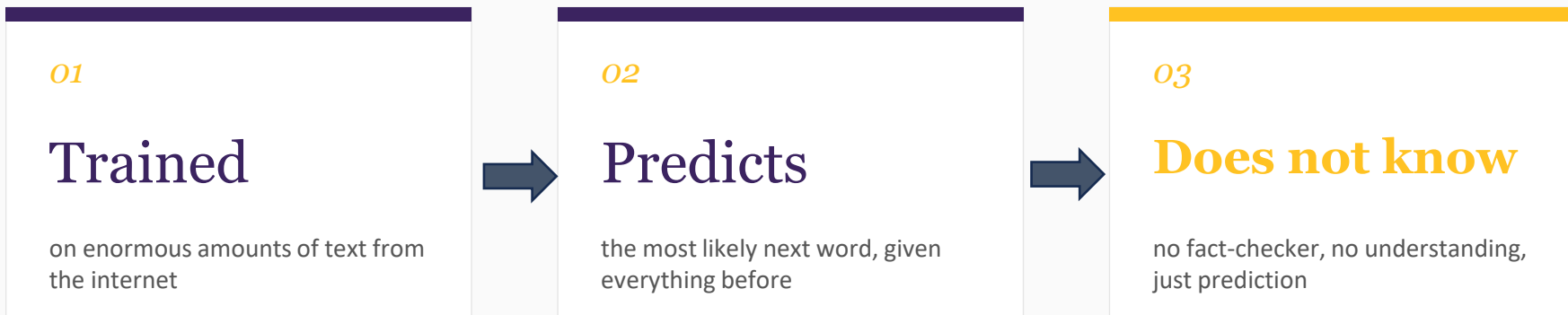


How LLMs work for educators.

*Plain language. Functional understanding
is enough for every pedagogical decision in this course.*

A probabilistic next-token predictor.

That is the whole mechanism.



Everything else in this lecture follows from that sentence.

It predicts probability, not truth.

CORRECT IN OUR WORLD

The capital of France is _____.

MODEL PROBABILITY

Paris 92%

Lyon 8%

Top prediction: Paris (92%)

high probability + factually correct

The model has seen this pattern many times in training data that matches the world.

FICTION, FLUENTLY DELIVERED

The capital of Wakanda is _____.

MODEL PROBABILITY

Birnin Zana 91%

Jabari 9%

Top prediction: Birnin Zana (91%)

high probability + factually false

The model has seen the fictional name in training, paired with its fictional capital.

The same scoring mechanism can make fluent fiction look almost as probable as fact.

Three practical consequences.



Hallucination

Plausible content that is not true.

- More likely on niche or recent topics
- Looks identical to true content on the surface



Bias

Training data carries the biases of its sources.

- Defaults skew toward dominant cultural assumptions
- Asking the model to be "unbiased" does not solve this



Overconfidence

Fluent prose reads as authoritative.

- Readers weight confidence as competence
- This is where student trust calibration breaks

What LLMs are actually good at.

STRONG

The course will use these

- + Summarizing long text
- + Translating between languages
- + Drafting initial structure
- + Asking questions on a topic
- + Generating examples and analogies
- + Rephrasing for accessibility

WEAK OR RISKY

Design around these

- Arithmetic at scale, exact recall
- Citation accuracy
- Recent or niche facts
- Stable identity or sustained reasoning
- Anything requiring values judgment

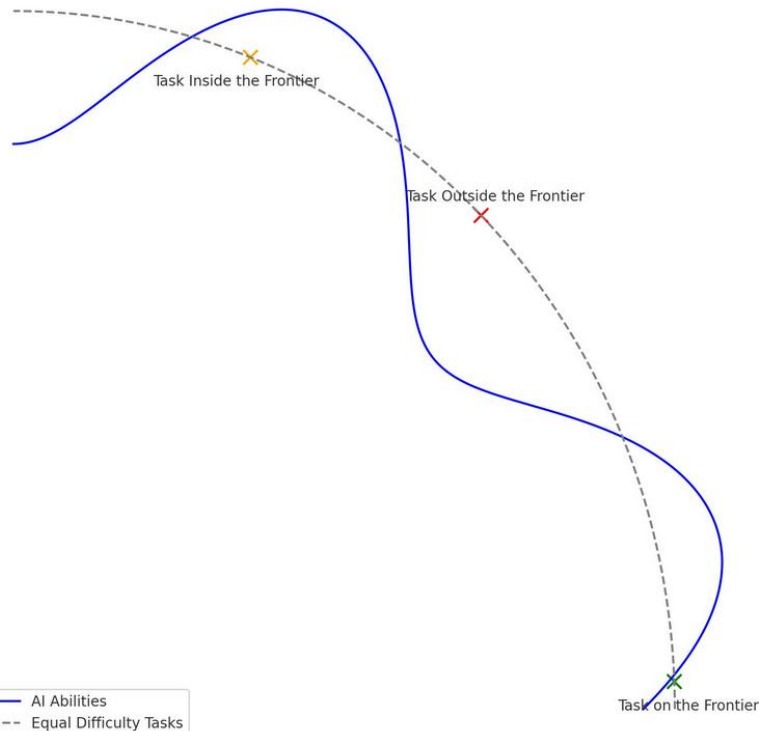
The jagged frontier.

There is no clean line between "AI can do this" and "AI cannot do this." Expect inconsistency.

INSIDE THE FRONTIER

Example: Draft a polished email from meeting notes.

Easy to check.



OUTSIDE THE FRONTIER

Example: Give legal advice from memory.

Requires verification.

What this course will not teach.

~~*How to use [any specific tool]*~~

Interfaces age out within months.

~~*How to write "the perfect prompt"*~~

There isn't one. We will review a framework, but think of it as a mental model on how to best interact with an AI.

~~*Whether you should adopt AI*~~

That is your decision, anchored to your Skeptic's Charter.

~~*How to detect AI use*~~

Detection is unreliable and inequitable. We will not rely on it.

Three takeaways.

1. An LLM predicts the next word from patterns in training data. It does not know.
2. Hallucination, bias, and overconfidence are not bugs to be patched. They are properties of how the system works.
3. The main question is what you want students to be able to do.



ELABORATION

In one sentence, explain why "fluency is not understanding" matters for grading.