUMER PRODUCTS

NYT, April 7, 2026

Estimated Hallucination Rate:
10%

The New York Times

◆ A.I. Overview

The **Neuse River** borders the western side of Goldsboro, North Carolina.

How Accurate Are Google's A.I. Overviews?

JOURNALISM

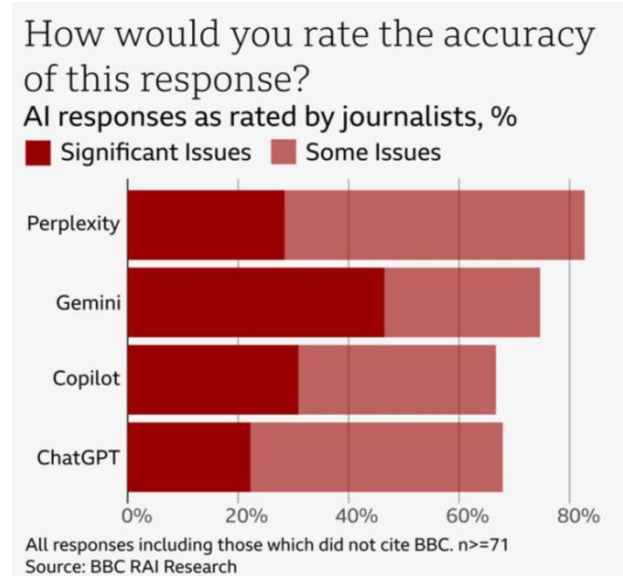
Perplexity, Gemini, Copilot, ChatGPT Hallucinate With RAG

BBC, Feb 11, 2025

51% of AI answers to news problems have significant issues.

19% of AI answers which cited BBC content introduced factual errors.

13% of the quotes sourced from BBC articles were either altered or didn't actually exist in that article.



LEGAL

Stanford, June 26, 2024

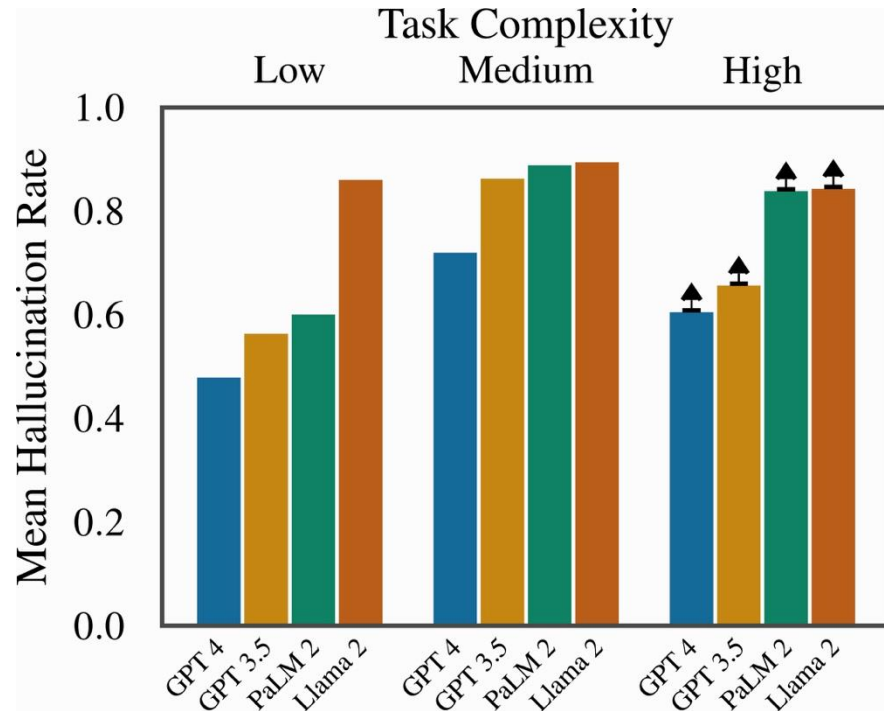
Large Legal Fictions: Profiling Legal Hallucinations in Large Language Models

Matthew Dahl ✉, Varun Magesh, Mirac Suzgun, Daniel E Ho Author Notes

Journal of Legal Analysis, Volume 16, Issue 1, 2024, Pages 64–93, <https://doi.org/10.1093/jla/liae003>

[jla/liae003](https://doi.org/10.1093/jla/liae003)

Published: 26 June 2024



Hallucinated citations are polluting the scientific literature. What can be done?

Tens of thousands of publications from 2025 might include invalid references generated by AI, a *Nature* analysis suggests.

By [Miryam Naddaf](#) & [Elizabeth Quill](#)

<https://www.nature.com/articles/d41586-026-00969-z>

MEDICINE

NBC News

May 13, 2026

ARTIFICIAL INTELLIGENCE

Most U.S. doctors are quietly using this AI tool. Few patients know about it.

OpenEvidence, an AI-powered medical search tool, has become a fast friend to America's doctors and is now used by nearly two-thirds of physicians.

<https://www.nbcnews.com/tech/tech-news/openevidence-ai-doctor-medical-physician-login-app-what-npi-uptodate-rcna341064>

Sept 2026

MedXPertQA:
Expert-Level Specialty Board Questions

Accuracy: 72.8%

FINANCE: AI Hallucination not Publicized

An enterprise:

“The bot says ‘the investment is made’ but it was not.”

A large consulting firm:

“Evaluation shows many hallucinations in an enterprise chatbot”

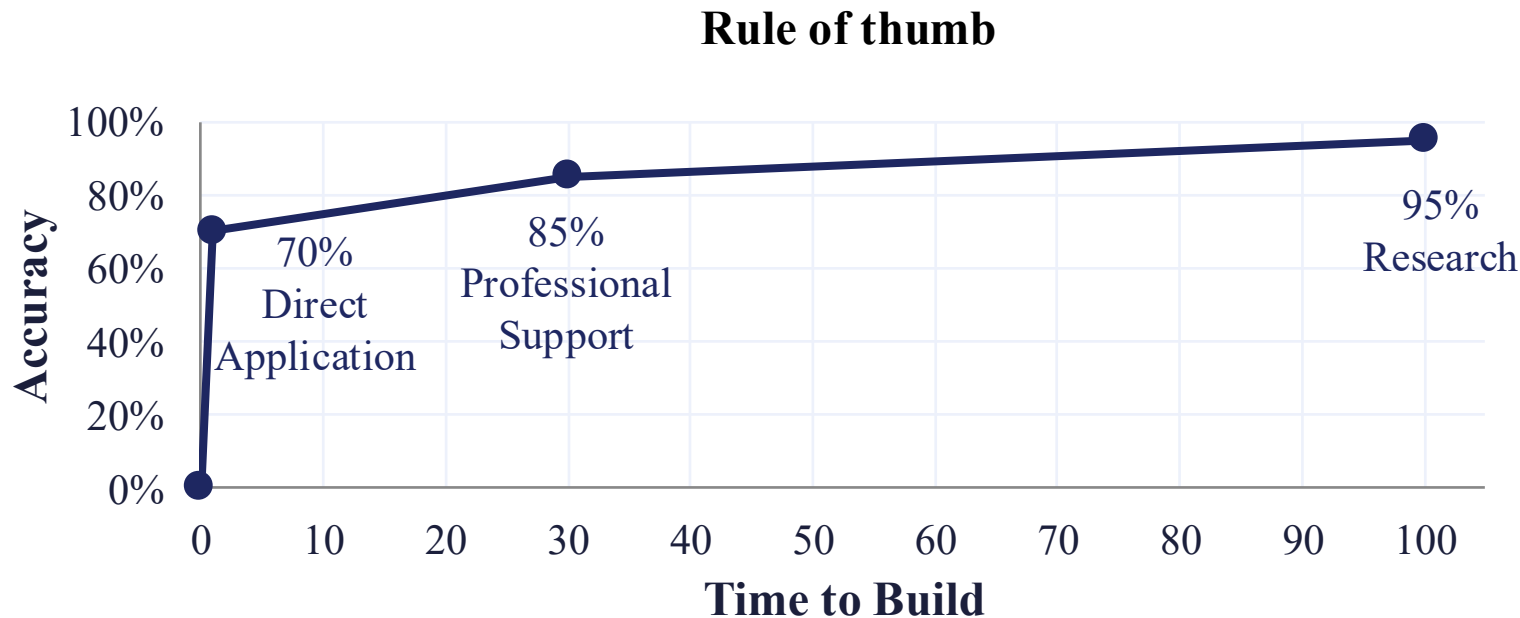
A popular chatbot vendor:

“Does not sell bots for high-stakes applications.”

A large enterprise infrastructure vendor:

“No enterprises deployed AI in high-stakes applications.”

High-Stakes LLM Agents are Hard to Build!



The Last 10% is 90% of the problem.

LLM Weaknesses

from Next-Word Prediction Training

Out-of-Distribution Tasks

- **New** information
 - **Long**-tail information
 - **Long** instructions
 - **Long** documents
 - **Long** horizon
 - **Long** contexts
- } RAG (Retrieval Augmented Generation)
- } Context management
e.g. automatic compression

FOCUS OF CS224V

TODAY'S AI AGENTS ARE NOT DEPLOYED IN HIGH-STAKES APPS
(MEDICINE, SCIENCE, FINANCE, LAW)

FUNDAMENTAL TECHNIQUES TO IMPROVE:
ACCURACY, ACCOUNTABILITY, LONG HORIZON, LONG CONTEXT

IMPACT
AFFORDABLE HIGH-STAKES SERVICES
KNOWLEDGE DISCOVERY ACCELERATION

Outline of This Lecture

1. Problem: LLM Hallucination in High-Stakes Applications
2. Basic concept: Decompose problems into reliable LLM tasks
3. Course information
4. Course syllabus preview
 - Stanford A-OS Agentic Operating System
 - Long-Horizon, Long-Context Custom-Facing Agents
 - Long-Horizon, Long-Context Researcher for Science
 - Accountable Decision Making

WHY IS THERE HALLUCINATION WITH RAG

1. RETRIEVAL DOES NOT RETURN THE ANSWER
2. LLMs SUPPLY THE ANSWER FROM MEMORY

WikiChat: Stopping the Hallucination of Large Language Model Chatbots by Few-Shot Grounding on Wikipedia

Sina J. Semnani Violet Z. Yao* Heidi C. Zhang* Monica S. Lam

Computer Science Department

Stanford University

Stanford, CA

{sinaj, vyao, chenyz, lam}@cs.stanford.edu

- Conversational Q&A in open domain with 97% accuracy
- Winner of the Wikimedia's Research of the Year Award, 2024
- In Proceedings of Association for Computational Linguistics: EMNLP, Singapore, December 2023.

WikiChat (7 Prompts)

Traditional (Factuality)

1. Formulate query from input
 - Retrieve documents
2. Filter each retrieved doc

LLM (Conversationality + Factuality)

3. Ask GPT to generate answer
4. Extract claims
5. Fact-check/remove each claim
 - Retrieve documents

-
6. Draft
 7. Refine

LLM calls: $5 + n + c$

n : # documents retrieved based on user queries

c : # claims generated

THE PROBLEM IS HARDER THAN EXPECTED

3 STUDENTS, 4 MONTHS

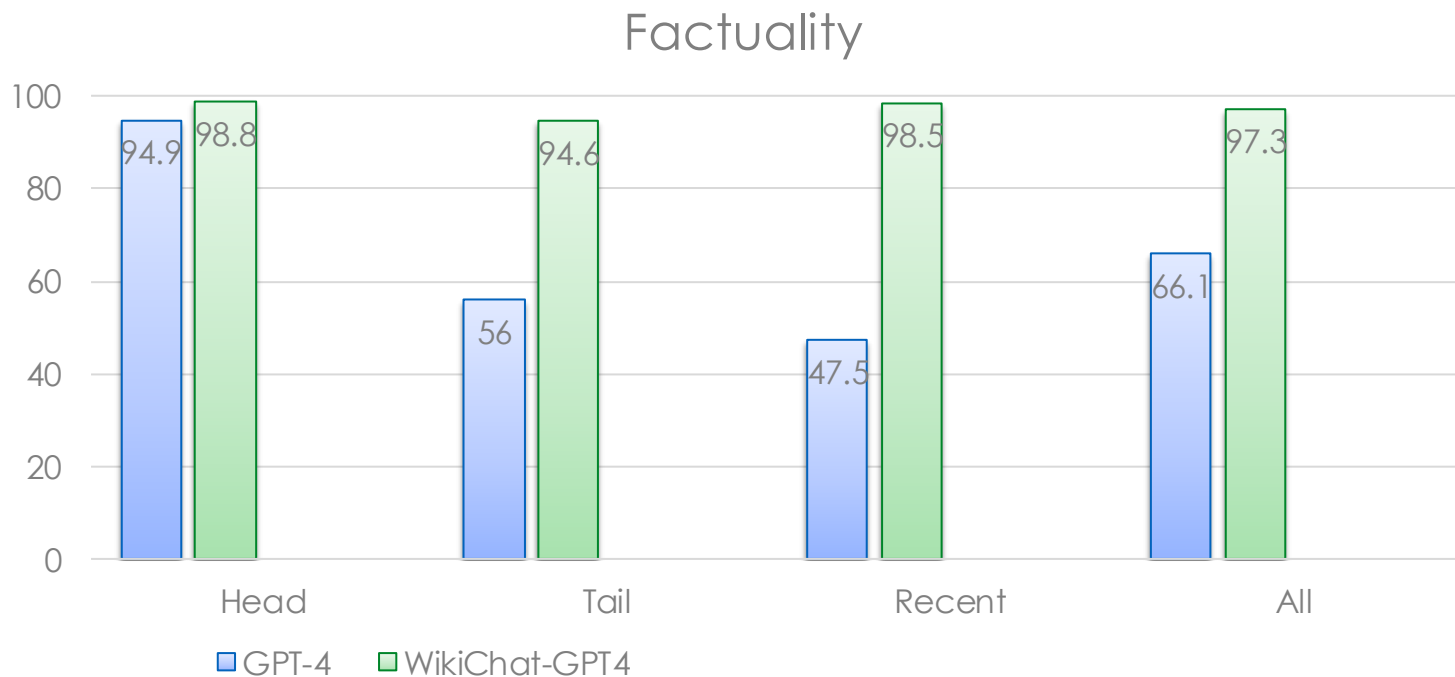
THE DEVIL IS IN THE DETAILS

DECOMPOSE PROBLEMS INTO RELIABLE LLM TASK

SPEED CAN BE REGAINED VIA DISTILLATION

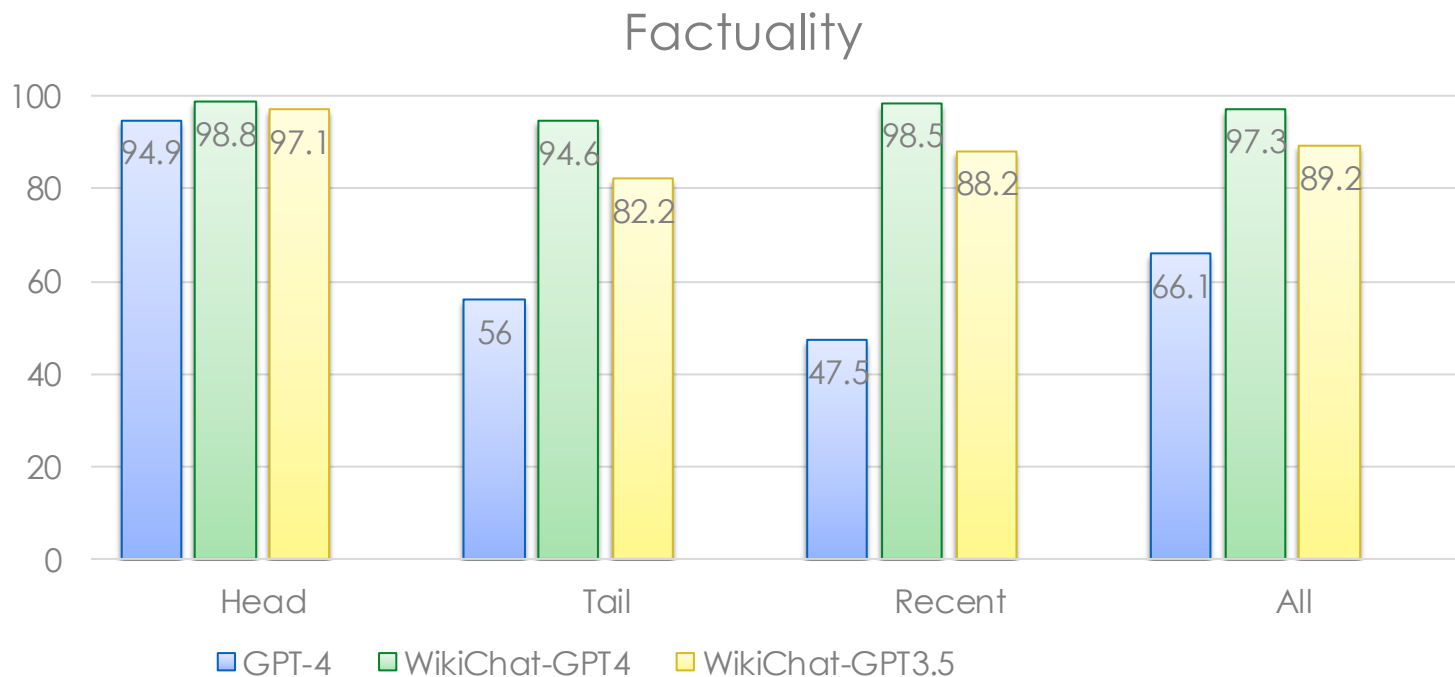
PUBLICLY AVAILABLE AS OPEN SOURCE

Simulated Dialogues: WikiChat-GPT4



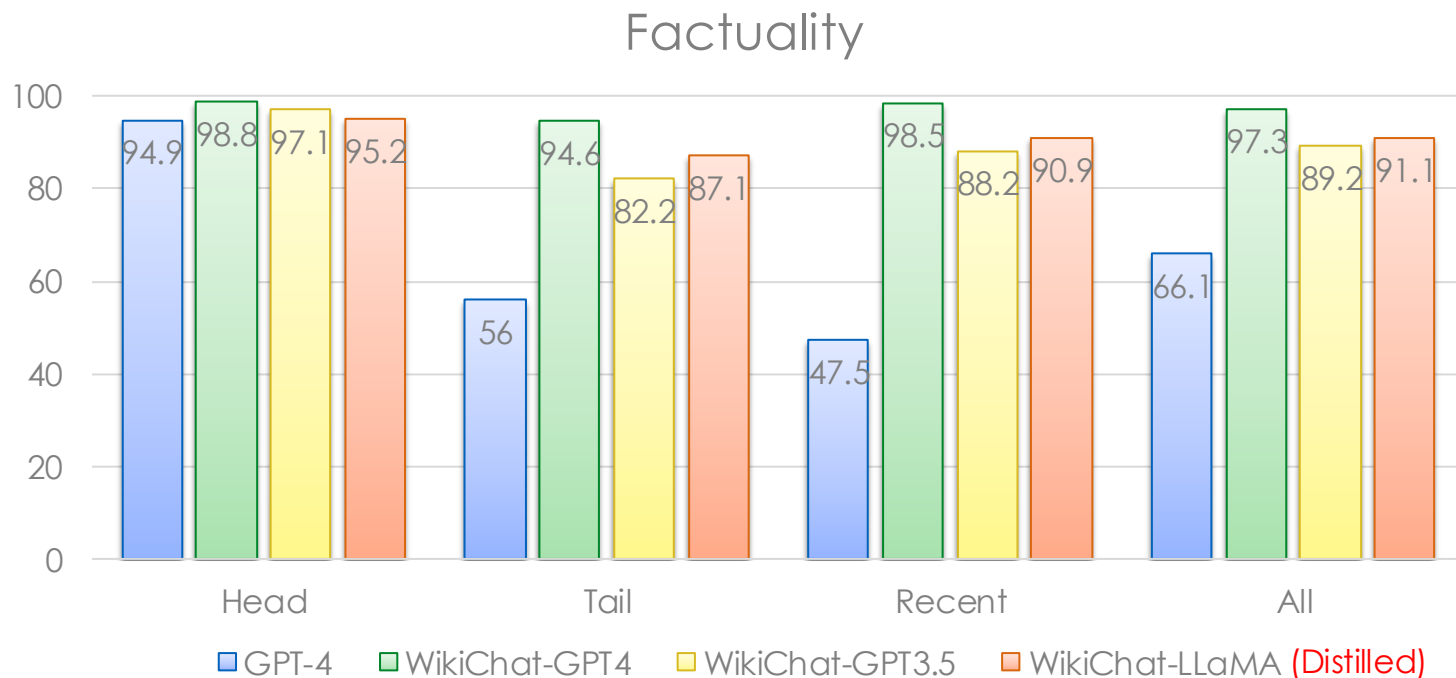
Semnani, Sina et al. WikiChat: Stopping the Hallucination of Large Language Model Chatbots by Few-Shot Grounding on Wikipedia, EMNLP Findings 2023

Simulated Dialogues: WikiChat-GPT3.5



Semnani, Sina et al. WikiChat: Stopping the Hallucination of Large Language Model Chatbots by Few-Shot Grounding on Wikipedia, EMNLP Findings 2023

Simulated Dialogues: WikiChat-LLaMA



Semnani, Sina et al. WikiChat: Stopping the Hallucination of Large Language Model Chatbots by Few-Shot Grounding on Wikipedia, EMNLP Findings 2023

Real User Evaluation

- Real user evaluation was seldom performed before LLM because of poor accuracy
- **Most important metric!**

User study: User reads the first sentence of a new Wikipedia page

Model	User Rating (out of 5)	Factuality
GPT-4	3.4	42.9%
WikiChat using GPT-4	3.8	97.9%

GPT-4: Users are not even aware that over half of the statements are false

FUTURE LLM MODELS WILL STILL HALLUCINATE
A HIGHER RATE FOR HARDER QUESTIONS

1. DO NOT TRUST DEMOS
2. UNDERSTAND THE UNDERLYING TECHNOLOGY
3. EVALUATE SYSTEMATICALLY

Outline of This Lecture

1. Problem: LLM Hallucination in High-Stakes Applications
2. Basic concept: Decompose problems into reliable LLM tasks

3. Course Information

4. Course syllabus preview

Stanford A-OS Agentic Operating System

- Long-Horizon, Long-Context Custom-Facing Agents
- Long-Horizon, Long-Context Researcher for Science
- Accountable Decision Making

History of this Course

- Actively create and advance new technologies
Adopted by real Users
- Course publications: 2023(2), 2024(5), 2025(6), 2026(2)
- Each CS224V class builds on previous year's research
 - Select class projects continue as research in 2026-2027
 - Results used in next year's CS224V

Secret to Success

Working software supports increasingly complex research

Live Demos: www.knowledge.org



Watch recordings from our 2025 Workshop

The Stanford Open Virtual Assistant Lab  presents a pilot of

A Public AI Assistant to World Wide Knowledge

Advancing global education and accelerating knowledge discovery with Artificial Intelligence

Access the WWKnowledge Assistant:

BROWSING

Collaboratively research the web with Storm. Create research articles with comprehensive and accurate bibliographies.

[Get your own article now >](#)



Chat with popular Wikipedias in different languages

Start chatting in English, Deutsch, Français, Español, 日本語, Русский, Português, Italiano, 中文, فارسی ... [>](#)



DEEP-DIVE

The African Times newspaper (1867 and 1910), available digitally for the first time.

[Chat with the African diaspora in that period >](#)



Co-hosted with Wikidata from Wikimedia, explore the world's knowledge

[Deep-dive into 15 billion facts and ask questions >](#)



Explore FEC campaign donation data (built in partnership with Big Local News)

[Chat and ask questions about campaign donations >](#)



Contributing to WWKnowledge:

Upload your Sources

We are always welcoming new data sources to include in WWKnowledge.

[Upload your CSV >](#)

[Upload your corpora >](#)

Learn the Genie Framework

WWKnowledge is developed on the Stanford OVAL Lab's Genie Agent Framework

[Learn to build with Genie >](#)

Learn at our Workshops

Learn about the Genie framework from its creators, and about funding and collaboration opportunities

[Watch the 2025 recordings >](#)

In Collaboration With



Course Objectives

Prepare you to push the limits of agentic AI

Lectures:

1. How to build a basic agent? (2 homeworks in first 2 weeks)
2. Challenges of today's LLM-based agents
3. Core techniques to build reliable agents
4. Advanced techniques to push the limit of LLMs

Hands-on quarter-long projects on topics of your choice

1. Create or improve general tools
2. Use any tools on any applications

Purpose of the Assignments

Goal: Prepare you for your project proposal

- Hands-on experiment with state-of-the-art tools
- An example of a compelling, challenging, yet doable project

Homework: Rare Diseases

- In partnership with Rare Genomics Institute
(A winner of the Anthropic's Rare Disease AI Grant Program)
- About 400 million people live with one of [more than 7,000 rare diseases](#)
- Goal of the homeworks:
 - Deep research on rare diseases
 - Extraction of diagnosis information
 - Diagnosis of rare diseases using patient records

Project Apprenticeship

- Assistance with project selection: Hardest part in research!
 - We have about 20 suggested projects on the website
 - Student-initiated projects are also welcome
- Weekly group mentorship meeting
 - We want to make you succeed!

Project Mentorship

All homeworks and projects are to be done in pairs

- Week 4: Project proposal, with a weekly plan
- Weeks 5-10 (excluding Thanksgiving break):
 - Submit a written weekend update (every Monday)
 - Group meeting with mentors during the week
- Week 11: Poster presentation (Dec 2)
- Final project report due Dec 8, 2026.

This Course

	Grade
Participation	15%
Assignment	25%
Final Project	60%

Participation includes

- Class attendance and participation
- Ed discussion
- Meetings with project mentors

Outline of This Lecture

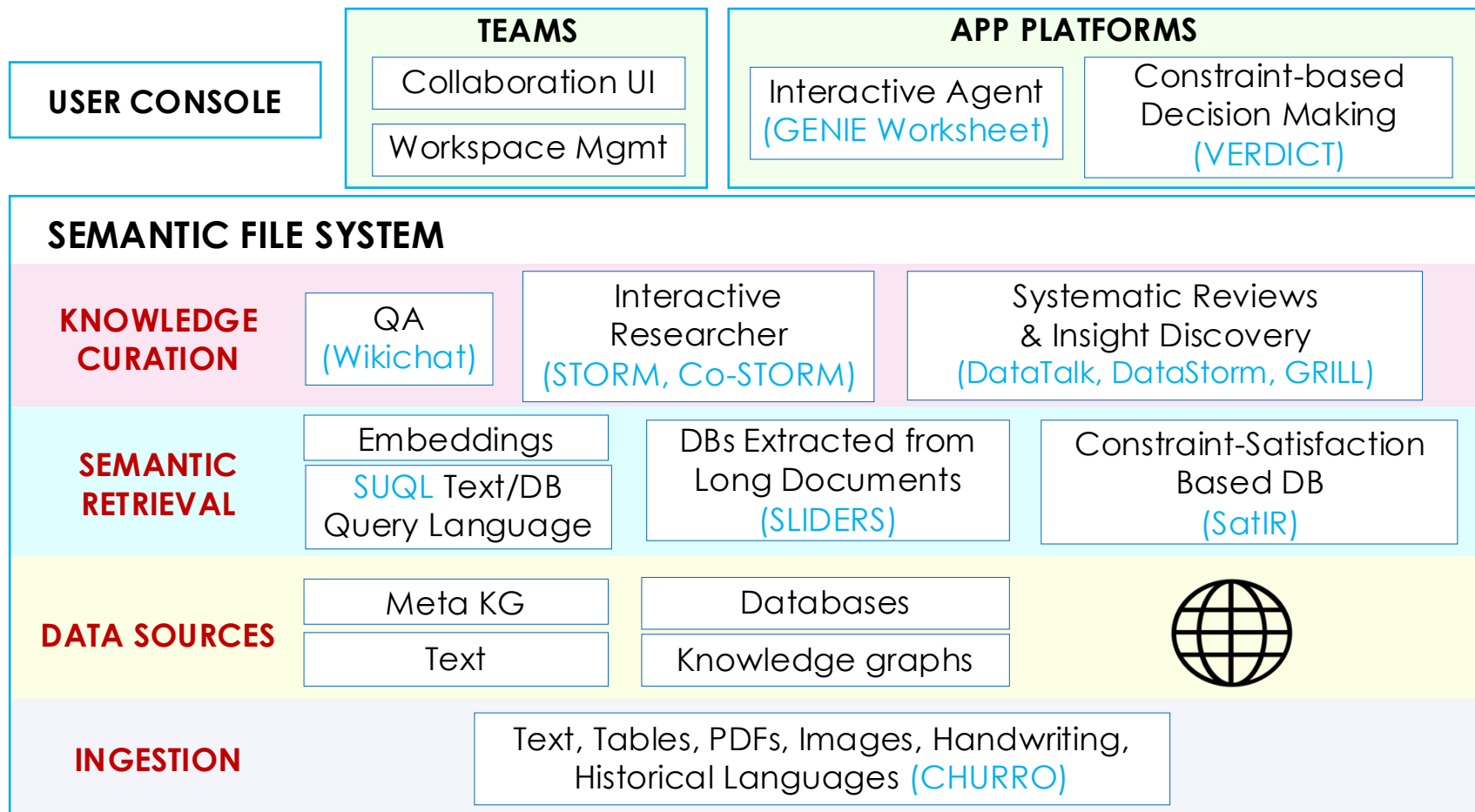
1. Problem: LLM Hallucination in High-Stakes Applications
2. Basic concept: Decompose problems into reliable LLM tasks
3. Course Information

4. Course syllabus preview

Stanford A-OS Agentic Operating System

- Long-Horizon, Long-Context Custom-Facing Agents
- Long-Horizon, Long-Context Researcher for Science
- Accountable Decision Making

Stanford AGENTIC OS (A-OS)



Basic Knowledge Access

Topic	Key Ideas	Significance
Transcribing historical texts: CHURRO VLM 2025	Fine-tune on manuscripts & prints (22 centuries, 45 languages)	vs. Gemini-Pro: 15.5x cheaper, similar accuracy history.genie.Stanford.edu
1-Pager Research Storm/Co-Storm 2024	Search-read-search loop. Interactive exploration into the Unknown-Unknowns.	1st deep researcher 1 Million organic users storm.genie.Stanford.edu
Accessing Texts, Databases, Knowledge graphs SUQL 2024	Structured/Unstructured Query Lang for hybrid DB + text. 1st use of agentic workflow for KGs.	Deployed in practice: DataTalk for journalists Wikidata agent datatalk.biglocalnews.org Spinach Wikidata bot

HIGH-STAKES AGENTS

ACCURACY, ACCOUNTABILITY

LONG-CONTEXT, LONG-HORIZON

1. **CUSTOMER-FACING AGENTS**

2. RESEARCHER FOR SCIENCE

3. DECISION MAKING

High-Stakes Customer-Facing Agents

- How to stop the agents from hallucinating?
- **Key Idea:** Full specifications are necessary
 - Not just functions to call!
- **Genie Worksheet: A Domain-Specific Language**
 - Specification of knowledge corpora + task workflow
 - Similar to workflow on webpages
 - Run-time to manage the long context

Real User Evaluation

Restaurant Reservation:

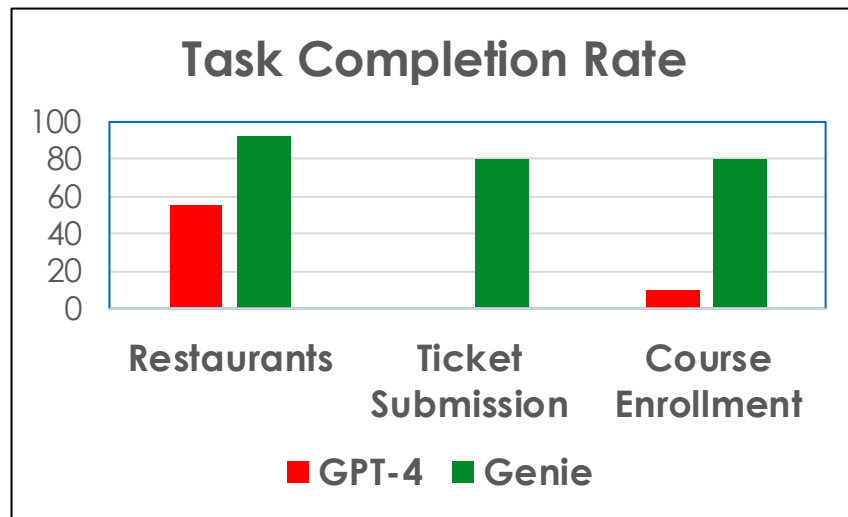
Real-life database from Yelp

Ticket Submission:

A subset of Service Now

Course Enrollment:

Multiple real-life databases of courses at Stanford University.



GPT-4 results are not acceptable

- Low completion rate: even 54.5% on restaurant is not good enough!
- Course: Hallucinates non-existent courses despite using knowledge base results

HIGH-STAKES AGENTS

HIGH-RECALL, LONG-CONTEXT, LONG-HORIZON

1. CUSTOMER-FACING AGENTS

2. **RESEARCHER FOR SCIENCE**

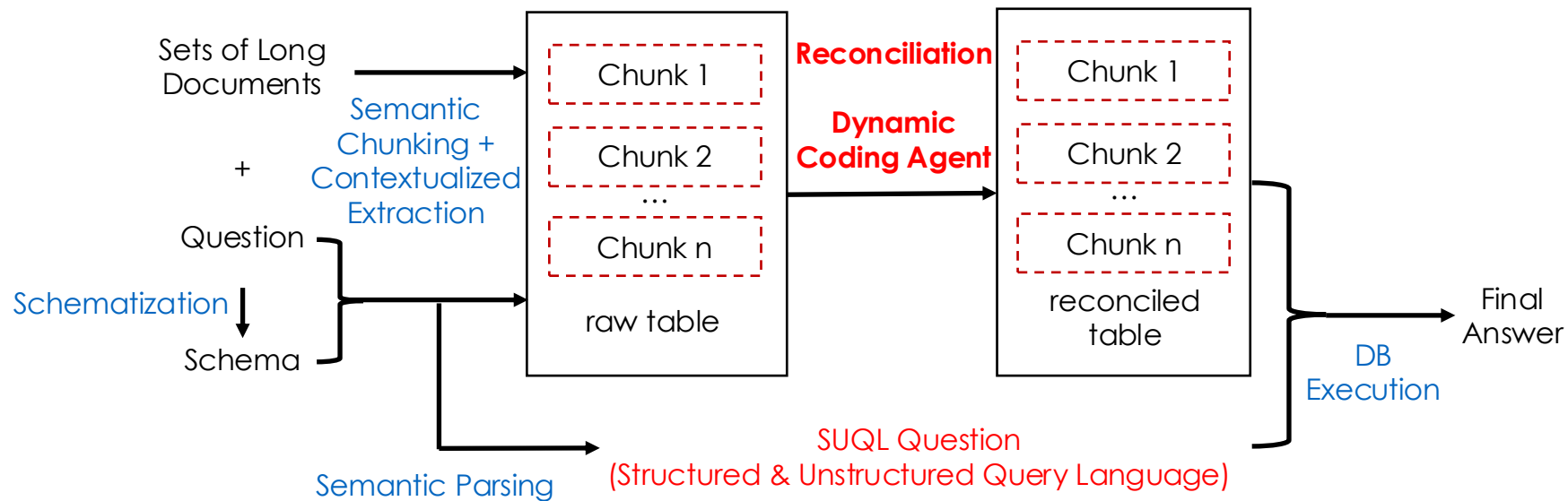
3. ACCOUNTABLE DECISION MAKING

CONTEXTS ARE NEVER LONG ENOUGH

ACCOUNTABILITY:
NEEDS TO TRACK THE EVIDENCE

ANALYZING LARGE SETS OF LONG DOCUMENTS

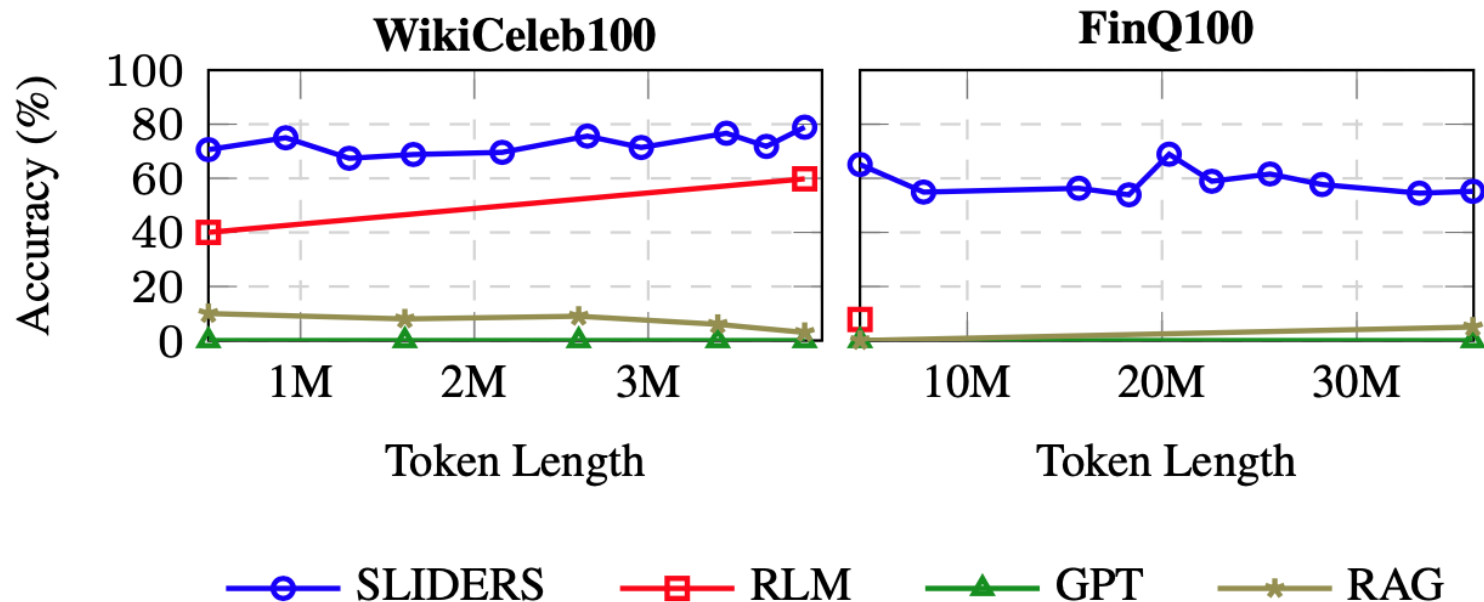
SYSTEMATIC REVIEWS



Key: Paragraph level provenance

ANALYZING LARGE SETS OF LONG DOCUMENTS

SYSTEMATIC REVIEWS



Long-Horizon Deep Research

Creates hypotheses, finds supports and refutations iteratively

RNA Sequencing Research

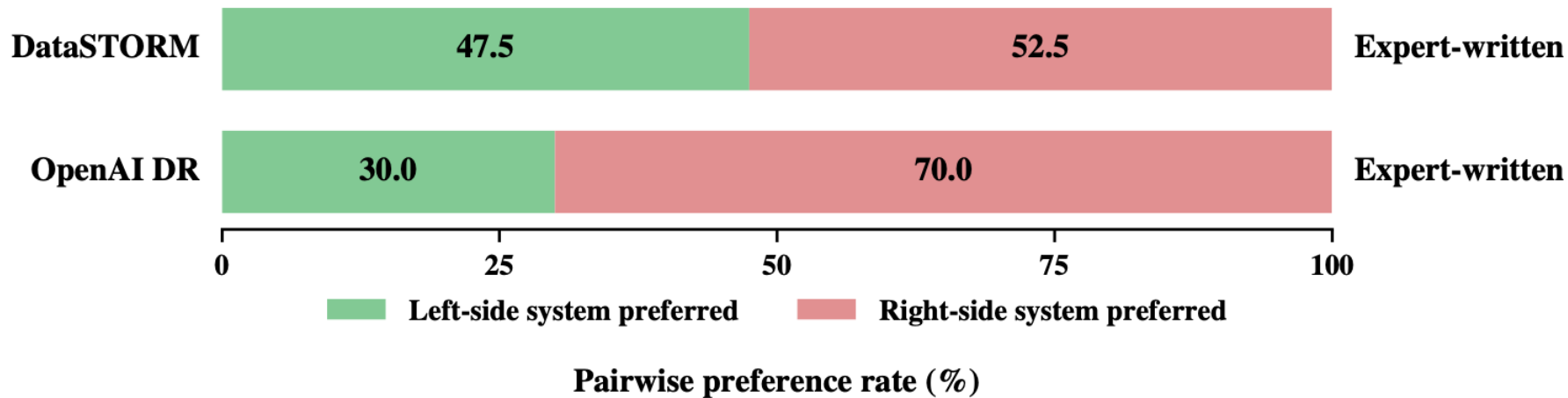
- *“What is the gene responsible for cancer drug resistance?”*
- Our results informed the direction of the experiment design

Rare Diseases

- *“For patients with genetic alterations in the Stxbp1 gene, what are the different clinical phenotypes associated with the gene, detailing the specific variations that has impact?”*
- Researcher: “Grill is an order of magnitude better than Codex”

USER EVALUATION

Automatically derive insights from databases
Insights on international conflicts



HIGH-STAKES AGENTS

HIGH-RECALL, LONG-CONTEXT, LONG-HORIZON

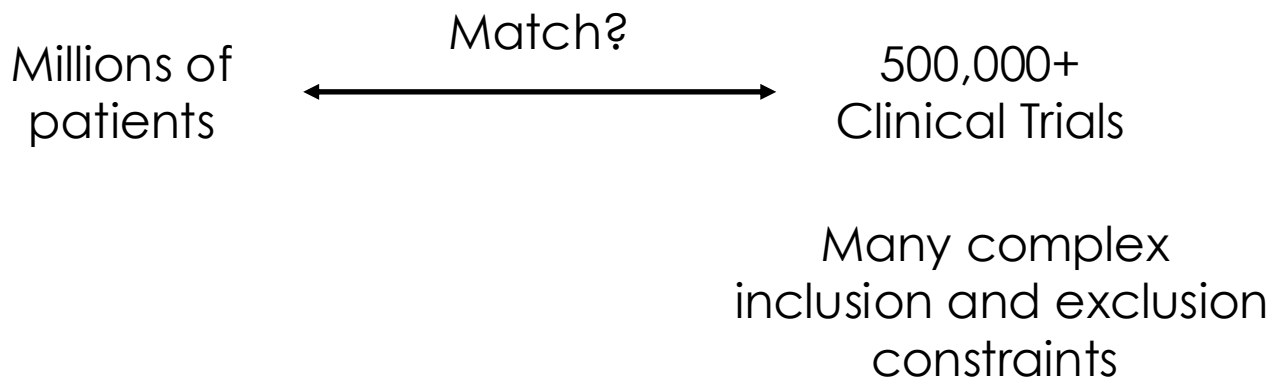
1. CUSTOMER-FACING AGENTS
2. RESEARCHER FOR SCIENCE
3. **ACCOUNTABLE DECISION MAKING**

High-Stakes Clinical Decisions

- Clinical trial matching
- Rare disease diagnosis
- Treatment guideline conformance

CLINICAL TRIAL MATCHING

Clinical Trials are critical for medical advances
80% Clinical trials fail for missing enrollment targets



Previous State of the Art

- Retrieval of eligible trials: Embedding-based retrieval
 - Low precision: returns many ineligible trials – wastes time
 - Low recall: misses potentially life-saving trials
- LLM-based trial matching: 90% Accuracy
 - Decision may still be wrong 10% of the time
 - Clinicians need information for continued health care
 - They need **accountability**

Accountability for High-Stakes Decisions

- **Verified:** Verified rationale
 - LLM: *the patient does not meet the eligibility criteria*
- **Faithful:** Explicit assumptions
 - LLM: *No clinically significant cardiovascular disease* satisfied by *controlled hypertension*
- **Consistent:** Policy adherence
 - LLM: Does not adhere to given policies:
Missing HIV status should remain unknown
- **Actionable:** Pivotal conditions
 - LLM: *Ineligible because ECOG > 1*
But changing *ECOG = 0* because of other conditions

Case-level accuracy.

Should we trust it?

Over-ridable.

There are lots of missing data:

Has it made wrong assumptions?

Steerable and Fair.

Ambiguity needs decision policies.

Are all patients treated equally?

Actionable.

Guarantees that decision will flip
if actions are taken.

Core Idea

	Language Understanding (LLM)	Decision Making (Theorem Solver)
Function	Clinical trials Patient records $\xrightarrow{\text{LLM}}$ SMT	$\text{SMT} \xrightarrow[\text{Solver}]{\text{SMT + MaxSAT}}$ Accountable Decision
Correctness	Formalization forces disclosure of ambiguity & missing values. SMTs may not be 100% faithful. But they are interpretable .	Verifiably correct with respect to SMTs Yes/No decision with A proof for the derivation Assumptions: all the residual constraints Adherence to policies is guaranteed! Pivotal conditions (MaxSAT)

Experimental Results (TREC 2021)

Backbone	CoT	LLMMATCH	ZSPM	VERDICT
GPT-5-mini	0.711	0.776	0.759	0.828
Claude Haiku 4.5	0.704	0.587	0.615	0.800
Qwen2.5-7B	0.663	0.649	0.695	0.738
Qwen2.5-7B (with distillation)				0.830

Clinician Study

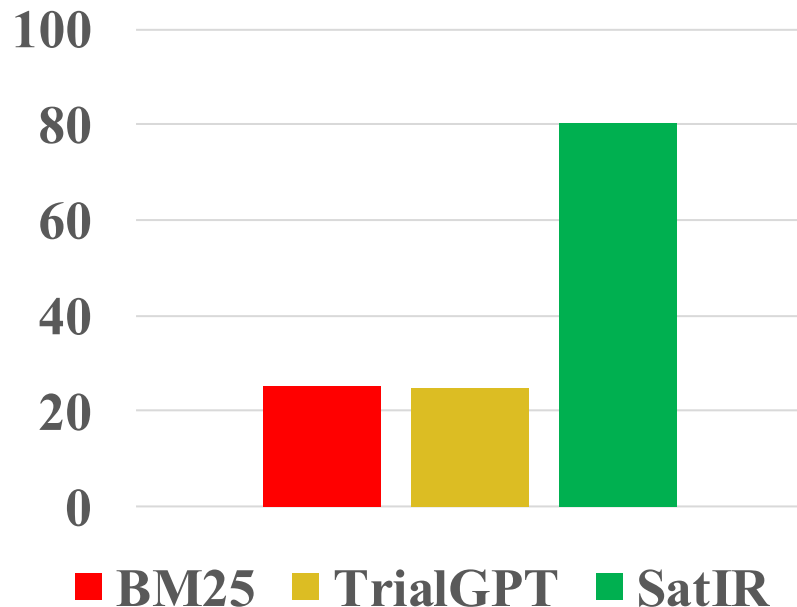


Trial Retrieval

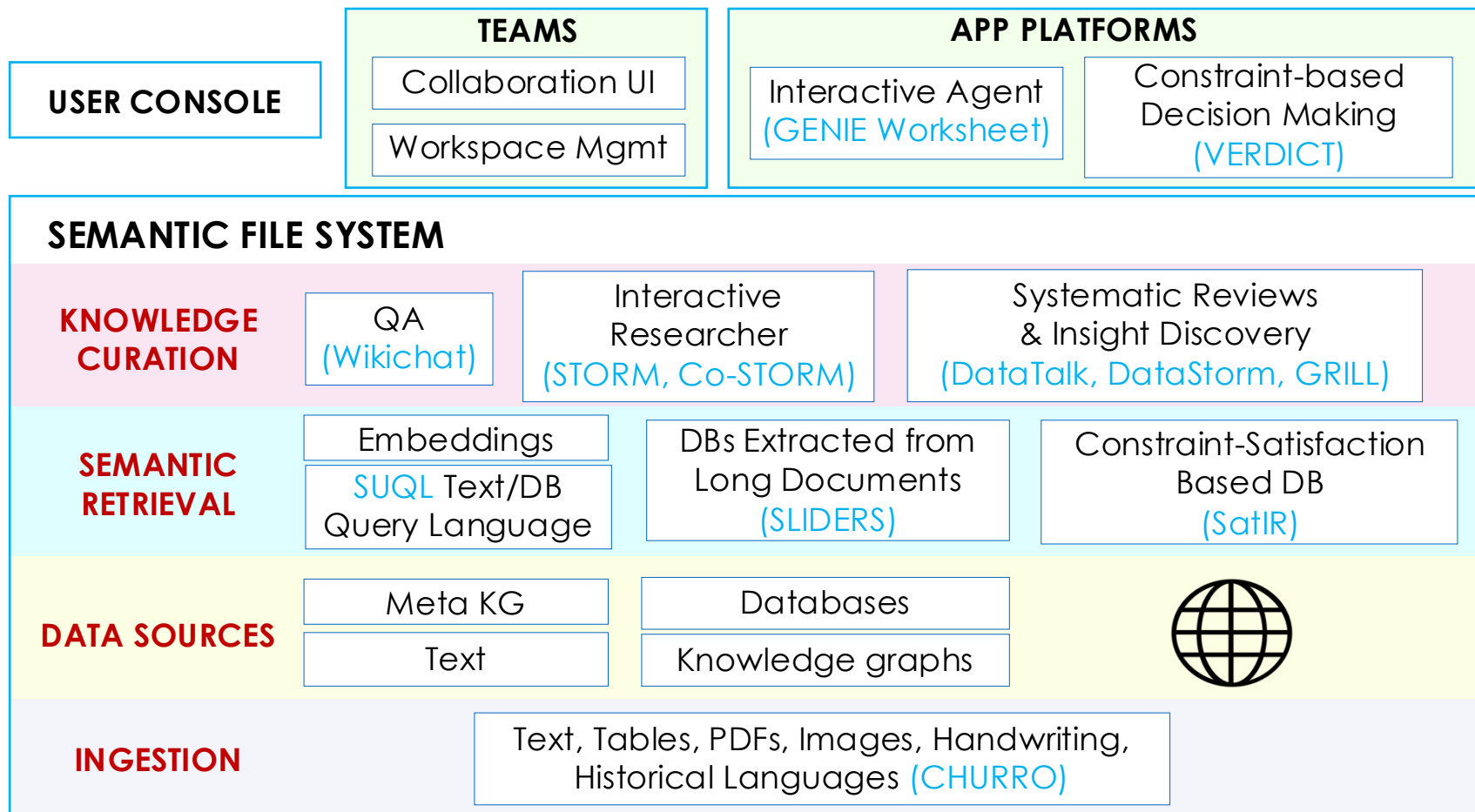


- Cost of matching patients with SMTs for all pairs is too high
- Project SMTs to DBs
- 0.15s/patient over 3,621 trials

TREC2022 Recall



Stanford AGENTIC OS (A-OS)



Course Schedule at a Glance

Dates	Lectures / Homeworks	Projects
9/23 - 10/ 5	Introduction; Hallucination-free chat, researcher 2 homeworks	Research Project Ideas
10/ 7 - 10/21	Structured and unstructured data retrieval and query language	Student-initiated ideas Project discussions Project proposals
10/26 - 11/2	Agentic workflow, task-oriented agents, Deep researcher.	Weekly meetings with mentor
11/ 4 - 11/11	Insight discovery, accountable decision- making. meta-optimizations	Weekly meetings with mentor
11/16 – 11/18	VLM:Historical Transcription; Training LLMs	Weekly meetings with mentor
11/23 - 11/25	<i>Thanksgiving</i>	
12/2		Final project posters (3:00-5:40)