

START HERE

AI, Safety and What Goes Wrong

The failure modes of AI, stated plainly, with the numbers.

First edition, September 2026

Addaly

Series editor: Dr Mustafa Mun

Draft, open for academic review

START HERE

AI, Safety and What Goes Wrong

The failure modes of AI, stated plainly, with the numbers.

Series editor: Dr Mustafa Mun

Addaly

First edition, September 2026 · Draft, open for academic review

AI, Safety and What Goes Wrong

© 2026 Addaly.

Licensed under Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International (CC BY-NC-SA 4.0). You may copy, print, share and adapt it for any non-commercial purpose, with credit, under the same licence.
creativecommons.org/licenses/by-nc-sa/4.0

First edition, September 2026, build 62a26075.

Written by AI models to a written standard, with Dr Mustafa Mun as series editor. Exam keys checked blind by a second model: 3,665 of 3,666 agreed. Academic review invited.

The manuscript was drafted with the assistance of a large language model and was edited and published under his name. That is stated plainly here rather than left to be discovered, because a reader is entitled to know how a book they are trusting was made.

How to cite: *Mun, M. (2026). AI, Safety and What Goes Wrong. Addaly. addaly.com/learn/ai-and-you.*

This book is the printed form of the course of the same name at addaly.com/learn/ai-and-you. The website is the living version and is corrected first. Where the two differ, the website is right.

Free to read, free to download, free to print, and free to teach from. There is no paid edition and there never will be.

Every diagram in it is drawn from data by code, not generated as a picture, so the numbers on the page are the numbers in the argument. The answers to every check, test and examination question are at the back.

9 modules · 73 lessons · 37 figures · 9 case studies · 51,639 words · about 10 hours

Brief contents

	Contents	6
1	How these systems fail	11
2	Bias, and how to measure it	49
3	Truth, sources and verification	89
4	Your data, and who ends up with it	129
5	Decisions made about you	167
6	Manipulation, attack and abuse	207
7	People, work and dependence	247
8	Ownership, law and the public record	287
9	Living with it: cost, governance and your own practice	327
	Examination	367
	Answers	397

1

MODULE 1

How these systems fail

Sort a failure into design, capability or misuse, read a benchmark or a demo for what it leaves out, and set how much checking a use needs by what it costs if wrong.

How these systems fail, before any particular harm: the three kinds of failure, what a model computes, where its data came from, what a benchmark score or a demo leaves out, why failures make no sound, why people believe the machine, who answers for it, and how much checking a use needs.

1	Design, capability or misuse	12
2	What the model computes	14
3	Where the data came from	18
4	What a benchmark score leaves out	21
5	Failure without an error message	24
6	How a demo is built	27
7	Why people believe the machine	30
8	Who is responsible when it is wrong	33
9	Triage: sorting risk by what it costs	36
	<i>Case study: Six cases that did not exist</i>	41
	<i>Module 1 test</i>	46

Lesson 1 · 3 min

Design, capability or misuse

THE ONE THING TO KEEP

Every failure is design, capability or misuse, and each needs a different fix.

In March 2023 two New York lawyers filed a brief in *Mata v. Avianca*, a personal injury claim against an airline in the federal court in Manhattan. Six of the decisions it cited did not exist. ChatGPT had written them, with judges' names, docket numbers and quotations. When one lawyer asked it whether a case was real, it said yes.

That is one failure. The news carries dozens a week, and they are not all the same kind. Before reading on, sort six real headlines into three bins.

Three bins, three fixes

Design. The system works as built, and the build is the problem. Amazon's experimental hiring tool learned from ten years of CVs, most of them from men, and marked down the word "women's". Nothing was broken. It learned what it was shown. An illustrative case makes the same point and runs through this course's film: a bank's loan model turns down a whole district, because few loans were ever made there, so few repayments were ever recorded. The fix is a different target, different data or a different threshold.

Capability. The system fails at the thing it claims to do. ChatGPT is not a legal database. It produces plausible text, and a plausible citation is what Steven Schwartz asked it for. The fix sits outside the tool: read every citation in a primary source before you sign.

Figure 1.1

Asked if its case was real, and what the court found

What ChatGPT told Steven Schwartz

- Varghese "does indeed exist"
- It is on Westlaw and LexisNexis
- The other cases are "real" authorities too

What the court found, 22 June 2023

- The Eleventh Circuit's clerk: no such ruling
- Its legal analysis is "gibberish"
- Six cited decisions did not exist
- A \$5,000 penalty on the lawyers and firm

From Judge Castel's opinion and order on sanctions. The reply about the citation came from the process that wrote the citation, so it checked nothing.

Misuse. The system works, and somebody points it at you. In 2024 a finance worker in Hong Kong paid out about \$25 million after a video call in which every colleague was a deepfake. The technology did its job. The fix is procedure: hang up and call back on a number you already hold.

Find the person who chose

When a system turns you down, no machine decided. People chose the model, the data, the target, the score at which a yes becomes a no, and whether anyone looks again. "The algorithm decided" moves those choices from people to software. It usually appears in a press statement.

In *Mata* the people are easy to name. One lawyer used a tool he had never used before. Another signed the brief without reading the cases. When the airline said the cases could not be found, both stood by them. On 22 June 2023 Judge P. Kevin Castel fined them and their firm \$5,000. He wrote that using a reliable AI tool is not improper in itself, and that the rules make lawyers the gatekeepers of what they file.

So ask two questions of any failure. Which bin is it in? Which human decision put it there? The whole Mata story, with the choice the firm faced in April 2023, is this module's case study.

BEFORE YOU MOVE ON

A delivery firm's routing model sends vans past a primary school at pick-up time, because it was told to minimise minutes and nothing else. Which kind of failure is this?

- A. Capability: the model failed to find the best route for vans
- B. Misuse: somebody aimed the routing tool at the schoolchildren
- C. Design: it met its target, and the target left out safety
- D. Capability: the map data must have been old or incomplete

The answer and why: p. 399

Lesson 2 · 4 min

What the model computes

THE ONE THING TO KEEP

A model writes likely next text, and nothing in that loop checks it is true.

The answer in Google's own advert

On 6 February 2023 Google launched its chatbot, Bard, with a short video. Someone asked what new discoveries from the James Webb Space Telescope they could tell a nine-year-old about. One of Bard's answers said the telescope took the very first pictures of a planet outside our solar system.

It did not. The first such picture was taken in 2004 by the European Southern Observatory's Very Large Telescope in Chile. Reuters reported the error, and on 8 February Alphabet lost about \$100 billion in market value.

A product from one of the best-funded research labs in the world said something false, fluently, in its own advert. Nobody outside Google knows exactly what produced that sentence, and Google's launch post said Bard drew on information from the web. The mechanism underneath still explains why a slip like this is ordinary.

Figure 1.2

Bard's answer, and the record

Google's Bard advert, February 2023

- Asked: new Webb telescope discoveries to tell a 9-year-old
- Answered: Webb took the very first pictures of a planet outside our solar system

The record

- First image of such a planet: 2004, ESO's Very Large Telescope
- September 2022: NASA announces Webb's own first direct image of one
- Alphabet's market value fell about \$100 billion on 8 February 2023

The true claim NASA made in September 2022 was about Webb's own first image. Bard's version drops one qualifier and becomes false.

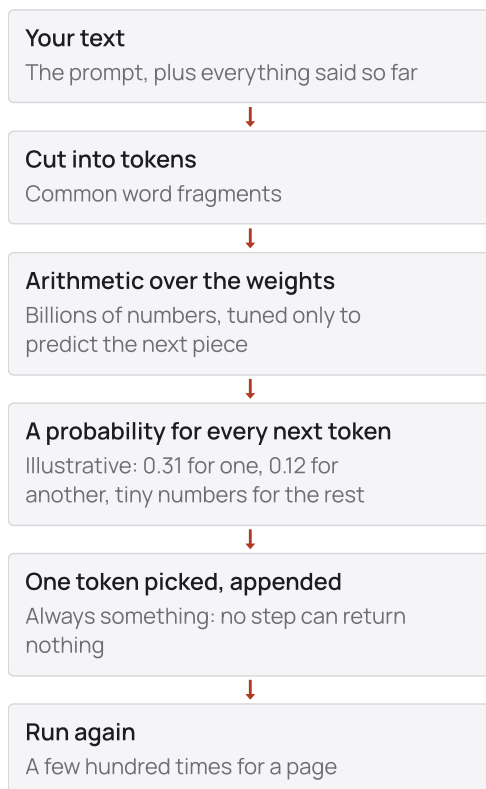
A pile of numbers that predicts the next piece

A language model is a very large set of numbers, called weights. Training adjusts them so that, given some text, the arithmetic predicts what comes next. That is the whole objective. Truth is not in it, and neither is caution.

Your text is cut into tokens, pieces of words. The model gives a probability to every possible next token, picks one, adds it, and runs again. A page of output is that loop a few hundred times.

Figure 1.3

The loop that produces a page of output



No step in this loop looks anything up. Facts, grammar and nonsense are held in the same weights, so a fabricated citation and a real one come from the same arithmetic.

Three things follow from that loop.

Nothing is filed as a fact. Facts, grammar and style all live in the same weights, in the same form. No step in the loop looks a fact up or checks that a sentence is true. A search step added outside the model can hand the loop a true page, and the loop still writes the sentence.

Nothing in the loop forces "I don't know". A model can learn to say it. In 2022 Saurav Kadavath and colleagues at Anthropic found that models carry usable signals of their own uncertainty. Saying so is still a trained habit,

not a built-in stop, and later training can wear it down. OpenAI's GPT-4 report found the model's confidence matched its accuracy well after pre-training, and less well after the training that followed.

Fluency is cheap and specifics are expensive. Grammar appears in every sentence the model trained on. A particular fact about a particular telescope appears rarely. The model is excellent at the first and much less reliable at the second, and the output looks the same either way.

Three stages, and the one that adds confidence

Pre-training is the costly stage: a vast amount of text and months of computing. It produces a model that continues text. Ask it a question and it may reply with five more questions, because lists of questions look like that.

Fine-tuning teaches the shape of an assistant from a much smaller set of example conversations.

Preference training pushes the model towards answers that human raters scored well. OpenAI's 2022 paper on InstructGPT describes the method. In 2023 Mrinank Sharma and colleagues at Anthropic found that human raters, and the preference models trained on their choices, sometimes preferred a convincing, agreeable answer to a correct one. In the same study, five assistants trained this way often gave up correct answers when a user pushed back.

Around the model sits a layer you never see: a hidden instruction, filters, sometimes a search step. When behaviour changes overnight and the model has not, that layer is usually why.

One question for any surprise

When an output surprises you, ask whether a system that predicts likely text, with no idea of truth, would produce it. A confident wrong fact: yes. Agreeing when you push back: yes. A famous fact right and a rare one wrong: yes. You seldom need a bigger explanation than that.

BEFORE YOU MOVE ON

A model answers a question about a small town's parking rules in fluent detail, and every specific is wrong. What does the mechanism suggest happened?

- A. Few examples existed, so it wrote a likely answer of that shape
- B. It retrieved the wrong entry from its internal database of rules
- C. It was unsure and should have errored, but error handling failed
- D. A safety filter removed the true answer and put in a vague one

The answer and why: p. 399

Lesson 3 · 3 min

Where the data came from

THE ONE THING TO KEEP

Training data is a stack of human choices, and each one shows in the answers.

Forty-five terabytes in, 570 gigabytes kept

In 2020 OpenAI published the recipe for GPT-3, the model family ChatGPT grew from. Its largest ingredient was Common Crawl, a free archive of the web. The team took 41 monthly snapshots from 2016 to 2019: 45 terabytes of compressed text. After filtering, they kept 570 gigabytes. Roughly one part in 80 survived.

"Trained on the internet" sounds like a fact of nature, like rain. It describes a chain of decisions, and the first one was which part in 80 to keep.

CASE STUDY · MODULE 1

Six cases that did not exist

A documented case. The sources follow the account.

Two lawyers at a New York firm answering an airline's motion to dismiss, in federal court in Manhattan, in the spring of 2023.

In February 2022 Roberto Mata sued the airline Avianca. He said a metal serving cart had struck his left knee on a flight from El Salvador to John F. Kennedy Airport in New York. He filed in a New York state court. Avianca moved the case to the federal court for the Southern District of New York, because claims about international flights fall under a treaty, the Montreal Convention.

Two lawyers at the firm Levidow, Levidow & Oberman worked on it. Steven Schwartz did the research and the writing. His practice had always been in the New York state courts, and he was not admitted to practise in this federal court. So his colleague Peter LoDuca put his own name on the filings.

In January 2023 Avianca asked the judge, P. Kevin Castel, to dismiss the claim. The argument was simple. The Montreal Convention gives a passenger two years to bring a claim, and Mata had waited longer. To keep the case alive, Mata's side needed court decisions showing that the clock had stopped, perhaps because Avianca had been through bankruptcy.

The firm's research service had limited access to federal cases. Schwartz had heard about ChatGPT from press reports and from his family. He had never used it. He took it to be a kind of search engine. He asked it to argue that the airline's bankruptcy paused the time limit. Then he asked it for case law, for specific holdings and for more cases. Each time it gave him what he asked for.

On 1 March 2023 the firm filed its answer to the motion. It cited and quoted decisions with names such as *Varghese v. China Southern Airlines*, each with a volume and page number in the Federal Reporter. LoDuca signed it without reading any of the cases. He relied on a colleague of more than twenty-five years.

On 15 March Avianca's lawyers replied. They said they could not find most of the cases, and the few they could find did not say what the filing claimed. The judge's own search failed too. In April the court ordered LoDuca to file copies of the decisions, and warned that failing to do so would end the case.

The decision

Put yourself at the firm in the middle of April 2023. The other side says the cases cannot be found. The court cannot find them either, and it wants copies. The only place the cases have ever appeared is a chat window.

There are two paths. The first is to tell the court now. The citations came from a chatbot, nobody at the firm can verify them, and the filing should be withdrawn. That is embarrassing. It weakens a client's case that was already in trouble, and it invites a sanction.

The second is to go back to the tool, ask again, and file what it gives. The tool is fluent and specific. It supplies judges' names and docket numbers. Asked directly whether *Varghese* is a real case, it says yes.

Before you open what happened, choose a path. Then name the human decision, at each step, that brought the firm to this point.

YOUR ANSWERS

1. Sort what went wrong in this case into the three kinds of failure: design, capability and misuse. Which kind was the fabricated citations, and what fix does that kind need?

2. Schwartz asked the chatbot whether Varghese was a real case, and it said yes. Why was that check worth nothing?

3. List the human decisions in this story, from the first prompt to the April affidavit. Which one did the court treat as the most serious, and why that one?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 44.

CASE STUDY · MODULE 1

How it played out

The firm took the second path. On 25 April 2023 LoDuca filed an affidavit with what were said to be copies or excerpts of the decisions. Schwartz had put it together from ChatGPT's answers, and LoDuca signed it without asking him a question. The court found the "decisions" were fabricated. The clerk of the Eleventh Circuit confirmed that *Varghese* was not a ruling of that court, and the judge called its legal analysis gibberish. Six of the cited decisions did not exist at all.

On 4 May the court ordered LoDuca to explain why he should not be sanctioned. The first admission that the cases came from ChatGPT arrived on 25 May, in an affidavit from Schwartz. Schwartz's filings included screenshots in which he had asked the chatbot whether *Varghese* was real, and it had said it was.

On 22 June 2023 Judge Castel sanctioned both lawyers and the firm. They were to pay a \$5,000 penalty, send the ruling to their client, and write to each real judge whose name had been attached to a fake opinion. The opinion said there is nothing inherently improper about using a reliable AI tool for assistance, but that the rules make lawyers the gatekeepers of what they file. The court also said the record would have looked quite different if the lawyers had come clean after the airline's reply or after the April orders. The bad faith it found lay in what came after the first filing: standing by the fake cases once they were questioned, and misleading statements to the court. Mata's claim itself was dismissed as out of time under the Montreal Convention.

The questions, discussed

1. Sort what went wrong in this case into the three kinds of failure: design, capability and misuse. Which kind was the fabricated citations, and what fix does that kind need?

The fabricated citations are a capability failure. The tool failed at the thing Schwartz believed it did, which was to find the law. A language model produces plausible text, and a plausible citation is exactly what he asked for. Nobody pointed the system at a victim, so this is not misuse in the sense the

lesson means. A capability failure cannot be fixed by the person using the tool, only worked around, so the fix is verification outside the tool: every citation read in a primary source before it is filed, and a rule that whoever signs a document has read what it cites.

2. Schwartz asked the chatbot whether Varghese was a real case, and it said yes. Why was that check worth nothing?

The same process that produced the citation produced the answer about the citation. A model asked whether its own output is true generates the most plausible reply, and after it has just supplied a case, the most plausible reply is that the case exists. A check only counts if it comes from somewhere independent of the thing being checked: the Federal Reporter, the court's own records, a legal database. Schwartz had that check within reach. He told the court he had looked Varghese up on a free site and could not find it.

3. List the human decisions in this story, from the first prompt to the April affidavit. Which one did the court treat as the most serious, and why that one?

Using a tool he had never used, for research in an area of law he did not know. LoDuca signing the March filing without reading the cases. Not acting on the airline's reply that the cases could not be found. LoDuca telling the court he was on holiday to win more time, which he admitted at the hearing was false. Filing the April affidavit instead of disclosing where the cases came from. The court weighed most heavily what happened after the warnings: the lawyers had been told the cases might not exist, and they doubled down rather than withdrawing the filing. The first mistake was a failure of care. Standing by it once questioned was a choice, and the rules on sanctions are written about choices.

Sources

1. Mata v. Avianca, Inc., No. 22-cv-1461 (PKC), Opinion and Order on Sanctions (S.D.N.Y. 22 June 2023)
2. Mata v. Avianca, Inc., No. 1:22-cv-01461 (S.D.N.Y.), docket and filings, CourtListener

MODULE 1 TEST

Check what stuck

6 questions, one right answer each. This one is for you, not for a record. The answers, and the reason behind each, are in the key on p. 398. If two go wrong, re-read the module before the exam.

- 1 A model answers a question about a small town's municipal rules fluently and wrongly, with no hedging anywhere in the answer. Which feature of the mechanism accounts for the missing warning?
 - A. It retrieved an entry from its internal store of facts, and that entry happened to be out of date and incomplete
 - B. A safety filter removed the uncertainty warning before the answer was displayed
 - C. Every prompt produces a probability distribution, and none of them means the model has nothing
 - D. Municipal rules fall after the training cutoff, and post-cutoff material is flagged as unreliable

- 2 A widely copied training pipeline kept crawled pages that a classifier judged to resemble the pages a particular English-speaking community links to. What does that choice do to everything downstream?
 - A. It defines good writing as resembling that community's references, and the model inherits the definition
 - B. It removes bias, because a quality classifier scores a page on how well it is written rather than on who wrote it or where they live
 - C. It affects only the model's grammar, since the classifier judges style rather than subject matter
 - D. It has no lasting effect once books and licensed archives are mixed into the corpus

- 3 A model was reported near the top of a professional exam's range, and a later re-analysis put the figure far lower. Nothing was fabricated in either report. What moved the number?
- A. The model was retrained between the two analyses and lost capability on the exam material
 - B. The re-analysis used a harder version of the exam released after the model's training cutoff, so contamination no longer helped it
 - C. The original score came from several attempts scored as correct if any attempt succeeded
 - D. The comparison group: the percentile was computed against a pool that included repeat takers
- 4 Six months after launch, a drafting tool's reported accuracy is unchanged, and support staff say they spend longer fixing each draft than they used to. Which measurement would show what is happening?
- A. The system's accuracy on its original held-out test set, re-run every month to confirm that the score has not moved
 - B. The edit distance between what the system proposed and what the human actually sent
 - C. The number of drafts generated per agent per day
 - D. The proportion of customers who reply to the message that was sent
- 5 A hospital reports that clinicians override its triage model in 0% of cases and offers the figure as evidence that the model is accurate. What does the figure more likely mean?
- A. The model has been calibrated to the clinicians' own historical decisions and now matches them exactly on every case
 - B. The clinicians lack the training needed to evaluate the model's recommendations
 - C. The review is not happening, because a reviewer who never disagrees is not forming an independent view
 - D. The model is operating inside its training distribution on every case seen so far

- 6 An airline's website chatbot described a refund policy the airline did not have, and the airline argued the chatbot was a separate entity responsible for its own statements. Why did that fail?
- A. The chatbot is part of the company's website, and a company is responsible for the information on it
 - B. The model provider had accepted liability for outputs in its terms of service, so the claim moved up the chain to the provider instead
 - C. The airline had not obtained the customer's consent before deploying an automated system
 - D. The chatbot's answer fell outside the disclaimer the airline published about automated tools

A worked example, with the arithmetic shown

Two groups, a thousand applicants each. The model is applied identically to both.

Group A. Approved and repaid: 630. Approved and defaulted: 70. Rejected who would have repaid: 90. Rejected who would have defaulted: 210.

Group B. Approved and repaid: 350. Approved and defaulted: 50. Rejected who would have repaid: 150. Rejected who would have defaulted: 450.

Figure 2.1

Group B: 1,000 applicants, one loan model

Approved	Rejected
Would have repaid	
350 true positives	150 false negatives
Would have defaulted	
50 false positives	450 true negatives

True positive rate $350 \div 500 = 70\%$, against 87.5% for Group A. Precision $350 \div 400 = 87.5\%$, against 90%. The same four boxes support 'an approval means the same thing for everyone' and 'a qualified Group B applicant is nearly three times as likely to be refused'.

Now four rates, each answering a different question. Do the division yourself as you read; the numbers are chosen to divide cleanly.

Approval rate — of everyone who applied, who got a loan. Group A: 700 of 1,000, so 70%. Group B: 400 of 1,000, so 40%.

True positive rate, also called recall or sensitivity — of the people who *would have repaid*, who got approved. Group A: 630 of 720, so 87.5%. Group B: 350 of 500, so 70%. This is the rate that matters to a creditworthy

applicant. Being in Group B and deserving a loan means a 30% chance of refusal, against 12.5% in Group A.

False positive rate — of the people who *would have defaulted*, who got approved anyway. Group A: 70 of 280, so 25%. Group B: 50 of 500, so 10%.

Precision, or positive predictive value — of the people the model approved, who actually repaid. Group A: 630 of 700, so 90%. Group B: 350 of 400, so 87.5%.

Stop and look at what just happened. On precision the two groups are almost identical: 90% against 87.5%. A lender could truthfully say "an approval means the same thing regardless of group". On the true positive rate they are not remotely identical: a qualified Group B applicant is nearly three times as likely to be turned away. Both statements describe the same four numbers.

This is not a hypothetical construction. It is the exact structure of the argument about the COMPAS recidivism tool, which Why you cannot have all three takes apart.

Which rate belongs to whom

The reason people talk past each other is that each rate answers the question of a different party.

The **institution** cares about precision and about the false positive rate, because those are its losses. The **individual who was rejected** cares about the false negative rate, because that is their harm. The **regulator** often looks at the approval rate, because that is what shows up as an outcome gap in the population. A **defendant** in a risk-scoring system cares about being wrongly labelled high risk, which is a false positive in that framing.

None of these people is confused. They are computing different fractions of the same table, and the fractions genuinely disagree.

EXAMINATION

AI, Safety and What Goes Wrong: final examination

This paper covers the whole course:

- how these systems fail
- bias and how to measure it
- truth and verification
- your data and who ends up with it
- decisions made about you
- manipulation and security
- people and dependence
- ownership and law
- living with it

64 questions, the whole pool. The site draws 25 for a sitting, pass mark 70%.
Work any 25 under your own conditions. Answers from page 449.

1 Module 1 · How these systems fail

A voice clone of somebody's daughter telephones asking for money, and the audio is convincing. Which of the three kinds of failure is this, and what does that imply about the fix?

- A. A capability failure: the system did not do the thing it claims to do, so it needs better training data
- B. Misuse: the technology worked, so the fix is procedural rather than technical
- C. A design failure: the system worked as designed and the design was the problem, so it needs a governance decision
- D. None of the three, because no model was involved in placing the telephone call itself

2 Module 1 · How these systems fail

A model becomes noticeably more confident and more agreeable after the stage of training that turns a base model into an assistant. Where did that come from?

- A. Fine-tuning added a rule instructing the model to sound authoritative when answering factual questions
- B. Pre-training on a larger corpus, which raises the probability of the top token in every distribution
- C. A product-layer system prompt that instructs the model never to express doubt to a user
- D. Alignment training pushed it towards responses human raters preferred, and raters prefer confidence

3 Module 1 · How these systems fail

A model scores markedly better on a benchmark published before its training cutoff than on an equivalent one published after. What should you reach for first?

- A. Contamination: the earlier benchmark's questions and answers were probably in the training data
- B. Saturation: the earlier benchmark has stopped discriminating between models
- C. Drift: the world changed between the two benchmarks being written
- D. A difference in scaffolding, since the later benchmark was probably evaluated with fewer attempts allowed

4 Module 1 · How these systems fail

A churn model built before a competitor entered the market now performs badly, although the input data looks much as it always did. What has changed?

- A. Data drift: the distribution of the inputs has moved away from the training distribution
- B. The model is out of distribution on every case, so it has stopped producing outputs at all
- C. Concept drift: the same input should now produce a different answer
- D. The training labels were wrong from the beginning, and the launch evaluation failed to detect the error

5 Module 1 · How these systems fail

You are shown a polished demonstration of a document-reading system. Which question carries the most information about whether it will work for you?

- A. Which model does it use underneath, and how many parameters does that model have
- B. What is the success rate on a random sample of real inputs, with the sample size
- C. How long did the team spend building it, and how many customers are live on it now
- D. Can it process the specific document shown in the demonstration faster than the version we saw

6 Module 1 · How these systems fail

Two tasks are equally reversible and equally repeated. One affects only you and one affects a tenant who cannot see the tool and has no route to complain. What does the triage say?

- A. They sit in the same tier, because reversibility is the axis that decides the tier
- B. The second sits lower, because a tenant can raise the matter through an ordinary complaints process
- C. Neither can be tiered without knowing which model was used and how accurate it has been measured to be
- D. The second sits a tier higher, because the person bearing the cost did not choose the risk

Module 1 • How these systems fail

The module starts on p. 11.

MODULE 1 TEST

Module 1 test, Q1, p. 46

A model answers a question about a small town's municipal rules fluently and wrongly, with no hedging anywhere in the answer.

Which feature of the...

Answer C: Every prompt produces a probability distribution, and none of them means the model has nothing

Why: There is no separate store of facts and no state that means empty. A question about a rule that exists and one about a rule that never existed both produce perfectly ordinary distributions, and both come out as fluent text. Fluency is cheap because grammar appears in every sentence of the training data; a specific municipal rule may appear four times in a trillion words.

Module 1 test, Q2, p. 46

A widely copied training pipeline kept crawled pages that a classifier judged to resemble the pages a particular English-speaking community links to. What does that...

Answer A: It defines good writing as resembling that community's references, and the model inherits the definition

Why: The filter is where the decisions live. Operationally defining quality as resemblance to one community's reading list makes that community's norms a property of the model's output, and adding books later does not undo the exclusion already applied to the crawl.

Module 1 test, Q3, p. 47

A model was reported near the top of a professional exam's range, and a later re-analysis put the figure far lower. Nothing was fabricated in...

Answer D: The comparison group: the percentile was computed against a pool that included repeat takers

Why: A benchmark result measures a specific task, under specific conditions, against a specific comparison, and the summary sentence drops all three. Here the comparison group was chosen, and choosing it moved the headline by tens of percentiles without anyone stating a falsehood.

Module 1 test, Q4, p. 47

Six months after launch, a drafting tool's reported accuracy is unchanged, and support staff say they spend longer fixing each draft than they used to...

Answer B: The edit distance between what the system proposed and what the human actually sent

Why: This is degradation the humans are absorbing. Reported accuracy stays flat because staff are patching the difference, and nobody has measured how much patching is happening. Edit distance and the override rate are the two numbers that make silent degradation visible; re-running the old test set cannot, because the test set has not drifted.

Module 1 test, Q5, p. 47

A hospital reports that clinicians override its triage model in 0% of cases and offers the figure as evidence that the model is accurate. What...

Answer C: The review is not happening, because a reviewer who never disagrees is not forming an independent view

Why: An override rate of zero is the signature of a rubber stamp. Oversight requires time, independent information, real authority to disagree and somebody counting how often it happens; where overriding is administratively expensive and agreeing is free, the recommendation becomes the default.

Module 1 test, Q6, p. 48

An airline's website chatbot described a refund policy the airline did not have, and the airline argued the chatbot was a separate entity responsible for...

Answer A: The chatbot is part of the company's website, and a company is responsible for the information on it

Why: The tribunal held that it makes no difference whether information comes from a static page or a chatbot; the company owns what its interface says. Providers in fact disclaim output heavily, so the chain does not move upward, and the responsibility lands where it would have landed if a member of staff had said it.

THE LESSON CHECKS**Lesson 1, p. 14**

A delivery firm's routing model sends vans past a primary school at pick-up time, because it was told to minimise minutes and nothing else. Which...

Answer C: Design: it met its target, and the target left out safety

Why: Every failure is design, capability or misuse, and each needs a different fix.

A Capability: the model failed to find...: The route looks like a mistake, so it feels as if the tool fell short. It found exactly the fastest route it was asked for.

B Misuse: somebody aimed the routing tool...: The harm lands on children who never chose the tool, which feels like being targeted. Nobody pointed it at anyone.

D Capability: the map data must have...: Old maps do cause bad routes. Nothing here says the map was wrong, and the route really was the fastest.

Lesson 2, p. 18

A model answers a question about a small town's parking rules in fluent detail, and every specific is wrong. What does the mechanism suggest happened?

Answer A: Few examples existed, so it wrote a likely answer of that shape

Why: A model writes likely next text, and nothing in that loop checks it is true.

B It retrieved the wrong entry from...: Assumes a look-up step inside the model. The weights hold patterns, not a table of entries to retrieve from.

C It was unsure and should have...: Treats "I don't know" as a built-in error state. A model can learn to say it, but nothing in the loop raises an error.

D A safety filter removed the true...: A filter leaves a refusal or a hedge. It does not produce confident, specific, wrong details.

START HERE

AI, Safety and What Goes Wrong

The failure modes of AI, stated plainly, with the numbers.

WHAT IT TEACHES	A working guide to how AI systems fail and what to do about it. Where bias actually comes from and why deleting the sensitive column does not remove it.
WHAT IS INSIDE	Nine modules and 73 lessons, 37 figures drawn from data, nine case studies, nine module tests, a 64-question examination, and every answer with the reason behind it.
LICENCE	CC BY-NC-SA 4.0. Copy, print, share and adapt it for any non-commercial purpose, with credit, under the same licence. There is no paid edition.
THE COURSE	Free to read, with the lessons, the films and the examination, at addaly.com/learn/ai-and-you

Addaly

First edition, September 2026 · build 62a26075



Scan to open the course