

GENERATIVE AI SERIES: DAY 23

Evaluating Your RAG System: Is It Actually Working?



You've built your brilliant RAG system. But... how do you prove it's good? Let's dive into the most critical step: measurement. 📈

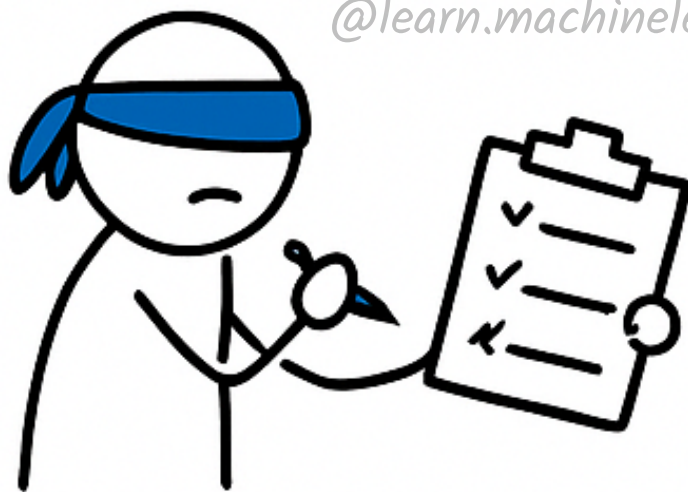
Without Data, You're Just Guessing!

Relying on a few test questions to see if your system works is like flying blind. For a reliable, professional system, you need objective numbers.

Guesswork

- The answers feel pretty accurate.

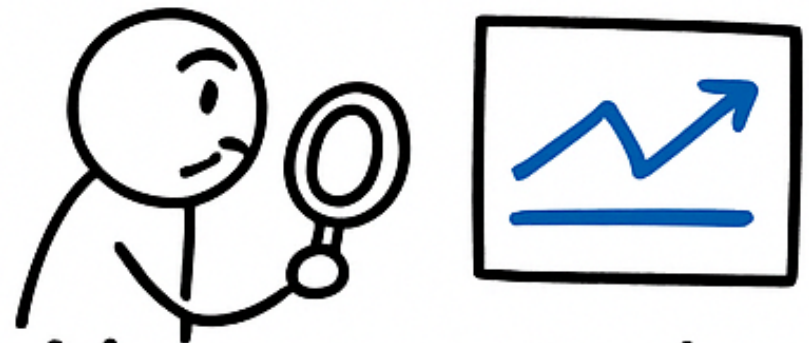
@learn.machinelearning



Guesswork

Measurement

- Our answer relevancy score is 0.92, an improvement of 15% after re-ranking.



Measurement

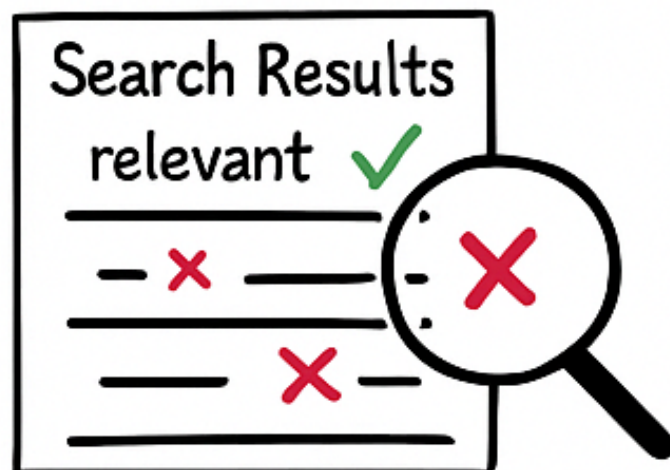
Evaluation turns your gut feeling into hard data.

Metrics 1 & 2:

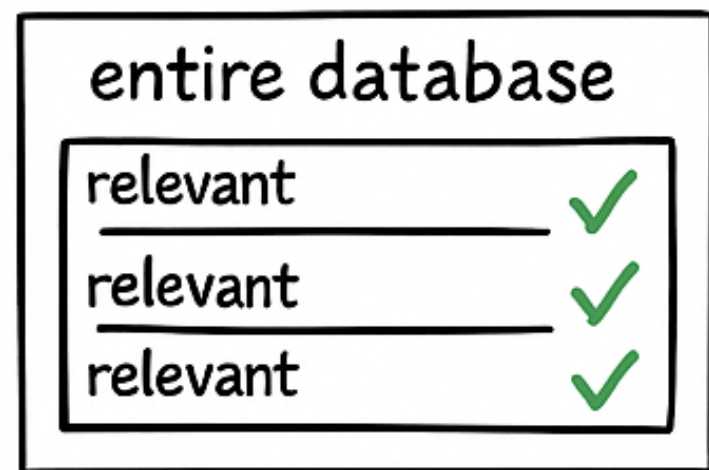
Context Precision & Recall

Before we even get an answer, we must check the quality of our retrieved context.

- **Context Precision:** Of all the chunks we found, how many were **actually** relevant?
(Are we finding junk?) *@learn.machinelearning*
- **Context Recall:** Of all the relevant chunks that exist in the database, how many did we **actually** find? (Are we missing important info?)



Context Precision



Context Recall

Goal: **High Precision AND High Recall**

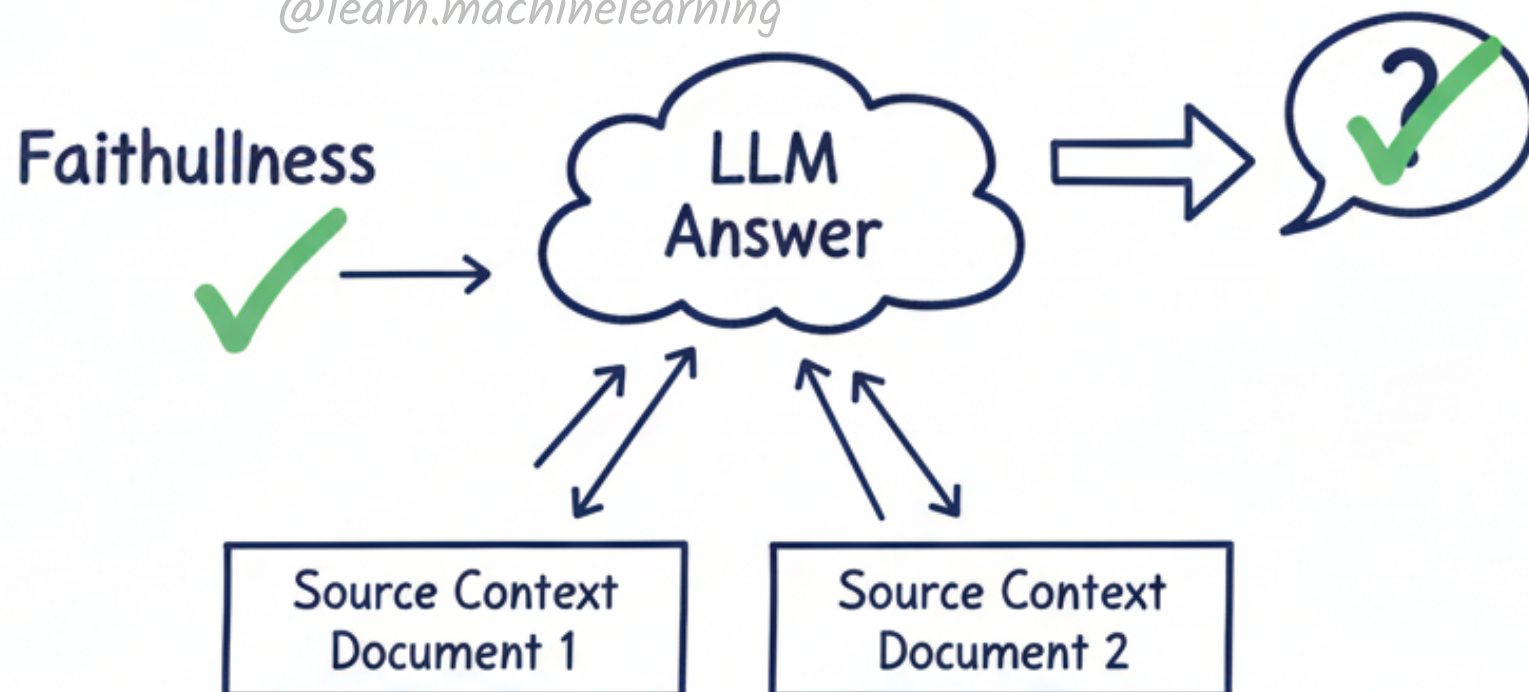
Metrics 3 & 4:

Faithfulness & Relevancy

Answer Faithfulness: Is the answer only based on the provided context? (Is it making stuff up / hallucinating?)

Answer Relevancy: Does the answer directly address the user's original question? (Is it on-topic?)

@learn.machinelearning



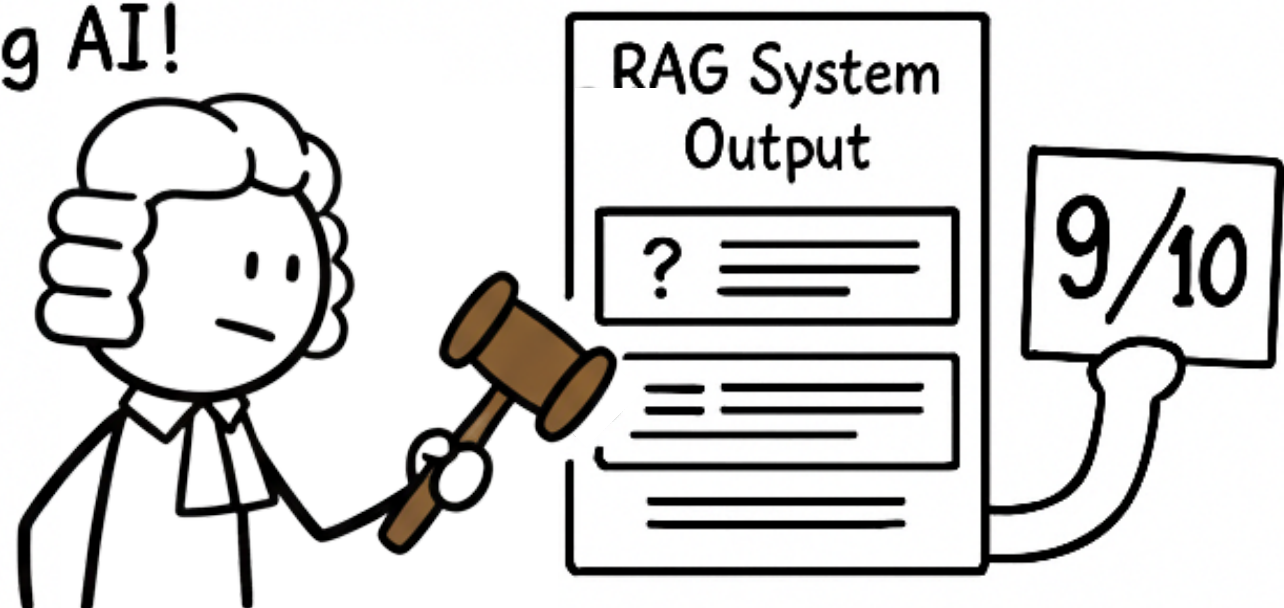
An answer can be factual but completely irrelevant!

How to Automate This?

Meet RAGAs!

- Checking these metrics manually is impossible at scale. So, we use modern tools!
- ✓ RAGAs (RAG Assessment): An open-source framework that uses LLMs to evaluate other LLMs
- ✓ How it works: You provide the question, the retrieved context, and the final answer. RAGs then uses clever prompting to automatically score your system on faithfulness, relevancy, and more
- ✓ It's AI judging AI!

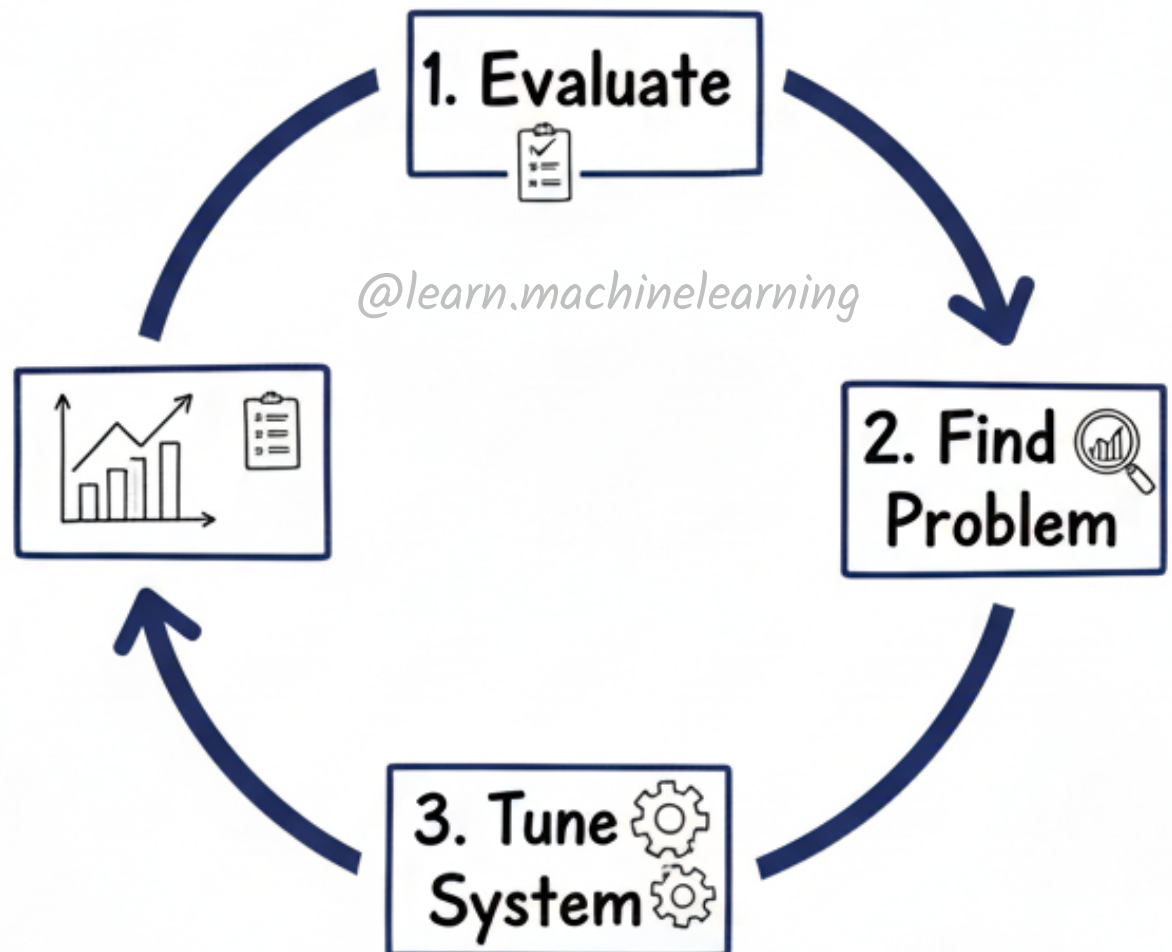
@learn.machinelearning



Build → Evaluate → Tune → Repeat

Evaluation isn't a one-time step;
it's a cycle.

Run Evaluation:
Identify Weakness:
Tune Your System:
Re-Evaluate



Day 23 Takeaway: What You Can't Measure, You Can't Improve!

Evaluating your RAG system with clear metrics (like Precision, Recall, Faithfulness, and Relevancy) is what transforms a cool experiment into a trustworthy, production-grade application. It is the most critical step for building reliable AI.

@learn.machinelearning



Next Up: Day 24: The Full RAG Blueprint Anatomy: Anatomy of a Production-Ready System

From your experience, what's the biggest sign that an AI-powered is untrustworthy? ↓
Follow @learn.machinelearning for the final day of our RAG Deep Dive! Let's discuss!
SAVE this post! These metrics are the key to building quality AI.  