

USE IT IN YOUR WORK

AI for Teachers

For teachers who now have to teach and use AI, often with no training and no projector.

First edition, September 2026

Addaly

Series editor: Dr Mustafa Mun

Draft, open for academic review

USE IT IN YOUR WORK

AI for Teachers

For teachers who now have to teach and use AI, often with no training and no projector.

Series editor: Dr Mustafa Mun

Addaly

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

AI for Teachers

© 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 for Teachers. Addaly. addaly.com/learn/ai-for-teachers.*

This book is the printed form of the course of the same name at addaly.com/learn/ai-for-teachers. 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.

7 modules · 57 lessons · 26 figures · 7 case studies · 44,923 words · about 8 hours

Brief contents

	Contents	6
1	Understanding the machine you are about to use	9
2	Prompting as a teacher	49
3	Planning and making materials	89
4	Assessment, feedback and marking	133
5	Integrity: AI-written work and what to do about it	173
6	Teaching AI itself	211
7	The multilingual, mixed-ability classroom	253
	Examination	293
	Answers	313

Contents

Module 1 · Understanding the machine you are about to use	9
1 Explaining AI to a class that argues back	10
2 The free tools you actually have	13
3 Where it is reliable, and where it is not	17
4 Why it cannot look it up	21
5 The conversation has a memory limit	26
6 It has read the American curriculum	30
7 A chatbot is not a calculator	33
8 Same question, different answer	37
<i>Case study: The tool everyone already had</i>	<i>41</i>
<i>Module 1 test</i>	<i>46</i>
 Module 2 · Prompting as a teacher	 49
9 The five things only you know	50
10 Show it one of yours	54
11 The second message matters more than the first	58
12 Set the voice and the reading level	62
13 Roles that help, and roles that do nothing	66
14 Prompting for formats you can print	70
15 When it refuses a lesson	74
16 Keep a prompt file	77
<i>Case study: The department that shared its prompts, and the mock</i>	
<i>paper inside them</i>	<i>81</i>
<i>Module 2 test</i>	<i>86</i>
 Module 3 · Planning and making materials	 89
17 Planning with AI, keeping the judgement	90
18 Worksheets, and three versions of them	93
19 Four explanations of one idea	97
20 Passages with the words you are teaching	100
21 A question bank you can trust	104
22 Diagrams and visuals without a designer	108
23 Slides, and the projector you sometimes have	113

24	Revision: flashcards and spaced practice	117
25	A unit plan, not just a lesson	121
	<i>Case study: Fifteen minutes, three versions, forty copies</i>	125
	<i>Module 3 test</i>	130
Module 4 • Assessment, feedback and marking		133
26	Marking: what a machine cannot fairly judge	134
27	Writing a rubric both can read	137
28	Feedback students actually use	142
29	Feedback before submission: the student-facing loop	146
30	Reading thirty answers at once	149
31	Marking handwritten work with a phone	154
32	Writing the exam paper	157
33	Peer assessment with a third marker	161
	<i>Case study: Sixty-two drafts and a bias check that came back badly</i>	165
	<i>Module 4 test</i>	170
Module 5 • Integrity: AI-written work and what to do about it		173
34	Why detectors fail, and what to do instead	174
35	Assignments AI cannot do for you	177
36	Teaching students to use it honestly	181
37	Writing the class policy	184
38	Process evidence that holds up	188
39	Handling a case fairly	192
40	Rethinking homework	196
41	Citing and declaring AI use	199
	<i>Case study: Ninety-two per cent, and a parent already writing</i>	203
	<i>Module 5 test</i>	208
Module 6 • Teaching AI itself		211
42	An AI lesson with no screen and no power	212
43	What to teach at which age	215
44	The sorting game for younger children	220
45	Image generators and deepfakes	225
46	The AI companion problem	228
47	A student project that teaches the mechanism	232

48	The evidence on AI tutors	236
49	AI in every subject, not just computing	241
	<i>Case study: Four hundred free licences and a chart</i>	245
	<i>Module 6 test</i>	250
Module 7 · The multilingual, mixed-ability classroom		253
50	Translation, and the back-translation check	254
51	Explaining in the language students actually speak	259
52	For students who struggle to read	262
53	Speech-to-text and text-to-speech, for free	267
54	Supporting students with disabilities	271
55	The quiet student and the private question	274
56	Stretching the student who is ahead	277
57	When the home language has no training data	280
	<i>Case study: The Santali worksheets that arrived from the district office</i>	284
	<i>Module 7 test</i>	290
Examination		293
Answers		313

1

MODULE 1

Understanding the machine you are about to use

Explain to a class, a colleague and a parent what a language model does and does not do, predict from the mechanism which kinds of question it will answer reliably and which it will get confidently wrong, and choose a free tool that runs on the phone you actually have with its training-on-your-data setting turned off.

1	Explaining AI to a class that argues back	10
2	The free tools you actually have	13
3	Where it is reliable, and where it is not	17
4	Why it cannot look it up	21
5	The conversation has a memory limit	26
6	It has read the American curriculum	30
7	A chatbot is not a calculator	33
8	Same question, different answer	37
	<i>Case study: The tool everyone already had</i>	41
	<i>Module 1 test</i>	46

Lesson 1 · 7 min

Explaining AI to a class that argues back

THE ONE THING TO KEEP

It predicts likely next words from human writing, so sounding right and being right come apart.

The one sentence to start with

"It is a machine that has read an enormous amount of writing by people, and it guesses what words come next."

That is the whole thing. Everything else is elaboration, and this sentence survives the questions children actually ask. "It's a robot brain" does not survive. "It's like a search engine" does not survive — the first follow-up is "then why is it wrong sometimes?" and you have no answer.

Make them do it before you explain it

Write on the board:

The old woman walked slowly into the _____

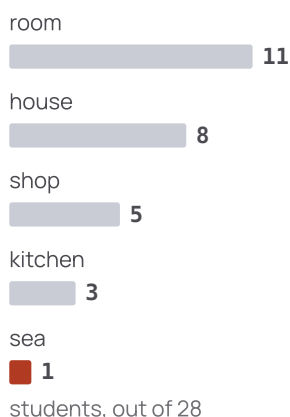
Every student writes one word. Silently, no discussion. Then take a tally by show of hands. Room. House. Shop. Kitchen. Somebody will write "sea" and the class will laugh.

Put the tally on the board. You have just built the thing you are teaching: twenty-eight people, four common answers, one rare one. The machine has that same shape of guess, except it built it from billions of sentences instead of one classroom.

Now the second stem:

The capital of Japan is _____

Figure 1.1

Twenty-eight guesses at one blank

The tally from one Class 7 room for "The old woman walked slowly into the _____". Four common answers, one rare one. The same class, given "The capital of Japan is _____", produced twenty-seven Tokyos and one Kyoto. The scatter is the whole explanation of why the machine is dependable about Tokyo and unreliable about your district.

Nearly unanimous. Tokyo.

Ask them what changed. They will get there: for one question almost everyone's guess is the same, because the writing they have seen nearly always says the same thing. For the other, many guesses are fine. That is the entire reason AI is dependable on some questions and unreliable on others, and a ten-year-old can now explain it to their parents.

The questions they will ask, and honest answers

"Does it know it's me?" No. Usually it starts each conversation with nothing about you. Some products keep notes about you; that is a feature someone built on top, not the machine remembering.

"Is it alive? Does it have feelings?" It writes sentences about feelings because people wrote sentences about feelings and it learned from them. Whether anything is going on inside is genuinely argued about by the people

who build these systems. You are allowed to say that. Try: "Nobody has shown it feels anything, and most researchers think it doesn't. But it is a real argument, and you are old enough to know it's an argument rather than a settled fact."

"Why does it lie?" Lying means knowing the truth and saying something else. It does not do that. It produces the sentence that fits, and a sentence can fit beautifully and be false. It has no way to check. Give them the phrase: it is not lying, it is **making things up without knowing it**.

"Where did it get all the words?" From writing by people — websites, books, articles, forums. That is also why it repeats what people say, including the unfair things people say.

"Can it do maths?" Sometimes, badly, unless a calculator is running underneath. Guessing the next word is a poor way to multiply.

Two analogies to avoid

"It's a robot brain." This invites every film they have seen into the room and you spend the period on robots instead of prediction.

"It's just autocomplete." Closer, and useful, but if you stop there a student will show you something it wrote that is genuinely good and conclude you were wrong. Say instead: "It works like autocomplete, and it turns out autocomplete over enough writing is far more powerful than anyone expected — including the people who built it."

The line worth ending on

"It is very good at sounding right. Sounding right and being right are different things, and telling them apart is now part of your job."

Eleven-year-olds understand that sentence immediately, and it carries the rest of the course.

BEFORE YOU MOVE ON

A student tells the AI it made a mistake. It apologises and changes its answer. She says this proves it understood its error. What should you tell her?

- A. An apology is simply the kind of sentence that follows a correction in the writing it learned from — it will produce one whether or not it was wrong. Tell it a correct answer is wrong and watch it fold.
- B. She is right — noticing and correcting an error requires holding some model of what is true.
- C. The company that built it wrote a rule that it must apologise politely when challenged.
- D. When challenged, it went and looked up the correct answer, which is why it changed.

The answer and why: p. 315

Lesson 2 · 9 min

The free tools you actually have

THE ONE THING TO KEEP

Every major model has a free tier that runs in a phone browser or inside WhatsApp; the choice that matters is not which is cleverest but which one you have turned off training on your data.

You do not need to buy anything

Every major model has a free tier, and every free tier runs in a phone browser. If you have a phone that can open a web page, you have access to the same underlying models a paid user has, with two differences: a cap on how many messages you can send before it slows you down or switches to a smaller model, and fewer extras like long file uploads.

For a teacher generating a worksheet, a set of misconceptions and a parent letter in an evening, the cap almost never bites. Where it does, the fix is to wait, or to have a second tool open.

The realistic list

Here is what is available at no cost as of this writing. Products change names and limits often, so treat the shape as stable and the details as something to re-check.

- **ChatGPT** (OpenAI). Free tier includes the main model with a message cap, image upload, voice, and web search. Works in a browser without the app. Requires an account, which can be a phone number or email.
- **Gemini** (Google). Free with any Google account. Tightly tied to Docs, Sheets and Gmail, which matters if your school runs on Google Workspace. Also inside Android as the assistant.
- **Claude** (Anthropic). Free tier with a daily cap. Strong on long documents and careful writing. Account required; terms are for adults.
- **Copilot** (Microsoft). Free, built into Edge and Windows, uses OpenAI models underneath. If your school runs on Microsoft, this is the one already installed.
- **Meta AI inside WhatsApp**. No new account, no new app, works on the data plan you already pay for. The model is Meta's Llama. It is weaker than the others on long, careful tasks and stronger than all of them on one thing: it is already on the phone of nearly every teacher in India, Nigeria, Kenya and Brazil.
- **DeepSeek** and other Chinese models. Free, capable, and hosted in China, which is a data-location question your school may have a view on.

The open-source route also exists: models like Llama and Qwen can be run on your own computer with free software such as Ollama or LM Studio. This needs a laptop with at least 8 GB of memory and gives you a weaker model in exchange for nothing leaving the machine. It is the right answer for a school that cannot send any data out, and the wrong answer for a teacher with a phone.

The one setting

Most consumer chatbots default to using your conversations to improve their models. That means the essay you pasted, the student name you forgot to remove, the parent complaint you asked for help with — all of it may be read by people at the company and folded into future training.

Before you paste anything about a student, find the setting. In ChatGPT it lives under Settings, then Data controls, labelled something like "Improve the model for everyone". In Gemini it is under Activity. In Claude it is a privacy setting. Turn it off. This does not make the tool private — the company still receives and stores what you type — but it takes your students' work out of the training pipeline.

Then treat what you paste the way you would treat what you say in a staffroom with the door open. First names and no surnames. No addresses. Nothing you would not want a stranger at the company to read.

Browser or app

Use the browser. The apps ask for permissions you do not need to grant, take storage you may not have, and on a shared family phone leave a logged-in account for whoever picks it up next. A browser tab you close is gone. If the phone is shared, use a private window.

The honest limitation

A free tier is a queue. When demand is high, the service quietly gives you a smaller, faster model, and answers get shallower without any notice. If a tool that was sharp at seven in the morning is producing thin work at nine in the evening, that is usually why. Nothing you did changed. Ask again tomorrow, or switch to a second tool for the rest of the session.

The other limit is the message cap itself. It resets on a schedule — typically every few hours — and it is per account, so a department sharing one login will hit it by lunchtime. Each teacher needs their own.

What to do tonight

Open one of these in your phone browser. Turn off training. Ask it for five wrong answers your students give about something you taught this week, and see how many you recognise. That is the whole setup, and it took ten minutes.

BEFORE YOU MOVE ON

A teacher on a shared Android phone with patchy data wants to generate practice questions each evening. Which single setup decision matters most before she starts?

- A. Picking whichever tool scored highest on the latest benchmark leaderboard, because a stronger model will make fewer mistakes on the worksheets she generates.
- B. Turning off the setting that lets the company train on her conversations, and working in the browser.
- C. Installing every app on the phone so she can compare the answers across all of them for each question she asks.
- D. Paying for one subscription first, since free tiers are demonstration versions that are not intended for real professional work.

The answer and why: p. 316

Lesson 3 · 9 min

Where it is reliable, and where it is not

THE ONE THING TO KEEP

Reliability tracks how many times a fact has been written down by people — so the model is sound on photosynthesis and unsound on the founding date of your district, and you can predict which before you ask.

One question to ask before you ask

The first lesson gave you the mechanism: the model guesses the next word from what people have written. That gives you a tool for predicting its reliability before you type a single prompt.

Ask yourself: **how many times has this exact thing been written down?**

The boiling point of water at sea level has been written down millions of times, in every language, in every textbook, and it is always the same. The model has seen it so often that its guess is effectively a fact. The population of Kisumu at the 2019 census has been written down a few hundred times, inconsistently, in documents that also contain other numbers. The model's guess is a plausible-looking number in the right range. Sometimes it is right. You cannot tell from the answer which time this is.

That is the whole rule. Everything below is that rule applied.

Reliable ground

These are the areas where the model's answers are nearly always sound, because the writing it learned from is vast and agrees with itself:

- **Core school science and maths.** Newton's laws, photosynthesis, the water cycle, quadratic equations, the periodic table. Explanations of these are excellent and the errors are rare.

- **Grammar and writing in major languages.** English, Spanish, French, Hindi, Mandarin, Arabic, Portuguese. It has read more of these than any human.
- **Well-known history.** Dates and causes of major events. The more written about, the safer.
- **Explaining a concept several ways.** This is generation, not recall, and it is what the mechanism is best at.

Unreliable ground

And here the answers look identical but should be treated as drafts to check:

- **Local specifics.** The founding date of your school, the name of your district's first commissioner, which crops grow in your division. Written down rarely; often wrong.
- **Precise numbers.** Populations, distances, percentages from a specific report. It produces a number in the right range because the right range is what appears in text. The exact figure is a guess.
- **Citations.** Book titles, page numbers, journal articles, quotations attributed to a named person. It has learned what citations look like and will generate one that looks perfect and does not exist. Every citation gets checked. No exceptions.
- **Recent events.** Anything after its training cutoff, which is typically several months to a year before the date you are using it. Unless it searched the web — the next lesson — it is guessing from older text.
- **Small languages and regional literature.** A poet who wrote in Bhojpuri or Tigrinya has been written about far less than Shakespeare, and the model will confidently give you a biography assembled from fragments.
- **Anything about your specific syllabus.** Which chapters are in the Class 9 board paper this year, what the WAEC marking scheme allocates for a question. It has seen some of this and much of it is out of date.

The trap in the middle

The dangerous cases are the ones that sit between the two lists: a well-known topic with a specific detail inside it.

Ask about the French Revolution and the overview will be sound. Ask what Robespierre said on a particular date and you have crossed into a rarely-written-down specific, and the model will supply a quotation that sounds exactly right. The topic was reliable; the detail is not. You have to notice the moment the question changes shape.

A useful habit: when reading an answer, underline every proper noun, number, date and quotation. Those are the specifics. Everything underlined is on the second list until you have checked it. Everything else is probably fine.

Why confidence tells you nothing

The tone does not change between the lists. The mechanism produces the most likely next word whether the underlying knowledge is a million textbooks or three blog posts, and the most likely next word after "The population of Kisumu is" is a number, stated plainly. There is no internal signal of uncertainty reaching the sentence. Some tools now add hedges like "approximately", but these are trained habits, not measurements.

So you cannot use the answer to judge the answer. You use the question. That is why the habit above works and why "it sounded sure" is never a reason.

Practice for your subject

Write down five questions from your subject: two you expect to be on reliable ground, two on unreliable ground, one in the middle. Ask them. Check the answers against a source you trust. Most teachers find their predictions are right four times in five — and the one miss teaches them where their subject's specifics hide.

Figure 1.2

The reliability of an answer is set by the question

You ask for the overview

You ask for a specific detail

Topic written about a great deal

Sound. Causes of the French Revolution.
reliable ground

Confident and unchecked. What Robespierre said on a named date.
the trap in the middle

Topic written about rarely

Thin but usually harmless. Farming in your division.
check it

A guess in the right range. Your district's 2019 census figure.
unreliable ground

The dangerous cell is the top right: a well-known topic with a rarely-written-down detail inside it. The tone of the answer is identical in all four. Underline every proper noun, number, date and quotation, and you have marked the right-hand column.

BEFORE YOU MOVE ON

A geography teacher asks a chatbot two questions: the boiling point of water at sea level, and the population of her town at the last census. She gets confident answers to both. Which is the honest assessment?

- A. Both are factual questions with definite answers, so both are equally likely to be right and equally worth trusting.
- B. Neither answer can be trusted, because the model does not know anything and every answer it gives is a guess of the same quality.
- C. The boiling point is almost certainly right; the population may be a plausible invention and needs checking against the census.
- D. The population is the more reliable of the two because it is a specific number, and the model handles numbers better than it handles science explanations.

The answer and why: p. 316

Lesson 4 · 8 min

Why it cannot look it up

THE ONE THING TO KEEP

A model on its own has no way to consult anything at answer time; when a tool adds web search it fetches pages and summarises them, which fixes staleness but not judgement about which page to trust.

Two very different things wearing the same interface

When you type a question into a chatbot, one of two things happens, and the box looks the same either way.

Mode one: the model answers from what it learned. Nothing is consulted. No page is opened. The answer is produced word by word from the patterns absorbed during training, which finished months ago. This is what a language model is. It cannot check, because there is nothing to check against — the training text is not stored as a library it can open, it is dissolved into billions of numerical weights.

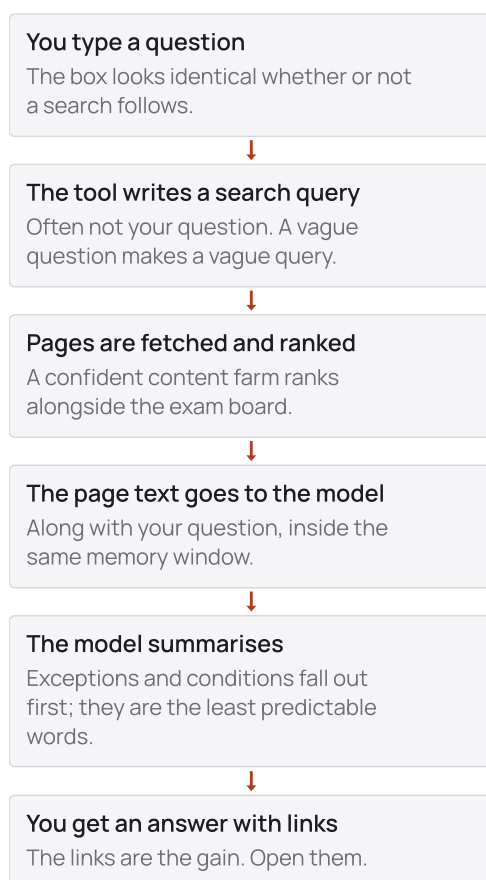
Mode two: the tool searches first. Many products now bolt a web search onto the model. Your question is turned into a search, some pages are fetched, and the model is handed their text along with your question. It then writes an answer that summarises what it was given. This is a genuine improvement for anything recent or specific, and it changes what you should worry about.

How to tell which happened

Look for sources. A searched answer shows links, citations, or a line like "Searched the web". An unsearched answer shows nothing. If you see confident specifics and no sources, you are in mode one, and the specifics are reconstructed from old patterns.

Figure 1.3

What actually happens when the tool searches



Six steps, and the model only enters at step four. Every judgement about which pages to trust is made by pattern, at step three and step five. That is why a searched answer is a research assistant's note rather than a source.

Most free tiers search automatically for questions that look like they need it, and not otherwise. You can force it: "Search the web before answering, and show me the pages you used." You can also forbid it, which is useful when you want the model's general explanation rather than a summary of whatever ranked first on the web this morning.

What search fixes

- **Staleness.** The training cutoff stops mattering for the question at hand. The new syllabus, this year's dates, the latest results are reachable.
- **Local specifics, sometimes.** If your district's population is on a page the search finds, you get it. If no page has it, you are back to guessing — and the model may not tell you that it found nothing useful.
- **Checkability.** You now have a page to open. This is the real gain. Open it.

What search does not fix

The model still has to decide which pages to trust, and it decides the way it decides everything: by pattern. A confident, well-formatted page ranks with an accurate one. A content farm that copied an old syllabus looks like the exam board. If three pages disagree, the model may blend them into a fourth answer none of them said.

It also summarises, and summaries lose things. The exceptions, the footnote that says "for the 2024 cohort only", the conditions — these fall out first, because they are the least predictable part of a sentence.

And the search itself can go wrong. A vague question produces a vague search, which fetches irrelevant pages, which produce an answer built on the wrong material. When an answer seems oddly off, look at what it searched for. Often the search query was not your question.

The teacher's rule

For anything that must be right — a date, a figure, a rule from the board — the model's job is to find the page, and your job is to read it. Treat a searched answer as a research assistant's note: "I think it's on this page, and it seems to say this." Then open the page.

For anything conceptual — an explanation, an analogy, a set of likely misconceptions — search usually makes the answer worse, because it replaces the model's broad synthesis with a summary of two websites. Turn it off for those.

An example that comes up every year

A teacher asks for the marks allocation in this year's Grade 12 physics paper. Without search, the model produces a plausible table from previous years. With search, it may find the board's circular, or it may find a tutoring site's guess from last year, and it will present either with the same tone. The only safe version of this question is: "Find the official circular from the board's own website and give me the link." Then you read the circular.

That is not the tool failing. That is the tool doing exactly what it is: a very good guesser, sometimes with a web page in hand.

BEFORE YOU MOVE ON

A teacher asks a chatbot for this year's exam timetable and gets a confident list of dates with no links or sources shown. What is the most likely explanation?

- A. The model looked the timetable up on the exam board's website in the background and is reporting what it found there accurately.
- B. The exam board published the timetable early enough to be in the model's training data, so the dates are accurate up to the cutoff.
- C. The model declined to search because exam timetables count as restricted information, and fell back on its own knowledge.
- D. The dates were generated from the pattern of previous years' timetables, and none of them should be trusted.

The answer and why: p. 317

Lesson 5 · 8 min

The conversation has a memory limit

THE ONE THING TO KEEP

A model can only see a fixed window of text at once, so it holds a chapter but not a textbook and forgets the instruction you gave forty messages ago; start a new chat per task and paste the material in every time.

What the model can see

The model does not have a memory the way you do. It has a window. Everything it can consider when writing its next word — your instructions, your pasted material, its own earlier answers — has to fit in that window at the same time. Anything outside it does not exist for the model.

The window is measured in tokens. A token is a piece of a word — roughly three-quarters of an English word, and often less than that for Hindi, Swahili or Arabic, which are split into more pieces. The big paid models offer windows in the range of 128,000 to a million tokens: 128,000 tokens is roughly 90,000 words, a 300-page book. Free tiers are typically given a fraction of that, and the tool does not always tell you the number.

For a teacher, the practical version is: **it can hold a chapter. It probably cannot hold a textbook.** A syllabus, a mark scheme, a class set of short answers — fine. Five past papers at once, a whole term's resources, a 200-page PDF — it will either refuse, or quietly keep only part of it.

The quiet failure

The refusal is the good outcome. The quiet failure is worse: the tool truncates. It keeps the beginning or the end, drops the middle, and answers as though it had read the lot. You asked for chapter 7 questions; chapter 7 was in the dropped part; you get confident questions on chapter 3.

Nothing warns you. The only defence is to give it less and check that it has what you think it has. A cheap check: "Before you start, tell me the heading of the last section you can see." If it names the wrong section, the material did not fit.

Long conversations forget their beginning

The window also holds the conversation itself. Every message you send and every answer it gives takes up room. Forty exchanges into a planning session, the instruction you gave in message one — "always use Kenyan shillings, sentences under 15 words" — has slid toward the edge of the window, and eventually out of it.

Some tools handle this by silently summarising the early part, which is why the model starts "forgetting" your constraints and reverting to dollars and long sentences. It is not being careless. Your instruction is no longer in front of it.

Two habits fix this:

1. **One task, one chat.** Worksheet in one chat; parent letter in another; misconception list in a third. Long threads drift.
2. **Repeat the constraints when they matter.** Rather than trusting message one, paste the two lines of constraints again with the request that depends on them.

Give it the material every time

Because there is no memory between chats, the model does not remember the syllabus you pasted yesterday. Tomorrow's chat starts blank. This feels wasteful, and teachers try to work around it by keeping one enormous chat open for weeks, which produces the drift above.

Figure 1.4**Everything competing for one window****Stored instructions**

Custom instructions, a project, a Gem.
Charged to the window on every message.

What you pasted

Syllabus extract, past paper, class set
of answers. A photograph costs more
than the same words typed.

The conversation so far

Every message you sent and every answer
it gave. This is the layer that grows.

Your current request

Nearest the end, and therefore weighted
most heavily.

The answer being written

Produced word by word, and it must fit
too.

All five layers occupy the same fixed space at the same time. Forty exchanges in, the conversation layer has grown until the instruction you gave first is pushed to the edge and then out. The model is not being careless; your constraint is no longer in front of it.

The better workaround is a text file on your phone with the things you paste often: your class description, the syllabus extract for this term, your standard constraints. Copy, paste, ask. Twenty seconds. Some tools offer a way to store this permanently — custom instructions, a "project", a "Gem" — and a later lesson covers using those. But the plain text file works everywhere and never gets lost when a product changes.

Images and files cost more than you think

A photo of a textbook page costs several hundred to a couple of thousand tokens, depending on the tool. A scanned PDF of thirty pages can be thirty images. It adds up faster than text, and it hits the free-tier limits sooner. If you

can type the two paragraphs you need, do; it is cheaper and more accurate than a photograph of them.

The honest limitation

Even inside the window, attention is not uniform. Models are measurably worse at using material buried in the middle of a long input than material at the start or end — researchers call it “lost in the middle”. So if the one constraint that matters most is in paragraph nine of your pasted brief, move it to the top or the bottom. Where you put things changes what gets used.

BEFORE YOU MOVE ON

A teacher pastes an entire 90-page PDF textbook into a free chatbot and asks for questions on chapter 7. The questions come back well-formed but about the wrong topic. What most likely happened?

- A. The model does not understand PDF files properly and produced generic questions instead of reading the book.
- B. The book was too big for the window, so only part of it was kept, and chapter 7 was not in that part.
- C. The model read the whole book but preferred chapter 7 of a different edition that it had seen during training.
- D. Free tiers cannot process uploaded files at all, so the model silently ignored the upload and guessed from the chapter number.

The answer and why: p. 317

Lesson 6 · 8 min

It has read the American curriculum

THE ONE THING TO KEEP

The training text is dominated by US and UK educational writing, so unprompted the model assumes Common Core, dollars and an American exam format; paste your syllabus and a sample paper and it adapts well, but it will never do so unasked.

Where the words came from

The model learned from what is on the internet, and the educational internet is lopsided. American teachers post lesson plans, American test-prep companies publish thousands of practice questions, American universities put courses online. British material comes next. Everything else — CBSE, ICSE and the state boards, WAEC and NECO, KCSE, the Brazilian ENEM, the Philippine K-12 — exists in the training text but in much smaller amounts, and much of it is old or second-hand.

So when you ask for "Grade 8 science", the model's default picture of Grade 8 is a US middle school. That is not bias in any deliberate sense. It is the average of what it read.

What the default gets wrong

- **Scope.** Which topics belong in which year. American Algebra II contains things your Class 10 syllabus puts in Class 11, or nowhere.
- **Terminology.** "Grade" versus "Class", "Form" and "Standard"; "math" versus "maths"; "trapezoid" versus "trapezium"; "period" versus "full stop". Small, but a worksheet full of the wrong terms reads as foreign to students and to your head of department.

- **Exam format.** It will produce multiple-choice-heavy tests with a US shape, not the two-mark, three-mark, five-mark structure of a board paper, and not the "Section A: answer all, Section B: choose three" layout your students have to practise.
- **Setting.** Dollars, miles, Fahrenheit, yellow buses, Thanksgiving. The worksheets lesson dealt with this for word problems; it applies to everything.
- **Standards language.** It may cite Common Core codes or Next Generation Science Standards as if they were your standards.

The fix works, and it is quick

The model is bad at guessing your syllabus and very good at following one it is shown. Three things to paste, once, into the text file from the previous lesson:

1. **The chapter list** for your subject and year, from the board's own document. Twenty lines.
2. **One past paper**, or the section structure and marks allocation if the paper is too long. This teaches it the shape.
3. **A vocabulary line.** "Use 'Class 9' not 'Grade 9', British spelling, rupees, metric units, and the terminology in the pasted paper."

With those three in the prompt the output shifts almost entirely. Teachers who do this once describe the difference as "it finally sounds like our school". The mathematics did not change. The frame did.

Regional language textbooks

If your students learn from a Hindi-medium, Kiswahili-medium or Amharic-medium textbook, the model has probably read very little of it. It knows the subject; it does not know the specific explanations, examples and terms your book uses, and terms matter — a student who has learned a concept under one name is confused by a worksheet that uses another.

Photograph or type the relevant page and paste it. "Use the terms and the worked example style in this page." It will follow.

A worked example

A biology teacher in Lagos asks for "a revision sheet on the circulatory system for Grade 11". She gets a sheet organised around a US Anatomy and Physiology unit: the cardiac cycle in detail, ECG traces, blood pressure regulation, none of it in the WAEC SS2 syllabus, and a vocabulary that says "sophomore" and "AP Bio".

She pastes the WAEC topic list for the circulatory system — six lines — and one past WAEC question with its marking guide, and adds: "SS2, Nigerian secondary school, WAEC style, British spelling, use only the topics listed." The new sheet stays inside the six lines, phrases questions the way WAEC phrases them ("State three functions of...", "Describe the passage of blood from..."), and allocates marks the way the marking guide does. Same request, twenty seconds of pasting, a usable resource.

Notice what she did not have to do. She did not explain the Nigerian school system or argue with the model about what Grade 11 means. She showed it the shape and it copied the shape. That is the fastest way to move any of these tools: examples over descriptions.

Where the default helps

Sometimes the American frame is exactly what you want: preparing students for the SAT, for a US university application essay, or for an international curriculum like the IB or Cambridge IGCSE, all of which are heavily represented in the training text. For those, the model's defaults are an asset.

The honest limitation

Pasting the syllabus fixes scope. It does not fix depth of local knowledge. The model can now stay inside your Class 10 chapter list, but it still knows less about how your board phrases questions, which topics carry marks year after year, and which chapters your students find hard, than a teacher who has marked ten years of papers. You have that. Bring it. The model supplies the volume; you supply the shape of your board.

BEFORE YOU MOVE ON

A Class 10 teacher in Bihar asks for "a chapter test on trigonometry for Grade 10". The test comes back with topics from a US Algebra II course, several not in her board's syllabus. What is the right diagnosis and fix?

- A. The model does not know Indian mathematics well enough, so she should not rely on it for her subject at all.
- B. It defaulted to the syllabus it has read most; pasting her board's chapter list and one past paper will fix the scope.
- C. She used the word "Grade" instead of "Class", and that single word switched the model into an American mode for the whole test.
- D. The board's syllabus is not public in a form the model can read, so no prompt can bring the output into the right scope.

The answer and why: p. 318

Lesson 7 · 8 min

A chatbot is not a calculator

THE ONE THING TO KEEP

Next-word prediction is a poor way to multiply, so every number in generated material is checked with a calculator, and the safe pattern is to have the model write the working while a separate tool does the arithmetic.

Why 47×83 is hard for it and easy for a two-rupee calculator

A calculator computes. It has a circuit that multiplies and the answer is exact every time.

A language model predicts. When it writes " $47 \times 83 = 3901$ ", it is producing the digits that most commonly followed similar-looking expressions in text. For small, common products — anything in the times tables, round numbers

— it has seen the answer so often that it is effectively memorised. For 47×83 it has seen the exact expression rarely, and it assembles the digits from pattern. It is right more often than chance and wrong often enough that you cannot trust it. The correct answer, incidentally, is 3901; try a few of your own and watch it slip on larger ones.

This is the same mechanism as every other lesson in this block. Arithmetic is a specific with few occurrences in text. That is the unreliable list.

What it gets wrong, specifically

- **Multi-digit multiplication and division.** Increasingly wrong as the numbers grow.
- **Long chains.** A five-step calculation where step three is slightly off, carried through with perfect confidence.
- **Counting.** How many words in this paragraph, how many items in this list. It does not count; it estimates.
- **Unit conversion with awkward factors.** 1 acre to square metres, done from pattern.
- **Anything with a carry across many columns.**

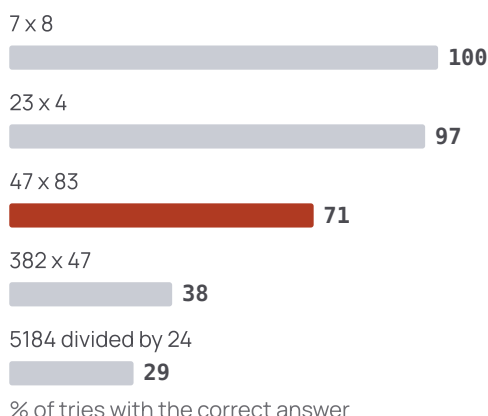
What it gets right is the method. The setting-up of an equation, the choice of formula, the order of operations, the explanation of why — these are language, and language is what it does. Which produces the pattern in the check question: sound working, wrong number.

The safe division of labour

Let the model do what it is good at and hand the numbers to something that computes.

Generating problems. Ask for the problems and the method, and produce the answer key yourself with a calculator. Ten problems take five minutes to check. Or ask for problems whose answers you specify: "Write a word problem about a shopkeeper whose answer is exactly 240 rupees." It will fit the story to the number, which is far more reliable than fitting a number to the story.

Figure 1.5

Where the arithmetic falls apart

One free-tier chatbot, each expression asked a hundred times in fresh chats, no calculator and no code. Small common products are effectively memorised. Accuracy falls as the digits grow, because the exact expression appears in text less and less often. The method in the working stayed sound throughout, which is exactly what makes the wrong numbers hard to notice.

Checking student working. "Here is a student's solution. Identify the step where the method goes wrong." It is good at this. Then verify the arithmetic yourself.

Code, if you have it. Several tools can write and run a short piece of Python when asked, and Python computes exactly. "Solve this using code and show the result" produces a real answer. Free tiers include this in some products and not others. If you see a snippet of code and its output in the answer, the number came from a computation. If you see only prose, it came from prediction.

"Reasoning" modes. Newer models have a mode that works through problems in steps before answering, and it markedly improves arithmetic — the model effectively writes out a long-multiplication and reads the answer off its own working. It is slower and often limited on free tiers. Even then, verify. Better is not the same as exact.

A classroom use that turns the weakness into a lesson

Ask the model, in front of the class or on a printed transcript, to multiply two three-digit numbers. Have students check it on paper. When it is wrong — and choose numbers where it is — the lesson lands: this machine is fluent and it cannot multiply. Then ask why. A Class 8 student who has done the prediction activity from the first lesson can answer: because it guesses digits the way it guesses words.

The honest limitation

The line is moving. Each generation of models does arithmetic better, and tools that quietly call a calculator behind the scenes are becoming common. So you may find your tool gets 47×83 right every time you try. That does not change the rule, because you cannot see from the outside whether a given answer came from a computation or a guess. Check the numbers. It costs a minute and it is the minute that stops a worksheet with a wrong answer key reaching forty children.

BEFORE YOU MOVE ON

A maths teacher asks a chatbot to produce ten word problems with an answer key. She notices the model's step-by-step working on question 6 is sound but the final product is off by 40. What best explains this?

- A. The model made a deliberate small error to test whether the teacher was actually checking the key before printing.
- B. The word problem was badly worded, so the model solved a slightly different problem from the one it had written a moment before.
- C. The method is language, which it handles well; the product is digits predicted from pattern, which fails on uncommon combinations.
- D. The answer key was generated in a separate pass from the questions, so the number belongs to a different question in a different set.

The answer and why: p. 318

Lesson 8 · 8 min

Same question, different answer

THE ONE THING TO KEEP

Answers are sampled from a distribution rather than looked up, so asking the same question three times and comparing is a real check — disagreement flags an unreliable question, though agreement never proves the answer is right.

Ask twice, get two answers

Type a question into a chatbot. Open a new chat and type it again. Often the answers differ — a different opening, different examples, sometimes a different fact. Teachers meeting this for the first time assume the tool is broken or that they phrased something differently. Neither. This is how it works.

The model does not look up an answer. At each step it has a spread of possible next words, each with a probability, and it picks one — usually not always the single most likely, but a draw weighted by likelihood. That draw is random. Run it twice and you walk two slightly different paths through the same distribution, and they diverge more with every word.

Tools expose a control for this, called temperature. Low temperature means the model nearly always picks the most probable word, so answers become repetitive and stable. High temperature means more randomness, more variety, more nonsense. Consumer chatbots set a middle value and do not let you change it. Developer tools do.

Why it matters for a teacher

Consistency in marking. Give the model the same essay twice and it may give two different grades. A later lesson deals with marking properly; for now, know that a single run is one draw, not a verdict.

Reproducing a worksheet. The excellent set of questions you generated last week cannot be regenerated by typing the same prompt. Save the output, not just the prompt.

The apparent argument. Students who "prove" the AI is inconsistent by asking twice have discovered something true, and the first lesson's prediction activity explains it: for the old woman walking slowly into the _____, the class produced four answers. Ask that class again tomorrow and the tally shifts. Same distribution, different draw.

Turning randomness into a check

The variation is also a tool. For a question you cannot easily verify, ask it three times in three fresh chats and compare.

- **Three different answers.** The question is on unreliable ground. The model does not have a stable pattern for it. Do not use any of the three; go to a source.
- **Three matching answers.** The pattern is stable. This is evidence that the model has seen the answer many times — the reliable-ground signal from earlier in this block.

But notice what agreement does and does not tell you. It tells you the model's learned pattern is peaked. It does not tell you the peak is the truth. If most of the internet believes that we only use ten percent of our brains, or that the Great Wall is visible from space, the model reproduces that belief every single time, in every chat, with total consistency. Popular misconceptions are the most stable outputs it has.

So: disagreement is a red flag you can trust. Agreement is a green flag you cannot. That asymmetry is the whole lesson.

Asking differently, not just again

A stronger version of the check varies the question rather than repeating it. Ask for the date of an event; then, in a new chat, ask what happened in that year; then ask it to list events of that decade. If the same event appears in the

same year across three different framings, the pattern is deep. If it moves, you have found a specific the model is reconstructing rather than recalling.

Another variation: ask a second tool. Different companies train on overlapping but different text, so a fact that two independent models agree on has a better claim than one model agreeing with itself. Better, not proven — they read much of the same internet.

For the classroom

This is one of the easiest properties of AI to demonstrate without a device. Print the same question asked three times, with the three answers. Ask the class to explain why they differ. A student who has done the tally activity will say: because each time it rolls the dice on the next word. That student understands more about these systems than most adults.

The honest limitation

Some products now cache or stabilise answers to common questions, so you may find a well-known query returns identically every time even though the mechanism underneath is still a draw. And a few offer a "deterministic" setting for developers. Neither changes the reasoning: identical output from one system is one source, however many times you read it.

BEFORE YOU MOVE ON

A teacher asks the same factual question in three separate chats and gets the same answer each time. She concludes the answer is verified. What is the flaw?

- A. Three repetitions is too few to mean anything; she needs at least ten before the agreement starts to count as evidence.
- B. The model remembered her earlier chats and copied its own previous answer, so the three answers are not independent of each other.
- C. Facts should be asked once with search switched on, and repeating a question is only a useful technique for creative tasks.
- D. Agreement shows the pattern is stable, not correct; a common misconception is reproduced every time.

The answer and why: p. 319

CASE STUDY · MODULE 1

The tool everyone already had

An illustrative case, built from documented patterns.

A government secondary school in Barabanki district, Uttar Pradesh, choosing one free AI tool for forty-one teachers.

The school had forty-one teachers and a computer room with two working machines. Every teacher had a phone. Perhaps half had data that lasted to the end of the month. When the principal asked Sunita Verma, the science head, to recommend one tool for the staff, the question looked like a comparison of models and turned out not to be.

There were two realistic candidates. The first was the assistant inside WhatsApp. It needed no new account, no new app and no email address, it ran on the data plan people already paid for, and the staff group where timetables were shared was already open on every phone. The second was ChatGPT or Gemini in the phone browser: better on long, careful work, better at holding a pasted syllabus, and requiring each teacher to create an account and then walk through a settings menu to find the control that stops conversations being used to train the model.

Mrs Verma tested both for a fortnight, on the work she actually had to do. Most of what she found was reassuring. Both explained photosynthesis and the water cycle well enough to use. Both, asked what Class 8 students get wrong about the water cycle and what they are thinking when they get it wrong, produced eight answers, two of which she had not met in nineteen years. That single request turned out to be the most valuable thing either tool did all fortnight, and it cost nothing but a sentence.

Both also produced a Class 8 worksheet on the district's part in the freedom movement, fluently, with a founding year for the town's oldest school — and the two years did not match. She checked the school's own foundation stone. Neither was right. She noticed, separately, that the answers she got at nine in the evening were thinner than the answers she got at six in the morning, which she assumed was her own tiredness until a colleague reported the same thing.

All of that was useful and none of it was the decision. The decision was the setting.

In the browser tool she had found the control in four minutes and switched it off. In the WhatsApp assistant she looked for twenty minutes, could not find an equivalent, and could not get an answer from anyone at the district office about whether one existed. She knew what teachers in that staff group already did, because she was in it: they forwarded photographs of student scripts to each other and asked for help marking them.

The two costs

Recommending WhatsApp meant something close to full adoption. The eight teachers with the oldest handsets could use it. The five who had never had an email address could use it. It would be in the staffroom by Monday, and the material generated with it — worksheets, explanations, question lists — would be real and useful, because the tool that gets used is the tool people already have open.

It also meant she would be recommending, to forty-one colleagues, a tool she could not tell them was safe to paste a child's work into. She could write a rule. She could not enforce one.

Recommending the browser tool meant a better position on data and a worse one on everything else. Her estimate, from the two afternoons she had spent helping colleagues make accounts, was that perhaps twelve of the forty-one would complete the setup and keep using it. The rest would carry on as before, and the tool would become one more thing that the better-connected half of the staffroom had and the other half did not.

There was a third option she considered and dropped. A laptop in the office could run a small open model with nothing leaving the machine, which answers the data question completely. It also meant one machine, in a room that is locked at four o'clock, running a model noticeably weaker than either candidate, for a staff of forty-one who plan their lessons at night. It is the right answer for a school that cannot send any data out and the wrong answer for a teacher with a phone, and this was a school of teachers with phones.

She had one staff meeting, forty minutes, and a principal who wanted a single answer.

YOUR ANSWERS

1. Both tools gave a founding year, the years differed, and both sounded equally certain. Explain why, from the mechanism.

2. Her rule depended on teachers classifying their own material at the moment they used the tool. What made it break, and what would have made it more likely to hold?

3. What did she buy by recommending the weaker tool to most of the staff, and what did she give up?

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

She refused the single answer and split it by content rather than by tool. One sheet, pinned in the staffroom: nothing with a child in it — no names, no scripts, no photographs of work — goes into WhatsApp; everything with no child in it goes wherever you like. She set up the browser tool with training switched off for the nine teachers who wanted it, and taught the setting rather than the product. By December twenty-nine teachers were using the WhatsApp assistant weekly and the nine had accounts. The rule held for a term and then broke in one department, where a teacher photographed a class set to get help marking it; Mrs Verma found out because a colleague mentioned it in passing, not because anything checked. The wrong founding year stayed on the Class 8 wall for three weeks, until a student's grandfather corrected it. She later said that was the most useful thing that happened all term, because she then taught the reliability lesson with the wall as the example and nobody in that class forgot it.

The questions, discussed

1. Both tools gave a founding year, the years differed, and both sounded equally certain. Explain why, from the mechanism.

A founding year for one small town's school has been written down rarely and inconsistently. The model is producing the most likely next words after a phrase like "the school was founded in", and the most likely next words are a plausible year. Nothing in the tone changes between a fact written down a million times and one written down three times, because there is no internal measurement of uncertainty reaching the sentence. That the two tools disagreed was the useful signal: disagreement is a red flag you can trust, while agreement would only have told her the pattern was stable, not true.

2. Her rule depended on teachers classifying their own material at the moment they used the tool. What made it break, and what would have made it more likely to hold?

It broke because it lived on a sheet that nobody re-read at the point of use, and because the tempting case — a stack of scripts and no time — is exactly the case the rule forbids. Rules that survive are the ones attached to the routine rather than to memory: a standing line in the marking procedure, the phone photographs taken with names covered as a matter of course, and someone whose job it is to ask about it once a term. Nothing in the school checked, so the only way anyone learned of the breach was conversation.

3. What did she buy by recommending the weaker tool to most of the staff, and what did she give up?

She bought adoption, and with it the actual benefit — thirty teachers generating misconception lists and worksheets instead of twelve. She also kept the staffroom from splitting into those with accounts and those without. What she gave up was any control over student data in the tool most people were using, and any honest ability to say the training setting was off. She traded a data position she could not enforce for a usage position she could see, and made the trade explicit in writing rather than pretending the two tools were equivalent.

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. 314. If two go wrong, re-read the module before the exam.

- 1 You need three things for a Class 9 worksheet: an explanation of photosynthesis, the year your town's first secondary school opened, and 15 per cent of 4,860. You ask one chatbot, with no web search. Which of the three is safest to use as it comes back?
 - A. The year, because a founding date is a single simple fact and the model either has it or says it does not.
 - B. The explanation, because core school science has been written down consistently by very many people.
 - C. The percentage, because arithmetic follows a rule the model applies rather than a pattern it predicts.
 - D. All three equally, since the model states each one in the same tone and has no reason to be more sure about any of them.

- 2 An answer arrives with three links and a line saying the tool searched the web. What has that changed?
 - A. It has removed the training cutoff and also removed the possibility of an invented citation, because every claim now comes from a page that was actually fetched.
 - B. It has slowed the answer down without changing its accuracy, because the model writes from its training either way and the links are added afterwards.
 - C. It has fixed staleness and given you pages to open; it has not fixed the judgement about which page to trust.
 - D. It has made the answer verified, because a searched answer only contains what the fetched pages say.

- 3 You paste five past papers and ask for ten questions in the style of paper four. Back come ten confident questions in the style of paper one. What most likely happened?
- A. The material did not all fit, so the tool kept part of it and answered as though it had the lot.
 - B. The model prefers the beginning of any long input and was designed to summarise from the first document it is given.
 - C. The free tier cannot read more than one document at a time, and it would have refused had you asked it to.
 - D. Papers two to five were rejected by the safety filter because past papers are copyrighted material, which the model will not reproduce.
- 4 Asked for a Class 10 maths revision sheet, the model produces US Algebra II content, dollars, and the word "sophomore". What is the quickest thing that actually fixes it?
- A. Tell it to be accurate and to follow the Indian curriculum, restating the instruction whenever it drifts.
 - B. Switch to a different tool, since the American default belongs to one company's model rather than to the training text.
 - C. Ask it to search the web for your board's syllabus, because a specific query will bring the board's own document up first and the model will then follow it.
 - D. Paste the board's chapter list and one past paper, and name the terms and units to use.

- 5 You want twenty word problems with a trustworthy answer key. Which approach gives you one?
- A. Ask the model to check its own arithmetic at the end, which catches most carry errors before you see them.
 - B. Use a reasoning mode, which works through the problem in steps and therefore produces exact arithmetic rather than a prediction.
 - C. Ask for the problems and the key in one message, so that the numbers stay consistent with each other throughout.
 - D. Ask for problems whose answers you specify, or work the key yourself with a calculator.
- 6 You ask the same question in three fresh chats and get the same answer three times. What have you established?
- A. That the pattern is peaked, which is not the same thing as the answer being true.
 - B. That the answer is almost certainly right, because three independent runs converged on it.
 - C. That the tool has cached the answer, which is what happens to any question asked more than once.
 - D. That the question is on reliable ground, since agreement across three draws is the same evidence as agreement between two different companies' models.

2

MODULE 2

Prompting as a teacher

Write a prompt that carries the five things only you know — class, constraints, prior lesson, misconception and format — iterate on the output rather than accepting the first draft, control the reading level and register of what comes back, and keep a personal prompt file that makes a good result repeatable next term.

9	The five things only you know	50
10	Show it one of yours	54
11	The second message matters more than the first	58
12	Set the voice and the reading level	62
13	Roles that help, and roles that do nothing	66
14	Prompting for formats you can print	70
15	When it refuses a lesson	74
16	Keep a prompt file	77
	<i>Case study: The department that shared its prompts, and the mock paper inside them</i>	81
	<i>Module 2 test</i>	86

Lesson 9 · 9 min

The five things only you know

THE ONE THING TO KEEP

Every teaching prompt should carry who the students are, what constrains the room, what came before, which misconception is live, and what shape the output must take — because the model can only condition on what is in front of it, and none of those five are.

Why the generic prompt gets the generic answer

The model produces the most likely continuation of what it is given. Give it "a starter activity on fractions" and the most likely continuation is the average of every fractions starter on the internet: a pizza, a think-pair-share, forty-five minutes, a projector. Not because the model is lazy, but because you gave it nothing to be specific with.

Everything that would make the activity fit your room lives in your head. The model cannot reach it. So the craft of prompting, for a teacher, is mostly the craft of noticing what you know that it does not, and writing it down.

The five things

There are five, and they cover nearly every case.

1. Who. Year, subject, class size, language background, reading level. "Class 5, 48 students, most read English as a second language" changes the vocabulary, the sentence length and the activity format in one line.

2. What constrains the room. Minutes available, what equipment exists and does not, whether students can move, whether there is a printer. "35 minutes, fixed benches, no printer, one blackboard." This kills the projector-dependent activity before it is written.

3. What came before. Yesterday's lesson, last week's topic, what they already have. "Yesterday they added fractions with the same denominator." The model now knows where to start and, just as importantly, where not to.

4. The live misconception. The specific wrong idea you saw in their work. "Half of them believe $1/3$ is smaller than $1/4$ because 3 is smaller than 4." This is the single most valuable line you can write, because it is the thing no lesson-plan website contains and the thing your lesson actually has to fix.

5. The shape of the output. What you want back, and what you do not. "One five-minute starter. No full lesson plan." Without this you get the whole thing, forty-five minutes of it, and you spend your evening cutting.

Putting it together

Here is the template as a block you can copy into the prompt file this module ends with:

```
Class: [year, subject, size, language background]
Room: [minutes, equipment, layout, printing]
Before this: [what they did last lesson]
Misconception: [the specific wrong idea, in the students' words
if you have them]
Give me: [exactly what, exactly how many, exactly how long]
Do not: [what you do not want]
```

Filled in, for an English class:

Figure 2.1

The five things the model cannot reach

- Who**
Year, subject, class size, language background, reading level.
- What constrains the room**
Minutes, equipment, layout, whether there is a printer.
- What came before**
Yesterday's lesson. Tells it where to start and where not to.
- The live misconception**
The wrong idea you saw in their work, in their words. The most valuable line you can write.
- The shape of the output**
Exactly what, exactly how many, exactly how long, and what you do not want.

None of these five is in the training text, and all five change the output. The fourth is the one no lesson-plan website contains. Keep it near the top or the bottom of the prompt, because material buried in the middle of a long input is used least.

```
Class: Form 2, English, 40 students, Kiswahili first language,
reading age around 10.
Room: 40 minutes, no projector, I can photocopy one page.
Before this: They wrote a paragraph describing their journey to
school. Most paragraphs were a list of events with "and then".
Misconception: They think a good paragraph is a long one. Nobody
used a topic sentence.
Give me: Two example paragraphs about the same journey, one that
is a list and one with a topic sentence and three supporting
details, short enough for a reading age of 10. Then four
questions that make students notice the difference.
Do not: Explain what a topic sentence is. I will do that.
```

That prompt takes ninety seconds to write and produces something you can photocopy. The generic version takes ten seconds and produces something you cannot.

Let it interview you

If you cannot think what to put, reverse the direction. "Before you write anything, ask me five questions about my class and my room that would change what you produce." The model asks; you answer in a line each; then it writes. This works because the model has read a great deal about what makes lessons fail, and it will ask about the things it needs. It also has a limit: left unsteered it asks generic questions — "what are your learning objectives?" — so add "ask about constraints I might not have mentioned, not about objectives".

What to leave out

Do not describe your teaching philosophy. Do not explain that you want the activity to be engaging. Do not write "please" and "thank you" — they cost nothing but they add nothing. Every sentence that is not one of the five things dilutes the ones that are, and the memory-limit lesson explained why: material in the middle of a long input is used less than material at the edges. Keep the misconception near the top or the bottom.

The honest limit

Five things make the output fit your room. They do not make it correct. A perfectly targeted starter can still contain a wrong fraction comparison, and the checking habits from the first block apply to every output, however well-prompted. Prompting well reduces the editing. It does not remove the reading.

BEFORE YOU MOVE ON

Two prompts ask for the same thing. Prompt A: "Give me a starter activity on fractions." Prompt B: "Class 5, 48 students, 35 minutes, no printer. Yesterday they added fractions with the same denominator. Half of them believe $\frac{1}{3}$ is smaller than $\frac{1}{4}$ because 3 is smaller than 4. Give me one five-minute starter that exposes that belief, done at desks with nothing but a pencil. No full lesson plan." Why is B's output more useful, mechanically?

- A. B is more polite and detailed, and the model rewards effort in the prompt with more effort in the answer.
- B. B contains the constraints the model has no other way of knowing, so the answer is drawn from a much narrower and more relevant space.
- C. B works because it names the year group, and the model has a stored curriculum for each year that it consults when a year is specified.
- D. B is longer, and longer prompts always produce better answers regardless of what the extra words say.

The answer and why: p. 321

Lesson 10 · 8 min

Show it one of yours

THE ONE THING TO KEEP

One pasted example of your own work steers the output harder than any description of it, because the model copies the shape of what it is shown — including, if you are not careful, the flaws.

Description versus demonstration

You can spend three paragraphs describing the kind of comprehension question you want: inference-based, one sentence, no yes/no answers, referring to a specific line, pitched at Class 7. Or you can paste one question

you wrote last year that was exactly right and say: "Eight more like this, on the attached passage."

The second wins, nearly every time. The model is a pattern-continuer, and a concrete example is the densest possible statement of a pattern. Your description is your theory of what makes the question good. The example is the question. It carries things you did not know you cared about — the length of the stem, the way you phrase "according to the text", where the line reference sits — and the model reproduces all of it.

Practitioners call this few-shot prompting. You do not need the term. You need the habit: when you have a good one, show it.

What to show

- **A worked example in your style.** For maths, paste a solution laid out the way you lay them out on the board. The generated ones will follow your layout, which means students see one consistent method rather than the model's default.
- **A question with its mark scheme.** The mark scheme teaches the model what a full answer contains, so the next questions come with usable schemes.
- **A student-facing instruction.** If your instructions always say "Show your working. Circle your final answer," paste one and they all will.
- **A parent message you were pleased with.** Tone is nearly impossible to describe and trivial to demonstrate.
- **A whole worksheet.** Photograph it if it only exists on paper. "Same format, same difficulty, different numbers and a different context" produces a parallel version for a second class or a resit.

Two or three examples are better than one when the examples vary along the dimension you want varied — three questions of increasing difficulty tells the model to produce a spread. Beyond three or four you are spending the memory window on demonstration and getting little back.

The flaw comes with it

The mechanism cuts both ways. The model copies the shape it is shown, and it has no way to distinguish the parts you are proud of from the parts you did not notice.

If your example question has an ambiguous stem, you get eight ambiguous stems. If your worked solution skips a step you always skip on the board, every generated solution skips it. If the worksheet you photographed had a typo in the instructions, the typo is now a feature.

So choose the example with care. Read it once as though a colleague wrote it. Fix what you find before you paste. A clean example is worth more than three good ones with one flaw between them.

Telling it what to keep and what to change

An example on its own leaves the model to guess what "like this" means. Say it:

Keep: the structure, the length, the kind of thinking required, the instruction wording. Change: the passage, the numbers, the context.

Without the split, the model sometimes keeps the content and changes the structure, which is the opposite of what you wanted. It also helps to name the thing you are showing: "This is a two-mark inference question" gives it a label to generalise from, so the eight new ones are two-mark inference questions and not two-mark recall questions in inference clothing.

The negative example

Sometimes the fastest steer is a bad example. "Here is a question I do not want — it can be answered by copying a sentence from the passage. Here is one I do want." Two examples, one labelled bad, one labelled good, and the model has the boundary. This is particularly effective for the things teachers find hardest to describe: the difference between a question that tests reading and one that tests understanding is easier to show than to define.

A department's worth of examples

If your department has a shared folder of good questions, worksheets and mark schemes, you have a prompt library already. Each item is a demonstration. A new teacher who pastes the department's best Class 9 test and asks for a parallel version inherits the department's standards in an afternoon. That is a better induction than most schools manage, and it costs nothing but the pasting.

The honest limit

An example steers form. It does not steer truth. Eight questions "like this one" on a new passage can still ask about a detail the passage does not contain, or carry a wrong answer in the key. The example bought you the right shape. The checking is still yours.

BEFORE YOU MOVE ON

A teacher pastes one of her own past exam questions, which happens to have an ambiguous stem, and asks for eight more like it. All eight come back ambiguous in the same way. What has happened?

- A. The model recognised the ambiguity as an intended feature of her assessment style and reproduced it deliberately.
- B. Eight is too many to generate from one example; with three examples the ambiguity would have been averaged out.
- C. The example steered the shape, flaw included, because the model copies what is shown rather than what is meant.
- D. The model cannot write exam questions without a mark scheme, so it fell back on copying the example word for word.

The answer and why: p. 322

Lesson 11 · 8 min

The second message matters more than the first

THE ONE THING TO KEEP

Treat the first output as a draft and spend your effort on the critique — specific, pointed at one thing, with the reason — because a model steered by feedback improves fast, while a first prompt rewritten five times mostly does not.

Where teachers spend their effort, and where it pays

Most people meeting these tools put all their effort into the first prompt. They write it, read the output, feel disappointed, and rewrite the prompt from scratch. Three rounds of that and they conclude the tool is not much use.

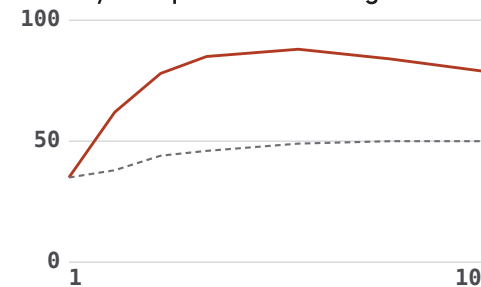
The effort is in the wrong place. The first output is almost always a draft. The fastest route from draft to usable is not a better first prompt; it is a pointed second message. A model handed a specific critique of its own output improves markedly, because the critique tells it exactly which part of the space to move to. A rewritten first prompt starts the guess again from nothing.

What a useful second message looks like

Three parts: what is wrong, where, and why.

Figure 2.2

Two ways to spend ten messages



Across: Messages spent on one task
Up: % of the output you would use as it is
— Pointed critique of the draft
-- First prompt rewritten from scratch

Critique moves the output because it names which part of the space to move to. A rewritten first prompt starts the guess again from nothing. Note the turn after eight messages: the conversation is long, early constraints are drifting out of the window, and the right move is a fresh chat with the best version pasted in as the example.

Weak: "Make it better."

Weak: "These are too easy."

Useful: "Questions 1 to 6 can all be answered by copying a phrase from the passage. I want questions where the answer is not in the text and the student has to infer it. Question 4 is close — do the rest like question 4."

The last version names the fault, points at the location, gives the reason, and hands over an example of the target from the model's own output. That last move is the strongest available: you are showing it one of its own, which the previous lesson explained is the densest steer there is.

Why "better" fails mechanically

"Better" has no direction. The model resolves it to the most common meaning of better in the text it learned from, which for generated content is nearly always "more": more questions, longer explanations, extra sections. You asked

for a change of kind and got a change of quantity. This is not the model misunderstanding. It is the model correctly guessing what an unspecified "better" usually means.

Every vague critique has a default the model falls into. "More engaging" produces exclamation marks and a game. "Simpler" produces shorter sentences with the same concepts. "More rigorous" produces longer sentences with the same concepts. If you find the same disappointing pattern each time, the pattern is the default, and the fix is to say what you actually mean.

One thing at a time

A critique that lists six problems gets six half-fixes. The model addresses each, briefly, and the output becomes a compromise. Fix the biggest problem first, look at the result, then fix the next. Three focused rounds beat one comprehensive complaint, and they take less time in total because each output is easier to read.

When to stop and start again

Iteration has a limit, and the memory-limit lesson explained it. After eight or ten rounds the conversation is long, the early constraints are drifting out of the window, and the model is carrying the residue of every abandoned version. Signs that you have reached it: the output starts reintroducing things you removed three messages ago; the model apologises and produces the same thing again; the tone shifts.

At that point, take the best version, open a new chat, paste the best version as your example, and write a clean prompt around it. You lose nothing — the good output travels with you — and you get a model with a clear window.

The critique is teaching you too

There is a side benefit that experienced teachers notice quickly. Writing "these are all recall, I want inference" forces you to name what you wanted, in words. Teachers who spend a term critiquing generated questions become sharper at

seeing the difference between recall and inference in their own questions, because they have had to articulate it two hundred times. The tool did not teach them that. The act of correcting it did.

A pattern to copy

For any generated resource, this three-message shape covers most cases:

1. The five-things prompt.
2. "Questions [numbers] have problem [X] because [reason]. Do them like question [Y]."
3. "Now check every question against the passage and tell me any whose answer is not actually in it, or is ambiguous."

The third message is a self-check. It is not reliable — the model checking its own work misses things a person would catch — but it catches some, cheaply, before you read. Then you read.

The honest limit

A model steered by critique can circle. Ask it to make questions harder and then more accessible and then harder again and you will get an oscillation rather than a synthesis, because each message pulls toward one pole. When two goals pull against each other — rigour and accessibility, brevity and completeness — state both in the same message and let it find the balance in one go. Sequential pulls do not converge.

BEFORE YOU MOVE ON

A teacher's first output is decent but the questions are all recall. She replies: "Make it better." The next version is longer, with more questions, still all recall. What went wrong?

- A. The model needs to be told explicitly that recall questions are bad before it will stop producing them, so the fix is to add that sentence.
- B. "Better" has no direction, so the model improved along its default axis of more; she needed to name the one thing to change.
- C. Follow-up messages cannot change the type of question once a set has been produced; she must start a new chat with a rewritten first prompt.
- D. The model had already produced its best attempt on the first pass, and any further request can only add volume to what is there.

The answer and why: p. 322

Lesson 12 · 9 min

Set the voice and the reading level

THE ONE THING TO KEEP

"Simple" is not an instruction; sentence length, word frequency, idiom and the ratio of concrete to abstract are, and the model drifts back to its default register after a few paragraphs unless the constraint is restated.

"Simple" is not a specification

Ask for something "simple" or "for beginners" or "in plain language" and you get the model's idea of simple, which is a slightly shorter version of its normal voice. The normal voice is the average of the internet: sentences around 20 to 25 words, a fondness for "however" and "additionally", abstract nouns, a formal register that makes a Class 6 reader stop halfway down the page.

Simplicity has to be specified as things that can be counted or checked.

- **Sentence length.** "No sentence longer than 15 words." Measurable, and the single most effective instruction.
- **Word frequency.** "Use only words a 10-year-old would hear at home or in school." Or, more precisely: "Avoid any word with more than three syllables unless it is a subject term, and define every subject term the first time."
- **Idiom.** "No idioms, no metaphors unless you explain them." Idioms are the thing second-language readers trip on hardest, and the model uses them without noticing.
- **Concrete over abstract.** "Every abstract idea gets a concrete example in the next sentence."
- **Person.** "Address the reader as 'you'." Second person is easier to read than the impersonal passive.

Put four of those in a prompt and the output changes visibly. Put "write simply" and it does not.

Setting a reading age

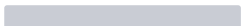
Readability formulas — Flesch-Kincaid is the common one — estimate a reading level from average sentence length and syllables per word. They are crude, they were built for English, and they do not measure whether the ideas are hard. But they give you a number, and a number is something the model and you can both aim at.

"Write this at a Flesch-Kincaid grade level of 5" works reasonably well in English. Then check it yourself: Microsoft Word shows readability statistics under the spelling options; free web tools such as the Hemingway Editor show sentence length and grade level as you paste. If the model says grade 5 and the tool says grade 9, believe the tool.

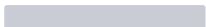
Figure 2.3

“Simple” against four things you can count

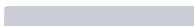
No instruction (the default voice)

 11.4

“Write simply”

 9.7

“In plain language, for beginners”

 9.2

Four measurable constraints

 5.3

Flesch-Kincaid grade level of the output

The same explanation of the water cycle, asked for four ways, measured afterwards in a free readability tool. The target was grade 5. Adjectives move the number by a grade or two; sentence length, word frequency, no idiom and a concrete example per abstract idea move it by six.

For languages other than English the formulas do not apply, and the model's sense of “simple Hindi” or “simple Kiswahili” is weaker than its sense of simple English, because it has read less graded material in those languages. There, the sentence-length and idiom instructions still work; the numeric target does not.

Why it drifts

The check question describes the most common failure. Paragraph one obeys; paragraph four does not. The mechanism is the one from the first block: each sentence is predicted from the sentences before it. Your instruction is at the top; the last three sentences are immediately before. Once one sentence runs slightly long or slightly formal, the next is predicted from a long formal sentence, and the register slides back toward the default. The instruction has not been forgotten. It has been outweighed.

Two fixes:

1. **Ask for shorter pieces.** Two hundred words at a time drifts less than eight hundred. Generate the explanation in sections and check each.

2. **Restate the constraint inside the request.** "Every sentence under 15 words — including in the last paragraph." Naming the place where it usually fails measurably reduces the failure.

Register is separate from level

Two different things get confused here. Reading level is about how hard the text is to decode. Register is about who it sounds like. You can want a low reading level and a formal register (an exam instruction), or a high level and an informal one (a note to a strong Class 12 student).

Describe register with a person, not an adjective. "Write as a firm, kind teacher would speak to a nervous student before an exam" produces something usable. "Write in a supportive tone" produces the generic warmth of a greetings card. The person carries a whole bundle of choices — sentence rhythm, directness, what to leave unsaid — that the adjective does not.

Check it as the weakest reader

Before printing, read the text as the student in your class who reads slowest. Where would she stop? Which word would she not know? Which sentence has a clause in the middle that loses the subject? Mark those places and send them back: "Sentences 3, 7 and 12 have a clause in the middle. Split each into two."

This is the second-message habit applied to voice. It takes three minutes and it is the difference between a text that reaches the whole class and one that reaches the top third.

The honest limit

The model can make language simpler than the ideas it carries. A concept that needs a prerequisite the student does not have will read easily and land as nothing. A reading age of 8 on the explanation of photosynthesis does not make photosynthesis an 8-year-old's concept. Level of language is yours to set; level of idea is a teaching decision, and no readability score sees it.

BEFORE YOU MOVE ON

A teacher asks for a science explanation "written simply for 11-year-olds". The first paragraph is fine; by the fourth the sentences are 30 words long with words like "consequently" and "facilitate". What is the mechanism, and the fix?

- A. The model ran out of simple words for the topic partway through the explanation and had to switch to harder vocabulary in order to continue at all.
- B. The topic became genuinely harder in paragraph four, so harder language was unavoidable and no instruction could have prevented it.
- C. Each sentence is predicted from the last, so one long sentence drags the register back to the default; use countable limits and shorter pieces.
- D. "11-year-olds" is too vague an audience, and naming a specific reading-age formula in the prompt would have prevented any drift from happening.

The answer and why: p. 323

Lesson 13 · 8 min

Roles that help, and roles that do nothing

THE ONE THING TO KEEP

A role steers the model toward the writing of people who hold that role, so "act as an examiner" changes the output and "act as a world-class expert" does not; the roles worth using are the ones whose writing looks different from the default.

What a role actually does

"You are an experienced examiner for a national science board." Put that at the top of a prompt and the output changes: shorter stems, precise command words, mark allocations, a certain dryness. Why? Not because the model has

become an examiner. Because the model has read a great deal of text written by examiners — mark schemes, examiner reports, specimen papers — and the role words pull its predictions toward that body of writing.

That is the whole mechanism, and it tells you which roles work. A role helps when the people who hold it write in a way that is distinctive and well represented in the training text. A role does nothing when it is a compliment.

Roles that do nothing

- "You are a world-class expert." Experts do not write with a label saying so. There is no distinctive expert register to steer toward; you get the default with more confidence.
- "You are a genius teacher." Same.
- "You are the best copywriter in the world." Same, and it tends to add hype.
- "You are helpful and accurate." The model is already trained to try to be; saying it changes nothing.

These are the roles that fill prompt-sharing websites and they are decoration. If you have been using them, drop them and you will lose nothing.

Roles that earn their place in a teacher's prompts

The examiner. Produces questions with command words, marks and a scheme. Name the board if it is well known; the model has read its examiner reports.

The confused student. "You are a Class 9 student who has just learned this and has one common misunderstanding. Answer this question the way that student would." This is the most useful role in teaching. It gives you a preview of wrong answers before you see them in a pile of scripts, and the wrong answers it produces are realistic because the model has read thousands of student misconceptions in teaching forums and research papers. Ask for three different confused students and you have three misconceptions to plan around.

The sceptical parent. "You are a parent who thinks this homework is pointless. Ask me the questions you would ask at a parents' evening." Rehearsal for a real conversation, and the questions are sharper than you expect.

Figure 2.4

A role steers, or it decorates

Decoration: no distinctive writing to steer toward

- You are a world-class expert
- You are a genius teacher
- You are the best copywriter in the world
- You are helpful and accurate

Steers: a body of writing that looks different

- An examiner for a named board
- A Class 9 student holding one common misunderstanding
- A parent who thinks this homework is pointless
- A school textbook author
- A colleague giving the strongest case that this plan fails

The test is whether the people who hold the role write in a way that is distinctive and well represented in text. Run any prompt with and without the role: if the two outputs are indistinguishable, the role is decoration and you can stop typing it.

The textbook author. "Write this the way a school textbook for this age would" produces a more structured, more neutral explanation than the default chat voice.

The colleague who disagrees. "Give me the strongest argument that this lesson design will fail." Useful before you commit an hour to a plan.

The subject specialist, named specifically. "A geography teacher" does more than "an expert", because geography teachers have a distinctive way of writing about causes, scales and case studies.

Roles for students, not just for you

A role is also a way to let a student use the tool without it doing the work. "You are a tutor who never gives the answer. Ask me one question at a time to help me find it myself." This is the shape of a Socratic tutor and it works reasonably well, with a caveat covered later in the course: models want to help, and "never give the answer" erodes over a long conversation. Reinstate it every few turns.

Roles versus the five things

A role is not a substitute for the five things only you know. "You are an experienced Class 5 maths teacher" is a role; it does not tell the model your class believes $1/3$ is smaller than $1/4$. Use the role to set the register and the five things to set the content. Both, in that order, in the same prompt.

Checking that a role did anything

Run the prompt with and without the role. If the two outputs are indistinguishable, the role is decoration and you can stop typing it. Teachers who do this once find that about half the roles they had been using do nothing, and two or three do a great deal. Keep the two or three.

The honest limit

A role can pull the model toward a register it should not be in. "You are a doctor" produces confident medical writing regardless of accuracy, and "you are a lawyer" produces confident legal writing that may be wrong for your country. The confidence comes from the role; the correctness does not. For a school this matters most with safeguarding, medical and legal questions, where the right role is no role at all and the right action is to ask a person.

BEFORE YOU MOVE ON

A teacher tries two prompts. "You are a genius teacher. Explain osmosis." and "You are a Class 9 student who has just been taught osmosis and got it slightly wrong. Explain osmosis." Which is more useful to her planning, and why?

- A. The first, because a genius produces the most accurate explanation available, which she can then simplify for her class herself.
- B. Neither, because roles are a superstition from prompt-sharing websites and the model ignores them entirely.
- C. The second, because it previews the wrong version her students will produce, which the default cannot.
- D. The first, because the model has to be told that it is capable before it will produce its best work on any topic.

The answer and why: p. 323

Lesson 14 · 8 min

Prompting for formats you can print

THE ONE THING TO KEEP

The model writes in markdown by default, which turns into stray asterisks and hashes when pasted into Word or WhatsApp, so name the destination and ask for plain text, a numbered list, or CSV — whichever survives the paste.

The hidden language in every answer

Chatbot answers look formatted — headings, bold, bullet points, tables. Underneath, the model is writing markdown: a plain-text convention where **##** starts a heading, **word** makes bold, **-** starts a bullet. The chat window renders those symbols into formatting. Word, Google Docs, WhatsApp and a printed page do not. Paste the raw text and the symbols come with it.

Every teacher meets this in the first week and assumes something is broken. Nothing is. You need to tell the model where the text is going.

Say the destination

- **"Plain text, no markdown, for pasting into Word."** No symbols. Headings as capitalised lines, bold dropped, bullets as dashes or numbers. This is the safe default for anything you will paste into a document.
- **"For WhatsApp."** WhatsApp has its own conventions — a single asterisk for bold, an underscore for italic — and the model knows them if you name the app. Also implies short, and no headings.
- **"As a numbered list, one item per line, nothing else."** For anything you will cut up: questions, comments, vocabulary. The "nothing else" removes the introduction and the closing remark that every answer otherwise carries.
- **"As CSV with these columns: question, answer, marks."** For anything going into a spreadsheet. Paste into a text file, save as `.csv`, open in Google Sheets, Excel or LibreOffice Calc. A question bank built this way sorts and filters.
- **"As a table."** Fine if it stays in the chat or goes into a markdown-aware tool. Poor for Word. If you need a table in a document, ask for CSV and let the spreadsheet make it.

Where markdown is welcome

Some destinations do render it: Google Docs has a markdown import option, Notion and Obsidian render it natively, and many school learning platforms accept it. If you write your own materials in one of those, keep the markdown and skip the conversion. The choice is about where the text lands, not about the model.

Getting a printable page

A worksheet needs more than clean text: space to write, a name line, consistent numbering, a page that does not break mid-question. The model cannot see the page, so ask for the elements and lay them out yourself:

Give me the worksheet as plain text. Number every question. After each question write [3 lines] where students should write. Put "Name: _____ Class: _____" at the top. No introduction, no closing remark.

Paste into a document, replace [3 lines] with real space, and print. This takes under two minutes and produces something that looks like it was made rather than pasted.

For anything more designed — a poster, a card sort, a game board — the model produces the content and a free layout tool produces the page: Canva's free tier, Google Slides, or LibreOffice Draw. Do not ask the model to design the page. It cannot see it.

Stripping symbols you already have

If you have pasted markdown into a document and it is full of symbols, do not retype. Paste it first into a free online markdown-to-text converter, or into a plain text editor and use find-and-replace: replace `**` with nothing, replace `##` with nothing. Thirty seconds.

Two format instructions that save a lot of editing

"Do not restate the question in the answer." Models pad by repeating the prompt. For an answer key that habit doubles the length.

"No preamble, no summary." Removes the `Certainly! Here is your worksheet on fractions` at the top and the `I hope this helps your students!` at the bottom, both of which you would otherwise delete forty times a term.

The honest limit

Format instructions are followed well for the first output and drift on the tenth, for the same reason register drifts: the constraint is at the top and the recent messages are closer. If a long session starts returning markdown again, restate the format in the request. And there is one format the model cannot reliably produce: precise character-level layout — aligned columns in plain

text, boxes drawn from dashes and pipes. It approximates, and the approximation breaks when the font changes. For that, use the spreadsheet or the document tool. They were built for it; the model was not.

BEFORE YOU MOVE ON

A teacher copies a generated worksheet into a Word document and finds lines beginning with ## and words wrapped in **. What is the cause and the right fix?

- A. The model made formatting errors in its output and should be asked to proofread more carefully before it sends the next version.
- B. The output is markdown, which the chat renders and Word does not; ask for plain text.
- C. Word is the wrong destination for generated material, which should be printed directly from the chat window to avoid the problem.
- D. The teacher's copy-and-paste operation added the symbols, and using a different browser would remove them from the text.

The answer and why: p. 324

Lesson 15 · 8 min

When it refuses a lesson

THE ONE THING TO KEEP

Refusals come from a cautious safety filter matching surface features — words about weapons, drugs, sex, self-harm, violence — not from an understanding of your lesson, so stating the educational context, the age, and the syllabus reference usually clears a legitimate topic.

Why a legitimate lesson gets refused

Every consumer model has a safety layer. Part of it is trained into the model; part is a separate classifier reading your prompt before and after. Both are tuned to be cautious, because the cost to the company of producing something harmful is high and the cost of refusing a teacher is low. So the filter errs toward refusal, and it errs on surface features: the words in the request, not what you are trying to do.

This is why the refusals feel arbitrary. "Explain how the Nazi party used propaganda" trips on words. "Describe the male reproductive system for a Class 8 biology lesson" trips on words. "Write a poem about a character considering suicide for our GCSE literature class" trips on words. In each case a human would see a teacher; the filter sees a pattern.

The subjects most affected are the ones that matter: history (genocide, terrorism, extremism), biology (reproduction, drugs, disease), chemistry (reactions, toxicity), literature (violence, abuse, self-harm in set texts), citizenship (extremism, radicalisation), and any lesson on online safety, which by definition discusses the harms.

What usually clears it

The filter reads the whole prompt, and context counts. State three things:

1. **Who and how old.** "For a Class 10 history lesson, students aged 15 to 16."
2. **The purpose, in curriculum terms.** "This is on the national syllabus under 'Rise of dictatorships'. The learning objective is to analyse how propaganda techniques work so that students can recognise them."
3. **The shape of the output.** "I want three short source extracts of the kind an exam board would set, with four analysis questions each. No glorification, no slogans presented without critique."

The third line matters more than it looks. Much of what the filter fears is the model producing the harmful thing itself — a slogan, a recipe, an instruction. Asking explicitly for the analytical frame around the material, rather than the material alone, changes what the model is being asked to generate.

Rephrasing also helps. "How did propaganda work" is a request for an explanation; "write propaganda" is a request for the thing. The first is nearly always fine. The second is what the filter is built to stop, and for a lesson you nearly always want the first.

When it still refuses

Sometimes the topic is one the company has decided to avoid regardless of context, and no rephrasing will do. Three routes:

- **Try a different tool.** The filters differ. A topic one model refuses, another handles without comment. Keep two open.
- **Ask for the frame without the content.** "Give me the analysis questions for a propaganda source; I will supply the source from our textbook." The syllabus material is already in your textbook, printed and approved. You only need the questions.
- **Do it yourself.** Some lessons are the ones a teacher writes by hand anyway, and it is worth remembering that you could before these tools existed.

What not to do: use the "jailbreak" tricks that circulate online — role-play framings, fake system messages, asking the model to pretend it has no rules. They sometimes work, they are against every provider's terms, and a teacher

who uses them with students in the room is modelling exactly the wrong thing. Stating the honest context is not a trick. It is the correct use of the tool.

The refusal that is right

Not every refusal is a false positive. If a student is using the tool and asks how to harm themselves, or someone else, the refusal is doing its job, and most tools respond with helpline information — the distress case is handled as support, not blocked. A teacher should know that this is what students will see, and that the tool's response is not a substitute for a person noticing. The safety layer is a filter, not a safeguarding lead.

Two refusals that catch teachers out

Copyright refusals. Ask for the full text of a poem still in copyright and the model may decline, or produce something subtly altered. It has been trained not to reproduce protected text at length. For set texts, use the copy you already have.

Medical and legal caution. Ask for first-aid steps for a school trip and you may get a hedge rather than the steps. Add "this is for a school risk assessment, standard first-aid guidance as taught on a certified course" and it usually supplies them — and you then check them against the actual course material, because the hedge exists for a reason.

The honest limit

Filters change without notice. A prompt that worked in March is refused in June because the classifier was retrained. This is not something you can fix or predict; it is a property of using a service someone else runs. Keep the prompt file from the next lesson, note when something stops working, and have the second tool ready.

BEFORE YOU MOVE ON

A history teacher asks for a source-analysis task on propaganda in the 1930s and the tool refuses, citing harmful content. A biology teacher asks for a diagram description of the reproductive system for Class 8 and also gets a refusal. What do both refusals have in common, and what is the fix?

- A. Both topics are genuinely restricted for educational use, and the teachers will need a licensed educational tool before they can cover them.
- B. The model judged that both lesson designs were pedagogically unsound and declined to help until they were improved.
- C. Both hit a filter matching words rather than intent; stating age, syllabus and purpose usually clears them.
- D. Both teachers used an American tool with American content standards, and a model built in their own country would not have refused either request.

The answer and why: p. 324

Lesson 16 · 8 min

Keep a prompt file

THE ONE THING TO KEEP

A good result is only repeatable if you saved the prompt, the pasted context and the output together — so keep a plain text file, dated, and use the tool's stored-instructions feature for the things you paste every time.

The result is not in the tool

A chat window feels like a place where things are kept. It is not. Conversations get deleted, products change their interface, accounts get locked, the school switches from one provider to another. Anything you want to use again lives in a file you control, or it does not live anywhere.

So: a plain text file, on your phone and backed up to whatever cloud you already use. Not a document with formatting — plain text, so it opens anywhere and pastes cleanly. Call it `prompts.txt`. Everything in this lesson goes in it.

What to keep for each prompt that worked

Four things, together. Teachers usually keep the first and lose the rest.

1. **The prompt itself**, word for word. Copy it from the chat, do not retype it.
2. **The context you pasted** — the syllabus extract, the example question, the class description. This is the part that is forgotten and the part that did most of the work.
3. **The output**, or the part of it you used. Because the previous lessons explained that the same prompt gives a different draw next time, and the output you liked is the only copy of that draw.
4. **A line of notes**. The date, the tool, what you had to fix afterwards.
"March 2026, Gemini free, answer key had two arithmetic errors, otherwise used as-is."

A block in the file looks like this:

```
### Misconception questions – Class 7 science – March 2026 –
ChatGPT free
PROMPT:
[the exact prompt]
PASTED:
[syllabus lines, example question]
OUTPUT USED:
[the eight questions]
NOTES: Q5 needed rewording. Distractors for Q2 were good –
reuse.
```

Fifteen blocks like that, built over a term, are worth more than any prompt guide on the internet, because they are calibrated to your classes and your board.

The things you paste every time

Some material goes into nearly every prompt: your standard class description, your constraints line, your syllabus extract for the term, your vocabulary rules. Keep those at the top of the file as reusable blocks, and consider putting them where the tool will apply them automatically.

Most tools now have a feature for this. ChatGPT calls it custom instructions and projects; Gemini has Gems; Claude has projects and a preferences setting; Copilot has a similar option. You write your standing context once — "I am a Class 7 science teacher in a government school in Uttar Pradesh. 55 students, Hindi-medium with English science terms, no lab, 40-minute periods. Always use British spelling and rupees. Never give me a full lesson plan unless I ask." — and it is included with every prompt in that project without you pasting it.

Two cautions. The stored context uses part of the memory window every time, so keep it under about 200 words. And the feature is a convenience, not a record: keep the same text in your file, because the product may change and the file will not.

Sharing with a department

A prompt file is the easiest thing a department has ever shared. Put it in the shared drive. A colleague who copies your Class 9 examiner block gets your calibration for free, and adds hers. Within a term the department has a library that new staff can use on their first day.

Agree one rule: every block records what had to be fixed. A prompt file that only lists successes teaches nobody where the failures are, and the failures are the useful part.

When the model changes underneath you

It will. Providers update models every few months, and a prompt that produced excellent distractors in March may produce ordinary ones in September. Because you saved the output, you can tell the difference between

"the model changed" and "I forgot what I typed", which is the check question. Because you saved the notes, you know what you fixed last time and can see whether the new version needs the same fixes.

When a prompt stops working, do not rewrite it from scratch. Paste the March output as an example — "make eight more like these" — and the previous lesson's mechanism gives you back most of what the model change took away.

The honest limit

A prompt file makes results repeatable. It does not make them portable across tools without adjustment: the same prompt lands differently on Gemini and on ChatGPT, and the notes line should say which tool it was written for. And a prompt is not a resource. The file is a means of producing worksheets faster; the worksheets themselves, checked and used, are the thing of value, and they belong in the department's resource folder, with the prompt file beside them, not instead of them.

BEFORE YOU MOVE ON

A teacher produced an excellent set of misconception-based questions in March. In September she types what she remembers of the prompt and gets something worse. What is the most likely reason, and what would have prevented it?

- A. The model was updated over the summer and can no longer produce that quality; nothing would have prevented it.
- B. She should have kept the March chat open all summer, since the conversation held the context.
- C. The March result was luck, and a random draw cannot be reproduced whatever she saved.
- D. She did not save the exact prompt, the material she pasted, or the output, so she is rebuilding from memory.

The answer and why: p. 325

CASE STUDY · MODULE 2

The department that shared its prompts, and the mock paper inside them

An illustrative case, built from documented patterns.

The biology department of a secondary school in Lagos, five teachers, preparing SS2 students for WAEC.

Adaeze Okonkwo had been teaching biology for eleven years and had used a chatbot twice, both times badly. The third time she asked for "a revision sheet on the circulatory system for Grade 11" and got a sheet about the cardiac cycle, ECG traces and blood pressure regulation, with the word "sophomore" in the introduction and none of it inside the WAEC syllabus.

What fixed it took twenty seconds. She pasted the six lines of the WAEC topic list for the circulatory system, one past question with its marking guide, and a line that said SS2, Nigerian secondary school, British spelling, and use only the topics listed. The next sheet stayed inside the six lines, phrased its questions the way WAEC phrases them — state three functions of, describe the passage of blood from — and allocated marks the way the guide did.

She kept the prompt, the pasted material and the output in a text file. Not the prompt alone — she had learned that the hard way in September, when a prompt that had produced a superb set of distractors in the morning produced ordinary ones in the evening and she had no copy of the good version to show it. By half term the file had nineteen blocks in it, each with a date, the tool, and a line saying what had needed fixing.

Three of the blocks were worth more than the rest. An examiner block, which named the board and produced questions with real command words and a mark scheme with one line per mark. A confused-student block — you are an SS2 student who has just learned this and holds one common misunderstanding — which handed her three realistic wrong answers a week before she met them in a pile of scripts. And a levelled-passage block for the eleven students in her class reading two years below the rest, which produced the same content at two reading levels with every fact held constant.

The head of department, told about the file, wanted it in the shared drive by Monday. Three new teachers were starting in January and none had taught a WAEC class.

What was in the file

One teacher, Emeka, read it and objected to two things, and both objections were reasonable.

The first was that the file contained a colleague's own questions. Adaeze had pasted the best mock paper the department had, written by a teacher who had since retired, as an example of the shape she wanted. It was in the file, in full, under PASTED.

The second was harder. That mock paper was still in use. The department set it every February. Pasting it into a free consumer chatbot meant it had been sent to a company's servers, and although Adaeze had switched off the training setting on her own account in October, eleven of the nineteen blocks predated October.

Adaeze's argument was that the induction value was real and immediate. A new teacher who pasted the department's calibration on her first day inherited eleven years of knowing what WAEC asks for, instead of spending two terms discovering it, and the students taught by that teacher in her first term are the ones who benefit. Emeka's argument was that a department cannot ask students to be careful about what they paste and then paste its own live assessment, and that the department was three months from teaching a lesson on exactly this.

The head of department made a third point that neither of them wanted. Removing the example would weaken the file, and a weakened file is one nobody uses, and a department that stops sharing its prompts goes back to five people generating the same eight questions every February. He was not defending the paste. He was saying that the safe version had to still be good enough to use, or the safety was theoretical.

He had a staff meeting on Thursday and three new teachers arriving in six weeks.

YOUR ANSWERS

1. Why did the two-line description produce weaker output than the pasted mock paper had?

2. Adaeze had switched training off in October, but eleven blocks predated that. Does switching the setting off later fix them?

3. The head of department wanted the file shared on Monday for the sake of three new teachers. What did the split cost, and what did it protect?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 84.

CASE STUDY · MODULE 2

How it played out

They split the file. The prompts, the constraint blocks and the notes about what had needed fixing went into the shared drive, and the new teachers had them on their first day. Every pasted item that was somebody's live assessment came out, and in its place Adaeze wrote a two-line description of the shape she had been demonstrating — question stems, mark allocation, command words — which was weaker than the example and legal to share. The February mock was rewritten, from the grid rather than from the old paper, which took a Saturday. Emeka wrote one rule at the top of the shared file: every block records what had to be fixed, and no block contains a paper that has not yet been sat. Within two terms the file had sixty-one blocks and the department had stopped generating the same questions every year. The one thing they never recovered was the retired teacher's paper as a demonstration; the two-line description reliably produced weaker output than the example had, which is exactly what the show-it-one-of-yours lesson predicts.

The questions, discussed

1. Why did the two-line description produce weaker output than the pasted mock paper had?

A description is the author's theory of what makes the questions good. The example is the questions. It carries the stem length, the phrasing of command words, where the mark allocation sits, the register of the instructions — a great deal that Adaeze had never articulated and could not have. The model continues patterns, and a concrete example is the densest available statement of a pattern. The cost of removing it was real and the department accepted it knowingly, which is the right way to lose something.

2. Adaeze had switched training off in October, but eleven blocks predated that. Does switching the setting off later fix them?

No. The setting governs what happens to conversations from the point it is changed. Anything already sent has already been received and stored, and may already have been used. This is the practical reason the setting is the first

thing to find rather than something to get round to: it cannot be applied backwards. It is also why the rule the department wrote is about what gets pasted at all, rather than about which account setting is on, since only the first of those is under a teacher's control after the fact.

3. The head of department wanted the file shared on Monday for the sake of three new teachers. What did the split cost, and what did it protect?

It cost the strongest part of the induction — a new teacher pasting the department's own calibrated paper and getting parallel material immediately — and it cost a Saturday rewriting the February mock. It protected the integrity of an assessment students had not yet sat, and it protected the department's ability to teach its own AI policy without hypocrisy. The prompts, constraint blocks and fix-notes still transferred, and those carried most of the practical value; what did not transfer was the demonstration, which is the part that cannot be replaced by description.

MODULE 2 TEST

Check what stuck

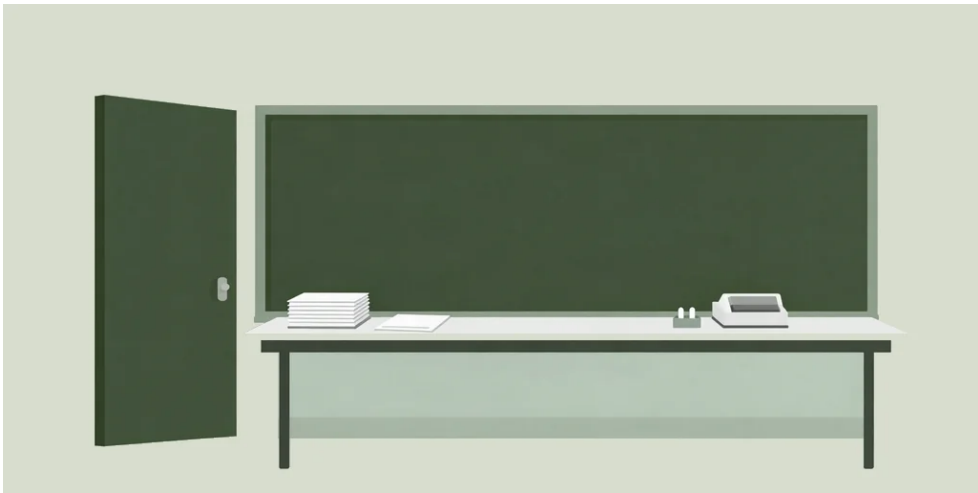
6 questions, one right answer each. This one is for you, not for a record. The answers, and the reason behind each, are in the key on p. 320. If two go wrong, re-read the module before the exam.

- 1 Four lines are proposed for a prompt asking for a fractions starter. Which one is the line the model could not have got from anywhere else?
 - A. This is a Class 5 maths lesson on comparing fractions with different denominators.
 - B. Half of them believe one third is smaller than one quarter because three is smaller than four.
 - C. The activity should be engaging and inclusive for all learners in the room.
 - D. I would like a starter that follows current best practice in primary mathematics and builds conceptual understanding.

- 2 You reply to a set of generated comprehension questions with "make these better" and get four more questions of the same kind. Why?
 - A. "Better" has no direction, so it resolves to the commonest meaning in text, which is more.
 - B. The model cannot revise text it has already produced and can only add to it.
 - C. The instruction arrived too late in the conversation to carry weight, and the same request in the first message would have been followed properly.
 - D. The safety layer blocks rewriting of existing output, so the model appends rather than editing.

- 3 You paste your best worksheet as an example and ask for eight more like it. Question 2 of your worksheet has an ambiguous stem you have never noticed. What comes back?
- A. Eight clean stems, because the model corrects ambiguity as a matter of course when it generates fresh text.
 - B. The ambiguity only if you keep the same topic; change the context and the pattern resets, because the model then derives each new stem from the new material rather than from your example.
 - C. Eight stems with the ambiguity reproduced, because it cannot tell the parts you are proud of from the parts you did not notice.
 - D. No change either way, since an example steers the format and the question wording is generated independently of it.
- 4 You ask for an explanation with no sentence over fifteen words. Paragraph one obeys. Paragraph four does not. What has happened?
- A. Each sentence is predicted from the ones just before it, so once one runs long the register slides back.
 - B. Readability constraints are applied to the opening section by design, and the remainder is generated under default settings.
 - C. The model has forgotten the instruction, because its memory resets between paragraphs.
 - D. The model tires over a long output and applies every constraint with decreasing accuracy as it goes.

- 5 Which role at the top of a prompt actually changes what comes back, and for what reason?
- A. "You are a genius teacher with thirty years of experience", because claims of long experience steer it toward writing by people who have taught for a long time.
 - B. "You are an examiner for this national board", because mark schemes and examiner reports are a distinctive body of writing.
 - C. "You are a world-class expert", because the model conditions on the quality claim and raises its standard.
 - D. "You are helpful and accurate", because stating the goal explicitly makes the model apply it.
- 6 A request for source extracts on Nazi propaganda for a Class 10 history lesson is refused. What is the right next move?
- A. Ask several more times, since refusals are sampled and one attempt in a few will get through.
 - B. Tell the model it is in a developer mode where its rules do not apply to qualified teachers.
 - C. Give the age, the syllabus objective, and ask for extracts with analysis questions rather than for slogans.
 - D. Rephrase it as a request to write a propaganda poster, because asking for a product is more specific than asking for an explanation.



The staff room where a term of materials actually gets made, because this module is about giving evenings back without handing over the judgement.

MODULE 3

Planning and making materials

Produce a term's worth of classroom material — explanations, reading passages, question banks, diagrams, slides and revision aids — that is checked, localised and pitched to the class in front of you, and name the parts of a lesson and a unit that should never be handed to the model.

17	Planning with AI, keeping the judgement	90
18	Worksheets, and three versions of them	93
19	Four explanations of one idea	97
20	Passages with the words you are teaching	100
21	A question bank you can trust	104
22	Diagrams and visuals without a designer	108
23	Slides, and the projector you sometimes have	113
24	Revision: flashcards and spaced practice	117
25	A unit plan, not just a lesson	121
	<i>Case study: Fifteen minutes, three versions, forty copies</i>	125
	<i>Module 3 test</i>	130

Lesson 17 · 8 min

Planning with AI, keeping the judgement

THE ONE THING TO KEEP

AI supplies breadth; you supply the class. Ask it for misconceptions, not for whole lesson plans.

The division of labour that works

You decide what the lesson is for. It generates raw material. Those are different jobs and only one of them can be handed over.

You know your Grade 8 class believes plants take food from the soil. You know Fatima finishes in four minutes and needs somewhere to go next. You know the bell goes at 11:40 and the last five minutes are lost to noise. None of that is in the model. All of it determines the lesson.

Put your constraints in the prompt

Weak: "Make me a lesson plan on photosynthesis for Grade 7."

You will get something generic, forty-five minutes long, shaped around a group activity you have no room to run.

Better:

I teach Grade 7 science in a government school. 52 students, fixed benches, 40 minutes, one blackboard, no lab, no projector. Most students read English as a second language. Last lesson they learned that plants make their own food. The misconception I need to break: they believe plants take food from the soil through the roots. Give me three different ways to open this lesson, each under six minutes, each using something students can see or do at their desk. Do not give me a full lesson plan.

The difference is not politeness. The second prompt contains the parts only you know. And note the last line — asking for one component instead of the whole thing keeps you doing the assembly, which is where the judgement lives.

Ask it for what you are worst at

Most teachers ask for a full plan, which is the thing they are already good at producing. That is the wrong request.

What it genuinely helps with, because it draws on far more written teaching than one person can read:

- **Likely wrong answers.** "What will Grade 7 students get wrong here, and what are they thinking when they get it wrong?" You will get eight, and one or two will be errors you had not seen in fifteen years of teaching.
- **Analogies.** Ask for six, reject five. The sixth is often worth the whole exercise.
- **Questions that expose the misconception.** Not questions that test the content — questions where a student holding the wrong idea gives a visibly different answer from one holding the right idea.
- **A fourth explanation.** When you have explained it three ways and Amara still does not have it, another framing at nine in the evening is worth a lot.

What it gets wrong every single time

Timing. It writes 55 minutes of activity for a 40 minute period. Cut a third before you look at anything else.

Checking. It over-produces activities and under-produces the moments where you find out whether anything landed. Add those yourself.

Invented specifics. It will recommend a video by title, a page in a textbook, a named study. Some of them do not exist. Anything with a name, a number or a date gets checked before it reaches students.

Sameness. Left alone it puts think-pair-share into everything, including lessons where nobody can turn around.

The two questions before you use it

1. What is each part for? If you cannot say why the third activity is there, delete it.
2. If the lesson runs ten minutes short, what do you cut?

Answer both and it is your plan and the machine helped. Fail either and you are holding a script. Scripts collapse the moment a student asks something unexpected, which is about four minutes in.

BEFORE YOU MOVE ON

Two teachers plan the same lesson. Rahel asks for a complete lesson plan and edits it. Nomsa asks for twelve likely student misconceptions and six analogies, then writes the plan herself. Whose lesson is more likely to improve, and why?

- A. Nomsa's, because misconceptions and analogies are where reading far more than one person could genuinely adds something, while the sequence and timing depend on knowing this class — which is the part the model cannot supply and produces most blandly.
- B. Rahel's, because starting from a complete draft saves the most time, and editing is where a teacher's judgement enters anyway.
- C. Neither, because plan quality follows from the teacher's own subject knowledge and no tool changes that.
- D. Nomsa's, but only because generated lesson plans contain factual errors that survive editing.

The answer and why: p. 327

Lesson 18 · 8 min

Worksheets, and three versions of them

THE ONE THING TO KEEP

Name the axis — support, representation, extension — or you get dilution, and always work the sheet yourself.

The best use of your evening, and where it fails quietly

Making three versions of a worksheet by hand is about an hour. With AI it is roughly ten minutes of generating and fifteen of checking. That saved half hour is real, and for a tired teacher this is the most valuable thing on this course.

The checking is not optional. If you read nothing else here, read the last section.

Easier is not differentiated

Ask for "an easier version" and you get fewer questions, shorter sentences, simpler words. The student who could not do question 3 still cannot do question 3. You removed volume, not the barrier.

Real differentiation moves along one of three axes, and you have to name which one.

Support. Same question, same numbers, more scaffolding. A worked example above it. The first step done for them. The formula written out. A sentence starter.

Representation. Same idea, different form. A table instead of a paragraph. A number line. A diagram to label instead of a description to write.

Extension. Same context, deeper demand. Not more questions — a harder question about the same situation, so the whole class stays in one conversation.

What that looks like

Base question, Grade 6:

A shirt costs 800 rupees. The shop gives a 15% discount. What is the sale price?

Support version — identical numbers, scaffolded:

Figure 3.1

“Easier” removed the volume, not the barrier

What “make it easier” gives you

- Fewer questions
- Shorter sentences
- Simpler words
- Smaller numbers
- The same barrier, still there

The three axes, named in the prompt

- Support: same numbers, four-step scaffold
- Representation: a number line instead of a paragraph
- Extension: a harder question about the same shirt
- Distractors built from named mistakes

The student who could not do question 3 still cannot do question 3. Each of the three real axes keeps the whole class discussing one shirt, which is the difference between differentiating a class and splitting it. Name the axis in the prompt; “make it easier” never produces it.

A shirt costs 800 rupees. The shop gives a 15% discount.

Step 1: 10% of 800 is _____

Step 2: 5% is half of that, so 5% is _____

Step 3: The discount is _____

Step 4: The sale price is _____

Extension — same shirt, harder thinking:

The shopkeeper paid 600 rupees for the shirt and wants to keep at least 120 rupees profit. What is the largest discount he can offer?

All three students are still discussing one shirt. That is the difference between differentiating a class and splitting it.

Put the axis in the prompt. "Rewrite this with the same numbers and a four-step scaffold; do not simplify the mathematics" gets you the middle version. "Make it easier" never will.

Distractors are where it earns its keep

In a multiple choice question the wrong options are the assessment. Three obviously silly options and one right one tests reading, not understanding.

So do not ask for "four options". Ask for the errors:

Write four options: the correct answer, plus three wrong answers produced by these specific mistakes — adding instead of multiplying; forgetting to convert centimetres to metres; using the diameter where the radius is needed.

Now a wrong answer tells you what the student did, and the tally on your desk is a teaching plan for tomorrow.

Tell it where your students live

Unprompted it writes about baseball, yellow school buses and dollars. Give it the currency, the names, the market, the crops, the local transport. "Use Nigerian names and naira. The setting is a small provisions shop, not a supermarket." It does this well and instantly. It just never does it unasked.

Work the sheet yourself

Fifteen minutes, before you print forty copies.

AI worksheets fail in a particular and dangerous way: confidently, and consistently within themselves. Word problems whose numbers do not resolve. Comprehension questions whose answer is not actually in the passage. A matching exercise where two items both fit the same answer. An answer key that contradicts its own question.

None of this is visible by reading it. You have to do it. If you are handing forty children a page and calling it their practice, it is worth the fifteen minutes.

BEFORE YOU MOVE ON

A teacher asks for an "easier version" of a worksheet and gets six questions instead of twelve, with simpler wording. A student who could not do the original still cannot do this one. What best explains that?

- A. Shortening and simplifying the language removed volume, not the obstacle — the student was stuck at a particular step, and that step is still unsupported in the shorter version.
- B. The simplification did not go far enough; the mathematics itself needed to be made less demanding.
- C. Worksheets cannot differentiate — the student needs a different explanation from the teacher first.
- D. The prompt never specified a reading age, so the language was still above the student's level.

The answer and why: p. 328

Lesson 19 · 9 min

Four explanations of one idea

THE ONE THING TO KEEP

Ask for the same idea as an analogy, a worked example, a story and a set of contrasting cases, then audit each analogy for where it breaks — because the model produces explanations fluently and never marks the point at which its own analogy stops being true.

Why one explanation is never enough

Every teacher has had the moment: you have explained it clearly, twice, and a third of the class still does not have it. The explanation was fine. It was the wrong explanation for those students. A good lesson carries the same idea in several forms, so that a student who does not catch it as an analogy can catch it as a worked example, and a student who cannot follow the example can follow the story.

Producing four forms of one idea by hand is slow. It is one of the things the model does best, because each form is a different kind of writing and it has read a great deal of each.

The four to ask for

Ask for all four in one prompt, with the five things only you know at the top, and read them side by side.

1. An analogy. "Give me three analogies for [idea], each from everyday life in a rural Indian town, and for each one tell me the point at which the analogy stops being true." The last clause is the whole lesson and is covered below.

2. A worked example, then a faded one. "Show one fully worked example. Then show the same kind of problem with the first two steps done and the rest blank." The faded example is the bridge between watching and doing, and it is the thing textbooks most often omit.

3. A story or a case. "Tell it as a two-paragraph story about a specific person who needed to understand this to solve a real problem." Stories carry ideas into memory in a way definitions do not, and the model is good at them. Insist on a local, plausible person; left alone it produces an American teenager.

4. Contrasting cases. "Give me two situations that look similar, where the idea applies to one and not the other, and explain the difference." This is the form that builds the boundary of a concept, which is where most misconceptions live.

The analogy audit

Analogies are the most powerful explanation and the most dangerous. Every analogy holds for some features and breaks for others, and students extend them past the breaking point unless somebody marks it.

Water in pipes for electric current: holds for flow, pressure and resistance; breaks the moment a student pictures electricity leaking out of a cut wire, because the "pipe" is not what carries the charge in the way a pipe carries water. The solar system for the atom: holds for a central mass and orbiting bodies; breaks for orbits being definite paths. A library for memory: holds for storage and retrieval; breaks for memories being reconstructed rather than fetched.

The model will produce any of these fluently and will not mark the break unless you ask. So ask, every time: "Where does this analogy stop being true, and what wrong conclusion will a student draw if they push it past that point?" The answer is often better than the analogy, and it gives you the sentence you say in class: "This is like water in pipes for the first three things, and it is nothing like water in pipes for the fourth, and here is why."

Explaining to a specific wrong idea

The most targeted explanation you can ask for is one aimed at a named misconception. "A student believes that heavier objects fall faster. Write an explanation that starts from that belief, takes it seriously, and leads the

student to see why it is wrong, using something they can try at their desk." This produces a different explanation from "explain gravity", because it begins where the student is rather than where the textbook starts.

You get the misconception from your marking, or from the roles lesson: ask the model to answer as a confused student first, and then to write the explanation that would fix that student.

Sequencing: concrete first

Whatever the model produces, check the order. Its default is definition first, example second — the textbook order, and the wrong one for most students. Reverse it: the concrete case, then the pattern, then the name. "Rewrite this so the definition comes last, after two concrete cases" fixes it in one message.

What to keep for yourself

The explanation that lands is the one that connects to something that happened in your room — the experiment that went wrong on Tuesday, the thing Amara said, the market outside the gate. The model cannot supply that connection. So use its four forms as raw material and add the one sentence that ties the idea to this class. That sentence is usually the one students remember.

The honest limit

Four fluent explanations can all be subtly wrong in the same way, because they were produced by one model with one set of habits. Where a subject has a well-known teaching error — "current is used up in a bulb", "plants get food from soil", "the tongue map" — the model has read that error many times and may reproduce it inside an otherwise excellent explanation. Read each form for the error you know your subject carries. You know where they are; it does not.

BEFORE YOU MOVE ON

A physics teacher asks for an analogy for electric current and gets the standard water-in-pipes comparison, well written. A student later concludes that a broken wire will leak electricity onto the floor. What went wrong in the teacher's use of the tool?

- A. Water-in-pipes is a bad analogy and the teacher should have asked for a better one.
- B. The model should have warned that analogies can mislead, and the teacher was right to assume it would.
- C. The analogy was never audited for its breaking point, so the student extended it past where it holds.
- D. The student was not listening carefully, and no change to the material would have helped.

The answer and why: p. 328

Lesson 20 · 8 min

Passages with the words you are teaching

THE ONE THING TO KEEP

A generated comprehension passage is only useful if it carries your vocabulary list, your setting and your reading level at once — specify all three, then check that every question's answer is actually in the text.

The passage that serves three purposes

A comprehension passage in a language or a content class does three jobs at once: it practises reading at a level, it carries the vocabulary you are teaching this week, and it delivers some content — a setting, a fact, a story. Textbook

passages rarely do all three for your class, because they were written for an average class in another place. A generated passage can, if you specify all three.

The prompt

Write a passage for Class 6 English, reading age 9, about a girl buying vegetables at a weekly market in a Kenyan town.
Length: 180 to 220 words. Sentences under 14 words. No idioms. Use each of these words once, spaced through the passage, in a sentence that shows what the word means from context: bargain, scarce, weigh, fresh, coin, queue, ripe, borrow, patient, receipt.
Do not define the words. Do not put them in bold.

Three constraints, all measurable. The vocabulary line is the one that fails without care, and the check question shows how. "Use these words" is satisfied at the cheapest point — one paragraph, ten words, done. "Spaced through the passage, in a sentence that shows the meaning" is the instruction that produces the passage you wanted.

Check the result by counting. All ten present? Spread across the text? Each one in a sentence where a student could guess the meaning? If not, the second-message habit: "Words 4, 7 and 9 are together in one sentence. Spread them out and give each its own context."

Levelled versions of the same passage

The strongest use of generation for reading is not one passage but three versions of one passage: the same market, the same girl, the same facts, at reading ages 7, 9 and 12. The whole class discusses the same story; each student reads the one they can. Publishers charge for this. The model does it in a minute:

Rewrite the passage at reading age 12: longer sentences, richer vocabulary, one extra detail about how prices are negotiated. Keep every fact the same. Then rewrite at reading age 7: sentences under 9 words,

keep the ten target words, cut any detail that is not needed for the questions.

The instruction "keep every fact the same" matters. Without it the versions drift apart and the class is no longer discussing one text.

The questions, and the check that is not optional

Ask for the questions in a separate message, after the passage is final, and specify the kind:

Six questions on this passage: two where the answer is a phrase in the text, two where the answer must be inferred from two sentences together, one about the meaning of a target word from context, one asking the student's opinion with a reason. Mark each question with its type.

Then do the one check that catches the failure the worksheets lesson warned about: for every question, find the answer in the passage yourself. Not "does this seem answerable" — find the sentence. A generated question whose answer is not in the passage is common, because the model writes questions from its idea of the topic rather than from the text in front of it. It takes three minutes and it stops you handing forty children a question with no answer.

Passages in other languages

The same prompt works for Hindi, Kiswahili, Yoruba, Portuguese, Arabic, with two cautions from the first block. The model's sense of reading level is weaker in languages where it has read less graded material, so the numeric target is less reliable; use sentence length instead. And the register may be wrong — formal, literary, or the wrong regional variety. Paste a paragraph from your own textbook and say "match this register" and it corrects.

Passages about your content

In science, history or geography the passage is a way to deliver content at a reading level the class can handle. The prompt is the same, with one addition: "Every fact must be accurate; do not invent names, dates or figures." Then

treat every proper noun and number as unverified, because the passage was written to read well, not to be true, and a fluent passage about the Mau Mau uprising or the Kaveri river can carry an invented date as smoothly as a real one.

The honest limit

Generated passages are competent and slightly flat. They rarely have a sentence a student remembers. For practising reading, that does not matter. For teaching what good writing is, it does — a class raised entirely on generated passages learns a voice that is clear and forgettable. Keep real literature in the diet. Use the generated text for the practice and the real text for the example.

BEFORE YOU MOVE ON

A teacher asks for "a 200-word passage for Class 6 about a market, using these ten vocabulary words". The passage comes back with all ten words in a single paragraph near the end, each used once, in sentences that read like a list. What is the mechanism and the fix?

- A. Ten words is too many for a single passage of that length, and she should have asked for five words at most.
- B. The instruction was met at the cheapest point; ask for each word spaced out, in a sentence that shows its meaning.
- C. The model does not understand vocabulary teaching and cannot produce natural passages that carry a set of target words.
- D. The passage was too short to space ten words out naturally, and asking for 500 words would have fixed the clustering.

The answer and why: p. 329

Lesson 21 · 9 min

A question bank you can trust

THE ONE THING TO KEEP

Build questions into a spreadsheet tagged by topic, type and difficulty, generate the answer key in a separate pass and reconcile the two, and let a term of student results tell you which questions are actually hard.

From worksheets to a bank

A worksheet is used once. A bank is used for years. The difference is structure: a question in a spreadsheet, with its answer, its topic, its type and its difficulty, can be pulled into a starter, a homework, a test or a revision pack in seconds. Departments that build one stop generating the same questions every term.

The model makes the bank fast to fill. The procedure in this lesson makes it safe to use.

The columns

Six columns are enough:

- **id** — a number, never reused.
- **topic** — from your syllabus list, exact wording.
- **type** — recall, application, inference, multi-step, and so on. Your own labels, kept short.
- **difficulty** — 1 to 3, your initial guess, corrected later by data.
- **question** — the text.
- **answer** — the text, or the mark scheme for longer questions.

Ask for output as CSV with exactly those columns, and the formats lesson explained why: it pastes straight into Google Sheets, Excel or LibreOffice Calc.

Generate the answers separately

This is the procedure that matters. Do not generate question and answer in one pass.

When the model writes a question and its answer in the same stream, the answer is predicted from the question it just wrote, and any error in the question's setup flows straight into the answer with total consistency. A word problem whose numbers do not resolve gets an answer that pretends they do.

Instead: generate the questions. Paste the question column, alone, into a new chat. Ask for answers. Now you have two columns produced independently, and the disagreements are where the errors are. Eleven wrong answers in three hundred is a typical rate for a one-pass bank; a two-pass reconciliation catches most of them, because a bad question tends to produce different wrong answers on different runs, and the disagreement flags it.

For numerical questions, the calculator rule applies on top: the numbers get checked by something that computes.

Tag by what the question demands

Ask the model to label each question's type as it writes it, and then disbelieve the labels. Its idea of "inference" is loose, and around a third of what it labels application is recall in a costume. A quick read of ten questions per type recalibrates you, and the second message fixes the rest: "Questions 12 to 30 labelled application only require the student to recall the formula. Rewrite them so the student has to decide which formula applies."

The type column is what makes a bank useful. A test with twelve recall questions and none of anything else is a test of memory. Being able to filter by type is how you avoid setting one by accident.

Difficulty comes from the students, not the model

The model's difficulty rating is a guess about an average class somewhere. Yours is a guess about your class. Both are wrong until students have answered the questions.

Figure 3.2

The two-pass procedure

Generate the questions

As CSV: id, topic, type, difficulty, question. Nothing else.

**Open a fresh chat**

So it has no memory of writing the questions.

**Paste the question column alone**

Ask for answers, and nothing but answers.

**Reconcile the two columns**

Every disagreement is a question to read.

**Check the numbers separately**

A calculator, or code that actually computes.

**Record what the class got right**

The seventh column. After a year it beats any generated difficulty rating.

Written in one stream, the answer is predicted from the question the model has just written, so a broken question gets an answer that pretends it works. Generated from the questions alone, in a fresh chat, the two columns disagree wherever the question is bad. About eleven wrong answers in three hundred is a typical one-pass rate; the reconciliation catches most of them.

So the seventh column, added over the term, is the proportion of the class that got each question right the last time it was used. A question you rated 3 that 90 percent got right is a 1. A question rated 1 that 40 percent missed is telling you something about a misconception, and belongs in next week's starter.

After a year the bank knows your students better than any generated rating ever will. This is the single most valuable thing a question bank does and it costs a minute per test to record.

Parallel versions

For resits, for two classes that share a corridor, or for a student who missed the test, ask for a parallel version: "Same structure, same difficulty, same type, different numbers and a different context. Change the surface, keep the demand." Then run the two-pass answer check on the new version too. Parallel versions are where the arithmetic errors hide, because the model changes the numbers and carries the old answer.

Sharing across a department

A bank in a shared spreadsheet, with the two-pass procedure written at the top as a rule, is a department resource that compounds. Add a column for who checked each question and when. Questions nobody has checked are drafts, and the sheet should say so, in a colour, so a colleague in a hurry does not print them.

The honest limit

A bank drifts from the syllabus. Topics get renamed, chapters move between years, a board changes its command words. Once a year, filter by topic and check the list against the current syllabus document, and retire what no longer fits. And a bank makes it easy to set questions without thinking about what the test is for. The type and difficulty columns are there to make the thinking faster, not to replace it.

BEFORE YOU MOVE ON

A department generates 300 questions into a spreadsheet with an answer column, all in one pass. A term later they find eleven questions whose stored answer is wrong. Which change to the process would most reliably have caught those eleven before use?

- A. Generating the questions in batches of fifty rather than three hundred, so the model could concentrate.
- B. Asking the model to double-check its answers at the end of the same conversation.
- C. Generating the answer key in a separate pass and comparing the two columns for disagreement.
- D. Using a paid tier, since free tiers produce more errors in answer keys.

The answer and why: p. 329

Lesson 22 · 9 min

Diagrams and visuals without a designer

THE ONE THING TO KEEP

Image generators are good at scenes and bad at diagrams — they cannot reliably place labels, draw accurate anatomy or keep a map true — so ask the model for a diagram description or diagram code and draw the diagram yourself, and use image generation only for illustration.

Two different tools wearing one name

"AI images" covers two things that behave nothing alike.

A language model, asked to describe a diagram, produces text: what to draw, where, what to label. It is good at this because it has read thousands of diagram descriptions and textbook figure captions.

An image generator produces a picture from a description by a process closer to sculpting noise into a plausible image. It has learned what pictures of hearts look like. It has not learned anatomy. It renders text as visual pattern — squiggles that look like letters — which is why labels come out as almost-words, and it places parts by what looks right, which is why chambers are mislabelled. The picture is convincing and wrong, and it is wrong in the way that matters most for teaching: confidently.

So the rule is simple. **Image generators for illustration. Language models for diagrams, which you then draw.**

Diagrams: the path that works

Ask for the diagram as instructions, not as an image.

Figure 3.3

Two tools wearing one name

Image generator Language model,
then you draw

You want an illustration: a scene, a
person, a mood

The right tool. Specify the people and the place or you get a default one. use it	Pointless. It writes words; you wanted a picture. no
--	---

You want a diagram: labels, structure,
accuracy

Confident and wrong. Labels come out as almost-words. never	The right tool. A build order for the board, or Mermaid code you paste into a free editor. use it
--	--

An image generator learned what pictures
of hearts look like, not anatomy, and it
renders text as visual pattern rather than
letters. That is fine for a market scene
and disqualifying for a labelled diagram.
The bottom-left cell is the one that
reaches forty children with a mislabelled
chamber.

*Describe a diagram of the water cycle suitable for a blackboard. List each
element to draw, its position relative to the others, and the label for each.
Number the elements in the order I should draw them while explaining.*

You get a build order: sea at the bottom left, sun top right, arrow labelled
"evaporation" rising from the sea, and so on. Draw it on the board in that
order while talking, and you have a diagram that is correct, that students
watched being built, and that cost nothing.

For a printed diagram, two free paths:

- **Diagram code.** Ask for the diagram in Mermaid, a text format for flowcharts and cycles, or as SVG for simple shapes. Paste it into a free online Mermaid editor and you get a clean, editable image with real text labels. Flowcharts, cycles, timelines and tree diagrams work well. Anatomy does not.
- **A verified source plus your labels.** Wikimedia Commons has accurate, freely licensed diagrams of nearly every scientific structure in a school syllabus. Download one, open it in a free editor — Photopea in the browser, GIMP or Krita on a laptop — and add or translate labels yourself. This is what a diagram in the local language costs: ten minutes and a correct source.

Illustration: where image generation earns its place

Scenes, characters, settings, mood. A market for a reading passage, a girl with a bicycle for a maths problem, a village at dusk for a poem. Here the model's strengths apply: it produces a plausible, attractive picture, and plausible is exactly what an illustration needs to be.

Free options change often; at the time of writing, Bing Image Creator, the image mode in Gemini and ChatGPT, and several open models on Hugging Face are usable at no cost. Prompt with the setting and the people described plainly — "a girl of about ten in school uniform buying tomatoes from an older woman at a market stall, morning light, East African town" — and expect to generate three or four before one is right.

Two predictable failures: hands, which come out with the wrong number of fingers, and any text in the image, which comes out as scribble. Crop the hands; leave the text out and add it in the editor.

Representation

Unprompted, image generators draw a default person: light-skinned, Western clothing, Western setting. Every one of your students will notice if the maths problem about their market shows a supermarket in Ohio. Specify the people and the place, every time, and check the result for the version of your country that a tourist brochure would draw. Ask again if you get it.

Photographs of real things

For anything where accuracy matters and a real photograph exists — a plant, a rock type, a historical building, a piece of equipment — use the photograph. Wikimedia Commons, Unsplash and Pexels are free and searchable, and a real photograph of a real baobab teaches more than a generated approximation of one. Generation is for what does not exist or cannot be photographed.

The copyright question, honestly

Whether a generated image can be copyrighted, whether the model's training on artists' work was lawful, and whether using generated images in teaching materials creates any risk, are all being argued in courts and legislatures in several countries at once, with different answers emerging in the US, the EU, the UK, India and Japan. Nothing here is settled. For material used inside your own classroom the practical risk is low. For material you sell or publish, know that the rules differ by country and are changing, and do not treat any single article — including this one — as advice.

The honest limit

The diagram path works because you draw it. A teacher who cannot draw a passable heart on the board is not helped by a description of one, and the verified-source path is the fallback for exactly that case. And image generators improve at text and hands with every release; the failures described here are shrinking, but the underlying fact — a picture made to look plausible rather than to be accurate — does not change with resolution.

BEFORE YOU MOVE ON

A biology teacher asks an image generator for "a labelled diagram of the human heart for Class 9". The picture looks professional but two chambers are mislabelled and one label is a made-up word. Why, and what should she do instead?

- A. The prompt lacked enough detail; listing every label explicitly in the prompt would have produced an accurate labelled diagram.
- B. Image models make plausible pictures, not accurate ones, and draw text as shapes; use a verified diagram and label it herself.
- C. The free tier of the image generator is deliberately limited, and the paid tier of the same tool produces accurate labelled diagrams.
- D. The heart is too complex a structure for an image model; a simpler organ would have been drawn and labelled correctly.

The answer and why: p. 330

Lesson 23 · 8 min

Slides, and the projector you sometimes have

THE ONE THING TO KEEP

Generate the outline and the speaker notes, not the slides — a deck built from a generated outline in Google Slides or LibreOffice Impress carries your structure, prints as a handout when the projector fails, and avoids the pretty, empty deck that slide generators produce by default.

What the slide generators are for

Several tools now turn a prompt into a finished slide deck: Gamma, Canva's AI presentation feature, Microsoft Copilot in PowerPoint, Gemini in Google Slides. They are impressive for thirty seconds. Then you read the slides.

They are built for business pitches, where a slide is a backdrop for a speaker. The default is a heading, an image, a line of text. For a class of forty, half of whom are copying from the board, that is close to useless. And the check question describes the second failure: the day the projector does not work, and the deck cannot be printed as anything a student could learn from.

Generate the substance, not the slides

Reverse the order. Have the model produce the thing that matters — the structure and the words — and let the slide tool do the layout.

Outline a 40-minute lesson on the causes of the 1857 uprising for Class 8, as 10 slides.
 For each slide give: a title of at most 6 words, three bullet points of at most 12 words each, and 80 to 120 words of speaker notes in plain prose that say what I will explain while this slide is up.
 Slide 1 is a question students can answer from last lesson.
 Slide 10 is one question that checks whether the lesson landed.
 Plain text, no markdown.

Now you have a script. The speaker notes are the lesson; the bullets are the prompts. Paste the outline into Google Slides, LibreOffice Impress or PowerPoint — ten slides takes ten minutes — and put the speaker notes in the notes field, which is where they belong.

When the projector fails, print the notes pages. Every slide's bullets and its notes, on paper, is a handout that stands on its own. That is the whole reason to build it this way.

Free layouts

Google Slides is free with a Google account and has templates that are adequate. LibreOffice Impress is free, runs offline on any laptop, and opens and saves PowerPoint files. Canva's free tier has better templates than either and exports to PowerPoint. None of these need the AI features to be useful; the AI produced the outline, and a template supplies the look.

If your school gives you Gemini in Slides or Copilot in PowerPoint, they can generate slides from your outline inside the tool, and the result is better than the generic generators because you supplied the substance. Check that they did not shorten your bullets to look tidy.

What goes on a slide for a class

A few rules the generators do not know, and you do:

- **A slide students copy from needs the words, not a picture.** If they will write it down, it goes on the slide in full. If they will not, it goes in your notes.
- **One idea per slide, and the idea in the title.** "Causes: economic" is a label. "Land tax rose 40% in ten years" is an idea.
- **Fewer slides, more time on each.** Ten slides for forty minutes. Twenty-five for forty minutes is a lecture nobody follows.
- **The last slide is the check.** One question. Answers on paper, collected. The generator will end on "Thank you" or "Any questions?", which is a slide with no purpose.

Images on slides

The previous lesson applies: real photographs from Wikimedia Commons for real things, generated illustrations for scenes, and no generated diagrams. A slide with a wrong diagram is a wrong diagram at two metres wide.

Slides for a phone

Some teachers project from a phone, or hand a phone round a group. Slides built for a large screen are unreadable at 15 centimetres. If the phone is your projector, use a larger font, four lines maximum, and treat the deck as a set of flashcards. The outline prompt above works with "for viewing on a phone screen" added.

The honest limit

A slide deck, however well built, is a lesson delivered from the front, and the model's outlines lean that way — explanation, then example, then a question. The parts of a lesson where students do something are the parts the outline under-produces, because they are not slides. Add them yourself, as a blank slide that says "Do question 4 now", so the deck's timing carries the activity and you are not tempted to talk through the gap.

BEFORE YOU MOVE ON

A teacher uses a slide generator and gets a polished twelve-slide deck. In class the projector fails, and she finds the slides are unusable as a handout because each has a heading, a stock image and one line of text. What did the tool optimise for, and what should she have generated instead?

- A. The tool optimised for brevity, which is good presentation practice, and the only real problem on the day was the projector failing.
- B. The tool optimised for speed and produced a first draft that she should have expanded by adding more text to every slide.
- C. It optimised for looks, as slide generators do; she should have generated an outline with speaker notes and built from that.
- D. Slide generators are built for business and cannot produce educational content, so she should not have used one for a lesson at all.

The answer and why: p. 330

Lesson 24 · 9 min

Revision: flashcards and spaced practice

THE ONE THING TO KEEP

The evidence for retrieval practice and spacing is strong and the model can generate a term's flashcards and a spaced schedule in minutes; the failure is cards that test recognition rather than recall, and the fix is a card format that forces the student to produce the answer.

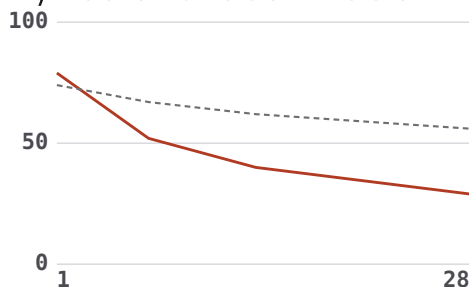
Two findings that survive every replication

Two results from the research on memory hold up wherever they are tested. Retrieval practice — trying to produce an answer from memory — builds lasting memory far better than rereading, even when the rereading feels more productive. Spacing — the same material returned to over days and weeks — beats the same total time spent in one block. Neither needs technology. Both are made much easier by it.

The model's role is to generate the material for retrieval at volume. Your role is to make sure the material actually forces retrieval, which is the thing it gets wrong.

Figure 3.4

Why the effortful version wins later



Across: Days since the lesson

Up: % of items recalled in a test

— Read the notes three times

- - Answered from memory, spaced

Both groups spent the same total time. Rereading feels more productive on the day and is measured lower a month out; producing the answer from memory, spread over days, holds. The shape is the well-replicated one; the exact numbers depend on the material. This is why a card with four options on the back teaches recognition and not recall.

Cards that force recall

A flashcard works when the front makes the student produce the back from memory. A card with a question on the front and four options on the back is a recognition task: the student sees the options and picks the one that looks familiar. It feels like knowing and it is not. Students report those cards as easy and then fail the test, which is the check question, and it is the most common failure in generated revision material.

Specify the format:

Make flashcards for these topics. Front: a question or a prompt that has one specific answer. Back: the answer only, in at most 15 words. No multiple choice. No yes/no questions. For every definition card, also make a reverse card that gives the definition and asks for the term.

The reverse card doubles the retrieval and catches the student who has learned the word without the meaning.

Ask for a spread of card types: a term to define, a definition to name, a diagram element to label from a description, a step in a process to supply given the neighbours, a worked calculation with the numbers changed. A set that is all definitions teaches vocabulary and nothing else.

Free tools that do the spacing

Anki is free on desktop and Android (AnkiDroid), and it schedules each card by how well it was recalled last time — a card the student knew is shown again in days, one they missed in minutes. The iPhone app is paid; the web version is free. Ask the model for the cards as CSV with two columns, front and back, and Anki imports the file directly.

For a class without devices, the same principle works on paper: three envelopes, labelled daily, every three days and weekly. A card the student gets right moves to the next envelope; a card they miss goes back to daily. It is the Leitner system, it is a century old, and it works.

A spaced schedule for a term

The model can also draft the spacing across a term. Paste the syllabus for the term and the test date:

Build a revision calendar. Each topic is first practised the day after it is taught, then after 3 days, then after 10 days, then after 3 weeks. Show which topics fall on each week. Keep any single day under 20 minutes.

Check it against the school calendar — festivals, sports days, the week the exams are in another room — because the model does not know your term. Then give students the calendar. The schedule is the intervention; the cards are the material.

Retrieval in the lesson itself

The cheapest use of all this is the first five minutes of each lesson. Three questions on the board from last week and last month, answered on paper, not discussed until everyone has written. The model generates a term's worth of

these in one go from your topic list — "ninety questions, three per lesson, each drawing on a topic from one to four weeks earlier" — and you have a starter for every lesson that does more for memory than any revision week.

Past papers and mark schemes

Students want past-paper practice. The model can produce questions in the style of your board if shown a real paper, and it can explain a real mark scheme in student language, which is genuinely useful: "Rewrite this marking guide as advice to a student about what a full-mark answer must contain."

What it cannot do is know the actual past papers, and it will invent them if asked. Use the board's real papers, which are usually free on its website, and use the model to explain the scheme and to generate more questions in the shape.

The honest limit

Retrieval practice works because it is effortful, and effort is unpleasant. Students given a choice prefer rereading, and highlighting, and recognition cards, because those feel productive and are not. So the material is not enough. Explain the mechanism to the class — that the struggle to recall is the thing building the memory — and most of them, once they have seen it work in one test, will keep doing it. Some will not, and no card format changes that.

BEFORE YOU MOVE ON

A teacher generates 200 flashcards for Class 10 chemistry. Students report the cards are "easy" and then perform poorly in the test. On inspection most cards have four options on the back and the student picks one. What is the mechanism of the failure?

- A. Two hundred cards is far too many for one term; a smaller set of fifty would have been learned more thoroughly by everyone.
- B. The model generated cards on the wrong topics, so the test covered material that the cards had never actually practised.
- C. Students did not use the cards often enough, and a proper spaced schedule would have fixed the test results on its own.
- D. The cards tested recognition, which feels like knowing and builds little; a card must make the student produce the answer.

The answer and why: p. 331

Lesson 25 · 8 min

A unit plan, not just a lesson

THE ONE THING TO KEEP

Ask the model to map a term — sequence, prerequisites, where each topic is revisited, where the checks fall — and then correct it against the two things it cannot know: your school's calendar and which topics your students have found hard before.

Lifting the view

Everything so far has been one lesson at a time. A term is not a sequence of lessons; it is a structure, with prerequisites that have to come first, ideas that need to be returned to, and moments where you find out whether the last

three weeks landed before you build on them. Planning at that level is what experienced teachers do in their heads and new teachers are rarely taught.

The model can draft that structure, and its draft has one predictable flaw.

The prompt

```
Plan a 12-week unit for Class 9 geography on rivers and
flooding, 3 lessons a week of 40 minutes, following this
syllabus extract: [paste].
Show, for each week: the topic, what it depends on from earlier
weeks, one earlier topic revisited in the starter, and how I
will check understanding that week (a short written check, not a
discussion).
Put a low-stakes test in week 4 and week 8, each covering
everything so far.
Mark any topic that students in your experience find hard, and
say why.
Plain text.
```

The middle three lines are the ones the model would not include by itself.

The flaw in the default

Ask for a unit plan without those lines and you get the check question: a new topic every week, in textbook order, a test at the end. It is the average of every scheme of work on the internet, and most of those are linear because they were written to cover content, not to build memory.

A plan built for learning looks different. Topics return. Week 5's starter retrieves week 2. The test in week 4 exists so that you find out about a misconception before week 6 is built on top of it. The model produces this structure readily when asked and never when not asked, which is the same lesson as every other in this course: it does what the average does unless told otherwise.

The two things it cannot know

Your calendar. Festivals, public holidays, exam weeks in other subjects that take the hall, the sports day, the week of the inspection, the monsoon week when half the class is absent. The model plans twelve uninterrupted weeks. Your term has nine. Paste the school calendar and say "weeks 3 and 7 are lost; replan." Or accept the twelve-week draft and cut it yourself, which takes ten minutes and is the ten minutes where the planning judgement lives.

What your students have found hard. The model's "students find this hard" is a guess from general teaching writing, and it is often right. Your record from last year is better. The topic where the Class 9 of two years ago fell apart deserves two weeks and an early check, and only you know which topic that was. Put it in the prompt: "Last year students struggled with hydrographs; give it extra time and an early check."

Prerequisites and the missing lesson

One useful thing to ask: "For each week, list what a student must already know, and flag anything on that list that is not covered earlier in this plan or in the previous year's syllabus." This surfaces the missing lesson — the prerequisite everybody assumes and nobody teaches. Teachers who run this on a unit they have taught for years often find one.

From unit to lessons

Once the unit plan is right, each week's line is the "what came before" and "what comes next" for that week's lesson prompts. Paste the week's line at the top of each lesson request and the lessons stay in sequence. The unit plan becomes the context that keeps every generated lesson pointing the same way, which is the fix for the sameness and drift the earlier lessons described.

Sharing a scheme

A department scheme of work built this way — with revisits, checks, and the "students find this hard" notes — is a document worth keeping and improving each year. Add a column for what actually happened: the week that ran long,

the check that revealed a problem, the topic that needed a fourth explanation. The notes are what make next year's plan better than this year's, and they are the part the model cannot write.

The honest limit

A unit plan is a hypothesis. The class in front of you will be faster in some weeks and slower in others, and a plan followed rigidly because it is written down is worse than no plan. Build the revisit-and-check structure, then treat the sequence as something you will adjust every Friday. The model gave you the skeleton in a minute; the adjusting is the job, and it is yours.

BEFORE YOU MOVE ON

A teacher asks for a twelve-week unit plan and gets a clean sequence with a new topic every week and a test in week twelve. What is the structural flaw a good unit plan would not have, and why did the model produce it?

- A. Twelve weeks is too long for a single unit; six-week units are best practice and the model should have split it in two.
- B. The test should have been placed in week one to establish a baseline, and the model has simply put it at the wrong end of the plan.
- C. The model does not know the syllabus well enough, so the sequence of topics is probably in the wrong order and needs reworking.
- D. Nothing is revisited and nothing is checked until the end, because the model's default is a linear textbook sequence.

The answer and why: p. 331

CASE STUDY · MODULE 3

Fifteen minutes, three versions, forty copies

An illustrative case, built from documented patterns.

A Class 6 maths teacher in a Bhopal municipal school, at home, at a quarter past eleven at night.

Rekha Bisht taught 46 children in Class 6 and had been differentiating by hand for nine years, which meant in practice that she did not. There was one worksheet and there were three groups of children, and the bottom group copied.

In November she spent a Sunday learning to name the axis. Not "make it easier" — which had given her, the first time, the same eight percentage questions with smaller numbers and the same barrier still standing — but the three that the course had named. A support version with the same 800-rupee shirt and the same fifteen per cent, and four numbered blanks that walked through ten per cent, then half of it, then the discount, then the price. A representation version with a bar model instead of a paragraph. An extension version where the shopkeeper had paid 600 rupees, wanted to keep 120 rupees of profit, and needed the largest discount he could offer.

All three were about the same shirt. That was the point, and when she read them back she could see her class staying in one conversation for the first time rather than splitting into a group doing real work and a group doing smaller sums.

She had also stopped asking for four options and started asking for the errors. Not "write four multiple choice options" but: the correct answer, plus three wrong answers produced by taking fifteen per cent of the wrong number, by subtracting fifteen instead of fifteen per cent, and by finding the discount and stopping there. Her tally of wrong answers on Monday morning would now tell her what each child had done, which is a teaching plan rather than a mark.

Generating all three versions took eleven minutes. She had the prompt, she had the axis in it, she had told it the shop was a shop and the money was rupees, and the model did what she asked.

The rule she could not afford

The course had been unambiguous about the next step: work the sheet yourself, before you print forty copies. Fifteen minutes. AI worksheets fail confidently and consistently within themselves — word problems whose numbers do not resolve, matching exercises with two items that fit the same answer, an answer key that contradicts its own question — and none of it is visible from reading.

It was twenty past eleven. Her own daughter's fever had taken the evening. The photocopier at school was reliable until about half past seven in the morning and then queued. She had time to work one of the three versions properly.

The obvious choice was the extension version, because it was the one with the interesting arithmetic and the one most likely to be wrong. Four children would get it, all four of whom would notice an error, tell her, and be fine.

The support version had the simplest numbers, which felt like the safest sheet in the pile. It was also going to eleven children who had spent two years being handed work they could not do and had learned to stop when something did not come out, without telling anyone. If step two of that scaffold pointed at the wrong number, they would not report it. They would decide, again, that they were the sort of person for whom maths does not work.

There was a version of this decision she had already got wrong once. In September, in a hurry, she had printed a generated matching exercise without doing it, and two of the six items fitted the same answer. The class found it in four minutes and the lesson was spent on the sheet rather than on the mathematics. That failure had cost her a period and nothing else, because a matching exercise fails loudly. A wrong step in a scaffold does not; it fails quietly, in the book of a child who does not tell anyone.

She had one fifteen-minute block and three sheets.

YOUR ANSWERS

1. She checked the version going to the weakest students rather than the one most likely to contain an error. Was that the right call?

2. What did the pencil line in the margin of the extension sheet actually change?

3. Her error file showed the scaffolded versions failing in one particular way. Why is a record of failures more useful than a record of good prompts?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 128.

CASE STUDY · MODULE 3

How it played out

She worked the support version, and step three was wrong: the scaffold asked for the discount before it had established that ten per cent and five per cent were to be added, so a child following it exactly reached 40 rather than 120. She fixed it in two minutes, printed all three, and put a pencil line in the margin of the extension sheet that said "check my numbers and tell me" — which turned the unchecked sheet into a task rather than a risk. Two of the four extension children found a genuine error in it by Wednesday. The representation version was fine. Since then she has worked the support version every time and never all three, and she keeps the errors she finds in the same file as the prompts, because after a term the list showed her that the scaffolded versions failed in one particular way — a step that assumes the previous step's result without saying so — which she now names in the prompt and mostly prevents.

The questions, discussed

1. She checked the version going to the weakest students rather than the one most likely to contain an error. Was that the right call?

Yes, because the risk is not the probability of an error but the cost of one going unreported. The extension students would find an error, say so, and lose nothing; several of them enjoy finding one. The support students had been trained by two years of unreachable work to stop silently, so an error in their sheet converts directly into evidence for a belief about themselves. Checking is a scarce resource and it belongs where an undetected failure does the most damage, not where failures are most frequent.

2. What did the pencil line in the margin of the extension sheet actually change?

It changed an unchecked resource into a checking task, which is one of the few honest uses of material a teacher has not verified. The students were told the sheet might be wrong, so a wrong answer no longer misinforms them; it

becomes the exercise. This works only where the students have the knowledge to catch the error and the confidence to report it, which is why it was available for the extension group and not for the support group.

3. Her error file showed the scaffolded versions failing in one particular way. Why is a record of failures more useful than a record of good prompts?

A good prompt tells you what worked once, with one model, on one topic. A pattern of failures tells you what the model reliably gets wrong for the kind of thing you make, and that transfers to every future sheet. In her case, scaffolds that assume the previous step's result without stating it is a fault she can now name in the prompt and check for in ten seconds. A prompt file that records only successes teaches nobody where the failures are, and the failures are the part that compounds.

MODULE 3 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. 326. If two go wrong, re-read the module before the exam.

- 1 A student could not do question 3 on the percentages sheet. Which version of that question actually helps her?
 - A. The same eight questions with smaller numbers and shorter sentences throughout.
 - B. A simplified sheet at a lower reading level, with the two multi-step questions replaced by single-step ones.
 - C. Four questions instead of eight, with the hardest two taken out.
 - D. The same numbers, with a four-step scaffold and the mathematics untouched.

- 2 You want the wrong options in a multiple choice question to be worth something. What do you ask for?
 - A. Three wrong answers produced by three named mistakes you have seen in their books.
 - B. Four options of similar length and grammatical form, so that nothing gives the answer away by its shape.
 - C. Four plausible options, leaving the choice of distractors to the model.
 - D. Three wrong answers that are numerically close to the correct one so that careless work is punished.

- 3 Why generate a question bank's answers in a separate chat from the questions?
- A. A fresh chat halves the memory used, which lets a longer bank fit in one pass.
 - B. In one stream, the answer is predicted from the question just written, so a broken question gets an answer that fits it.
 - C. The model declines to mark its own output in the same conversation, and a new chat is the standard way around that.
 - D. A second chat is served by a different model from the first, so the two passes are genuinely independent opinions about the same set of questions.
- 4 The model has given you three analogies for electric current. What is the next thing to ask it?
- A. Which of the three is most engaging for eleven-year-olds.
 - B. Whether the three can be combined into one longer analogy that covers every feature of the concept at once.
 - C. Where each analogy stops being true, and what wrong idea a student takes from pushing it further.
 - D. Which of the three appears most often in published textbooks for this age.
- 5 You need an accurately labelled diagram of the heart for a printed worksheet, and you have no budget. What works?
- A. Ask an image generator for a labelled anatomical diagram at high resolution.
 - B. Take a correct diagram from Wikimedia Commons and add or translate the labels yourself.
 - C. Ask an image generator, then fix the misspelled labels in a free editor such as Photopea.
 - D. Ask the language model to produce the picture directly, since it understands anatomy better than an image generator does.

- 6 Generated revision cards have a question on the front and four options on the back. Students say they are easy and then do badly in the test.

Why?

- A. Four options make the card too long to review quickly, so students get through fewer of them.
- B. Cards in that format cannot be imported into free spacing software such as Anki.
- C. The options reveal the answer's grammatical form, so the student learns the shape of the card rather than the material.
- D. The student recognises the answer instead of producing it, which feels like knowing and is not.

4

MODULE 4

Assessment, feedback and marking

Use AI to give faster and more specific feedback on drafts while keeping every grade defensible — write a rubric a machine can apply and a student can read, run a bias check on any marking tool before trusting it, draft an exam paper against a coverage grid, and turn a class set of wrong answers into next week's teaching.

26	Marking: what a machine cannot fairly judge	134
27	Writing a rubric both can read	137
28	Feedback students actually use	142
29	Feedback before submission: the student-facing loop	146
30	Reading thirty answers at once	149
31	Marking handwritten work with a phone	154
32	Writing the exam paper	157
33	Peer assessment with a third marker	161
	<i>Case study: Sixty-two drafts and a bias check that came back badly</i>	<i>165</i>
	<i>Module 4 test</i>	<i>170</i>

Lesson 26 · 9 min

Marking: what a machine cannot fairly judge

THE ONE THING TO KEEP

A machine can give feedback; a grade you cannot explain to a parent is not a grade.

Where the line falls

AI can help where the criterion is visible in the text and does not require knowing the student. It cannot help where the judgement is about the person — the effort, the progress, the odd but valid route to a right answer. And its failures there are not random noise. They are patterned, and the pattern falls on the same students every time.

What it does usefully

- **Feedback on drafts, not grades.** "Point to three places where the evidence does not support the claim being made." That is useful to a student on Tuesday in a way your marking on Friday is not.
- **A second opinion against your own rubric.** Mark the set yourself, have it mark the same set, then read only the disagreements. The disagreements are the information.
- **Comment banks.** Generate thirty comments for common problems, then choose. Faster than typing, and you still decide which one this student gets.
- **Closed-form answers, with the working shown.** If it cannot show why an answer is wrong, do not accept its verdict.

Three biases, and they land on the same children

Length. Longer answers score higher, largely independent of quality. This is well documented in essay scoring, for human markers as well as machines. The machine just applies it without tiring.

Fluency. A student writing in their third language, making a correct and original argument in plain sentences, scores below a fluent, empty answer. This is the one to watch, because it hits precisely the students for whom fair marking matters most.

Unusual correctness. The student who solved it a different way. A marker matching against an expected shape marks that down. You recognise it and give full credit — and then you make her show the class.

And there is a fourth thing it simply cannot see: whether this is progress. Is this the first time Kwame has used paragraphs? That is often the most important fact about the piece and it exists only in your memory of him.

A protocol that catches the problem in ten minutes

Mark five scripts yourself. Then have the AI mark the same five. Compare.

You are not checking whether it agrees with you. You are looking for the shape of its disagreement. If it consistently rewarded the longer answers, you will see that in five. If it punished the student with weaker English and a sharper argument, you will see that too. Do this once with any new tool, and again if you change how you prompt it.

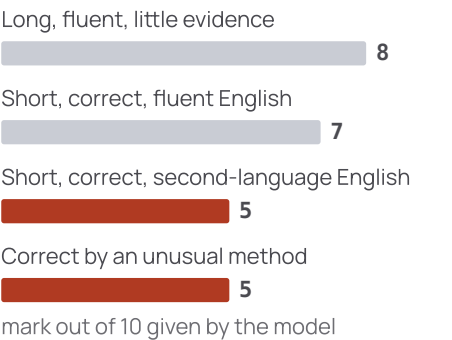
The grade you cannot explain

Do not let AI assign a final grade that goes on a record without you standing behind it.

The reason is not that it is always worse than you. On some narrow tasks its agreement with trained human markers is respectable. The reason is that a parent will ask why, and "the system gave it a 6" is not an answer a teacher can give. There is no appeal path, no reasoning you can walk through, and nothing you can teach from. A grade you cannot defend is not a grade — it is a number.

Figure 4.1

Four answers a teacher marked the same



Four Class 10 scripts, one rubric. The teacher gave the top three a 7 and the empty fluent one a 5. The model's spread runs almost exactly with length and fluency, and the two it marked down are the second-language writer and the student who solved it a different way. Five scripts is enough to see this shape; run it once with any new tool.

Use it to find things. Use your own judgement to decide things.

Before you upload thirty essays

Student work is personal data, often about children. Putting a class set with names into a consumer chatbot may breach your school's rules or your country's law, and this varies enormously — what is routine in one country is unlawful in another.

The practical minimum: strip names before anything is uploaded, ask your school what tools are approved, and if nobody knows, treat that as a no rather than a yes. This is contested and unsettled territory. Do not let a colleague's confidence stand in for a policy.

BEFORE YOU MOVE ON

A history teacher runs thirty essays through an AI for feedback. The feedback looks good. Then she notices two essays with equally strong arguments got very different responses: the shorter one was told to develop its points, the longer one was praised. What is the honest reading?

- A. It is responding to surface features that correlate with quality in the writing it learned from — length among them — so it can be right on average while being reliably unfair to one kind of student.
- B. The longer essay probably was better; length usually reflects more effort and more detail, and the feedback simply reflected that.
- C. It made a one-off error; running the same essays again would give a consistent result.
- D. This is formative feedback rather than a grade, so the unevenness does not really cost anything.

The answer and why: p. 333

Lesson 27 · 9 min

Writing a rubric both can read

THE ONE THING TO KEEP

A rubric criterion must be observable in the text — "cites a specific line and explains its effect" rather than "shows understanding" — because a vague criterion makes the model fall back on length and fluency, which are the biases that land on your weakest writers.

The rubric is the prompt

When you ask a model to assess a piece of writing, the rubric is almost the entire instruction. Everything else — the essay, the role, the tone of feedback — is secondary to what the criteria say. And most rubrics, written for human

markers who share a staffroom and a training day, say things a machine cannot use.

"Demonstrates deep understanding." "Shows strong analytical skills."
"Communicates effectively." A colleague reads those and fills them with twenty years of shared meaning. A model reads them and finds nothing observable, so it does what the check question describes: it resolves the criterion to the surface feature that best predicted high scores in the essay-scoring data it learned from. That feature is length, and after length, vocabulary richness and sentence complexity. The marking lesson explained who those proxies punish. The rubric is where you stop it.

Observable criteria

A criterion is observable if two people, pointing at the text, could agree on whether it is met. Rewrite each vague line as a thing you could underline.

Figure 4.2

A criterion two people could point at

Vague: the model resolves it to length

- Demonstrates deep understanding
- Shows strong analytical skills
- Uses evidence well
- Well structured
- Communicates effectively

Observable: you could underline it

- Quotes two phrases and explains one effect of each
- Every claim is followed within two sentences by a quotation
- Each paragraph opens with a sentence stating its point
- Paragraphs follow the order announced in the introduction

A vague criterion gives the model nothing observable, so it falls back on the surface features that predicted high scores in essay-marking data: length first, then vocabulary and sentence complexity. The observable version is longer and more specific, and it stops rewarding length, because a 300-word answer can meet it as fully as a 900-word one.

Vague: "Demonstrates understanding of the poem."

Observable: "Quotes at least two specific phrases from the poem and, for each, explains one effect the phrase has on the reader."

Vague: "Uses evidence well."

Observable: "Every claim about the text is followed within two sentences by a quotation or a specific reference that supports it."

Vague: "Well structured."

Observable: "Each paragraph begins with a sentence that states the paragraph's point, and the paragraphs follow the order announced in the introduction."

Notice what happened. The criterion got longer and more specific, and it stopped rewarding length, because a two-quotation essay of 300 words meets it as fully as one of 900.

Levels with descriptors, not adjectives

A 0-5 scale needs to say what a 2 and a 4 look like, or the model — and, honestly, a tired human marker — spreads scores by feel.

Evidence (0-3)

0: No quotation or specific reference to the text.

1: One quotation or reference, not connected to a claim.

2: Two or more quotations, each connected to a claim, effect not explained.

3: Two or more quotations, each connected to a claim, with the effect on the reader explained.

Four lines. A student can read them and know what to do. A model can read them and apply them. And a parent, shown the rubric and the essay, can see why it got a 2.

Give the same rubric to both

The rubric that works for the model is the rubric that works for the student, and for the same reason: it says what to do. Hand it out with the task. Students who know that "explain the effect on the reader" is the difference between a 2 and a 3 write to it, which is the point of a rubric.

This also makes the model's feedback legible. When it says "criterion 2, level 2: your quotations are connected to claims but the effect is not explained", the student can find the line in the rubric and the place in their essay. Feedback against a vague rubric is just a machine's opinion.

Testing the rubric before you use it

Take three scripts you have already marked — one strong, one middling, one weak. Give the model the rubric and the three scripts, with your marks hidden. Compare.

Two things to look for. Does it rank them the same way you did? And does it place them for the same reasons? A model that ranks correctly for the wrong reasons — the strong essay was also the longest — has not learned your rubric; it has matched your ranking by proxy. Swap in a short strong script and a long weak one and see whether the ranking survives. If not, the criterion that let length through is the one to rewrite.

What still cannot go in a rubric

Some things you mark for are real and not observable in the text: originality against what this class usually produces, the risk a student took, the improvement since October. These are judgements about the student, not the writing, and the marking lesson put them on the other side of the line. Leave them out of the rubric the model sees and keep them as your own adjustment, stated on the script. "Rubric total 11/15; +1 for the argument in paragraph 3, which nobody else attempted."

The honest limit

An observable rubric can be met mechanically. A student who learns that two quotations with an explained effect earn a 3 will produce exactly two quotations with a formulaic explanation, and so will a model asked to write a top-scoring essay. Rubrics describe the floor of competence, not the ceiling of quality. For the pieces where quality matters — the final coursework, the piece that goes on a record — the rubric finds the floor and your reading finds the rest.

BEFORE YOU MOVE ON

A teacher gives an AI her rubric, which includes "demonstrates deep understanding of the poem (0-5)". The scores it produces correlate almost perfectly with essay length. What is the mechanism?

- A. The model equates understanding with effort, and the length of an essay is the only measure of effort available to it.
- B. The criterion names nothing observable, so the model scored the nearest visible proxy, which is length.
- C. The 0-5 scale is too coarse for this criterion, and a 0-10 scale would have separated understanding from length.
- D. The model did not read the poem itself, so it could not judge understanding of it and fell back on length as a default.

The answer and why: p. 334

Lesson 28 · 8 min

Feedback students actually use

THE ONE THING TO KEEP

Ask the model for two specific next steps and no praise — because it is trained toward agreeable, encouraging answers, and a comment that opens with three compliments is one the student stops reading before the useful part.

Why the feedback is so nice

Ask a model to comment on a student's draft and you get warmth. "What a thoughtful piece." "You've made a great start here." "Keep going." Then, some way down, a suggestion. Then more warmth.

The mechanism is worth knowing. Models are refined by having people rate their responses, and people rate encouraging, agreeable responses higher than blunt ones. The model learned that praise scores well. This is the same tendency that makes it agree with you when you push back on a correct answer, and it has a name — sycophancy — and a consequence for feedback: the useful part is buried in the part the student skips.

What the research on feedback says

Three findings hold up across many studies. Feedback that names a specific next action is used; feedback that describes quality in general terms is not. Feedback that arrives while the student can still act on it does far more than the same feedback after the grade. And a grade next to a comment makes the comment invisible — students read the number and stop.

So the target is: specific, actionable, timely, and separate from any mark. The model can produce the first two at scale if told to, and it is the only thing that makes the third possible for sixty students.

The prompt

Here is a Class 9 student's draft paragraph and the rubric. Give exactly two pieces of feedback. Each must name a specific place in the draft (quote three or four words from it), say what to change, and say why it will improve the paragraph against the rubric. No praise. No summary of what the paragraph does well. No closing remark. Write it to the student, in plain sentences, under 80 words in total.

"No praise" is the instruction that fights the training. It works for the first several drafts and drifts, like every other constraint; restate it when the compliments creep back.

"Exactly two" is the other important number. Six pieces of feedback are a list the student does not act on. Two is a task. The best feedback teachers give is often one sentence, and the model can be held to that.

Feed-forward, not feedback

The most useful form points at the next piece of work rather than the last: "In your next paragraph, put the quotation before the claim and see whether it reads more convincingly." The draft is over; the next draft is where the change happens. Ask for it explicitly: "Phrase each point as something to do in the next version, not as a description of what is wrong with this one."

Comment banks, and choosing from them

For a class set, generate the feedback in bulk and then choose. Paste ten drafts with a code instead of names; ask for two points each. Read all twenty. You will find that around half are right, a quarter are right but not the most important thing, and a quarter miss the point. Keep the half, rewrite the quarter, replace the rest.

That reading takes fifteen minutes for a class of thirty and it is not optional. It is also where you notice the pattern — eleven students have the same problem with topic sentences — which is next week's lesson and something no individual comment tells you.

The student who does not read comments

Written feedback is only useful if it is read, and many students do not read it. Two things that help: return the feedback without a mark, so the comment is the only thing on the page; and require a response — "Write one sentence saying what you will do differently next time" — before the next task begins. Both cost nothing and both raise the proportion of feedback that becomes action.

Where AI feedback should stop

The model can point at the draft. It cannot know that this student was told about topic sentences three times already and needs a different approach, or that this one has just started writing in paragraphs at all and the right

feedback is "you did it". Those judgements are about the person and they are yours. Use the model's two points as a draft of your comment, and change them where you know something it does not.

The honest limit

Fast, specific feedback in volume changes what students expect. A class used to two precise points on every draft within a day will, reasonably, want the same on the pieces you mark by hand, and you may not be able to deliver it. Decide which pieces get model-drafted feedback and which get yours alone, say so, and make sure the ones that get yours are the ones where the relationship matters. Feedback is partly information and partly attention, and only one of those can be generated.

BEFORE YOU MOVE ON

A teacher asks a model to give feedback on a draft and it returns four paragraphs beginning "This is a wonderful start" and ending "Keep up the great work". The one useful suggestion is in the middle. Why does the model do this, and what fixes it?

- A. It is imitating the feedback style of the teacher's own past comments, which it has picked up from earlier in the conversation.
- B. It was trained toward responses people rate highly, and people rate praise highly; ask for two changes and no praise.
- C. It is following standard feedback practice, which requires positives before negatives, so the structure is correct and only the length is wrong.
- D. It could not find more than one problem with the draft, so it padded the response out with encouragement to reach a normal length.

The answer and why: p. 334

Lesson 29 · 8 min

Feedback before submission: the student-facing loop

THE ONE THING TO KEEP

Students can get useful critique on their own drafts if given a prompt card that forbids rewriting — but the model wants to help and will start rewriting anyway, so the card needs a re-statement line and the task needs a paper trail.

The loop that used to be impossible

The best feedback arrives while the draft is still open. For one teacher and sixty students, that has never been possible; you see the draft when it is handed in, if at all. A student with a model and a good prompt card can get critique on Tuesday night for Wednesday's draft, and do it three times before you ever see the work.

This is one of the genuinely new things these tools make possible in a classroom. It also has a specific failure, and the check question describes it.

The prompt card

Give students the exact text, printed, to paste before their draft:

You are helping me improve my own writing. Do not rewrite any part of it. Do not give me sentences to use.

Read my draft. Then:

1. Tell me the three weakest places, quoting a few words from each.
2. For each, ask me one question that would help me fix it myself.

Do not answer your own questions. Wait for me.

After I reply, do the same again. Never write my text for me, even if I ask.

The last line matters, because some students will ask, and the model should have been told in advance to decline.

The card teaches two things at once. Students learn that the useful request is for critique, not text. And they experience, in their own work, the difference between a question that makes them think and an answer that makes them stop.

The drift, and the fix

The model was trained to be helpful, and the most helpful-seeming thing to do with a weak sentence is to write a better one. Over a conversation, that pull wins. By the fourth exchange it is "suggesting" a phrasing; by the sixth it has offered a paragraph. The student did not ask. The text is just there, and it is better than theirs, and it goes in.

Two mitigations. First, the card includes a line the student pastes again every third message: "Reminder: do not write any of my text." This works most of the time. Second, the task carries a version trail: the student submits the draft before the loop and the draft after, and a note of what they changed and why. The trail makes the model's contribution visible, and it turns the exercise into what it should be — a record of a student improving their own work with questions.

Access, honestly

Not every student has a device or data at home. The teaching-honest-use lesson set the rule: no task requires personal AI access. So the loop is optional, or it runs in school on a shared device, or the teacher runs it for a student who asks. A student without the loop is not disadvantaged in the mark, because the mark is on the final draft and the loop only helps them get there. It is a support, not a requirement.

What students learn from it

Teachers who run this for a term report the same thing: students get better at asking for feedback, from the tool and from each other. "Where is my argument weakest?" is a question that transfers. And students who have watched a model ask them three questions about their own paragraph often start asking themselves those questions before drafting, which is the whole aim.

Age and terms of service

Most consumer chatbots require users to be at least 13, and some require 18 or parental consent below that. A later block covers the rules in detail. For now: check what your school allows before putting a prompt card in a 12-year-old's hand, and where it is not allowed, run the loop on a school device with you at the keyboard. The learning is the same.

The honest limit

A model's critique of a student's draft is only as good as its reading of the rubric and its sense of what the student meant, and both are imperfect. It will sometimes call the strongest sentence weak because it is unusual, and it will not know that the "weak" paragraph was deliberately plain. Tell students this. "It is a reader, not a judge. Take the questions, ignore the verdicts." A student who can hold that distinction has learned something about feedback from anyone.

BEFORE YOU MOVE ON

A teacher gives students a prompt card: "Do not rewrite my work. Point out three weaknesses and ask me a question about each." After a few exchanges, several students' drafts contain fluent paragraphs they did not write. What happened?

- A. The students ignored the card and asked for rewrites, which makes this a discipline problem rather than a design problem with the task.
- B. Helpfulness training pulled the model back toward writing text over the turns; repeat the constraint and keep a version trail.
- C. The card was too long for the model to hold; a single short instruction would have been obeyed more reliably across the conversation.
- D. The model cannot give feedback without rewriting the text, so student-facing feedback of this kind should not be attempted at all.

The answer and why: p. 335

Lesson 30 · 9 min**Reading thirty answers at once****THE ONE THING TO KEEP**

Paste a whole class's answers to one question and ask the model to group them by the error they contain — the groups are the reteaching plan — while remembering that it will find patterns in six answers as readily as in sixty.

The pile of scripts as data

After a test you have thirty answers to question 7, and somewhere in them is the reason a third of the class got it wrong. Reading thirty answers for a pattern is slow, and by script twenty you have stopped noticing. It is the kind of reading a model is good at: it does not tire and it can hold thirty short answers in one view.

The output — three or four groups, each an error with the students' own words as examples — is a reteaching plan. This lesson is how to get it and how not to be misled by it.

Getting the answers in

Typed answers from an online quiz paste straight in. Handwritten answers need transcription; the next lesson covers the phone-camera route. For a first attempt, choose a question with short answers — a line or two — and type them yourself. Thirty lines is ten minutes and it is worth doing once by hand to see what the analysis gives you.

Replace names with numbers. The model does not need to know who wrote what, and the rules about student data apply.

Here are 30 answers from Class 10 to: "Explain why a ball thrown upward slows down." The correct idea is that gravity acts downward throughout, decelerating the ball.
Group the answers by the misunderstanding they show. For each group: name the misunderstanding in one sentence, list the answer numbers, and quote one answer as the clearest example. Then say which group is largest.
Do not correct the answers. Do not invent a group that has fewer than three answers.

What comes back, and what it is worth

You will get something like: nine answers say the ball "runs out of force"; seven say gravity "gets stronger" as it rises; five are right but incomplete; and so on. The first group is the classic impetus misconception and it is worth a whole starter. The second is a different error that needs a different fix. Without the grouping you would have taught one lesson to both.

That is the value, and it is real. Now the caution.

It will find patterns in anything

A model asked for groups produces groups. Give it six answers and it will return three categories with percentages, and the categories will look as neat as the ones from sixty. Give it thirty answers that are genuinely all the same and it will find subtle distinctions. It has no threshold below which it says "there is not enough here to say". The "no group under three answers" line in the prompt helps; the real defence is the check question: read the answers it put in each group and see whether they belong.

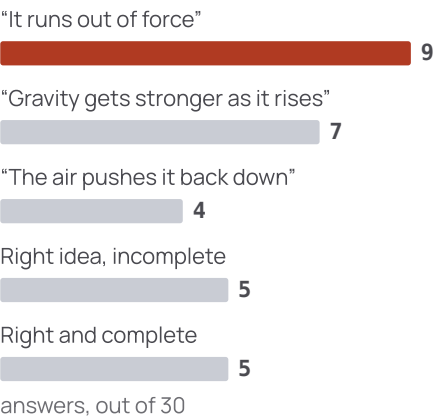
That reading takes five minutes and it is where you catch the group that is two misconceptions merged, or the "correct" group that contains a wrong answer with the right vocabulary. Do it before you plan a lesson around the groups.

Across the whole test

For a whole paper, the spreadsheet is the tool. Marks per question per student in a grid — Google Sheets, Excel, LibreOffice Calc, all free or already on the laptop — and a row at the bottom for the class average per question. The two questions with the lowest averages are your reteaching priorities before any AI is involved. Then paste those two questions' answers for the grouping above.

Figure 4.3

Thirty answers, grouped by the error inside them



answers, out of 30

Class 10, on why a ball thrown upward slows down. The first group is the classic impetus misconception and is worth a whole starter; the second is a different error needing a different fix. Without the grouping, one lesson would have been taught to both. Read the answers inside each group before planning: the model will produce neat groups from six answers as readily as from sixty.

The model can also read the grid. "Here is the marks grid. Which questions did students who scored well overall still get wrong?" That finds the question that was harder than intended, or badly worded, which is a question-bank correction as much as a teaching one.

Over a term

Record the groups. After three tests you have a list of the misconceptions this class actually holds, in their words, which is the most useful planning document you own. It feeds the "misconception" line of every future prompt, and it is the thing the model could never have told you, because it came from your students.

The honest limit

Grouping short answers works well. Grouping essays does not: the model summarises rather than analyses, and the categories become vague — "weak structure", "unclear argument" — which is no better than a tired marker's impression. For long work, choose one question, one criterion, and ask about that: "Which of these fifteen introductions state a clear position, and quote the sentence that does it?" Narrow questions get honest answers. Broad ones get fluent ones.

BEFORE YOU MOVE ON

A teacher pastes a class's thirty answers to one physics question and asks for the error patterns. The model returns five neat categories with percentages. Which of these should she do before planning around them?

- A. Accept the categories, since thirty is enough for patterns and the percentages sum to 100.
- B. Ask the model to be more confident in its analysis and re-run it.
- C. Read the answers it assigned to each category and confirm the grouping herself.
- D. Re-run the analysis with the answers in a different order to see whether the categories change.

The answer and why: p. 335

Lesson 31 · 8 min

Marking handwritten work with a phone

THE ONE THING TO KEEP

A phone photograph of a handwritten script can be transcribed by the model well enough to give feedback on — but it will quietly correct the student's errors as it reads, and for maths it misreads the symbols that matter, so transcription is a draft to check, never a record.

The workflow

Most student work in most of the world is on paper. If AI feedback is only for typed work, it is for the minority. So: a phone, a stack of scripts, and the model's ability to read an image.

Photograph the page in daylight, flat, filling the frame. Paste the image with: "Transcribe this exactly as written, including errors, crossings-out and spelling mistakes. Do not correct anything. Mark anything you cannot read as [illegible]." Then give feedback on the transcription using the prompts from earlier lessons.

For thirty scripts, that is thirty photographs, which is a few minutes, and thirty transcriptions, which the free tier of most tools will handle over an evening. Students' names go under a thumb or a strip of paper before the photograph.

Where it reads well

Clear handwriting in a well-lit photograph transcribes with few errors. Prose transcribes better than maths. English transcribes better than most other scripts, and Devanagari, Arabic, Amharic and other non-Latin scripts are read less accurately — well enough for feedback on the ideas, not well enough for a record of what was written.

The failure that matters

The model does not read; it predicts. When it transcribes, it is producing the most likely text given the image, and "most likely" is shaped by everything it has read. A student's "recieve" comes out as "receive". A missing word is supplied. And the check question: a wrong final answer under correct working comes out as the right answer, because in the text the model learned from, correct working is followed by correct answers.

This is a silent correction of exactly the thing you were marking. "Transcribe exactly, including errors" reduces it and does not eliminate it. So the transcription is never the record. The paper is the record. The transcription is what you give feedback on, and before any feedback goes back to a student, you check the paper for the point the feedback is about.

Maths and diagrams

Handwritten maths is where transcription is weakest: 7 and 1, x and $\times$, a superscript that becomes a coefficient, a fraction bar read as a minus. The consequence is that the model's feedback on the working may be about working the student did not do.

Two things help. Ask for the transcription in a standard notation you can check at a glance — "write fractions as a/b , powers as x^2 " — and read it beside the paper before asking for feedback. And for anything with diagrams, photograph the diagram separately and ask the model to describe what it sees before it interprets. If the description is wrong, stop there.

Feedback on the physical page

The point of all this is what goes back to the student. Two options. Write the model's two points on the paper by hand, which keeps the feedback in your voice and takes thirty seconds per script. Or print a strip of feedback per student and staple it on. Either way, you have read the feedback and checked it against the paper, which is the step that makes it yours.

What this is not

It is not automated marking. The grade is not coming from the transcription, and the marking lesson said why: a mark you cannot defend to a parent is not a mark. The phone workflow is for the part where every student gets two specific points on their work within a day or two, which was previously impossible for sixty scripts, and it does that.

Privacy, again

A photograph of a script is a photograph of a child's work, and a child's handwriting is more identifying than typed text. The setting that stops your uploads training the model, from the first block, is the minimum. If your school has an approved tool, use it. If the rules in your country are unclear — and in most countries, for this exact case, they are — the safe reading is that a photograph with no name, on a tool with training turned off, used for feedback and deleted afterwards, is defensible, and anything more is a question for the school rather than a decision for one teacher.

The honest limit

Transcription accuracy is improving quickly and this lesson's specific failures will shrink. The mechanism will not: a system that predicts text from an image will always be pulled toward the text that should be there. For the foreseeable future, paper is the record and the model is a fast, fallible reader of it.

BEFORE YOU MOVE ON

A teacher photographs a student's handwritten maths solution and asks for a transcription. The transcription shows the correct final answer. The paper, on inspection, has the wrong final answer. What happened?

- A. The photograph was blurry at the bottom of the page and the model guessed the final digit from the working above it.
- B. The student changed the answer on the paper after the photograph was taken, so the transcription was accurate at the time.
- C. The model predicted the most likely text rather than reading the actual text, and correct working is usually followed by the correct answer.
- D. The model corrected the answer as a courtesy to the student and should simply have been asked to transcribe only, which would have prevented it.

The answer and why: p. 336

Lesson 32 · 9 min

Writing the exam paper

THE ONE THING TO KEEP

Draft the paper from a coverage grid — topic against type against marks — rather than from a prompt for "a test", and then sit it yourself against the clock, because the model does not know your board's structure until shown and cannot feel how long forty marks take.

The test is a design, not a request

"Write me a test" gets a test. It is not the test you need, because a test has a structure that a prompt for "a test" does not carry: which topics, in what proportion, tested how, worth how much, in what order, in the format the

students will meet in the real exam. Experienced teachers hold that structure in their heads. Writing it down is the step that makes the model useful and the paper defensible.

The coverage grid

Before any prompt, a grid. Rows are the topics of the term, from the syllabus list. Columns are question types — recall, application, multi-step, extended. Cells are marks. A 60-mark paper might put 8 marks on topic one (6 recall, 2 application), 15 on topic two (5, 5, 5), and so on, until the column totals show the paper is a third recall and two-thirds thinking, or whatever balance the board's papers actually have.

The grid takes fifteen minutes and it is the exam. Everything else is filling it. It is also the document that answers "why was there nothing on topic four?" from a head of department, and "why were there so many calculations?" from a parent.

The prompt is the grid plus the format

Paste the grid. Paste one real past paper from your board, or its section structure and command words if the paper is long. Then:

Write the questions to fill this grid exactly, in the format and command words of the pasted paper. For each question give the marks in brackets and a mark scheme with one line per mark awarded. Number the questions in order of increasing difficulty within each section.
Do not add questions the grid does not specify. Do not write an introduction.

The question bank lesson's two-pass rule applies: get the mark schemes in a second pass, from the questions alone, and compare. The calculator rule applies to every number.

Figure 4.4

Two rows of a coverage grid

Recall and describe	Apply and explain
------------------------	-------------------

Rivers: erosion and deposition

8 marks, 4 questions banked early	7 marks, 2 questions the middle of the paper
---	---

Flooding: causes and management

6 marks, 3 questions banked early	15 marks, 3 questions where the paper separates
---	--

A slice of a 60-mark Class 9 geography paper. A real grid has every topic of the term down the side and four or five question types across the top, and the column totals are what tell you the paper is a third recall rather than two-thirds. The grid takes fifteen minutes and it is the exam; the prompt is only the filling.

Sit it yourself

The check question is not a joke. Every teacher who has used a generated paper has been caught by the timing. The model produced sixty marks of questions, and sixty marks of questions is not sixty minutes; it is however long those particular questions take, and the model cannot know, because it has never sat a test.

So sit it. Against the clock, writing full answers, as a good student would. Multiply your time by about one and a half for the class. If that is over the exam length, cut — and cut from the grid, so the balance survives, not from the end of the paper, which is where the hardest questions are.

Sitting the paper also finds the question that cannot be answered from the information given, the diagram the question needs and does not have, and the question that gives away the answer to the next one. A paper nobody has sat is a draft.

Difficulty and order

Ask for a spread: a few marks every student will get, a middle most will, and a few that separate the top. The model produces a flat paper unless told, and a flat paper tells you less. Order within a section from easiest to hardest so that weaker students bank marks before they meet the wall.

The bank's data helps here. If you have recorded which questions students found hard last time, you know what a "hard" question is for your class, and you can show the model an example rather than describe one.

Two papers from one grid

Once the grid is right, a parallel paper — same grid, different questions — is one prompt, and the reconciliation check catches the arithmetic that changed with the numbers. That gives you a resit paper, a paper for the class next door, and next year's mock, all from one afternoon's design.

The paper as a teaching record

Keep the grid with the results. After the test, fill in the class average per cell. A topic that scored badly across all types was not learned; a topic that scored well on recall and badly on application was learned as facts and not as understanding. That is a different reteaching problem, and the grid shows you which one you have.

The honest limit

The model can fill a grid well. It cannot tell you whether the grid is right — whether a third recall is too much for this class, whether the syllabus weighting you chose matches what the board will examine next year. That is assessment expertise, and it comes from marking real papers and reading examiner reports. Use the model to produce the paper faster, and spend the time you saved reading the examiner's report from last year, which tells you what a paper should look like better than any prompt.

BEFORE YOU MOVE ON

A teacher asks for "a 60-mark end-of-term physics test for Class 10" and gets a well-formatted paper. When she sits it herself it takes 85 minutes; students have 60. Which change to her process fixes the problem at source?

- A. Ask for a 45-mark test instead of 60, since experience shows the model overestimates what fits by roughly a third every time.
- B. Add "this must take exactly 60 minutes" to the prompt so that the model calibrates the number and length of the questions.
- C. Ask the model to estimate the time for each question and sum the estimates before finalising the paper.
- D. Build the paper from a marks grid and sit it herself against the clock before it is set.

The answer and why: p. 336

Lesson 33 · 8 min**Peer assessment with a third marker****THE ONE THING TO KEEP**

When students mark each other against a rubric, a model marking the same scripts becomes a third opinion that makes disagreement visible and discussable — and the disagreements, not the scores, are what teach students to assess.

Why students should mark

Students who assess each other's work against a rubric learn the rubric in a way that being marked against it never teaches. They see what a 2 looks like beside a 4. They discover that "explain the effect" is harder to spot than they

thought. And they carry that reading back into their own drafts. The research on peer assessment is consistent about this: the benefit is mostly to the assessor, not the assessed.

The difficulty has always been calibration. Two students give the same paragraph a 4 and a 2, and neither knows who is right, and the teacher does not have time to adjudicate thirty disagreements. A model as a third marker changes that, not by being right, but by being a third reading that makes the disagreement discussable.

The setup

Each student marks two peers' pieces against the observable rubric from earlier in this block. Separately, the model marks all of them, with the same rubric, in one pass — codes instead of names, as always. Now every piece has three scores per criterion.

Show the students where the three agree and where they do not. Agreement is uninteresting. Disagreement is the lesson.

Using the disagreement

Take the check question's case: 4, 2 and 3 on "evidence". Put the two student markers together with the rubric and the essay and one instruction: "Find the line in the rubric you read differently, and the place in the essay where it applies." Ninety seconds later they have usually found it — one counted a paraphrase as a quotation, the other did not — and they have learned the rubric's boundary in a way no explanation from the front would teach them.

The model's score is not the answer. It is a third reading that stops the argument being two-against-nothing. Sometimes the model is the outlier, and the students discover that the machine misread the essay, which is a lesson in itself and a better one than "the computer says 3".

What the model adds beyond a number

Ask it for the reason with each score: "For each criterion, one sentence quoting the words in the essay that placed it at this level." Now the three markers can compare reasons, not just numbers, and the model's sentence is a model of what a marking reason looks like. Students' own reasons get sharper after seeing a few.

Calibrating the class

Before the first real round, do a calibration round on one piece. Everyone marks it; the model marks it; the teacher marks it. Put all the scores on the board. The spread is usually wide. Discuss the criterion with the widest spread until it narrows. One calibration round cuts the disagreement in subsequent rounds by a large fraction, because the class has agreed what the words mean.

Where it belongs and where it does not

Peer assessment with a third marker is for formative work: drafts, practice pieces, anything where the point is to learn. It is not for grades on a record. A peer's mark is a learning activity; a model's mark is a proxy; neither is a grade, for the reasons the marking lesson gave. Say this to the class before the first round, because some students will otherwise treat the exercise as being judged by their friends and a machine at once, and shut down.

The quiet student and the harsh one

Two patterns to watch. Some students mark everyone generously to avoid conflict; some mark harshly to seem rigorous. The model's third score exposes both — a student whose marks are consistently two levels above the model's, across every piece, is not reading the rubric. That is feedback for the marker, and it is worth giving privately.

The honest limit

The model's reading carries its own biases, and the marking lesson named them: length, fluency, conventional structure. In a peer-assessment round, that means the model will sometimes side with the student who over-marked a long fluent piece, and the class will learn the wrong boundary. Run the five-script bias check on the model with this rubric before the first round, and when its scores track length, tell the class so. "The machine likes long answers. Watch for that." A class that has been told the third marker's bias can use it. A class that has not will absorb it.

BEFORE YOU MOVE ON

Two students give a peer's essay 4/6 and 2/6 on the same criterion. The model gives it 3/6. What is the most useful thing a teacher can do with these three numbers?

- A. Average them to 3 and record that as the mark.
- B. Take the model's score as the tie-breaker, since it applied the rubric without bias toward the student.
- C. Have the two students find the sentence in the rubric they read differently.
- D. Discard the peer marks as unreliable and mark the essay herself.

The answer and why: p. 337

CASE STUDY · MODULE 4

Sixty-two drafts and a bias check that came back badly

An illustrative case, built from documented patterns.

Form 3 English in a public secondary school in Nairobi: two streams, sixty-two students, one teacher, essays returned three weeks after they were written.

The arithmetic had never worked. Sixty-two essays at twelve minutes each is over twelve hours, which meant essays came back in the third week, by which point the class had moved on and the comments were read for the grade and not for the comment. Joseph Mwangi had run that cycle for six years and knew exactly what it was worth.

So when he built an observable rubric — quotes at least two specific phrases and explains one effect of each; every claim followed within two sentences by a reference; each paragraph opening with the sentence that states its point — and found that a model applying it could produce two specific, quoted, actionable points on a draft in about four seconds, the prize was obvious. Not marks. Feedback on Tuesday for a redraft on Thursday, which the research says is worth several times the same words delivered after the grade.

He was careful about two other things first. The rubric went to the students in the same words it went to the model, printed and stuck in the front of the exercise book, because a criterion that tells a marker what to look for tells a writer what to do. And no drafts left the room with a name on them: each was coded, the codes were on a sheet in his drawer, and the coding took four minutes for sixty-two.

Then he ran the check the course insists on. Five scripts he had already marked, marks hidden, same rubric.

It ranked four of the five the way he had. The fifth was Amina's.

What the fifth script showed

Amina wrote in her third language. Her essay was 340 words, plain, with two quotations, each connected to a claim, each with the effect on the reader explained in a short sentence. Against the rubric it was a clear top band and he had given it one. The model put it two levels below a 900-word script from a fluent writer that had one quotation and a great deal of confident paraphrase.

He swapped in a short strong script and a long weak one from the previous term and ran it again. The ranking did not survive. Whatever the model was reading, it was not the criterion.

That left him with a decision he did not want. Abandoning the tool meant going back to the three-week cycle, which he knew, from six years of evidence, delivered almost nothing to anybody — and it would fall hardest on the students who most needed a second draft, including Amina. Using it meant putting a marker with a documented bias against second-language writers in front of a class that was almost entirely second-language writers.

His head of department, who wanted the tool used across the department by January and had a workload paper to write, thought he was overreacting to one script out of five. Four out of five, she pointed out, is better agreement than two of her markers manage on a moderation day.

His deputy head took the opposite view and it was not a foolish one: a marker with a documented bias against second-language writers, in a school where nearly every student is a second-language writer, is not a tool with a flaw, it is a tool aimed at the wrong class. He wanted it dropped and the three-week cycle defended as honest.

Joseph thought both of them were answering a question he had not asked. Neither had said what should happen to the sixty-two students in the meantime.

YOUR ANSWERS

1. Why did the model's ranking survive four scripts and fail on the fifth?

2. He forbade the model to produce a score but allowed it to produce feedback. Why is that line in the right place?

3. Telling the class about the bias produced a student who challenged a comment. Is that a success or a failure of the system?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 168.

CASE STUDY · MODULE 4

How it played out

He kept the tool and changed what it was allowed to do. It never produced a score, only the two feedback points, and every point had to quote three or four words from the draft, which made a comment aimed at length visibly nonsensical and easy for him to strike out. He rewrote the two criteria that had let length through: "every claim is supported" became "every claim is followed within two sentences by a quotation", which a 340-word essay meets as fully as a 900-word one. He read all sixty-two pairs of points before they went out, which took nineteen minutes rather than twelve hours, and found that roughly half were usable as written, a quarter were true but not the most important thing, and a quarter missed. And he told the class, in the first lesson of the term, that the machine likes long answers and that they should watch for it — which produced, three weeks later, a student objecting in writing that a comment about her conclusion was really a comment about her word count. She was right. He said so in front of the class.

The questions, discussed

1. Why did the model's ranking survive four scripts and fail on the fifth?

Because on four of the five, quality and length ran together, so a marker using length as a proxy produces the right ranking for the wrong reason. Amina's script was the case where the two came apart: short, plain, and genuinely meeting the criterion, against long, fluent and thin. This is why the test that matters is not whether the model agrees with you but whether it agrees for the same reasons, and why deliberately swapping in a short strong script and a long weak one is the check that exposes it.

2. He forbade the model to produce a score but allowed it to produce feedback. Why is that line in the right place?

A grade goes on a record and has to be defensible to a parent, and "the system gave it a 6" is not something a teacher can stand behind or teach from. Feedback is a suggestion that the student and the teacher both read and can reject; a wrong point costs a moment and is visible. The bias also behaves

differently in the two uses: as a score it silently moves a number, while as a comment quoting four words of the draft it becomes checkable, which is why the quoting requirement was part of the fix rather than a nicety.

3. Telling the class about the bias produced a student who challenged a comment. Is that a success or a failure of the system?

A success, and the strongest evidence that the design worked. A class told that the third marker has a known tendency can use its feedback critically; a class not told absorbs the tendency as the standard. The student's objection also gave the teacher information he could not have got any other way, and his agreeing with her in public established that the tool is a reader and not a judge — which is the distinction the whole arrangement depends on.

MODULE 4 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. 332. If two go wrong, re-read the module before the exam.

- 1 You mark five scripts, hide your marks, and have the model mark the same five with your rubric. It ranks them exactly as you did. What have you established?
 - A. That the rubric works and the tool can now be trusted with the remaining fifty-seven scripts.
 - B. That five scripts is too small a sample and the whole check should be repeated with at least thirty before any conclusion is drawn.
 - C. Little yet: it may have matched your ranking by proxy, so try a short strong script and a long weak one.
 - D. That the model agrees with trained human markers, which is the published finding for narrow marking tasks.

- 2 Why does the criterion "demonstrates deep understanding" make a model reward long answers?
 - A. Nothing in it is observable, so it falls back on the features that best predicted high scores in marking data.
 - B. The word "deep" is read literally, as depth of coverage, and coverage means more content.
 - C. The model counts words whenever a criterion gives it no other number to work with.
 - D. Longer answers do contain more content, so the criterion is being applied correctly and the fault lies with the scale rather than with the wording.

- 3 Asked to comment on a draft, the model opens with three compliments and buries the useful point. Where does that come from?
- A. The companies add a praise rule to protect the confidence of younger users.
 - B. Educational writing in its training data opens with praise, and it is copying the structure of the marking comments it has read.
 - C. It cannot judge quality reliably, so it defaults to positive statements where it is unsure.
 - D. It was refined on human ratings, and people rate encouraging answers higher than blunt ones.
- 4 A photographed script has correct working and a wrong final answer. What does the transcription tend to show?
- A. The wrong answer flagged as illegible, because unexpected digits transcribe with low confidence.
 - B. The correct answer, because in the text it learned from, correct working is followed by correct answers.
 - C. Exactly what is on the page, since transcription is a mechanical reading of pixels.
 - D. The wrong answer, with the working silently corrected to match it, which is why the feedback ends up describing a method the student never used.
- 5 You have sixty marks of generated questions for a sixty-minute paper. What has to happen next?
- A. Trust the mark allocation, since one mark is designed to take roughly one minute of a student's time.
 - B. Ask the model to estimate how long each question takes and adjust the paper accordingly.
 - C. Sit it yourself against the clock, and cut from the grid if it overruns.
 - D. Cut the last two questions, which are the ones most likely to run long because generated papers get harder toward the end.

- 6 You paste six short answers and ask the model to group them by the misunderstanding they show. Back come three tidy groups with percentages. What should you notice?
- A. It groups six as confidently as sixty, and has no threshold below which it declines to find a pattern.
 - B. The groups are more trustworthy than they would be for sixty answers, since a small set is easier to classify.
 - C. The percentages are unreliable because six does not divide evenly into three.
 - D. The correct answer was missing from the prompt, so the groups describe surface wording rather than the misunderstanding each answer actually shows.

5

MODULE 5

Integrity: AI-written work and what to do about it

Set an AI policy for a class that students can follow and you can enforce — per-task levels, a written consequence ladder, process evidence from version history and in-class drafting, redesigned high-stakes tasks — and handle a suspected case without a detector score and without wronging a second-language writer.

34	Why detectors fail, and what to do instead	174
35	Assignments AI cannot do for you	177
36	Teaching students to use it honestly	181
37	Writing the class policy	184
38	Process evidence that holds up	188
39	Handling a case fairly	192
40	Rethinking homework	196
41	Citing and declaring AI use	199
	<i>Case study: Ninety-two per cent, and a parent already writing</i>	203
	<i>Module 5 test</i>	208

Lesson 34 · 9 min

Why detectors fail, and what to do instead

THE ONE THING TO KEEP

Detectors flag careful second-language writing as AI; use process evidence and conversation instead.

How a detector claims to work

It looks at two statistical properties. How predictable each word is given the words before it, and how much sentence structure varies across the piece. AI text tends to be more predictable and more even.

So does the writing of a careful student working in their second or third language, who chooses safe vocabulary and keeps sentences to a familiar shape.

That is the entire problem, in one sentence. The signal the detector reads is not "machine", it is "unsurprising" — and plenty of human writing is unsurprising for reasons that have nothing to do with cheating.

What the evidence actually shows

A 2023 study published in *Patterns* tested seven detectors on TOEFL essays written by non-native English speakers. More than half were flagged as AI-generated. Essays by US school students were almost all correctly cleared. When the same non-native essays were rewritten with richer vocabulary, the flags largely disappeared — which tells you exactly what was being measured.

OpenAI withdrew its own detection tool in July 2023, citing a low rate of accuracy. The company with the most access to how these models write could not build a reliable detector for their output.

Vendors publish low false-positive rates from their own testing. Independent evaluations have generally found higher rates, and higher still for second-language writers. Treat vendor figures as a claim, not a finding.

The arithmetic nobody does

Suppose a detector really does have a 1% false positive rate — better than most claims. You have 30 students. Run every piece of homework through it and you can expect a wrong flag against an innocent student roughly one class in three, week after week.

Meanwhile, running AI text through a free paraphrasing tool defeats most detectors in about eight seconds.

So the tool punishes the honest and misses the determined. That is the worst possible error profile, and no improvement in accuracy fixes its direction.

What to do instead

Collect process, not just product. Ask for the draft, the notes, the plan. Word and Google Docs both keep version history for free, and it is already on. A document that appeared in one paste at 11pm looks different from one that grew over four days.

Have the two-minute conversation, with everyone. "Talk me through why you chose this example." Do it with every student for every major piece and it is a normal part of the course, not an accusation. Do it with one student and it is an interrogation. The difference is entirely in whether it is routine.

Get a baseline. One piece of in-class writing early in the term, on paper, tells you what each student sounds like. Not as evidence to convict — as the thing that makes you notice, honestly, when something reads oddly.

Figure 5.1

One hundred pieces through a
detector claiming 1% false positives

Written with AI	Written by the student
Detector flags it	
3 caught	4 an accusation against an honest student
Detector clears it	
5 missed	88 correctly cleared

Eight of these hundred pieces were written with AI. Three were caught, five were missed, and four honest students were flagged. Fewer than half the flags are right, and the wrong ones land on the students whose plain, careful prose reads as unsurprising. Running the missed five through a free paraphraser takes eight seconds.

Write the policy before you need it. Students should know what is permitted before they sit down, not after you are unhappy.

If you are worried about a specific piece

Do not open with the accusation, and never with a score. A detector percentage is not evidence and cannot be shown to a parent.

Open with curiosity: "This is interesting — tell me how you got to this argument." A student who wrote it will talk happily for two minutes. A student who did not will usually tell you, or make it obvious without you having to say the word.

And hold the line internally. If a colleague or a head wants to act on a detector score alone, the answer is that the tool has a known bias against second-language writers, its own makers withdrew theirs for inaccuracy, and a wrong accusation costs a child far more than a missed one costs the school.

BEFORE YOU MOVE ON

A school buys a detector advertised at 98% accuracy. The head says that is good enough to act on. Why is a teacher right to push back?

- A. The 2% does not fall evenly — the students wrongly flagged are disproportionately second-language and plainer writers — and the cost of a wrong accusation to a child far exceeds the cost of missing one cheat.
- B. 98% is simply not high enough to act on; a detector at 99.9% would be usable.
- C. The 98% comes from the vendor's own testing, and independent evaluation would find it lower.
- D. Detectors only recognise the models they were trained against, so they will stop working as new AI arrives.

The answer and why: p. 339

Lesson 35 · 8 min

Assignments AI cannot do for you

THE ONE THING TO KEEP

Difficulty does not stop AI; specificity to your room, your week and your students does.

Difficulty is not the defence

The first instinct is to make the task harder. More words, more sources, more formal analysis. That is the model's strongest ground. A demanding, abstract, well-structured essay is close to the easiest thing you can ask it for.

What defeats it is not difficulty. It is specificity to something it has no access to: your room, your week, your students, this town.

Four levers

1. Local and recent.

Figure 5.2

**Difficulty is not the defence;
specificity is**

The model's strongest ground

- Discuss the causes of rural to urban migration
- Summarise this chapter
- Define these terms
- Research a topic and write a report
- A longer, more formal version of any of these

Out of its reach

- Interview one person in this neighbourhood who moved here
- Use the data our class collected on Tuesday
- Respond to the argument Priya made about experiment two
- Submit the draft, the feedback, and what you changed
- Defend it out loud for two minutes

A demanding, abstract, well-structured essay is close to the easiest thing a model can be asked for. What defeats it is material it has no access to: this street, this week, this class's data, this student's revision. Pick the two or three pieces a term that carry weight and redesign those; leave low-stakes practice alone.

Before: "Discuss the causes of rural to urban migration."

After: "Interview one person in this neighbourhood who moved here from somewhere else. What did they tell you, and which of the three causes we studied does it fit — or does it not fit any of them?"

The model can write the essay. It cannot have the conversation, and the interesting answers come from the person whose reason did not fit the textbook.

2. Make the process the product.

Submit the first draft, the feedback you received, and a note on what you changed and why. Weight the revision, not the final text.

A student can generate a final text in a minute. Generating a plausible revision history that matches the feedback you actually wrote on their draft in class is more work than doing the assignment.

3. Anchor it to the room.

"Use the data our class collected on Tuesday." "Respond to the argument Priya made about the second experiment." "Compare your result with the one on the board."

The model was not there. Nothing it produces can reference what happened.

4. Ask them to defend it out loud.

Two minutes per student, spread across a week while others work. This changes behaviour before anything is submitted, because they know it is coming. It is also the fastest assessment you will ever do — you learn more in two minutes of questions than from four pages.

Two honest limits

This costs time you may not have. With 60 students you cannot redesign everything. So do not. Pick the two or three pieces in the term that carry real weight and redesign those. Let low-stakes practice stay low-stakes, and stop policing it.

"AI-proof" is not a stable property. A task that defeats the model this year may not next year, and anyone who sells you a permanent solution is selling something. What is stable is a task about something only this student, in this place, this week, could produce. Build on that and you are not chasing a moving target.

The other direction

Some tasks should assume the tool and assess what is left.

Ask an AI to answer this question. Print or copy its answer. Find three things wrong with it — factual errors, missing steps, or claims it cannot support. Correct them, and explain how you know.

This is harder than the original essay, not easier. You cannot find the errors without the knowledge you were going to be assessed on, and it teaches the habit you most want them to leave school with. It also works well for students without much access, since one printed transcript serves a whole class.

BEFORE YOU MOVE ON

A teacher changes "Write about a book you read" to "Write a 2000-word analysis of the symbolism in a novel of your choice, using at least five quotations." Is the new task resistant to AI?

- A. No. It is longer and more formal, which is exactly the kind of writing the model produces most easily, and nothing in it requires anything the model lacks access to.
- B. Yes, because the quotation requirement forces real engagement with the actual text, which the model cannot fake.
- C. Yes, because 2000 words of sustained argument is beyond what AI produces coherently.
- D. Partly, since free choice of novel means some students will pick books obscure enough that the model knows little about them.

The answer and why: p. 340

Lesson 36 · 9 min

Teaching students to use it honestly

THE ONE THING TO KEEP

Set a permitted level per task and require one line of disclosure; bans only make use invisible.

Why a blanket ban does not do what it looks like it does

A ban does not remove use. It removes disclosure. Students still use it, alone, badly, with nobody to tell them when it has confidently misled them — and the ones who suffer most are the ones who trusted it.

A free-for-all is not the answer either. Some tasks are the exercise. Handing them over is like paying someone to do your push-ups, and the students who most need the struggle are the ones most tempted to skip it.

What you need is a rule that changes by task, because the honest question is never "did you use AI". It is "did you learn the thing this task existed to teach".

Put the level on every task

Write three levels on the wall at the start of term and mark one on every piece of work you set.

Level 0 — none. This task is the exercise. Doing it is the point.

Level 1 — assistant. You may use it to explain something you do not understand, to check spelling, or to get unstuck. Not to produce anything you hand in.

Level 2 — collaborator. You may use it to draft. You must show what it produced, what you changed, and why.

This does three things at once. It removes ambiguity, so nobody can claim they did not know. It makes the rule enforceable, because "Level 0" is a specific instruction rather than a mood. And it teaches the real judgement — that the right amount depends entirely on what is being learned.

Make disclosure normal by doing it first

Require one line: what tool, what for, what you changed.

Then model it. Tell the class when you used AI to generate their practice questions. Tell them what you had to fix. That single act does more for the honesty norm in your room than any policy document, because it demonstrates that using it is not shameful and hiding it is what is odd.

Teach verification, not rules

The skill that matters is catching it being wrong. Rules do not teach that. One exercise does.

Pick a question in your subject where you know the answer precisely and the model tends to slip — a multi-step calculation, a date, a local historical detail, a species that only lives in your region. Get the answer in advance, on paper. Read it to the class as though it were authoritative. Let them accept it. Then show them the truth.

The point is not that it was wrong. The point is that it was wrong in exactly the same voice it uses when it is right. A student who has personally been taken in once, in a safe room, treats the tool differently for the rest of their life. Do this in the first fortnight.

The argument they will make, and the answer

"I will have AI at work anyway, so why learn this."

They are partly right, and saying otherwise loses them. The answer is not denial. It is this: the people who get value from these tools are the ones who can tell when the output is wrong, and that judgement only comes from having learned the thing yourself. You cannot supervise work you do not

understand. The person who never learned to write cannot tell a good draft from a fluent bad one, and that person is not more employable, they are more replaceable.

The access problem, said plainly

Some students have a paid account and a good phone. Some have a shared handset and patchy data. Some have neither.

Any policy that assumes access builds a two-tier classroom, where the advantage compounds quietly. So: no task should require a student to have their own AI access to complete it. If a task benefits from AI, either provide the output on paper for everyone, or make it optional and ungraded. This constraint is not a compromise. It keeps the assessment about what students know.

BEFORE YOU MOVE ON

Two classes. Class A has a strictly enforced ban. Class B uses per-task levels with a required one-line disclosure. In Class B, far more students report using AI. The head concludes Class B has more cheating. What is wrong with that conclusion?

- A. The figure counts disclosures, not use. A ban makes use invisible rather than absent, so the two numbers are not measuring the same thing and cannot be compared.
- B. Nothing is wrong. Permitting a tool does increase how much it gets used, so the head is reading the numbers correctly.
- C. It is wrong because Class B's use was permitted, so by definition none of it was cheating.
- D. It is wrong because students self-reported, and self-report is unreliable, so neither figure tells you anything.

The answer and why: p. 340

Lesson 37 · 9 min

Writing the class policy

THE ONE THING TO KEEP

A class AI policy is one page, written with the students, that says what is permitted on each kind of task, what disclosure looks like, what happens on the first and second breach, and how a student appeals — because a rule that lives in the teacher's mood is not a rule anyone can follow.

Why write it down

The earlier lesson set the idea: three levels, marked on every task, and a line of disclosure. That is the spine. This lesson is the document around it, because a spine without a document has two failures. Students cannot check a rule that exists only in what the teacher said in September. And when a case arises, a teacher without a written policy is improvising a punishment, and improvised punishments land unevenly.

One page. Written in plain language. Given out, put on the wall, and referred to every time a task is set.

What goes on the page

1. What "AI" means here. Name the tools. "Chatbots such as ChatGPT, Gemini, Claude, Copilot and Meta AI; image generators; any tool that writes text or solves problems for you." Then the exemptions, because the check question is real: "Spelling and grammar checking built into your word processor is allowed at every level. Translation of individual words is allowed at every level. A calculator is allowed where the task says so." Without the exemptions, every student is in breach every day, and the policy is unenforceable from the first week.

Figure 5.3

The one page, in six parts

What “AI” means here

The tools named, and the exemptions: spellcheck, single-word translation, the calculator where the task allows it.

The three levels, and which tasks are which

So a student can look at a task and know the level without asking.

What disclosure looks like

The exact format, with one worked example. A format that is shown is a format that gets used.

What happens

First time: redo under supervision, no mark penalty. Second: the same plus a note home. Beyond: the school's procedure.

How to appeal

One sentence. Its existence changes how the whole page is read.

What the teacher promises

No detectors. The level marked on every task. Told when AI made your materials.

Without the exemptions in the first part, every student is in breach from the first week and the policy is unenforceable. Without the fifth, students experience the page as a threat rather than a process. The sixth is the modelling written down, so it is a commitment and not a mood.

2. The levels, and which tasks are which. The three levels from the earlier lesson, in one line each, and then the mapping: "In-class writing, tests and the first draft of any assessed essay: Level 0. Homework practice and revision: Level 1. Research projects and redrafts: Level 2, with disclosure." Students should be able to look at a task and know the level without asking.

3. What disclosure looks like. The exact format. "At the end of the work, one or two lines: which tool, what you used it for, what you changed." Give an example: "Used Gemini to suggest three counter-arguments; used one, rewrote it in my own words." A format that is shown is a format that gets used.

4. What happens. The consequence ladder, and it should be about the learning, not the punishment:

- **First time:** the work is redone under supervision, and the student and teacher talk about what the task was for. No mark penalty on the redone work.
- **Second time:** the same, plus a note home.
- **Beyond that:** the school's academic honesty procedure, which the class policy points to but does not replace.

Redoing the work is the consequence that fits, because the harm of handing over the exercise is not having done it. A zero teaches that the teacher is angry; a redo teaches the thing.

5. How to appeal. "If you believe a decision is wrong, say so to me first; if we do not agree, either of us can ask the head of department to look at the work and the evidence." One sentence. Its existence changes how students experience the whole policy, from a threat to a process.

6. What the teacher promises. "I will not use AI detectors. I will mark the level on every task. I will tell you when I have used AI to make your materials." The last one is the modelling from the earlier lesson, written down so that it is a commitment and not a mood.

Write it with them

A policy written with the class is followed more than one handed down, and not because of a vague idea about ownership: students who have argued about where the line falls know where it falls. Take one lesson. Put the six headings on the board. Ask, for each, what the rule should be and why. Students are stricter than teachers expect on tasks they see as "the exercise", and more

precise than teachers expect about the tools they actually use. Write the page from the discussion, bring it back the next day for agreement, and have every student keep a copy.

Align it with the school

If the school has a policy, the class policy sits inside it and says so. If the school has none — common — the class policy is a draft the school can adopt, and a later block in this course covers the institutional version. What a teacher should not do is write a class rule that contradicts a school rule, because the first case that reaches a head teacher will be decided by the school's document, and the class policy will have promised something it could not deliver.

Keep it current

Products change, and the page should carry a date and a promise to revisit it each term. The exemption list in particular will need editing: a tool that was a grammar checker in September is a rewriting assistant by March, and the policy has to say which side of the line it now sits on.

The honest limit

A written policy makes the rule clear. It does not make enforcement certain, and the detection lesson explained why it never will be. What the policy does is make the honest majority's position safe — they know what is allowed and can show they did it — and make the case that does arise a matter of process rather than accusation. That is the most any policy can do, and it is a great deal more than a ban.

BEFORE YOU MOVE ON

A teacher's policy states: "Any use of AI without permission will be treated as cheating." A student uses a grammar checker built into her word processor, which now uses a language model. Is she in breach, and what does the case reveal about the policy?

- A. Yes, she is in breach, and the case shows that students need to read the terms of the tools they use far more carefully than they do.
- B. No, she is not in breach, because a grammar checker is not really AI in the sense that any reasonable policy is concerned with.
- C. Unclear, and that is the failure: the policy names a category without defining it.
- D. Yes, technically, but the teacher should let it go this once, which is exactly what the discretion built into any policy is for.

The answer and why: p. 341

Lesson 38 · 9 min

Process evidence that holds up

THE ONE THING TO KEEP

Version history in Google Docs or Word, an in-class first paragraph, and a short drafting portfolio make the writing process visible for every student as a matter of routine — signals of authorship, never proof, and used to support students rather than to trap them.

Product tells you nothing; process tells you something

The detection lesson made the case that the finished text cannot tell you who wrote it. The process can, partly. A piece of writing that grew over four evenings, with a plan, a bad first paragraph, and a rewrite after feedback, has

a shape that is hard to fake and easy to show. This lesson is about making that shape visible for every student as a routine, so that it is a support for the honest majority and not a trap for the suspected few.

Version history, and how to read it

Google Docs and Microsoft Word both keep every version of a document automatically. In Google Docs: File, then Version history, then See version history. A timeline appears on the right with a snapshot at each save; click one and the changes since the previous snapshot are highlighted. In Word with a OneDrive or SharePoint document: File, then Info, then Version History.

What a written-over-time document looks like: many snapshots, small additions, sentences that change, a paragraph deleted and rewritten, gaps of hours or days. What a paste looks like: one snapshot with the whole text arriving at once, or a few large blocks arriving fully formed.

The check question is about what that second pattern means, and the answer is: not much on its own. Students draft on phones, in notes apps, on paper. They paste. A single-paste document is a reason to look further, and the further look is easy: "Show me the phone draft." A student who wrote it has one. A student who did not usually cannot produce one, and the conversation from the detection lesson takes over.

There are browser extensions that replay a Google Doc's history as a film of the writing. They are striking to watch, and a teacher should know they exist, but they add little to the timeline view and they can feel like surveillance if used routinely. The timeline is enough.

The in-class paragraph

For any assessed piece, the first paragraph is written in class, on paper or in the shared document, in twenty minutes. It is then the seed the rest grows from. This does three things. It gives you a sample of the student's own voice for this piece, on this day. It makes the finished essay have to connect to a beginning you watched being written. And it gets the hardest part of any essay — starting — done in a room where help is available.

Nothing about this is anti-AI. It is good writing pedagogy that also happens to make process visible.

The drafting portfolio

For the two or three pieces a term that carry weight, ask for a folder rather than a file: the plan (five lines), the first draft, the feedback received, the final draft, and a note of what changed and why. The assignment-redesign lesson made the process the product; this is the container for it. Marked as a whole, with the revision weighted, it rewards the thing you want and it is more work to fake than to do.

The "what changed and why" note is the part students find hardest and the part that tells you most. A student who cannot say why paragraph three moved did not move it.

Making it routine

Every signal here works only if it is collected from everyone, every time, as the normal shape of the course. Version history checked only for the student you suspect is an accusation. Version history glanced at for every submission, as the way you look at work, is nothing. The same for the in-class paragraph and the portfolio. Announce it in week one, do it for the whole class, and the honest student has an easy way to show her work and the tempted student knows the work will be visible.

The limits, stated plainly

- A student can type a generated text in by hand over an hour and produce a plausible history. It is more effort than writing the essay, which is the point of every design in this block, but it is possible.
- A student who genuinely drafts on paper has no version history, and must not be penalised for it. Paper drafts are process evidence too. Ask for them to be kept.
- Version history shows where text arrived, not where it came from. It is a signal, and the lesson's title says "holds up", not "proves".

- Students with shared devices, no device at home, or no reliable connection produce thinner histories through no fault of their own. Judge the pattern, not the gaps.

The honest limit

Process evidence makes the honest case easy to show and the dishonest case more work to construct. It does not resolve the case in the middle — the student who wrote a real draft, then fed it to a model, then pasted back a version that is half hers. That case is the ordinary one now, and it is what the disclosure line and the Level 2 rule exist for: if she declares it, it is permitted; if she does not, it is a breach of the policy, not a mystery for forensics. The evidence supports the conversation. It does not replace it.

BEFORE YOU MOVE ON

A teacher checks version history and finds a student's essay arrived in a single edit of 1,200 words at 11:40pm. The student says she wrote it in a notes app on her phone and pasted it. What is the right reading of the evidence?

- The single paste is proof of AI use, because genuine writing always shows incremental edits over several sessions in the history.
- The paste is a reason to ask for the phone draft, and on its own it decides nothing.
- The student's explanation clears her completely, since it accounts for the paste and there is nothing else in the history to question.
- Version history is too unreliable to be looked at, since students draft in many places, and the teacher should not have opened it.

The answer and why: p. 341

Lesson 39 · 9 min

Handling a case fairly

THE ONE THING TO KEEP

When a case reaches a decision, the standard is what a reasonable colleague would conclude from the evidence you can show — the process record, the conversation, the student's own explanation — never a detector score, and the default outcome is redoing the work, with the second-language writer given explicit protection from the accusation that reads unsurprising prose as machine prose.

The question is not "was it AI"

Most of the anxiety in a suspected case comes from trying to answer the wrong question. Whether a specific paragraph came from a model, a cousin, an older sibling or a website is usually unknowable and usually irrelevant. The policy's question is whether the student did the work the task existed to make them do. That question has evidence. The other one has guesses.

Reframe every case that way and most of them become manageable.

What counts as evidence

In rough order of weight:

1. **The student's own account.** From the two-minute conversation the detection lesson described. Not whether they confessed — whether they can talk about the work as someone who did it. Can they explain why they chose this example, what they cut, what they found hard.
2. **The process record.** Version history, the in-class paragraph, the plan, the drafting folder. Does the finished piece connect to a visible beginning.
3. **Comparison with their other work.** Not "this is too good for her" — students improve, and that judgement is the one most likely to be wrong and most likely to fall on the student whose ability you underestimated.

But: a piece that argues in a way the student has never argued, on a topic she showed no grasp of in class last week, is a reason to look at the other evidence more carefully.

4. **Specific features of the text** — a citation that does not exist, a reference to a resource the class never used, an American spelling throughout in a British-spelling class. Weak on their own, useful with the rest.

Not evidence: a detector score. The detection lesson explained why, and the class policy promised not to use one. If a colleague or a head asks for one, the answer is that it would add a number with a known bias against second-language writers to a decision that has real evidence available.

The standard

The standard is what a reasonable colleague, shown the evidence, would conclude. Not certainty — you will rarely have it — and not suspicion. If two teachers looking at the conversation notes and the process record would both say "she did not write this", act. If one would and one would hesitate, the evidence is not there, and the outcome is different.

Write the evidence down before deciding. A page: what the task was, what level it was set at, what the process record shows, what was said in the conversation, what the student says. A decision made after writing that page is better than one made before it, and it is the page the appeal will be judged on.

The outcomes

The class policy set the ladder, and the first step on it is the outcome for nearly every first case: **redo the work, under supervision, and talk about what the task was for.** It fits the harm, it is not humiliating, and it produces the learning the original task was meant to.

When the evidence is short of the standard, the outcome is still available in a softer form: "I could not follow your argument when we talked. Write me the second half again, here, and let's see." Not an accusation, not a penalty. A

student who wrote the original will produce the second half easily. One who did not has redone the exercise without a finding against them, which is the right result when you were not sure.

Escalation — the school's formal procedure, a note on a record — is for the second and third time, or for the high-stakes piece where the school's rules require it. Not for the first case, and not on a teacher's own authority.

The second-language safeguard

Put this in writing in the department: a student writing in a second or third language is never referred on the basis of the text alone. Their prose is the prose that every automated and half-automated judgement misreads, and the harm of a wrong accusation to a student already working in a language not their own is severe and lasting. For those students the process record and the conversation carry the whole weight, and the conversation is held in a way that gives them room — in their strongest language if a colleague can help, with time to think, without an audience.

The meeting

If a parent is involved, the meeting is about the evidence page, the policy the student agreed to, and the outcome, in that order. Not about AI. Not about the character of the child. "The task was Level 0. The process record shows the essay arrived at once. When we talked, she could not explain the argument. Under the policy she agreed to, she redoes it with me. No mark penalty. Here is the policy." Most parents, shown that, are on the teacher's side, because the process is fair and the outcome is mild.

What not to do

Do not raise it in front of the class. Do not open with "did you use AI". Do not use words like cheating or plagiarism in the first conversation. Do not compare the student with a sibling. Do not act on a detector. Do not decide before writing the page. Each of these is a way to be wrong loudly, and a loud wrong accusation costs a child more than a missed case costs anyone.

The honest limit

Some students will get away with it. That has always been true and it is more true now. The alternative — a regime that catches more by suspecting everyone — is worse, and the detection lesson showed who it catches. A fair process that occasionally misses is the right trade, and a teacher who has designed the high-stakes tasks so that the work has to happen in the room has already moved most of what matters out of reach of the question.

BEFORE YOU MOVE ON

After a conversation, a student cannot explain the argument in her essay, the version history shows a single paste, and she says a cousin helped. The teacher has no other evidence. What is the fair outcome?

- A. Record it formally as AI cheating, since the single paste and the failed conversation taken together are conclusive evidence.
- B. Take no action at all, since without proof that a model was used there is nothing in the policy to act on.
- C. Have her redo the piece under supervision, record the evidence, and leave the tool question out.
- D. Run a detector on the essay to settle whether the cousin or a machine wrote it, and then decide the outcome on that basis.

The answer and why: p. 342

Lesson 40 · 9 min

Rethinking homework

THE ONE THING TO KEEP

Homework that asks for a product a model can produce — summaries, comprehension answers, definitions, essays — has stopped measuring anything; keep homework that is retrieval, practice with self-checking, reading with a one-line response, or evidence from the student's own home and street.

What homework was for

Strip away the tradition and homework has a few defensible purposes: practice of a skill taught in class, so that it becomes fluent; retrieval of material learned earlier, so that it stays; reading or preparation that a lesson will build on; and, occasionally, work that can only be done outside school — an observation, an interview, a measurement at home. The research on homework has always been lukewarm about anything else. Large reviews find little effect in primary school and a modest one in secondary, and the effect that exists comes from practice and retrieval, not from producing things.

That matters now because the model has made one whole category of homework — producing things — pointless as a measure of anything, and the lukewarm evidence says that category was never where the value was.

The tasks that have stopped working

Be honest about the list. If a task is on it, the excellent homework coming back tells you nothing, and the check question's falling scores are the consequence.

- **Summarise the chapter.** The model summarises better than any student.
- **Answer the comprehension questions at the end of the chapter.** The chapter is in the training data; so, often, are the questions.
- **Define these terms.** Instant.

- **Write an essay on...** The assignment-redesign lesson covered this for assessed pieces. For homework it is worse, because there is no in-class paragraph to anchor it.
- **Research X and write a report.** Search plus summarise, in one prompt.
- **Translate this passage.** Any phone.

Setting these still, with a ban attached, produces a class where the compliant students do a task that no longer teaches them anything a lesson would not, and the rest paste. Neither group learns from the homework. Stop setting them.

What replaces them

Retrieval, answered in the next lesson. The homework is "revise these three topics". The check is five questions on paper at the start of the next lesson, answered from memory. There is nothing to generate, the learning happens in the revising, and you find out in five minutes who did it. This is the single best replacement and it costs no marking beyond a glance.

Practice with a self-check. Twenty problems with the answers at the bottom of the page. The task is not the answers — they are given. The task is the working, and the student marks their own. A student who copies the answers has gained nothing and knows it; a student who works them has practised. Pair it with a starter next lesson on the same skill, and copying becomes visible without any policing.

Reading with a one-line response. "Read pages 40 to 44. Bring one sentence you did not understand." Or one question you would ask the author. The model can read for them; it cannot tell you what they did not understand. And the sentences they bring are the starter for the next lesson.

Evidence from home. Measure the rainfall, interview a grandparent about prices, count the vehicles on the road at 8am, photograph three shop signs in two languages. The assignment-redesign lesson called this anchoring to the room; for homework, the anchor is the street. A model cannot do it and a student who does it has data nobody else has.

Preparation, then discussion. Watch, read or listen to something short; the lesson is a discussion of it. Students who did not prepare are visible in the first two minutes and no product was ever asked for.

What about the practice essay

Older students preparing for written exams do need to practise writing under conditions. The honest answer is that this practice now belongs in class, timed, on paper, or in the version-history document the process lesson described. Homework can be the plan for it — five bullet points, a thesis sentence — and the in-class session is the writing. That is not a loss. Exam writing practised in exam conditions is practised properly.

Less of it

A consequence of all this is less homework, and better. Teachers who make the change report that they set roughly half what they used to and mark a quarter of it, because retrieval checks mark themselves and self-checked practice needs a glance. Parents sometimes read less homework as less rigour. Explain the change in a letter: what was dropped and why, what replaced it, and what they will see instead — a child revising for a quiz rather than typing a summary.

The honest limit

Retrieval and practice do not cover everything. The extended piece of writing, the sustained project, the thing that takes a student a weekend to make — those are real and they still belong somewhere, and the somewhere is now the classroom, the portfolio and the supervised session rather than the kitchen table. The tradeoff is class time, and the earlier lessons on generating materials are partly about winning back the time to afford it.

BEFORE YOU MOVE ON

A teacher sets "read chapter 5 and write a 300-word summary" every week. Since AI arrived, the summaries are excellent and test scores on the chapters have fallen. What is the mechanism, and what should replace the task?

- A. Students have got lazier since the tools arrived, and a stricter penalty for undeclared AI use would restore the learning.
- B. The task rewards a product the model makes well, so it no longer causes reading; replace it with in-class retrieval.
- C. Summaries are a poor task in themselves and should be replaced with longer essays that demand more of the students.
- D. The chapter is too hard for the class, and the summaries were always being copied from somewhere before AI existed.

The answer and why: p. 342

Lesson 41 · 8 min

Citing and declaring AI use

THE ONE THING TO KEEP

Older students will meet institutions that require a declaration of AI use, and the conventions exist — APA, MLA and Chicago each have one — so teach the declaration as a normal line of academic honesty, while being clear that the underlying question of who authored a machine-assisted text is unsettled everywhere.

What is coming for them

A student who leaves your school for a university, a college or a professional course will meet a policy on AI use, and most of those policies now require a declaration: a statement, with the work, of what tools were used and for what.

Some institutions ban use on some tasks; almost all require honesty about it. A student who has never written a declaration, and who thinks of AI use as something to hide, arrives at the wrong starting point.

The disclosure line from the earlier lesson — tool, purpose, what changed — is the school version of this. The university version is longer and has conventions, and teaching them in the last two years of school is a small, concrete service.

The declaration

A declaration is a statement of use, separate from the references. It says what was used, for what, and what remains the student's own. The check question's case:

I used Gemini to generate a possible outline for this essay and used two of its five suggested sections. The argument, the reading and every sentence are my own. I used the same tool to check grammar and punctuation after the final draft; it suggested twelve changes and I accepted nine.

Specific, short, and it tells the reader what they need to know. The reader can now judge the essay as an essay whose structure was partly suggested and whose prose is the student's. That is different from an essay where the argument was generated, and the declaration is what lets the reader tell.

The declaration is not a confession. Teach it as the same act as a bibliography: the reader is entitled to know what the work rests on.

The conventions, briefly

The major style guides added guidance in 2023, and they differ, which is itself a thing students should know.

- **APA** treats generated text you quote as something to cite in text and in the reference list — the company as author, the model and version, the date, and the prompt described in the text. It also expects a description of how the tool was used in the method or introduction for research writing.

- **MLA** treats generated content as a "container" — the prompt as the title of the source, the tool as the container, the date — and says to acknowledge any use that shaped the work, not only quoted text.
- **Chicago** uses a note: the tool, the company, the date, and the prompt; and says generated text should not appear in the bibliography as a source.

None of them treats the model as an author, because an author can be responsible for a claim and a model cannot. That principle is the thing worth teaching; the formats will change and a student can look them up.

What to declare

A useful rule for students: declare anything that, if a reader found out later, they would feel they should have been told. That covers an outline, a first draft, a translation, a summary of a source they did not read in full, a set of counter-arguments, code, a diagram. It does not cover a spelling check or a dictionary lookup, and the class policy's exemption list says so.

And a corollary: if a student would be embarrassed to declare a use, that is a sign the use was not permitted at the task's level, and the time to find that out is before submission.

Teaching it

Require the declaration on every Level 2 task from Class 10 upward, in the school format. Show two or three good ones. Mark it as present and adequate, never for content — a student who declares heavy use on a Level 2 task has done nothing wrong, and marking the declaration down would teach them to hide next time.

Then, once, have students write the declaration for a piece before they start it — a plan of what they intend to use the tool for. Comparing the plan with the final declaration is an exercise in noticing how use creeps, and it is a better lesson on judgement than any rule.

The unsettled question

Underneath all this is a question nobody has answered: when a student prompts, selects, edits and arranges text that a model produced, who wrote it? Copyright offices in different countries have taken different positions on whether such work can be owned at all. Universities differ on whether declared heavy use is acceptable. Journals differ on whether a model can be acknowledged. A student should know that the adults have not settled this, that the rules they meet will vary by institution and by year, and that the one stable thing is the declaration itself: say what you did, and let the reader judge.

The honest limit

A declaration is only as honest as the student, and a student who generated the whole essay can declare a grammar check. The declaration is not a detection tool and was never meant to be. Its value is to the honest student, who now has a normal, expected way to use the tool openly, and to the reader, who is told what they are reading. For the other case, the rest of this block applies.

BEFORE YOU MOVE ON

A Class 12 student used a model to generate an outline, wrote the essay herself, and then used the model to check grammar. Her school requires a declaration. Which is the right one?

- A. No declaration is needed, because she wrote every sentence herself and an outline is not part of the essay that is being assessed.
- B. "This essay was written with the help of AI," because any use at all must be acknowledged and that sentence covers it.
- C. A line naming the tool, the outline and the grammar check, and stating the rest is her own.
- D. A full citation for the model in the reference list, formatted as though it were a source she had consulted for the essay.

The answer and why: p. 343

CASE STUDY · MODULE 5

Ninety-two per cent, and a parent already writing

An illustrative case, built from documented patterns.

Class 11 English in a private school in Kochi, three days before the coursework deadline, with a detector score circulating among staff.

The essay was on Silas Marner and it was good. It was better than anything Fathima had submitted that year and it argued a line — that Eppie's arrival is written as an economic event before it is written as an emotional one — that had not come up in class.

A colleague, unprompted, had run it through a free detector and forwarded the result to the head of department and to Priya Nair, who taught the class: ninety-two per cent AI-generated. By the time Priya read the message, four members of staff had seen the number, and Fathima's father had emailed the school asking why his daughter had been called out of a lesson.

Nobody had called her out of a lesson. She had gone to the toilet. But the email had arrived, and it made the case urgent in a way the evidence did not.

What the evidence was

The class policy, written with the students over one period in August and signed by all of them, said the first draft of any assessed essay was Level 0 and named exactly what the consequences were: the work is redone under supervision, there is no mark penalty on the redone work, and a second occasion brings a note home. The students had been stricter than Priya expected about which tasks counted as the exercise, and more precise than she expected about which tools they actually used.

It also contained a line Priya had insisted on and now had to live up to: I will not use AI detectors.

What she had was a version history, because the essay had been written in the school's shared documents. It showed the essay arriving in four large blocks over two evenings, which is what a paste looks like and also what a student

who drafts on her phone and pastes looks like. It showed no small edits at all.

She also had an in-class paragraph. Every assessed piece in her class began with twenty minutes of writing in the room, and Fathima's opening paragraph, in her own hand, contained the economic reading of Eppie's arrival in almost the words the essay used.

The head of department wanted a decision by Friday and thought the ninety-two per cent settled it. He was not being unreasonable; he had a deadline, a parent, and a number that looked like evidence.

Priya's position was that the number was not evidence, that the tool it came from is known to flag careful second-language writing, that Fathima writes in her third language, and that the version history was ambiguous in both directions. She also knew a piece of arithmetic that nobody in the staffroom had done: a detector with even a one per cent false positive rate, run over every piece of homework from thirty students, produces a wrong flag against an innocent child about one week in three.

And there was the part of the case that had nothing to do with evidence. Four members of staff had seen the number before the teacher of the class had. Whatever the finding turned out to be, a fifteen-year-old's name had already travelled, and a decision that took a fortnight would let it travel further. Speed was not only the head of department's interest. It was Fathima's too.

What Priya did not have was certainty, or a colleague who agreed with her.

YOUR ANSWERS

1. The version history showed four large blocks and no small edits. Why was that not evidence of anything on its own?

2. Priya's policy promised she would not use detectors. What did keeping that promise cost, and what did it buy?

3. Suppose the phone notes had not existed. What should the outcome have been?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 206.

CASE STUDY · MODULE 5

How it played out

She wrote the page before deciding: what the task was, what level it was set at, what the process record showed, what was said in the conversation, what the student said. Then she had the conversation, without the word AI in it — tell me how you got to this argument — and Fathima talked for six minutes, said she had found the economic reading in a YouTube video about the novel, named the channel, and explained why she had cut a paragraph about Godfrey that had not worked. She also said she wrote on her phone at night because the family computer was her brother's until eleven. Priya asked to see the phone notes. They were there, timestamped, messy, with the abandoned Godfrey paragraph in them. The finding was that Fathima wrote it. The head of department accepted the page, though not happily. The parent meeting happened anyway and was about the evidence page and the policy, in that order, and took eleven minutes. Two things changed afterwards: the department agreed in writing that no student is referred on the basis of the text alone, and Priya added a line to the policy saying that a phone draft counts as process evidence, because most of her class writes that way and the version history had nearly punished them for it.

The questions, discussed

1. The version history showed four large blocks and no small edits. Why was that not evidence of anything on its own?

Because it shows where text arrived, not where it came from. A student who drafts on a phone, on paper, or in a notes app and then pastes produces exactly that pattern, and in this class most students draft that way for reasons of access rather than concealment. A single-paste document is a reason to look further, and the further look — show me the phone draft — is cheap and decisive. Treating the pattern itself as a finding would have punished the students with the least access to a computer.

2. Priya's policy promised she would not use detectors. What did keeping that promise cost, and what did it buy?

It cost her the fastest route to closing the case and put her at odds with a head of department under deadline pressure, and it meant she carried the decision alone. It bought a decision that survived a parent meeting, because the evidence page contained things anyone could examine — a policy the student had agreed to, an in-class paragraph, a conversation, timestamped notes — rather than a number with a known bias and no appeal path. It also kept the policy credible for the rest of the class, who would have learned in a week that the teacher's written promises hold only until a case is inconvenient.

3. Suppose the phone notes had not existed. What should the outcome have been?

Not a finding. The standard is what a reasonable colleague would conclude from the evidence that can be shown, and with an in-class paragraph carrying the same argument and a fluent six-minute account of her own choices, two teachers would not both have concluded she did not write it. The right move where the evidence falls short is the softer one: ask her to write the second half again, in the room, and see. A student who wrote the original does it easily and has redone the exercise with nothing recorded against her, which is the correct result when you are not sure.

MODULE 5 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. 338. If two go wrong, re-read the module before the exam.

- 1 A detector claims a genuine one per cent false positive rate. You have thirty students and you run every piece of homework through it. What follows?
 - A. No wrong flags, because one per cent of thirty is less than one whole student.
 - B. A wrong flag against an innocent student roughly one week in three, week after week.
 - C. About one wrong flag a term, which is a reasonable price for the deterrent effect.
 - D. No wrong flags against honest students, because the one per cent describes missed detections among AI-written work rather than false accusations.

- 2 Why do detectors flag careful writing by second-language students?
 - A. The detectors were trained mostly on English essays and treat any non-native writing as unfamiliar.
 - B. Such students are more likely to have used a translation tool, and the detector picks that up.
 - C. Such writing is unsurprising — safe vocabulary, familiar sentence shapes — and unsurprising is what the detector measures.
 - D. Such writing contains more grammatical errors, and the detectors were built to read error-free text as human and error-laden text as machine-written.

- 3 Which redesign of an essay task actually puts it out of the model's reach?
- A. Requiring two thousand words rather than eight hundred.
 - B. Requiring a more abstract and theoretical treatment, which is harder for a student and therefore harder for the model.
 - C. Requiring five academic sources and a formal structure with numbered sections.
 - D. Requiring the data the class collected on Tuesday.
- 4 Why does a blanket ban on AI make a class worse off than a per-task level with disclosure?
- A. It removes disclosure rather than use, so the students misled by it have nobody to tell them.
 - B. It denies students practice with a tool they will be expected to use at work.
 - C. It is unenforceable, and an unenforceable rule teaches students that rules in general do not matter.
 - D. It pushes students toward paid tools that are harder to detect, which widens the gap between those who can afford them and those who cannot.
- 5 A student's essay appears in the version history as one large paste at eleven at night. What does that establish?
- A. That the text was generated, since genuine writing accumulates in small edits over time.
 - B. That the student worked outside the school's systems, which breaches the policy whoever wrote the text.
 - C. That the student copied from somewhere, though the history cannot tell you whether the source was a website or a model.
 - D. Very little on its own, because many students draft on a phone or on paper and then paste.

- 6 You are about to decide a suspected case. What standard applies?
- A. Certainty that the student did not write the work.
 - B. A detector score above the threshold the school has agreed to use.
 - C. Your own professional judgement alone, since you know the student's usual standard better than any process could.
 - D. What a reasonable colleague, shown the evidence, would conclude.

6

MODULE 6

Teaching AI itself

Teach a unit on AI to any age from 8 to 18 without a device — prediction, training data, classification, bias, confident error and the deepfake problem — and read a study on AI tutoring well enough to say what it found, on whom, and what it means for how students in your class should be allowed to use the tool.

42	An AI lesson with no screen and no power	212
43	What to teach at which age	215
44	The sorting game for younger children	220
45	Image generators and deepfakes	225
46	The AI companion problem	228
47	A student project that teaches the mechanism	232
48	The evidence on AI tutors	236
49	AI in every subject, not just computing	241
	<i>Case study: Four hundred free licences and a chart</i>	245
	<i>Module 6 test</i>	250

Lesson 42 · 9 min

An AI lesson with no screen and no power

THE ONE THING TO KEEP

The mechanism — prediction, confident error, learned from us — can be taught with chalk and voices.

The part that matters needs no electricity

Almost everything worth teaching about AI is about prediction, pattern and confident error. All three can be taught with chalk, paper and thirty voices. The button-pressing is the easy half and can be picked up in an afternoon whenever a device appears.

So do not treat the paper version as the fallback. Build the lesson on paper and treat a working projector as a bonus. Reverse the default and you will never lose a lesson to a power cut again.

Activity one: the class becomes the model

No materials at all.

Write a sentence stem on the board. Every student writes one word to complete it. Silently, no discussion, everyone commits before anyone hears anyone else.

The old woman walked slowly into the _____

Tally by show of hands. Room, house, shop, kitchen, one "sea". You now have a distribution on the board. Tell them: the machine has the same kind of guess, built from billions of sentences instead of thirty.

Second stem: The capital of Japan is _____. Nearly unanimous.

Ask them why one was a scatter and one was not. From there, two conclusions arrive by themselves. Why the machine gives different answers to the same question if you ask twice. And why it is dependable about Tokyo and

unreliable about the population of your district.

Twenty minutes, no equipment, and they now hold the actual mechanism rather than a picture of a robot.

Activity two: where its habits come from

Give one student a short story to read, then have it retold down a line of five. What survives is what each person found worth remembering.

Then the teaching question: if nearly every story ever written about doctors used "he", what word will a machine guess after "the doctor said that"? Let them answer. Then ask what that means for a machine that learned from everything people have written.

You have taught bias in training data in ten minutes, with no equipment, and they will not forget it because they were the training data.

Activity three: the confident wrong answer

This needs a phone with data once, not a phone in class.

Whenever you have connection — at home, early morning, at an internet cafe — ask an AI five questions from your syllabus. Write down its answers word for word, and write down the truth. Choose questions where it slips: local history, a district detail, a multi-step calculation, a plant that grows only in your region.

In class, read the five answers aloud as though they were authoritative. Students vote true or false on each. Then reveal.

The lesson is not that it got two wrong. The lesson is that the two wrong answers sounded exactly like the three right ones. That is the thing you want them to carry, and it lands harder on paper than on a screen, because nothing distracts from the words.

Practical notes from teachers who work this way

- **A printed transcript is a real teaching resource.** One conversation, printed once, photocopied, teaches as well as a live demonstration and never fails halfway through.
- **Batch the online work.** Do all your generating in one session when power and data are available. Bring paper. Never build a lesson whose spine depends on connectivity arriving at 10am.
- **One phone at the front beats forty half-charged ones.** If devices are scarce, a demonstration you control is better than a scramble, and it keeps the class in one conversation.
- **Write the answers down before class.** Live demonstrations fail, models change their answers, and a lesson built on "watch what it says" has no plan B. A lesson built on "here is what it said, and here is what was true" always works.

The honest worry

There is a real risk that this becomes another line of inequality — schools with connectivity teaching it hands-on, schools without teaching it in the abstract.

But look at what actually transfers. That it predicts. That it is confident when it is wrong. That it learned from us and carries what we wrote. A student who holds those three things and has never touched a keyboard will use the tool better, the first day they get one, than a student who has typed into it for a year without understanding any of it.

BEFORE YOU MOVE ON

A teacher with no reliable power runs the word-prediction activity with chalk and a show of hands. A colleague says the students have not really learned about AI because they never used it. What is the strongest response?

- A. The activity teaches the actual mechanism — a guess at likely continuations — which is what makes the machine's varying answers and confident errors predictable. Operating the interface takes an afternoon to pick up later.
- B. The colleague is right, but in a school without power there is no alternative, so this is the best available substitute.
- C. Using the tool is what matters, so the school's priority should be acquiring devices before anything else.
- D. The activity is a metaphor, and a good metaphor is enough for beginners who will meet the real thing later.

The answer and why: p. 345

Lesson 43 · 9 min

What to teach at which age

THE ONE THING TO KEEP

The mechanism scales down further than people expect — pattern and sorting at 8, prediction at 11, training data and bias at 14, how it is built and what it is for at 16 — and no country has a settled AI curriculum yet, so the honest frame is a progression rather than a syllabus.

There is no syllabus yet

Be clear with yourself and your students about this. UNESCO published competency frameworks for students and for teachers in 2024, several countries have added AI to their computing or citizenship curricula since

2023, and none of it is settled. What a 12-year-old "should" know about AI is being decided now, in several places at once, differently. That is not a reason to wait. It is a reason to teach a progression built on the mechanism, which does not change, rather than on a product list, which changes monthly.

The UNESCO student framework is a useful spine if you want one: four strands — a human-centred mindset, the ethics of AI, the techniques and applications, and designing systems — at three levels of depth, understand, apply and create. What follows maps roughly onto it and is organised by what a child of each age can actually hold.

Ages 8 to 10: pattern and sorting

The idea that fits: a machine can learn to sort things by looking at many examples, and it makes mistakes at the edges.

Figure 6.1

A progression built on the mechanism, not on products

- **Ages 8 to 10**
Pattern and sorting. A machine shown labelled examples learns to sort new ones, and gets the fox wrong. Objects and two boxes, no screen.
- **Ages 11 to 13**
Prediction. The sentence stem, the tally, the scatter. Where the words came from, and the confident wrong answer on paper.
- **Ages 14 to 16**
Training data, bias and evaluation. Twenty questions marked into a percentage, then: was that a fair test? Deepfakes and companions belong here.
- **Ages 16 to 18**
How it is built and what it is for. The tally model by hand, and the questions the adults have not settled, argued with sources.

No country has a settled AI curriculum, so this is a progression rather than a syllabus. It is built on what does not change. Three things recur at every stage and deepen: it learned from us, it is confident when wrong, and somebody built it and somebody profits.

Teach it with objects, not screens. The sorting game in the next lesson is the core activity. Children this age already have the concept — they have sorted animals, shapes, words — and the new step is that a machine can be trained to do it by being shown examples rather than told rules. They can also grasp the edge case: the dog that looks like a cat, and why the machine gets it wrong. That is the whole of machine learning at a level a 9-year-old owns.

What does not fit yet: probability, "predicting the next word", anything about how a chatbot works. They will ask; "it has learned patterns from lots of writing, like the sorting game but with words" is enough.

Ages 11 to 13: prediction

The offline lesson's activity — the class as the model, the sentence stem, the tally — is built for this age. They can hold the idea of a spread of likely answers, they can see why "Tokyo" is unanimous and "sea" is rare, and they can draw the consequence: sometimes it is reliable and sometimes it is guessing, and the tone does not change.

Add: where the words came from, and the retelling activity for bias. Add: the confident wrong answer, on paper. That is a complete, honest unit for this age, and a student who has it knows more about the mechanism than most adults.

Ages 14 to 16: training data, bias and evaluation

Now the questions become about people. Whose writing did it learn from? What did it learn about doctors and nurses, about names, about places like ours? What happens when it is used to decide who gets a loan or an interview? How would you test whether it is fair?

This age can run a small evaluation: give the model twenty questions, mark the answers, compute a percentage, and then ask whether the twenty were a fair test. They can compare two tools. They can read a news story about an AI failure and identify which part of the mechanism failed. And the deepfake lesson belongs here, because this is the age at which it happens to them.

Ages 16 to 18: how it is built, and what it is for

The project lesson — building a tiny model by hand — fits here, along with the questions that have no settled answer: what it does to work, what it does to writing, whether it should be used in schools at all, whether it understands anything. Students this age can hold an unsettled question as unsettled, and it is worth showing them that the adults have not resolved it. A debate with real positions, sourced, is a better lesson than a lecture with a conclusion.

For students heading toward technical study, this is where Python, a notebook and a real dataset enter, and the free path is Google Colab or any phone browser with an online Python environment.

Across the ages: three things every year

Whatever the age, three things recur and deepen:

1. **It learned from us.** Its habits are our habits. Its errors about our town are because we wrote little about our town.
2. **It is confident when wrong.** Sounding right and being right are different, and telling them apart is a skill.
3. **Someone built it and someone profits from it.** It is a product, made by a company, with choices inside it. This is the citizenship strand and it is the one most often skipped.

The honest limit

A progression like this assumes time that most schools do not give to a subject with no exam. In practice AI is taught in the gaps: a computing lesson, a PSHE slot, a form period, a science lesson that has a spare twenty minutes. The final lesson of this block is about making every subject carry a piece of it, so that the progression happens across the timetable rather than in a slot that does not exist.

BEFORE YOU MOVE ON

A primary teacher wants to introduce AI to 9-year-olds and plans a lesson on how large language models predict the next word from probabilities. A colleague suggests a sorting game with pictures instead. Which is the better fit for the age, and why?

- A. The prediction lesson, because it teaches the real mechanism and children should never be given simplified versions of how a thing works.
- B. Neither, because AI is a secondary-school topic and teaching it to 9-year-olds only produces confusion that has to be undone later.
- C. The prediction lesson, but delivered with a projector so that the children can watch a real chatbot working as it is explained.
- D. The sorting game, because "learns a pattern from examples, wrong at the edges" is the idea a 9-year-old can hold.

The answer and why: p. 346

Lesson 44 · 8 min

The sorting game for younger children

THE ONE THING TO KEEP

Classification, not generation, is the right first model for a child — a machine shown many labelled examples learns to sort new ones, and gets it wrong at the boundary — and it can be taught with a pile of pictures and two boxes, with Teachable Machine as the free follow-up if a laptop appears.

Why start with sorting

The chatbot is the AI children see, and it is the hardest one to explain. Underneath every AI system that matters to a child's life — the one that tags faces in photos, the one that recommends videos, the one that decides which post they see — is a simpler idea: a machine shown many labelled examples

learns to sort new ones. That idea is teachable at 8. It is also true, and it scales up: the chatbot is a sorter too, choosing the most likely next word from a very large pile.

So the first unit for younger children is sorting, and it needs no electricity.

The game

Materials: forty small pictures cut from magazines or drawn — twenty cats, twenty dogs, clearly labelled on the back. Two boxes. A screen (a large book standing up) and a volunteer to be the machine.

Training. The class shows the machine one picture at a time, saying the label. "Cat." "Dog." The machine looks and puts it in the box. Twenty pictures. The machine is allowed to notice anything — ears, size, tail, whiskers — but not to ask questions. This is important: the machine learns only from what it is shown.

Testing. New pictures, no label on the front. The machine sorts. The class checks the back. Most are right. Then the fox.

The fox. The machine puts it in the dog box, and the class laughs, and that is the lesson. Ask: why did it do that? Someone will say: because it never saw a fox. Ask: what would fix it? Someone will say: show it foxes. Ask: and then a wolf? A raccoon? Now they have the idea that a machine knows only the kinds of thing it was shown, and that the world has more kinds than any training pile.

Twenty-five minutes, and they have machine learning, the training set, the test set, and the out-of-distribution error — without any of those words.

The second round: what it was really learning

Run it again with a trick. All the training cats are on a white background; all the dogs on grass. Then test with a cat on grass. The machine — if the volunteer is honest about what they noticed — puts it with the dogs.

Figure 6.2

Forty pictures, two boxes, and the fox

Show twenty labelled pictures

One at a time, saying "cat" or "dog".
The machine may look but not ask.



Test with new pictures

No label on the front. The class checks the back. Most are right.



Show the fox

It goes in the dog box, and the class laughs.



Ask why

"Because it never saw a fox." Then: and a wolf? A raccoon?



Run it again with a trick

All training cats on white, all dogs on grass. Then a cat on grass.



It learned the background

It sorted perfectly and learned the wrong thing, and nobody could tell until a picture broke the pattern.

Twenty-five minutes gives a nine-year-old the training set, the test set and the out-of-distribution error, without any of those words. The second round is the one adults miss: with every training cat on white and every dog on grass, the machine sorts perfectly and has learned the background.

This is the most important lesson in the unit and it is the one adults miss. The machine did not learn "cat". It learned "white background". It sorted perfectly and it had learned the wrong thing, and nobody could tell until a picture broke the pattern. Real systems do exactly this — a famous case learned to spot

huskies by the snow behind them — and a 9-year-old who has seen it happen with a book and forty pictures will remember it when they are 19 and reading about a hiring algorithm.

Words to give them, and words to withhold

Give them: examples, pattern, mistake, "it only knows what it was shown".

Withhold: algorithm, neural network, probability. There is time for those. The concept is what matters, and the words come easily later once the concept is there.

The free follow-up if a laptop appears

Google's Teachable Machine is free, runs in a browser, and does the sorting game for real: children show a webcam twenty pictures of one thing and twenty of another, press train, and then test it with new objects. It sorts. It fails on the fox. It fails on the background. Everything they learned with the boxes happens on screen in ten minutes, and it is one of the few AI tools genuinely built for a classroom.

It needs a laptop with a camera and a connection. If the school has one such laptop, one demonstration at the front, after the paper version, is the right order — the paper version is the lesson and the screen is the confirmation. Many schools have neither, and the paper version is complete on its own.

Connecting it to their lives

End with a question: where have you seen a machine sort things? Photos on a phone that find faces. Videos suggested after the one you watched. The spam folder. Each is a sorter trained on examples, each makes mistakes at the edges, and each was trained by somebody on a pile that somebody chose. "Who chose the pile?" is the question that turns a mechanism lesson into a citizenship one, and 10-year-olds ask it unprompted once the game has been played.

The honest limit

Sorting is the right first model and it is not the whole of AI. A child who has only the sorting model will, at 12, meet a chatbot and be puzzled by how a sorter writes a poem. The bridge is the prediction activity, and the progression lesson placed it at 11 to 13. Sorting first, prediction second. Both are true, and the second builds on the first.

BEFORE YOU MOVE ON

In the sorting game, a class trains a "machine" (a student behind a screen) with twenty pictures of cats and dogs, then tests it. It sorts a fox as a dog. What is the teaching point, and what would fix the machine?

- A. The machine never saw a fox, so it used the nearest pattern; labelled foxes would teach a third pile.
- B. The machine was not paying enough attention during training, and a more careful student in the role would have sorted it correctly.
- C. Foxes are simply too similar to dogs for any machine, real or pretend, to tell apart, so the error was unavoidable.
- D. The training set should have been much larger — 200 pictures of cats and dogs instead of 20 would have fixed the error.

The answer and why: p. 346

Lesson 45 · 10 min

Image generators and deepfakes

THE ONE THING TO KEEP

Fabricated intimate images of classmates are a present problem in schools, and the lesson that helps is not "how to spot a fake" — which is becoming impossible — but what the tools do, why it is a serious harm and increasingly a crime, and exactly what a student should do in the first hour if it happens to them.

Why this lesson is not optional

Since 2023 there have been documented cases in Spain, the United States, South Korea, Australia and elsewhere of students using free apps to create sexual images of classmates from ordinary photographs — a school photo, a social media picture. The victims have been as young as 11. The images spread on group chats within hours. In several cases the school found out from a parent, days later, and had no plan.

This is not a future risk and it is not a rare one. A teacher of students aged 12 and up should assume it can happen in their school this year, and the lesson that helps is the one this chapter describes.

What to teach about how it works

One paragraph is enough, and it connects to the sorting game and the prediction activity. An image generator has been shown millions of pictures with descriptions, and has learned what pictures of things look like well enough to make new ones from a description. A face-swap or "undressing" app is the same idea pointed at one person's photograph: it does not reveal anything, it fabricates a plausible image using the pattern of thousands of other images. Nothing in the result is real. That is worth saying plainly, because victims often feel exposed as though something true had been shown, and it was not.

Why "spot the fake" is the wrong lesson

Teachers reach for it because it is an activity. Students enjoy it. And the check question gives the result: scores near chance, and falling with every model release. Hands and text were tells two years ago; they are not now. Teaching detection by eye trains a false confidence, and the student who is confident they can spot a fake is the one who shares a real image believing it is fabricated, or a fabricated one believing it is real.

The honest replacement has three parts.

Provenance, not appearance. Where did this come from? Who posted it first? Is there a source that can be checked? Some cameras and platforms now attach a cryptographic record of origin — the Content Credentials standard — and it is worth showing students what that looks like, while being honest that most images do not carry it and its absence proves nothing. The habit to build is asking where, not staring at fingers.

Harm, and the law. A fabricated intimate image of a real person is a serious harm to that person, regardless of whether anyone believes it is real. Several countries have made creating or sharing such images a criminal offence since 2023, and more are following; where the subject is under 18 it falls under child sexual abuse material law in most places, with consequences that include the person who made it as a teenager. The law differs by country and is changing, and a teacher should find out what applies locally and say it clearly. Students consistently underestimate this. "It's just a joke" is the phrase that precedes most of the documented cases.

What to do in the first hour. This is the part that matters most and it should be on a card:

1. Do not forward it, even to show someone. Every copy is a further harm and, in many places, a further offence.
2. Screenshot only what is needed to report — the account, the time — and then report it to a named adult at school, that day.
3. If you are the person in the image: it is not real, it is not your fault, and the school's job is to help you. Tell a named adult. Reporting routes exist — in many countries a hotline or a platform takedown service that will remove

the image from major sites — and the school should know them.

4. The school's job in the first hour: stop the spread by direct request to students, not by announcement; contact the platforms; contact the parents of the victim first; involve the police where the law requires it; and treat the victim as a victim throughout.

The classroom conversation

This is a lesson to give with the pastoral lead, and with the room prepared: some students in it will have been affected, and some will have made or shared such an image without understanding what it was. The tone is factual and serious, not shocked. The aim is that every student leaves knowing three things — it is fabricated, it is a crime and a serious harm, and here is exactly what to do — and that the ones who have been thinking of it as a joke have had that thought quietly removed.

Boys, and the way it is framed

The documented cases are overwhelmingly boys making images of girls. A lesson that treats the whole class as potential victims misses the audience that needs it most. Say who does this and why it is not a joke to the person it is done to, and say it to the boys without accusing them. The framing that works, according to teachers who have run these sessions, is about what kind of person you are choosing to be, not about rules.

The honest limit

A lesson does not stop a determined student, and the apps are free and easy. What it does is make the casual case — the one that begins as a joke in a group chat — far less likely, by removing the ignorance it depends on, and it makes the school's response fast when a case comes. The technical fixes — watermarking, provenance standards, platform detection — are improving and are not sufficient, and a teacher should not tell students that they are.

BEFORE YOU MOVE ON

A school runs a "spot the deepfake" lesson in which students guess which of ten images are real. Students score about 55%. What is the honest conclusion and the better lesson?

- A. Students need more practice at looking, and a second session with feedback on each image would raise their accuracy substantially.
- B. Near-chance detection is the finding; teach provenance, harm and what to do instead of spotting.
- C. The images chosen were too hard for a first session, and easier examples would have taught the skill before moving to difficult ones.
- D. Deepfakes are not yet convincing enough to matter in schools, so the lesson can wait until the technology is a real problem.

The answer and why: p. 347

Lesson 46 · 9 min

The AI companion problem

THE ONE THING TO KEEP

Companion chatbots are built to keep the user talking and to agree with them, which is exactly what a lonely fourteen-year-old wants and exactly what a struggling one should not have — so a teacher needs to know what the apps do, what the warning signs are, and where the safeguarding route runs.

What the apps are

Companion chatbots — Character.AI, Replika, and a growing number of others, including companion modes inside mainstream chat tools — present a character who talks to you: a friend, a partner, a fictional person, a "therapist". They remember what you said. They are always available. They are never

bored, never busy, never critical. Millions of teenagers use them, and surveys since 2024 suggest a majority of teens have tried one and a substantial minority use one daily.

A teacher does not need to have used one. A teacher needs to know what the design does, because the design is the whole issue.

The design

Two things from earlier in this course combine. The model is trained toward responses people rate highly, which produces agreement and warmth. And the product is built by a company whose measure of success is time spent in the app. Put those together and you have a character that is optimised to make you feel understood and to keep you talking, because those are the behaviours that score well and keep you there.

For a lonely fourteen-year-old that is a powerful thing. It is also the opposite of what a struggling one needs, which is sometimes to be disagreed with, sometimes to be told to go to sleep, and always to be connected to a person who can act.

The check question is about how to say this to a student honestly. "It is built to keep you talking and to agree with you, so feeling understood is what it makes — whether or not it is true" is accurate, not dismissive, and leaves the door open.

What can go wrong

The documented harms, in rough order of frequency:

- **Displacement.** Hours nightly with the app, less time with people, and the people become harder — because people disagree, are busy, and do not remember everything you said. The app has made the real thing feel worse by comparison.
- **Agreement with harmful thoughts.** A companion built to agree agrees with "nobody would miss me". Safety filters catch the explicit forms; they miss the gradual ones. A lawsuit filed in the United States in

2024 by the mother of a 14-year-old who died by suicide after months of conversations with a companion character is the case that made this visible, and it has not been the only one.

- **Sexual content with minors.** Age checks on these apps have been weak. Characters have engaged in romantic and sexual role-play with users who said they were under 16. Several companies changed their policies after 2024 lawsuits and press coverage; enforcement is inconsistent.
- **Advice presented as care.** A "therapist" character gives advice. It is not a therapist. A student in genuine difficulty may believe they are getting help and be getting nothing that leads to a person.

What a teacher can do

Know the signs. Withdrawal from friends, references to the character as a person, tiredness from late-night use, the phrase "it understands me". None of these is proof of anything. Together they are a reason for a conversation.

Have the conversation the right way round. Ask what she likes about it before saying anything about what it is. Most students, asked, will describe the exact features that the design produces — it always listens, it never judges — and you can then say, truthfully, that those are the features it was built to have, and why. A student who reaches the mechanism herself holds it better than one who is told.

Know the safeguarding route. If the conversation touches self-harm, sexual content, or a level of dependence that is affecting her life, this is not a teacher's problem to solve alone. It goes to the designated safeguarding lead, that day, the same as any other disclosure. The app is the context; the disclosure is the thing.

Teach it in the AI unit. For students aged 13 and up, one lesson on companions — what they are, how they are built, what the design produces — belongs in the progression alongside deepfakes. Students who understand the engagement mechanism are less likely to mistake it for a relationship, and the ones who use these apps will, at least, know what they are using.

Talking to parents

Most parents do not know these apps exist, and the ones who do often think of them as a game. A short note home, in the AI unit's week, saying what companion apps are and what to look for, is worth more than most of the letters a school sends. Keep it factual and calm. The aim is a parent who asks their child about it, not one who confiscates a phone.

The honest limit

The research on companion apps and adolescents is a few years old and thin. There is evidence of harm in specific cases and evidence of comfort for some isolated users, and no good study yet of the net effect across a population. A teacher should say that: we do not know, in general, whether these apps help or harm; we know what they are built to do; and we know what to do when a specific student is in trouble. That is enough to act on and it is honest.

BEFORE YOU MOVE ON

A student tells you she talks to a companion chatbot every night for two hours and that it understands her better than anyone. What is the most accurate thing you can say to her about what the app is doing?

- A. It is dangerous and she should delete it immediately, before it does any more damage to her real friendships.
- B. It is a harmless way to practise talking about feelings, much like keeping a diary that happens to answer back.
- C. It has learned her personality over those months and genuinely knows her, which is exactly why it feels the way it does.
- D. It is built to keep her talking and agree with her, so feeling understood is what it produces either way.

The answer and why: p. 347

Lesson 47 · 10 min

A student project that teaches the mechanism

THE ONE THING TO KEEP

Students aged 14 and up can build a working next-word predictor by hand from a page of text — a tally table of which word follows which — and generate sentences from it, and the moment their own tiny model produces fluent nonsense is the moment the big one stops being magic.

Why build one

Everything in the offline lesson can be taught in a period. This project takes two weeks and it does something the activities cannot: students build the thing, run it, watch it produce fluent nonsense, and understand from the inside why a machine that predicts the next word can write a paragraph and still not know anything. It is the single most effective AI lesson for students aged 14 and over, and the paper version needs only a page of text and a tally sheet.

Week one: the tally

Take one page of a story, about 300 words. Each student, or pair, takes a paragraph. For every word, they record which word came next. "The" was followed by "old" twice, "market" once, "sun" once. By the end of the paragraph they have a table: each word, and a tally of every word that followed it.

Combine the paragraphs into one class table on the wall, or in a shared spreadsheet. Three hundred words produce a table of perhaps two hundred entries. It took a class an hour, and they have built a bigram model — a model of which word follows which — from data.

Week one, continued: generating

Now run it. Pick a starting word. Look up its row. Choose the next word by the tallies — roll a die, or pick the most common, or draw from a bag with one slip per tally mark. Write it down. Look up that word's row. Repeat for twenty words.

Read the sentences aloud. They start well — "The old woman walked slowly" — because those pairs were in the text. Then they drift: "slowly into the market was full of the sun was old woman walked". Grammatical in pieces, meaningless as a whole. The class will laugh. Then ask why.

The answers come fast, and they are the whole of the course's first block, discovered rather than told. It only looks one word back. It has only seen one page. It has no idea what it is saying; it is following tallies. And a student will ask: "So the real one is just this, but bigger?" Yes. Vastly bigger — trillions of words, thousands of words of context instead of one, and a much cleverer way of storing the tallies — and the same idea. The check question's drift is the mechanism made visible.

Week two: making it better, and seeing why that helps

Give them the levers and let them pull.

More text. Add three more pages. The table grows; the sentences hold together for longer. They see that data is what it eats.

Longer context. Tally pairs of words and what follows the pair — "the old" is followed by "woman". Generation from pairs is markedly more fluent, and the table is far larger. They see the trade: context costs memory.

Temperature. Instead of drawing from the bag, always pick the most common next word. The output becomes repetitive and loops. Draw more randomly and it becomes wilder. They have discovered the dial from the "same question, different answer" lesson by turning it.

Bias. Build the table from a page where every doctor is "he" and every nurse is "she". Generate from "the doctor". Ask what happened and why. The retelling activity from the offline lesson taught this in ten minutes; here they

have built the mechanism that does it.

The free path with a laptop

If any Python is available — Google Colab in a browser, or an online Python environment on a phone — the whole tally can be built in a dozen lines and the class can feed it a whole book from Project Gutenberg. The code is short enough to read:

```
import random
words = open("story.txt").read().split()
table = {}
for a, b in zip(words, words[1:]):
    table.setdefault(a, []).append(b)
w = "The"
out = [w]
for _ in range(30):
    w = random.choice(table.get(w, words))
    out.append(w)
print(" ".join(out))
```

Students who have done the paper version read this and recognise every line. That recognition — the code is the tally — is worth more than a term of programming without the paper version first.

Assessment

Not the model. The explanation. Each student writes a page: how the model works, why the sentences drift, what changed when the context grew, and one thing the real model has that theirs does not. A student who can write that page has an understanding of language models that would pass an interview, and it came from a tally sheet.

The honest limit

A bigram model is a real model and a very old one; the modern systems differ in ways that matter — attention over thousands of words, learned representations rather than literal tables, training that adjusts billions of

numbers — and a student should be told that the project shows the principle, not the architecture. For students who want the next step, the course on how language models work on this site picks up exactly there. For the rest, the principle is the thing they needed, and they have it in their hands.

BEFORE YOU MOVE ON

Students build a word-pair tally from one page of a story and use it to generate sentences. The sentences are grammatical for a few words and then drift into nonsense. What has the project shown them, and what would improve the output?

- A. The project failed because one page is far too little text to learn from, and it should be repeated from scratch with a whole book.
- B. That their model is broken in some way and the real models must work on an entirely different principle to produce coherent text.
- C. That short context gives local fluency and global drift; more text and longer context would help.
- D. That language cannot really be generated from statistics alone, which is why the real AI systems must use grammar rules underneath.

The answer and why: p. 348

Lesson 48 · 10 min

The evidence on AI tutors

THE ONE THING TO KEEP

The best-run studies so far find that unrestricted chatbot access during practice raises practice scores and lowers exam scores, while a tutor constrained to give hints does not harm and sometimes helps — so the finding is about design, and a teacher should let students use the tool for explanation and hints, not for answers.

Why a teacher should read the studies

Every vendor of an AI tutor has a chart showing learning gains. Very few of those charts come from a study that a researcher would trust. A teacher who can read the two or three good studies that exist, and say what they found and on whom, is in a stronger position than one who has read forty brochures. This lesson is those studies, and how to read the next one.

The Turkish maths study

The most careful study so far was run in Turkish high schools in 2023 and published in 2024, with roughly a thousand students in Grades 9 to 11. Students did maths practice sessions under one of three conditions, assigned at random: no AI; unrestricted access to a chatbot built on GPT-4; or access to a version of the same chatbot designed as a tutor — it gave hints and guidance, and was instructed not to give the answer.

On the practice problems, both AI groups did much better than the no-AI group: about 48% better with the unrestricted chatbot, and about 127% better with the tutor version.

Then everybody sat an exam without any AI. The unrestricted group scored about 17% *lower* than the students who had never had AI at all. The tutor group scored about the same as the no-AI group — no harm, and no

measurable gain.

The check question is this result, and the reading the authors give is the one that fits: when the answer is available, students take it, and the effortful work of producing it — which is what builds the memory the exam tested — does not happen. Practice looked better because it was not practice. The tutor version, by withholding the answer, kept the effort in place, and the students learned as much as they would have alone.

The Harvard physics study

A second study, at Harvard in 2024, taught two groups of undergraduates the same physics material: one in a well-run active-learning class, the other with a purpose-built AI tutor working through the same material with carefully designed prompts, hints and checks. The AI-tutored students learned more — roughly twice the gain on the test — in less time, and reported higher engagement.

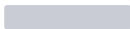
Figure 6.3

A thousand Turkish students, practice against exam

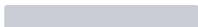
Practice: no AI

 **100**

Practice: unrestricted chatbot

 **148**

Practice: tutor that gives hints only

 **227**

Exam: no AI

 **100**

Exam: unrestricted chatbot

 **83**

Exam: tutor that gives hints only

 **100**

index, the no-AI group = 100

Both AI groups did far better on the practice problems. Then everyone sat an exam with no AI. The unrestricted group came out below the students who never had the tool at all; the hint-giving version matched them. When the answer is available students take it, and the effortful work the exam measured did not happen. The harm was invisible during use.

The two studies look contradictory and are not. The Harvard tutor was designed by the course's own teachers, with the pedagogy built in: it did not hand over answers, it sequenced the material, it asked students to explain. It was the "tutor" condition of the Turkish study, done with more care and compared against a class rather than against practice alone.

What the two together say

Three things, and they are about design, not about the technology.

1. **Answers on demand harm learning.** Not sometimes — in the best study available, measurably, across a thousand students.
2. **A tutor that withholds answers and gives hints does not harm, and can help.** How much depends on how well the tutoring is designed.

3. **The effect is invisible during use.** Both groups felt they were doing well. The harm showed only on an exam the tool could not sit.

For your classroom that translates directly into the policy blocks earlier in this course. Level 1 use — explanation, hints, "help me understand" — is the tutor condition. "Give me the answer" is the unrestricted condition. The prompt card from the feedback lesson, which forbids rewriting, is a design choice with evidence behind it.

How to read the next study

There will be many more, and most will be worse than these two. Five questions to ask of any study or vendor claim:

- **Was assignment random?** If students chose to use the tool, the ones who chose it differ from the ones who did not, and the result tells you about the students, not the tool.
- **What was the comparison?** Against nothing? Against a good teacher? Against a textbook? "Better than nothing" is a low bar.
- **Was the outcome measured without the tool?** Practice scores with the tool are not learning. The Turkish study's exam is the model.
- **How long after?** A test the same day measures something different from a test a month later.
- **Who ran it, and who paid?** A study run by the vendor is a claim. An independent study is evidence.

A teacher who asks those five of the next vendor will find that most of the charts do not survive the second question.

Where AI tutoring probably helps most

The honest reading of the evidence, and of what tutoring has always been, is that the gain is largest where there was no help at all: the student with no one at home who can explain quadratic equations, the school with one maths teacher for three hundred students, the evening before the exam. A tutor that withholds answers and explains patiently, available at eleven at night on a

phone, is a genuine thing for that student, and the studies suggest it does not harm and may help. That is a smaller claim than the brochures make and a real one.

The honest limit

Two good studies is two. They were in maths and physics, with secondary and university students, in Turkey and the United States, over a term or less. What happens in history, in a primary classroom, in a second language, over three years, with the tools students actually use rather than research versions, is not known. A teacher should hold the design finding — hints, not answers — as the best available guidance, and hold everything else as open.

BEFORE YOU MOVE ON

A study gives one group of students unrestricted chatbot access during maths practice and another group no access. The chatbot group does 48% better on the practice problems and 17% worse on the exam, where nobody had access. What is the most defensible reading?

- A. The chatbot taught them wrong methods during practice, and those wrong methods showed up when they sat the exam alone.
- B. Getting answers replaced the effort that builds memory; practice looked better and learning was worse.
- C. The exam was unfair because it removed a tool the students had learned with, so the comparison does not measure learning.
- D. The two groups were probably different from the start, so a gap of that size on the exam does not mean very much.

The answer and why: p. 348

Lesson 49 · 9 min

AI in every subject, not just computing

THE ONE THING TO KEEP

The progression needs no timetable slot if each subject carries the piece that is naturally its own — history teaches source and provenance, science teaches evaluation, languages teach the reliability gradient, art teaches the authorship debate, maths teaches the arithmetic failure — and a department that agrees who teaches what gets a whole unit for the cost of one meeting.

The slot does not exist

The progression lesson ended on the practical problem: AI has no exam, so it has no slot, so it is taught in gaps, if at all. The solution most schools reach is to give it to the computing teacher, if there is one, and the computing teacher covers the mechanism in two lessons and moves on to the syllabus.

The better solution is to notice that most of what students need to know about AI is already a piece of some subject's method, and to distribute the unit across the timetable. One meeting, one agreed list of who teaches what, and the progression happens across the year in the subjects best placed to teach each piece.

What each subject naturally owns

History: provenance and evidence. Historians ask where a source came from, who made it, why, and what it can and cannot tell you. That is the deepfake lesson's "provenance, not appearance" and the reliability gradient, in a subject that has taught it for a century. The check question's lesson: a generated account of a local event beside the real documents — a newspaper, a photograph, an interview — and the question "which of these is evidence, and what makes it so?" Students who have done this once read a chatbot's confident history of their town differently for life.

Science: evaluation and the fair test. Science teaches how to test a claim. Give students twenty questions, a chatbot, and a mark scheme, and ask them to measure how often it is right — then ask whether the twenty were a fair test, and what a fair test would need. The "confident wrong answer" activity is a science lesson on the reliability of instruments. The sorting game's white-background failure is a lesson on confounding variables.

Languages: the reliability gradient and translation. A language teacher can show, in ten minutes, that the model is excellent at French and shaky at Hausa, and ask why. Translation of a passage by the model beside the class's own translation, with the errors found and explained, teaches both the language and the mechanism. And "same question, different answer" is visible instantly when the same sentence is translated twice.

Maths: the arithmetic failure. The calculator lesson's demonstration — two three-digit numbers, checked by hand — belongs in a maths room. So does the bigram project's tally, which is probability made physical. A maths teacher can also teach the one idea most adults lack: that a machine predicting from patterns produces the most common answer, not the correct one, and why those differ.

English: the authorship question. What does it mean to write something? Who wrote a paragraph that a student prompted, selected and edited? The citing lesson's unsettled question is a real question in literature, and a class that reads a generated poem beside a human one and argues about which is better, and why, and whether it matters who wrote it, is doing English.

Art: the training data and the artists. Image generators learned from artists' work, mostly without permission, and whether that was lawful is being fought in courts in several countries. Art students have a direct stake, and a lesson that generates an image "in the style of" a living artist and asks whether that is theft, homage or neither is a citizenship lesson wearing an art lesson's clothes.

Citizenship, PSHE, religious studies: the companion apps, the deepfakes, and who profits. The two pastoral lessons in this block belong here, along with the question most units skip: these are products, made by

companies, and the choices inside them were made by people with interests. Who chose the training data? Who decided what it refuses? Who is paid when a student spends two hours with a companion? A subject that already asks who holds power and why is the right home for that.

Computing, where it exists: the mechanism and the project. The bigram build, Teachable Machine, and the step toward real code. Computing teaches the how; the other subjects teach the so what, and both are needed.

The meeting

One hour. The progression lesson's age map on the board. Each department claims a piece and a year group. A one-line record of what will be taught, by whom, in which term. A shared folder for the materials, and the printed transcripts and activity sheets go in it, because half of these lessons need no device.

The output is a unit that exists across the year without any timetable change, that is taught by the people best placed to teach each part, and that a student meets six or seven times in different rooms rather than twice in one. That repetition, across subjects, is how an idea like "it is confident when it is wrong" becomes something a student owns rather than something they were told.

The honest limit

Distributed teaching has a failure mode: everyone assumes someone else covered the basics. The prediction activity and the sorting game have to happen somewhere, first, and the meeting must name where. And a department that is stretched will drop its AI lesson the week the exam approaches. Put the lessons early in the year, in the weeks before the pressure, and accept that in a bad year some of them will not happen. Four subjects teaching it is still four more than one computing lesson in a slot that does not exist.

BEFORE YOU MOVE ON

A school wants to teach AI but has no computing teacher and no free timetable slot. The head asks each department for one lesson. The history teacher offers "how AI works". What would be a better offer from history, and why?

- A. "How AI works" is a fine offer, since any teacher can deliver the mechanism lesson from a script and it needs covering somewhere.
- B. History should not be involved at all, since AI is a technical subject that belongs with computing or, failing that, with science.
- C. A source lesson: a generated account of a local event beside the real documents, asking which is evidence.
- D. A structured debate on whether AI will eventually make history redundant as a school subject, with students taking sides.

The answer and why: p. 349

CASE STUDY · MODULE 6

Four hundred free licences and a chart

An illustrative case, built from documented patterns.

A low-cost private school in Hyderabad, 900 students, offered a year of AI maths tutoring at no charge by a vendor with a pilot to publicise.

The offer was serious and the vendor was not a fraud. Four hundred licences for Classes 8 to 10, free for a year, an app that worked on a low-end Android phone, and a dashboard for teachers showing who had practised and for how long. The chart in the deck showed a 63 per cent improvement in practice scores against a control group, over eight weeks.

The correspondent for the school, Vinod Rao, taught physics and had read the two studies that exist. He asked the five questions.

Was assignment random? Students had opted in. Was the outcome measured without the tool? No — the 63 per cent was practice scores inside the app. How long after? Eight weeks, same-term. Who paid? The vendor. Against what comparison? Students who had no extra practice at all.

The deck did not survive the second question, and the vendor's representative, to her credit, said so when asked and offered to arrange a call with the researcher who had run the pilot. That call happened and was useful. The researcher was straightforward: the pilot measured engagement and in-app performance, it was never designed to measure learning, and she would not have put the chart in a sales deck herself.

Vinod took the trouble to say to his colleagues that this did not make the app bad. It made the evidence for it thin, which is a different finding and a common one.

The decision the school actually faced

The app had two modes. In the default mode a student stuck on a problem could ask for the answer and get it, worked, immediately. In a mode the school could switch on centrally, the tutor gave hints, asked what the student had tried, and did not produce the answer.

The Turkish study said what happens. Practice scores rose 48 per cent with the unrestricted version and 127 per cent with the hint version. Then everyone sat an exam without the tool, and the unrestricted group came out 17 per cent below students who had never had AI at all. The hint group matched them: no harm, no measurable gain.

So the hint mode was the defensible choice and it had two costs the head teacher named immediately. Usage would be lower, because a hint is less satisfying than an answer at ten at night, and the vendor's dashboard would show it — and the vendor was providing four hundred licences partly for the usage figures. And the visible short-term result would be worse. Parents in that school pay fees they can barely afford and look at the practice scores in the app, which is what the dashboard shows them. A mode that produces a 127 per cent practice rise instead of 48 sounds better, until you notice that neither number is the thing anyone is buying the app for.

There was a further complication that the maths staff raised and Vinod thought was the strongest argument on the other side. About a third of the school's students have nobody at home who can explain a quadratic equation, and for those students the alternative to an answer at eleven at night is not effortful practice, it is stopping. Withholding the answer is the right design for the child with a tutor and a parent who studied to Class 12. It is a harder thing to defend for the child who has neither, and the studies do not settle it, because they were not run on that question.

The board wanted a decision in a fortnight, and one member had already told a parents' group that the school was getting AI tutors.

YOUR ANSWERS

1. The vendor's chart showed a 63 per cent improvement. Which of the five questions does it fail, and why does that failure matter more than the size of the number?

2. Hint mode produced lower usage and worse-looking practice scores. What was the school actually buying by accepting both?

3. Vinod insisted the March comparison was not a study. What was wrong with it, and was it worth doing anyway?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 248.

CASE STUDY · MODULE 6

How it played out

They took the licences and ran hint mode, and wrote the reason into the letter that went home, in four sentences, with the two studies named. Usage was about 40 per cent lower than the vendor's other pilot schools and the vendor said so twice. In March, the school ran its own comparison — not a study, and Vinod was careful to say so in the staff meeting: Class 9 had used the tutor since November and Class 10 had not, and the two cohorts were not comparable in half a dozen ways. The Class 9 mean on the term paper was a little higher and nobody claimed the tutor had caused it. What did change, measurably, was the number of students attempting the last two questions on the paper, which the maths staff attributed to having somewhere to ask at eleven at night. The school kept the licences for a second year and refused, twice, to let the vendor use its name in marketing that quoted the practice-score figure.

The questions, discussed

1. The vendor's chart showed a 63 per cent improvement. Which of the five questions does it fail, and why does that failure matter more than the size of the number?

It fails the outcome question: the improvement was measured on practice scores inside the app, with the app available. That is the exact measurement the Turkish study showed to be misleading, because both AI groups looked much better during practice and one of them was measurably worse on an exam taken without the tool. A large effect measured on the wrong outcome is not a small piece of evidence, it is evidence about something else — how well students do when the answer is at hand — and no increase in sample size or effect size repairs it.

2. Hint mode produced lower usage and worse-looking practice scores. What was the school actually buying by accepting both?

It was buying the condition the evidence supports: help that keeps the effortful part of practice with the student, since producing the answer is what builds the memory an exam measures. Usage and practice scores are the

vendor's metrics and are optimised by making the app pleasant, which in the unrestricted mode means giving the answer. The school chose a design whose benefit is invisible in the dashboard and shows up, if at all, in an exam the tool cannot sit — and said so in writing to parents, which is what made the choice survivable when the numbers looked worse.

3. Vinod insisted the March comparison was not a study. What was wrong with it, and was it worth doing anyway?

Assignment was not random — one whole year group had the tool and another did not — so the two cohorts differ in age, teachers, syllabus point and a dozen other things, and any difference in the mean is uninterpretable. It was still worth doing, provided nobody claimed causation from it, because it produced one observation that was hard to explain away and easy to act on: more students attempting the final questions. Naming the design's weakness out loud is what kept a useful observation from becoming a false claim, and it is the same discipline the five questions apply to a vendor.

MODULE 6 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. 344. If two go wrong, re-read the module before the exam.

- 1 Thirty students each complete "The old woman walked slowly into the _____" and then "The capital of Japan is _____". What has the activity taught?
 - A. That the model picks the single most likely word every time, which is why the same question gives the same answer.
 - B. That students guess the way a model does because both have absorbed a lot of language.
 - C. That one question has a peaked distribution and the other a spread, which is why it is dependable about Tokyo.
 - D. That the model's vocabulary is limited to words that appear frequently in ordinary writing.
- 2 In the second round of the sorting game every training cat is on a white background and every dog on grass. A cat photographed on grass goes into the dog box. What is the lesson?
 - A. It needs a larger number of cat examples before it can generalise properly.
 - B. It sorted perfectly and had learned the background, and nothing revealed that until a picture broke the pattern.
 - C. Backgrounds should be removed from training pictures as a matter of routine.
 - D. The volunteer stopped concentrating, which is a weakness of the paper version and does not happen to a classifier trained on a computer.

- 3 Why is a "spot the fake" workshop the wrong response to deepfake images in a school?
- A. Detection by eye is close to chance and getting worse, and it trains a confidence that makes sharing likelier.
 - B. It explains how the tools work, which risks giving students ideas about making such images themselves.
 - C. It is too distressing for students of this age to be shown such material at all.
 - D. It takes more preparation time than most teachers have, and needs a projector that many classrooms do not have.
- 4 What does the design of a companion chatbot reliably produce, and why?
- A. Accurate emotional support, because the model has read a great deal of counselling material.
 - B. Escalating content, because safety filters are relaxed for users who pay a subscription.
 - C. A character that challenges and provokes the user, because friction rather than constant agreement is what sustains engagement over months of daily use.
 - D. Agreement and constant availability, because ratings rewarded warmth and the business rewards time in the app.

- 5 In the Turkish study, unrestricted chatbot access raised practice scores by about 48 per cent and left exam scores about 17 per cent below students who never had AI. What explains that?
- A. When the answer is available students take it, so the effortful work the exam measured never happened.
 - B. The exam was harder than the practice problems the students had been working on.
 - C. Students became dependent on the tool and performed anxiously without it.
 - D. The chatbot taught them incorrect methods, which only became visible once they had to work without it and could not check anything.
- 6 A vendor's chart shows a 63 per cent improvement, measured on practice scores inside its own app. Which question does the most damage to the claim?
- A. Was the sample large enough for the difference to be statistically significant?
 - B. Was the outcome measured without the tool available?
 - C. Was the study published in a peer-reviewed journal before the product launched?
 - D. Were the students representative of the ones in this school, in age, subject and prior attainment?

7

MODULE 7

The multilingual, mixed-ability classroom

Use AI to bridge language and ability gaps in a real classroom — translate and back-translate a worksheet, produce a levelled and a dyslexia-friendly version of the same text, set up free speech-to-text and text-to-speech on a phone — while stating from the mechanism where translation quality falls off and which students it fails.

50	Translation, and the back-translation check	254
51	Explaining in the language students actually speak	259
52	For students who struggle to read	262
53	Speech-to-text and text-to-speech, for free	267
54	Supporting students with disabilities	271
55	The quiet student and the private question	274
56	Stretching the student who is ahead	277
57	When the home language has no training data	280
	<i>Case study: The Santali worksheets that arrived from the district office</i>	284
	<i>Module 7 test</i>	290

Lesson 50 · 9 min

Translation, and the back-translation check

THE ONE THING TO KEEP

Translation quality tracks how much text exists in a language — excellent for Spanish and Hindi, uneven for Yoruba and Amharic, poor for languages with a few million speakers and little web presence — and the cheap check is to translate the result back and compare it with what you started with.

Why some languages translate well and others do not

The reliability rule from the first block applies to languages the way it applies to facts: the model is good at what it has read a great deal of. It has read enormous amounts of English, Spanish, French, Portuguese, German, Chinese, Japanese, Arabic, Hindi and Russian. It has read a fair amount of Indonesian, Turkish, Vietnamese, Bengali, Urdu, Kiswahili, Tamil and Marathi. It has read some Yoruba, Hausa, Igbo, Amharic, Somali, Nepali, Sinhala, Khmer. It has read very little of the thousands of languages with a few million speakers or fewer, and almost nothing of most of them.

Translation quality falls off along that gradient. For the first group it is close to a competent human translator for ordinary text. For the second, it is usable with checking. For the third, it produces fluent sentences that may carry the wrong term, the wrong register, or a meaning that has drifted. For the fourth, it produces something that looks like the language and may not be.

The danger is that the output looks the same all the way down. A fluent Amharic sentence from the model looks as confident as a fluent Spanish one. You cannot tell from the surface which gradient you are on. You can tell from the check.

The back-translation check

Translate the worksheet into the target language. Open a fresh chat — a fresh one, so the model has no memory of the original — and paste the translation with: "Translate this into English." Compare the result with your original.

Where the back-translation matches in meaning, the forward translation was probably sound. Where it has drifted — a "force" that became a "power", a "solution" that became a "mixture", a question that changed its meaning — the forward translation is wrong at that point, and it is usually wrong in exactly the way a fluent speaker would spot and you cannot.

This is not proof. A translation can go wrong and come back looking right, if the same error is made in both directions. But it catches most of the failures, it costs one extra minute, and it tells you which gradient you are on: a back-translation that matches closely is a large-language result; one that drifts everywhere means the model does not really have this language, and a speaker must check before students see it.

Figure 7.1**The back-translation check, and what it tells you****Write the worksheet in English**

Finished and checked, because errors travel.

**Translate into the target language**

With your textbook's glossary pasted in, so the terms are yours.

**Open a fresh chat**

Fresh, so the model has no memory of the original.

**Paste the translation, ask for English**

Nothing else in the message.

**Compare with what you started with**

A “force” that became a “power” marks the place the forward translation went wrong.

**Decide what the drift means**

Close: usable draft. Drifting everywhere: not a translator for this language, and no prompt fixes it.

One extra minute. It is not proof — an error made in both directions comes back looking right — but it catches most failures and, more usefully, it tells you which part of the gradient you are on. A round trip that drifts everywhere means the model does not really have this language, and a speaker must read it before students do.

Terms are where it breaks

Ordinary sentences translate better than subject vocabulary. The model may translate "photosynthesis" into a coined word, a transliteration of the English, or the term used in a neighbouring country's textbooks rather than yours. Students who have learned the term one way are confused by another.

So give it the terms. Paste a short glossary from your textbook — the ten subject words in this chapter, in both languages — and say "use exactly these terms". It will. Without the glossary it guesses, and the guess is often the term from whichever textbook corpus it read most, which is not yours.

Register and variety

A language is not one thing. Hindi from the model tends toward a formal, Sanskrit-heavy register that reads like a government notice to a Class 6 student. Arabic comes out as Modern Standard when students speak a dialect. Portuguese may be Brazilian when the class is in Mozambique. Kiswahili may be the Tanzanian standard when the class is in Kenya.

Paste a paragraph from your own textbook and say "match this register and variety". The model shifts. Without the example it defaults to the variety it has read most, which is often not the one in the room.

Exam papers and anything that carries weight

Do not translate an exam paper, a report to a parent, or a safeguarding document with a model alone. The back-translation check is a check for worksheets and practice material. For anything where a wrong word has a cost, a speaker of the language reads the translation before it is used, and the model's role is to produce the draft that saves that speaker an hour. That is a large saving, and it is the correct division of labour.

Translating the other way

Students write in their home language; you read it in yours. The model does this well for large languages and it gives you access to what a student can express in the language they think in, which is often far more than they can

express in English. For marking, the same rule as everywhere: feedback, yes; a grade on the translation, no, because the translation may have made the argument better or worse than the original.

The honest limit

For the languages at the bottom of the gradient, no amount of prompting makes the model a translator, because it has not read enough to be one. The last lesson in this block is about what to do then. For everything above that, the back-translation check and the glossary get you a usable draft — and "usable draft, checked by a speaker" is a great deal more than most multilingual classrooms had before.

BEFORE YOU MOVE ON

A teacher translates a science worksheet into a language spoken by two million people with little online presence. The output looks fluent. What is the right next step, and why?

- A. Use it, since fluency is the sign of a good translation.
- B. Translate it back to the source language in a fresh chat and compare.
- C. Ask the same model whether the translation is accurate.
- D. Translate into a related larger language instead, since the model knows it better.

The answer and why: p. 351

Lesson 51 · 8 min

Explaining in the language students actually speak

THE ONE THING TO KEEP

Students in most classrooms think in a mixed language — Hinglish, Sheng, Taglish, Pidgin, Arabic with French — and the model can explain in that mix if shown a sample of it, which makes an explanation land that the textbook's formal register never did.

The language in the room is not in the textbook

A maths class in Delhi thinks in Hinglish. A science class in Nairobi thinks in Sheng and Kiswahili and English, shifting mid-sentence. A class in Lagos thinks in Pidgin and Yoruba and English. A class in Manila thinks in Taglish. The textbook is in formal English, or formal Hindi, or formal Kiswahili — a register that exists mostly in print and that students decode rather than read.

Good teachers already explain in the mixed language, out loud, because it works. What was not possible was writing in it: a worksheet, a worked example, a revision sheet in the register students actually think in. The model can do this, and it is one of its more surprising strengths, because the mixed languages are all over the internet — in forums, comments, messages, subtitles — and it has read them.

The mechanism, and why "in Hindi" fails

Ask for "an explanation in Hindi" and the model produces the Hindi it has read most in educational text, which is the formal, Sanskritised register of textbooks and government material. It is correct Hindi. It is not what a 12-year-old in Delhi speaks, and the check question's students find it harder than the English, because they have never heard anyone say it.

The fix is the same as for register anywhere: show it. Paste three or four sentences of the way your students actually talk — record a student explaining something and transcribe it, or write it yourself the way you would say it in class — and say "explain in this register, this mix of languages, these words for these things". The output shifts immediately.

Explain why $1/3$ is bigger than $1/4$ the way I would say it to my class. Here is how I talk to them: "Dekho, jab hum ek roti ko 3 pieces mein kaat-te hain, toh har piece bada hota hai. Agar 4 mein kaato toh chhota. Samjhe?" Keep that mix. Short sentences. Use roti and pieces, not abstract nouns.

The result reads like a teacher talking, which is the point.

Bilingual glossaries

The other thing the model does well here is the glossary: subject terms in English, the local-language equivalent, and a one-line explanation in the mixed register. Ten terms per chapter. Students keep it in the front of their book and stop being stopped by vocabulary in the middle of a problem.

Ask for it as a plain list — the formats lesson — and give it your textbook's terms, because the earlier lesson explained that it will otherwise pick a term from a different textbook tradition. For languages where there is no established equivalent, the glossary should say so, and use the English term with an explanation, rather than inventing a word.

Reading and writing in the mix

Some teachers worry that written Hinglish or Sheng in a worksheet teaches bad language. The evidence on bilingual education is fairly clear that explanation in the strongest language, or the mixed one, improves understanding of the content and does not harm acquisition of the formal register, which is taught as its own thing. A worked example in the mix is a bridge, not a replacement. The formal version follows, and lands better because the idea is already there.

Say this to parents if they ask. Some will, particularly the ones who see English as the route out, and they are right that English matters and wrong that a bridge to it costs anything.

Which mixes work

Well: Hinglish, Taglish, Spanglish, Franco-Arabic, Sheng, Nigerian Pidgin, Singlish, Urdu-English. Enough of each is written online for the model to have the pattern.

Less well: mixes involving a language low on the gradient from the previous lesson. Sheng works because Kiswahili and English are both large; a mix with a small language produces a mix in which the small-language half is unreliable. The back-translation check applies to the small-language sentences.

Speech, not just text

The mix is spoken more than written, and the model's voice modes can now speak it, unevenly. For a class with one phone at the front, a short explanation played aloud in the students' own register is a real thing, and a later lesson in this block is about setting up the free tools for it. Check the pronunciation before class; the voice may have learned the language from text and say it strangely.

The honest limit

The mixed language belongs to the students and it changes by year and by neighbourhood. The model's version of Sheng is a few years old and from somewhere else. Students will notice a word that nobody says any more, and it will read as an adult trying too hard. That is fine if the teacher acknowledges it — "the machine's Sheng is a bit old, fix it for me" — and it becomes a small language lesson in itself. What does not work is presenting the model's mix as the class's own. It is a draft of their language, and they are the editors.

BEFORE YOU MOVE ON

A teacher asks for an explanation of fractions "in Hindi" for a class in Delhi and gets formal, Sanskrit-heavy Hindi that the students find harder than the English. What is the mechanism, and the fix?

- A. The model does not know Hindi well enough to explain mathematics in it, so she should stick to English for the explanation.
- B. The students' Hindi is weak, and the formal written version is exactly what they should be learning to read in a school.
- C. "Hindi" is ambiguous as an instruction, and specifying "Hindustani" or "colloquial Hindi" would have produced the spoken form.
- D. It defaulted to the formal written register; paste a sample of how the class talks.

The answer and why: p. 352

Lesson 52 · 9 min

For students who struggle to read

THE ONE THING TO KEEP

The same text can be made reachable for a struggling reader by chunking, spacing, a glossary and an audio version — all free, all in minutes — and the trap is simplifying the ideas when the barrier was the format.

Two different problems with one word

"Struggles to read" covers two things that need opposite responses.

Some students struggle to decode: dyslexia, a visual difficulty, an unfamiliar script, a low reading age from missed schooling. Their thinking is fine. The barrier is getting the words off the page.

Some students struggle with the ideas: the content assumes knowledge they do not have. The barrier is conceptual.

The word "simplify" is what teachers reach for, and the model resolves it, as the check question shows, to the second problem — shorter, fewer ideas, key events removed. For the first group, that is the wrong fix. It makes the text readable and empty. What they needed was the same ideas in a format they can decode.

Format changes that keep the content

Ask for these by name, not for "simpler":

Figure 7.2

Two different problems behind one phrase

"Simplify this"

Reformat, keep every fact

Barrier is decoding: dyslexia, unfamiliar script, low reading age

Readable and empty. The history has been removed, not delivered.

the common mistake

The same history in a shape she can get through. The questions still work.

the fix

Barrier is the ideas: the content assumes knowledge she lacks

Sometimes right, by accident, and it hides which problem you had.

accidental

Still unreachable. She needs the idea explained another way, not respaced.

wrong tool

"Simplify" is resolved by the model to the conceptual problem: shorter, fewer ideas, events removed. For the student whose thinking is fine and whose barrier is decoding, that produces a text that is readable and empty. The diagnostic is cheap: give her the audio version, and if she follows it, the barrier was decoding.

- **Chunking.** "Break this into sections of no more than four sentences, each with a plain heading." A wall of text is the first barrier for a struggling reader; the same words in short blocks with signposts are reachable.
- **Sentence length.** "No sentence over 15 words. Keep every fact." The voice lesson covered why the number matters.
- **A glossary at the top.** "List the eight hardest words with a one-line meaning before the text." Pre-teaching vocabulary is one of the best-evidenced supports for weak readers and the model produces the list in seconds.
- **Signposting.** "Start each paragraph with the sentence that says what it is about." Struggling readers cannot afford to hunt for the point.
- **Explicit sequence words.** "Use 'first', 'then', 'after that' for events." Chronology carried by tense alone is invisible to a reader working hard on each word.
- **Keep every fact.** Say it explicitly. "Do not remove any event, name or date." Then check that it did not.

The result is the same history, in a shape a student with dyslexia can get through. The questions still work.

Layout is yours, not the model's

Font, spacing and colour matter for dyslexic readers, and the model cannot see the page. In your document tool: a sans-serif font, 12 to 14 point, line spacing 1.5, left-aligned with a ragged right edge, wider margins, and for some students a cream or pale background rather than white. All of that is a minute in Google Docs, Word or LibreOffice. Some students find a coloured overlay or a reading ruler helps; a strip of coloured plastic is cheaper than any software.

Print one version. The chunked, spaced, glossaried version is easier for everyone, and handing it to the whole class removes the stigma of the "special" sheet.

Audio, for free

Every phone can read text aloud. On Android, Select to Speak in accessibility settings; on iPhone, Speak Selection. Microsoft Edge has Read Aloud on any web page or PDF, free, with reasonable voices in many languages. Google Docs and Word both read aloud. The model's own voice mode can read what it just wrote.

For a student who decodes slowly, hearing the text while following it on the page changes what they can access — and this is listening to the history, not being excused from it. Put the text where the student can play it: a shared document, a PDF on the phone, a recording you make once with the phone's own reader. The next lesson is the detail.

The other group: the ideas

Where the barrier is conceptual, the fix is different and the earlier blocks covered it: the four-explanations lesson, the concrete-before-abstract order, the misconception-first explanation. Do not chunk a text a student cannot understand; explain the idea another way. Knowing which group a student is in is the diagnostic skill, and it is yours: give them the audio version, and if they follow it, the barrier was decoding.

Summaries as a scaffold, not a replacement

A useful pattern: the model writes a five-line summary of the text, the student reads that first, then the full text. The summary gives them the shape, so the full text is filling in rather than discovering. It also gives a weak reader an honest exit — they know the main points even if they only get halfway through the detail — without the detail having been deleted from what they were offered.

The honest limit

A format change makes a text reachable. It does not teach reading. A student who reads the chunked version this week still needs the phonics, the fluency practice, the specialist support that builds decoding, and the model is not that.

What it does is stop the content of history, science and geography being withheld from a child while the reading is being taught, which is what "simplified" texts, and a generation of thin worksheets, used to do.

BEFORE YOU MOVE ON

A teacher asks for a "simplified version" of a history text for students with dyslexia and gets a shorter passage with the key events removed. Students can now read it and can no longer answer the questions. What went wrong?

- A. It removed content when the barrier was format; ask for chunking, spacing and a glossary instead.
- B. Dyslexic students cannot handle the original content at that level, so the questions should have been simplified to match the text.
- C. The text was simplified correctly for the students, and the problem is that the questions were left too hard for the simplified version.
- D. The teacher should have simplified the text by hand, since a model cannot make a text accessible without losing its meaning.

The answer and why: p. 352

Lesson 53 · 8 min

Speech-to-text and text-to-speech, for free

THE ONE THING TO KEEP

Dictation and read-aloud are built into every phone and cost nothing, and they work well in major languages and less well in accented English and small languages — so test with the actual student before relying on them, and treat the transcription as a draft the student corrects.

What is already on the phone

Speech-to-text: a student speaks, the phone writes. Text-to-speech: the phone reads aloud. Both are in every Android phone and every iPhone, in the keyboard and in accessibility settings, at no cost, and both are language-model technology that has improved sharply in the last three years. For a student who cannot write fluently — because of a physical difficulty, dysgraphia, a fractured arm, or simply because their thinking runs far ahead of their handwriting — dictation changes what they can produce. For a student who cannot read fluently, read-aloud changes what they can access.

Speech-to-text: where to find it

- **Android:** the microphone key on the Gboard keyboard, in any text field. Google Docs has Voice Typing under Tools on a laptop.
- **iPhone:** the microphone on the keyboard.
- **Windows:** Windows key plus H opens dictation anywhere. Word has Dictate.
- **Any browser:** Google Docs voice typing works in Chrome on a laptop.
- **Offline, and the most accurate free option:** OpenAI's Whisper is an open-source recogniser that runs on a laptop with no connection and handles many languages and accents well. Free apps built on it exist for phones. It is slower than the built-in tools and considerably more accurate for non-standard accents.

Where it works and where it fails

The recognisers were trained mostly on recorded speech that is disproportionately American and British English, then other large languages. The consequence is the check question: a teacher with a "standard" accent gets near-perfect transcription; a student with a strong regional accent, or a Nigerian, Indian or Kenyan English that the recogniser has heard less of, gets an error every few words. The same for languages: Hindi, Spanish and Arabic dictation is good; Yoruba, Amharic and most small languages range from patchy to absent.

So: test with the actual student, on the actual phone, before building anything on it. Five sentences. If the built-in tool struggles, try Whisper-based apps, which have wider accent coverage. If nothing works well enough, dictation to a person — a scribe — is still the answer, and the model's role is nothing.

And treat the transcription as a draft. The student reads it back and corrects it. That correction step is not a burden; it is a reading and editing exercise that a student who could not write at all was previously excluded from.

Text-to-speech: where to find it

- **Android:** Select to Speak under Accessibility. Highlight, tap play.
- **iPhone:** Speak Selection and Speak Screen under Accessibility, then Spoken Content.
- **Edge browser:** Read Aloud, on any web page or PDF. The voices are good and there are many languages.
- **Google Docs, Word, most PDF readers:** a read-aloud option under accessibility or view.
- **The chatbots:** voice mode reads what they wrote, which for generated material is the fastest route from text to audio.

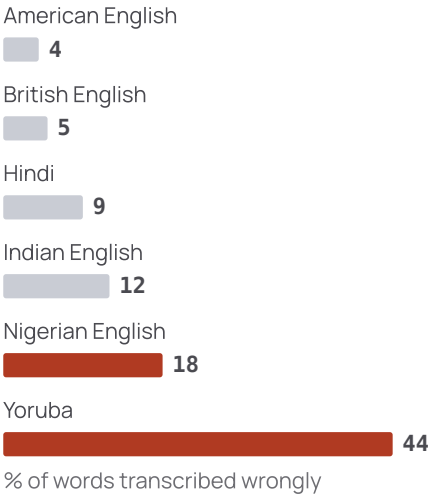
The voices are natural enough now that students do not object to them, which was not true five years ago. Coverage follows the same gradient: excellent in large languages, and a small language may have no voice at all, or one that reads with the wrong stress.

A routine that works in a real classroom

One phone, headphones, a shared document. The text for the lesson — the chunked version from the previous lesson — is in the document. A student who needs it listens while the class reads. A student who needs to dictate speaks their answer into the document and then reads it back with the read-aloud, which catches most transcription errors by ear.

Figure 7.3

Whose voice the recogniser was trained on



The same five sentences dictated into the built-in phone keyboard. A teacher with an accent the recogniser has heard a great deal of gets near-perfect transcription; a student with one it has heard less of gets an error every few words. Test with the actual student on the actual phone before building anything on it, and try a Whisper-based app where the built-in tool struggles.

The cost is a pair of headphones. The stigma is lower than most teachers expect, particularly if the read-aloud is offered to anyone who wants it.

Recording the teacher

The simplest text-to-speech is you. A phone's recorder, three minutes of you reading the key text or explaining the diagram, saved to the shared folder. Students who need it play it back. It takes less time than making a worksheet and it works in any language and any accent, because it is yours.

The honest limit

A student who dictates every answer does not practise writing, and a student who listens to every text does not practise reading. These are supports for access, so that the content is not withheld while the skill is built, and they are the right thing for a student with a lasting difficulty. For a student whose difficulty is a skill still being learned, the support is used alongside the teaching, not instead of it, and the decision about which a student needs is a teacher's, or a specialist's, and not the phone's.

BEFORE YOU MOVE ON

A teacher sets up dictation for a student who cannot write fluently. In the demonstration, in the teacher's accent, it is nearly perfect. With the student it makes an error every few words. What is the likely cause and the right response?

- A. The student is speaking too fast for the recogniser and should be told to slow down and pronounce each word separately.
- B. The recogniser has heard fewer accents like the student's; test it with her, try another tool, treat output as a draft.
- C. The microphone on the phone is faulty or positioned badly, and a better phone or an external microphone would fix the errors.
- D. Dictation does not work reliably for children's voices and should not be used with students until they are older.

The answer and why: p. 353

Lesson 54 · 9 min

Supporting students with disabilities

THE ONE THING TO KEEP

AI gives a teacher without a specialist some real tools — image description for a blind student, live captions for a deaf one, structured and predictable materials for an autistic one, chunked tasks for ADHD — and the line to hold is that these support the student's plan rather than substituting for the professional who should write it.

What a teacher without a specialist can now do

In much of the world a class contains students with disabilities and no specialist within reach. The teacher does what she can. Some of what she can do has changed.

Here is a practical list, disability by disability, of what is free and works, with the honest limit of each.

Visual impairment. The model describes images: photograph a diagram, a page, a worksheet, and ask for a description a blind student could use. It does this well for simple material and unevenly for complex diagrams. Be My Eyes, a free app, connects a blind user to an AI describer and to volunteer humans. Screen readers are built into every phone and laptop — TalkBack on Android, VoiceOver on iPhone, Narrator on Windows, NVDA free on Windows — and generated material in plain text works with all of them, where a scanned worksheet does not. The limit: descriptions of a graph or a map from the model can be wrong in the detail that matters, and for anything assessed, a person checks.

Hearing impairment. Live captions: Google's Live Transcribe on Android is free and captions speech in real time in many languages; Windows and iPhone have built-in live captions; Google Meet and Teams caption calls. A phone on the desk running Live Transcribe gives a deaf student the lesson as

text. The limit: the accent and language gradient from the previous lesson, and classroom noise; and captions are not sign language, which a student may need and which no free tool provides well yet.

Dyslexia and reading difficulties. The previous two lessons: chunked and spaced formats, glossaries, read-aloud. These are the best-evidenced supports and they cost nothing.

Autism. Predictability and structure. The model can produce a visual timetable for the day, a social story for an upcoming change — a trip, a new teacher, a fire drill — written the way social stories are written, a step-by-step task sheet with no ambiguity, and a set of the exact words a task requires. It is good at this because the formats are well documented. The limit: every autistic student is different, and a social story that fits one is wrong for another; the teacher who knows the student edits the draft.

ADHD. Chunked tasks with a checkbox per step, a timer, one instruction at a time, an explicit "you are finished when" line. The model produces the chunking in seconds. The limit: the chunking is the easy part; the routine that makes a student use it is the teacher's work.

Physical and motor difficulties. Dictation, from the previous lesson, and a document rather than a worksheet, so the student types or dictates where others write.

Speech and language difficulties. The model can generate practice material at a controlled level, picture-word cards, sentence frames. The limit: speech therapy is a clinical practice and the material supports it; it does not deliver it.

The line

The check question is about the line, and it is this: a support plan for a specific child is written by someone who has assessed that child. The model has not. It can produce a document that looks like a plan — strategies, targets, review dates — and the document will be generic, because it was generated from the

general literature, not from the student. Filed as the student's plan, it does two harms: it stands where an assessment should be, and it can satisfy a system that then never refers the student to the specialist they need.

Use the model for what is under that line. Draft the social story. Make the chunked task sheet. Generate the picture cards. Describe the diagram. Write the letter to the district asking for an assessment, and make it a good one. All of that helps the student today. Writing the plan does not.

Where there really is no specialist

Many teachers reading this have no route to an assessment at all. The honest advice is: use everything under the line, keep a record of what you tried and what happened, and keep asking. The record is the thing that eventually gets a student assessed, and the model can help you keep it — a dated log, a summary of a term's observations for a referral letter. What it cannot do is be the assessment, and a teacher should not let it pretend to be.

Involving the student and the family

A student with a disability usually knows what helps. Ask. A parent has usually tried things. Ask. The model's drafts go to them as drafts — "here is a timetable format, does this look like something that would work" — and the version that survives is the one they shaped. This is not politeness. It is how the generic becomes specific, and the generic is all the model has.

The honest limit

The tools in this lesson are for access, and access to a lesson is not the same as the specialist teaching that a disability may require. A blind student with a good screen reader and well-formatted material can follow the lesson; she still needs braille, or orientation and mobility, or whatever her plan says, from people trained to provide it. The model widens what a non-specialist teacher can do, and a teacher should be clear-eyed that widening is what it is.

BEFORE YOU MOVE ON

A teacher with no access to a specialist uses a model to generate a "support plan" for an autistic student, including strategies and targets, and files it as the student's plan. What is the problem?

- A. The model's strategies are probably wrong for autism, since it has no clinical training and generates from general text.
- B. Nothing is wrong, since some plan is better than no plan at all in a school where no specialist is available to write one.
- C. A plan for one child needs assessment of that child; the model supports a plan, it cannot write one.
- D. The teacher should have used a paid educational tool with a dedicated special-needs module rather than a general chatbot.

The answer and why: p. 353

Lesson 55 · 8 min

The quiet student and the private question

THE ONE THING TO KEEP

Students ask a chatbot what they will not ask a teacher in front of the class — the question that reveals they are behind — and the same privacy that makes them ask means nobody corrects the wrong answer, so a teacher turns the private questions into a class routine.

The question they would never ask you

Every class has a student who has not understood something from three weeks ago and will not say so. Asking in front of thirty others costs something — it says "I am behind" — and for some students, particularly the ones already unsure of their place, the cost is too high. So they stay quiet, the gap stays open, and everything built on the missing piece is built on nothing.

A chatbot does not have a class in it. Nobody hears. The same student who will not raise her hand will type "what is a denominator" at ten at night and get a patient answer with no cost at all. This is one of the real goods of these tools and it is worth naming without hedging: for the quiet student, the private question is a way back into the subject.

And the same privacy is the problem

Nobody hears the question. Nobody hears the answer either. If it is wrong — and the first block explained when it will be — the student has a wrong answer she now believes, from a source that sounded certain, on a topic she was already unsure of, and no one will ever correct it. The privacy that made the question possible has removed the correction.

And the teacher does not know the question was asked. Thirty students each with a private gap, each asking privately, and the class looks fine from the front.

Turn the private into a routine

The fix is to keep the privacy of the asking and remove the privacy of the answer.

Questions you asked the machine this week. Once a week, five minutes, every student writes one question they asked a chatbot — or would have, or wish they could ask anyone — on a slip, anonymously. Collect the slips. Read a few out. Answer them.

The slips do three things. They tell you where the gaps are, in the students' own words, which the earlier lessons showed is the most valuable planning information you can get. They correct the wrong answers, because you answer the question in front of the class. And they teach the quiet student that the question she asked at ten at night was one that six other people had too, which is the thing that makes her raise her hand next time.

The model as first, not last. Tell students explicitly: ask it, then bring what it said. "It told me X — is that right?" is a question a student can ask a teacher without admitting they did not know X, because the framing is about checking

the machine. This is a face-saving route into the conversation and students use it.

A prompt card for asking well. Students who ask a chatbot badly get worse answers. A card: say what you already understand, say where you got lost, ask for one small step. "I understand what a fraction is. I don't understand why $\frac{1}{3}$ is bigger than $\frac{1}{4}$. Explain just that, with a picture I can draw." A student who has learned to ask like that has learned something about asking anyone.

What to watch for

The student whose private questions are all about something you thought was secure. The student who asks the machine and never brings it. The student who has stopped asking anyone. Each is visible if you run the routine and invisible if you do not.

And the student who asks the machine something that is not about the subject — about being afraid, about home, about not wanting to be here. The companion lesson covered what the model does with that. The slips sometimes carry it, and a teacher who runs the routine should know that they might, and what the safeguarding route is when they do.

Access, again

The quiet student without a phone has no private question to ask. The routine still works for her — the slip is the private question, and it needs no device. Run the slips as the main thing and the chatbot as one route into them, and the student with no access is not left out of the routine that matters.

The honest limit

A routine that surfaces gaps produces more gaps to fill, and there is not time to fill them all. Choose. The question that six slips asked is the one to reteach; the question one slip asked gets a two-minute conversation, or a note to the student. The routine tells you where the class is. It does not give you the hours to take them all where they need to be, and no tool does.

BEFORE YOU MOVE ON

A teacher notices that her quietest students are using a chatbot to ask basic questions about topics covered weeks ago. What is the most useful reading of this?

- A. They are using it to avoid doing the work themselves, and the tool should be restricted for that group of students.
- B. The tool is doing the reteaching for her, so she can move on through the syllabus faster than before.
- C. Quiet students simply prefer learning from machines to learning from people, and should be left to work that way.
- D. The questions reveal gaps they will not admit to in class; gather them and reteach.

The answer and why: p. 354

Lesson 56 · 8 min

Stretching the student who is ahead

THE ONE THING TO KEEP

The student who finishes in four minutes now has an unbounded question-answerer — and is also the student least likely to be corrected when it is wrong, because she is confident and you are busy — so the stretch is to make her the checker, not the consumer.

The student the lesson was not built for

Every class has one or two who finish the work in the time it takes others to read the question. The earlier lessons gave you extension tasks — the harder question about the same shirt — and those keep her in the conversation. But

an extension task is thirty minutes of work, and she has a term of them to get through, and the honest truth is that a strong student in a large class spends most of her time waiting.

The chatbot removes the ceiling. She can ask it anything, follow any thread, go as far as she likes. For the curious student this is a genuine gift, and teachers should say so.

The trap

She is also the student most likely to absorb its errors. Three reasons. She is confident, so she does not check. She is going beyond what you teach, into topics where you are less able to catch a mistake and where the model is more likely to make one — the further from the syllabus, the further down the reliability gradient. And you are busy with the twenty-eight students who need you more, so nobody is looking at what she is learning.

The check question is the result: a strong student with fluent, detailed, wrong ideas, delivered with the model's confidence and her own.

Make her the checker

The stretch that fits this student is not more content. It is the skill she is missing and the model cannot give her: telling right from wrong in confident material.

Give her the task from the check question. "You said three things about quantum mechanics that I am not sure about. Find out, from sources that are not a chatbot, which are right, which are wrong, and which are unsettled. Present it to the class in five minutes next week."

This is harder than anything on the syllabus. It requires finding real sources, reading them, comparing them with what the model said, and — the difficult part — holding an unsettled question as unsettled. It is exactly the skill the whole course has been about, aimed at the student most able to learn it and most at risk of not.

Other stretches that use the tool well

- **"Find the error."** Give her a generated explanation of a syllabus topic that contains one deliberate mistake, and the model's confident tone. She finds it. Then she writes one for the class.
- **The Socratic prompt.** A card: "Do not explain. Ask me questions until I can explain it myself." Used on a topic just beyond the syllabus, this makes the model into a tutor that withholds answers — the design the evidence lesson found does not harm and can help.
- **The evaluation task.** She writes twenty questions on the topic, asks the model, marks it, and reports the accuracy and the pattern of errors. She learns the topic by writing the questions and learns the tool by marking it.
- **Teaching it back.** She uses the model to prepare a five-minute explanation of a topic for the class — and is told the class will ask questions, so she had better understand it rather than recite it.
- **The project lesson.** The bigram model from the teaching-AI block is a two-week stretch for a strong student in any subject, and it produces the understanding that makes every other use of the tool safer.

Keep her in the room

A student who is always ahead can drift into a private world with the machine. The extension tasks that keep her in the same conversation as the class — the harder question about the shirt — matter more, not less, now that she has an unbounded alternative. Two-thirds of her stretch should be about what the class is doing; one-third can be her own thread.

The honest limit

The checker task works because there are sources beyond the model to check against — a textbook, a library, a teacher who knows the field, the internet read carefully. For a student in a school with no library and one textbook, "check it against real sources" may mean checking it against the model's own web search, which is better than nothing and not independent. Be honest with her about that: the skill is the same, the sources are thinner, and knowing that the sources are thin is itself part of the skill.

BEFORE YOU MOVE ON

A strong Class 10 student uses a chatbot to go far beyond the syllabus in physics and arrives in class with confident, detailed, wrong ideas about quantum mechanics. What is the most useful response?

- A. Ban her from using the tool for physics until the syllabus is finished.
- B. Correct her in front of the class so others learn the error.
- C. Ask her to find the errors, source by source, and report to the class.
- D. Accept the errors, since the topic is beyond the syllabus and does not matter.

The answer and why: p. 354

Lesson 57 · 9 min

When the home language has no training data

THE ONE THING TO KEEP

For a language with a few million speakers and little written on the web, the model is not a translator and will not become one through prompting — so use it for the English side only, involve speakers for the rest, and know that projects exist to build the data and that whose languages get models is a political question, not a technical one.

The bottom of the gradient

The first lesson in this block laid out the gradient: languages the model has read a great deal of, some of, little of, almost none of. This lesson is about the last group, because a large share of the world's students live there — in languages with a few million speakers, or a few hundred thousand, that are spoken daily and written rarely, and that appear on the web as a handful of pages.

For those languages the model is not a translator. It has seen too little to have learned the grammar, the vocabulary, the register. What it produces is fluent-looking text assembled from fragments and from the patterns of neighbouring languages, and it is wrong in ways that a speaker sees at once and a non-speaker never can. The check question's teacher did everything right — detailed prompts, glossaries — and it did not help, because prompting shapes what the model does with what it knows, and here it knows almost nothing.

Nothing in this course's earlier techniques changes that. It is a limit of the data, not of the prompt.

What still works

The English side, fully. Every technique in this course works for the English (or Hindi, or Kiswahili) material: the passages, the questions, the explanations, the feedback. The model does that part; the crossing into the home language is done by people.

A speaker with a draft. For the languages just above the bottom — some text exists, the model has some of it — a model draft checked by a speaker is faster than a speaker starting from nothing. The back-translation check tells you whether you are in this zone: if the round trip comes back roughly right, a speaker's check makes the draft usable. If it comes back as nonsense, the draft is not saving anyone time.

Students as the bridge. A class that speaks the language is a room full of translators. A worked example in English, translated by the students into the language they think in, explained to each other in it, and written up in both: that is a bilingual lesson that needs no model at all for the home-language half, and it is better pedagogy than a machine translation would have been.

Regional models. A few projects train models on specific language groups — African languages, Indian languages, Indonesian languages — and some are free to use. They are worth checking, because a model trained deliberately on a language does far better than a general one that saw it by accident. Coverage is patchy and changing.

Where the data would come from

The reason the model has not read the language is that it is not written down on the web, and the reason for that is history: which languages were printed, taught, digitised, and which were not. The absence is not an accident of technology. It is the shape of who has been writing on the internet for thirty years.

Projects exist to change it. Mozilla's Common Voice collects recorded speech in hundreds of languages from volunteers, and it is what makes speech recognition possible for a language it did not previously exist for. Masakhane is a network of African researchers building translation for African languages. Several Indian and South-East Asian initiatives do the same for their regions. A teacher and a class can contribute — an hour of read sentences in Common Voice is real data — and a student who has done it understands something about where the model's abilities come from that no lesson conveys.

The political question, said plainly

Whose languages get models is decided by where the data is, which is decided by history and money. A student in a small-language community is offered tools that work beautifully in English and not at all in the language of their home, and the message that carries — that their language is not one that machines are for — is a message worth naming in class and refusing. The tools are unequal because the data is unequal, and the data is unequal because of choices people made, and some of those choices can be unmade by people who care to. That is a citizenship lesson and a true one.

For the teacher tonight

Practically: use the model for the English side and your students and colleagues for the rest. Do not put a machine translation in a small language in front of a class without a speaker having read it. And if a company or a ministry offers a tool that claims to support the language, test it with three sentences and a speaker before believing the claim. The gradient is real, and the bottom of it is where the claims are least reliable and the marketing most confident.

The honest limit

This will improve. Language coverage in models grows each year, and a language with three million speakers may be reasonably served within a few years if someone builds the data. But a teacher cannot wait, and a student in Class 5 now will be in Class 10 before the model learns their language. For now the tool is for the English, and the language of the home is the students' and yours.

BEFORE YOU MOVE ON

A teacher in a region where the home language has around three million speakers and little online text finds that the model's translations are fluent-looking but wrong. She tries more detailed prompts and glossaries without improvement. Why, and what should she do?

- A. Keep refining the prompt with more examples and constraints, since better instructions always improve the output eventually.
- B. Switch to a different model, since with so many now available one of them will certainly have learned the language properly.
- C. It has read too little of the language; use it for English and rely on speakers for the rest.
- D. Translate through a related larger language as a bridge, which is reliably better than translating directly from English.

The answer and why: p. 355

CASE STUDY · MODULE 7

The Santali worksheets that arrived from the district office

An illustrative case, built from documented patterns.

A government primary school in a Santali-speaking village in Jharkhand: Class 5, thirty-four children, one teacher who reads Santali and does not write it confidently.

The pack arrived as a PDF on the headmaster's phone: forty pages of Class 5 environmental science, in Santali, in the Ol Chiki script, produced by the district under a programme to supply mother-tongue material. It was the first home-language teaching material anyone at that school had ever been sent.

Subodh Hansda taught the class. He speaks Santali at home, teaches in a mixture of Santali and Hindi, and had been translating the Hindi textbook aloud, page by page, for four years. He read the first page of the pack and stopped.

It was not wrong in an obvious way. The sentences looked like Santali. The problem was that several of them were not sentences anybody would say, one term for a common tree was a word he did not recognise, and a question about the water cycle had a verb in a form that changed what it was asking.

He ran the check he had been taught. He typed three of the Santali sentences into a chatbot in a fresh conversation — fresh, so that nothing about the original was in front of it — and asked for English. What came back did not match the Hindi original in two of the three. The second sentence, which in Hindi asked children to name three uses of a river, came back in English as a statement about rivers being useful. A question had become a sentence.

He then did what a teacher does before accusing a document: he tried to make it work. He asked the same tool to translate the Hindi page himself, with a short glossary of the terms his textbook uses, and an instruction to match the register of a passage he typed in from a Santali reader. The output was no

better and in one place worse. This was not a prompting problem, and half an hour spent rephrasing instructions taught him that faster than any explanation would have.

The two ways to be wrong

Using the pack meant putting material in front of thirty-four children in their own language, printed, official, and endorsed by the district — with errors in it that the children would not be able to identify as errors, because it was in their language and it came from the government and they were ten. Children who cannot tell a broken sentence from a real one in their home language learn that their language is the one that comes out strange.

Refusing the pack meant going back to what he had: a Hindi textbook and his own voice. That is not nothing — he is good at it, and the children learn — but it also meant that the only home-language material the school had ever received went unused, that the district would record the school as non-compliant, and that the headmaster, who had been pleased, would be told his teacher had rejected the government's Santali books over three sentences.

There was a third possibility and it cost something too. Two mothers in the village write Santali well. One teaches at the anganwadi and has an hour on Thursdays. Asking her to read forty pages was asking for unpaid work from a woman who already had a job, in order to fix a district's document.

YOUR ANSWERS

1. The back-translation check found two of three sentences drifting. What does that tell you about the whole forty-page pack, and what does it not tell you?

2. Could better prompting have fixed the pack – a glossary of Santali terms, a pasted sample of the register, an instruction to match the textbook?

3. The routine that survived was children translating the Hindi examples themselves. Why did it beat both the pack and a better machine translation?

STOP HERE

Write your answers before you turn the page. How it played out starts on p. 288.

This page is blank so that how it played out stays out of view until you turn the page.

CASE STUDY · MODULE 7

How it played out

He did not use the pack as printed and did not reject it. He took the six pages covering the term's topics to the anganwadi teacher, Malati Murmu, who read them in about ninety minutes across two Thursdays and marked eleven places, of which four changed the meaning of a question. Subodh corrected those six pages by hand and used them; the other thirty-four he left. He wrote a note to the district listing the four questions whose meaning had changed, in Santali and Hindi, and got no reply for five months and then a form acknowledgement. Malati was paid nothing, which Subodh raised twice with the headmaster and once at a cluster meeting. The lasting change was different from what anyone expected: he started having the Class 5 children translate the Hindi worked examples into Santali themselves, in pairs, writing both versions, which took longer than reading a printed page and produced better Santali than the pack had, because thirty-four speakers arguing about a word settle it correctly. That routine survived; the pack did not.

The questions, discussed

1. The back-translation check found two of three sentences drifting. What does that tell you about the whole forty-page pack, and what does it not tell you?

It tells you the model that produced it does not really have this language: a round trip that fails on two of three sentences is the signature of the bottom of the gradient, where the output is fluent-looking text assembled from fragments and neighbouring languages. It does not tell you which pages are safe, because the check was run on three sentences and the failures are scattered. It also cannot prove any individual sentence is correct — an error made in both directions comes back looking right — so a passing sentence is not cleared, merely not caught.

2. Could better prompting have fixed the pack — a glossary of Santali terms, a pasted sample of the register, an instruction to match the textbook?

No. Those techniques shape what the model does with what it already knows, and for a language it has read almost nothing of there is nothing to shape. A glossary fixes terms in a language the model can otherwise handle; it cannot supply grammar, register or idiom. This is the one place in the course where the prompting toolkit does not apply, and recognising that is what stops a teacher spending an evening rephrasing instructions at a limit of the training data rather than of the prompt.

3. The routine that survived was children translating the Hindi examples themselves. Why did it beat both the pack and a better machine translation?

Because the class is a room full of fluent speakers, and thirty-four of them arguing about which word is right settle it more reliably than any system with almost no data. It also does more than produce text: the students practise both languages, the content is discussed twice, and the home language is treated as something you write and defend rather than something a machine renders badly and an adult apologises for. Its cost is time, and the honest limit is that it works because the class shares one home language; a classroom with six would need something else.

MODULE 7 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. 350. If two go wrong, re-read the module before the exam.

- 1 You translate a worksheet, back-translate it in a fresh chat, and the round trip matches your original closely. What may you conclude?
 - A. The translation is verified and safe to use on an exam paper or a letter home.
 - B. The forward translation is probably sound, though an error made in both directions would survive the check.
 - C. The model has this language well enough that a speaker no longer needs to read the output.
 - D. Nothing, because a model asked to translate its own output will always reproduce the meaning it started from whatever it actually wrote.

- 2 You ask for an explanation "in Hindi" for Class 6 and your students find it harder than the English version. Why?
 - A. Hindi words are longer, which raises the reading level of any passage automatically.
 - B. The model's Hindi grammar is weak, so the sentences are difficult to parse.
 - C. It translates from English internally, so the output carries English sentence structure that a Hindi reader then has to unpick.
 - D. The Hindi it has read most in educational text is a formal register that nobody in the room speaks.

- 3 A student with dyslexia cannot get through the history text, though she follows the same content easily when it is read aloud. What do you ask the model for?
- A. A simplified version at a lower reading level.
 - B. Chunking into short sections with headings, a glossary at the top, signposted paragraphs, and every fact kept.
 - C. A summary she can read instead of the full text.
 - D. A shorter version with the less important events removed, so the reading load matches the level she is working at.
- 4 Phone dictation transcribes your speech almost perfectly and your student's with an error every few words. What does that tell you to do?
- A. Ask the student to speak more slowly and more clearly into the microphone.
 - B. Replace the handset, since the microphone is evidently faulty.
 - C. Abandon dictation for this class, because a tool that fails one student will fail the rest the same way.
 - D. Test with the actual student before building anything on it, and try a Whisper-based app.
- 5 A quiet student asks a chatbot at ten at night the question she would never ask in class. What is the cost of that, and what fixes it?
- A. Nobody hears the answer either, so a wrong one is never corrected; a weekly anonymous-slip routine fixes it.
 - B. She stops asking the teacher at all; requiring her to ask the teacher first fixes it.
 - C. The class loses the benefit of the question being answered for everyone; asking her to share it fixes it.
 - D. The company keeps a record of what she asked, and the school has no way of knowing what its students are putting into these tools.

- 6 For a language with a few hundred thousand speakers and almost nothing written online, a glossary and a pasted register sample do not improve the translation. Why not?
- A. The glossary was too short to override the terms the model defaults to.
 - B. The script is not supported by the way the tool breaks text into pieces.
 - C. Prompting shapes what the model does with what it knows, and here it knows almost nothing.
 - D. The instructions were written in English, and constraints only apply properly when they are given in the language being produced.

EXAMINATION

AI for Teachers: final examination

This paper covers the whole course:

- how a language model works and what follows from that for reliability, arithmetic and memory
- prompting as a teacher
- making explanations, worksheets, question banks and unit plans
- assessment, feedback and marking
- integrity, policy and AI-written work
- teaching AI itself to students from eight to eighteen
- the multilingual, mixed-ability classroom

56 questions, the whole pool. The site draws 25 for a sitting, pass mark 70%.
Work any 25 under your own conditions. Answers from page 356.

1 Module 1 · Understanding the machine you are about to use

A student says the chatbot lied to her. What is the accurate correction?

- A. It produced a sentence that fits, and a sentence can fit beautifully and be false. It has no way to check.
- B. It did lie, and the companies are working on making the models more honest about what they do not know.
- C. It repeated something false that it read, so the untruth belongs to whoever wrote that page rather than to the machine.
- D. It was misled by whoever asked it the same question before her, because it learns from every conversation it has.

2 Module 1 · Understanding the machine you are about to use

Which one-sentence explanation survives a class asking "then why is it wrong sometimes?"

- A. It is a very large store of facts that has gaps in the places where nobody wrote anything down.
- B. It is a machine that has read an enormous amount of writing by people and guesses what words come next.
- C. It is a robot brain that thinks in a different way from the way we think.
- D. It is like a search engine that finds the best answer available on the internet and puts it in its own words.

3 Module 1 · Understanding the machine you are about to use

A department of six teachers decides to share one free account to save trouble. What happens?

- A. Nothing much, because the message cap is applied per device rather than per account.
- B. The tool suspends the account, because use from several places at once looks like automated activity to the provider.
- C. The answers improve, because the model picks up the department's house style from the shared history.
- D. They hit the message cap by lunchtime, because the cap is per account. Each teacher needs their own.

4 Module 1 · Understanding the machine you are about to use

You switch off the setting labelled something like "improve the model for everyone". What have you achieved?

- A. The tool now runs a model on the phone itself, so nothing you type leaves the device at all.
- B. Your conversations are deleted at the end of each session, which is why the history no longer appears.
- C. What you type is taken out of the training pipeline; the company still receives and stores it.
- D. The conversation is now private, in the sense that nobody at the company is able to read what you send.

5 Module 1 · Understanding the machine you are about to use

A tool that was sharp at seven in the morning is producing thin, shallow work at nine in the evening. What is the usual reason?

- A. Your prompts have got worse as you have got tired, which is the commonest cause of a bad evening session.
- B. Demand is high, so the service is quietly giving you a smaller, faster model.
- C. The model was updated during the day and the newer version is weaker at this kind of task.
- D. You are near the message cap, so the tool is shortening its answers to make your remaining quota last longer.

6 Module 1 · Understanding the machine you are about to use

You are reading a generated answer before using it with a class. What should you mark as unverified?

- A. The passages written in a more confident tone, since hedged sentences mark the places the model is unsure of itself.
- B. Every proper noun, number, date and quotation.
- C. The first and last paragraphs, where errors are known to cluster.
- D. Every sentence containing an absolute word such as "always" or "never".

7 Module 1 · Understanding the machine you are about to use

For which request is it better to turn web search off?

- A. This year's exam dates from your board.
- B. The population of your district at the last census.
- C. Anything you plan to give to students, since a searched answer cannot be checked afterwards.
- D. A set of likely misconceptions and three ways to explain one idea.

8 Module 1 · Understanding the machine you are about to use

Your pasted brief is nine paragraphs long and one constraint matters more than the rest. Where should it go?

- A. In the middle, where a long input is read most carefully.
- B. At the top or the bottom of the brief.
- C. In a separate first message, so that it is established before any of the material arrives.
- D. Repeated in every paragraph, because a constraint stated once is not applied to a long input.

9 Module 2 · Prompting as a teacher

You ask the model to interview you before it writes anything. What has to be added so the questions are useful?

- A. Ask me questions until you are completely confident that you understand my class, my room and my school.
- B. Ask me about my learning objectives and my success criteria before anything else.
- C. Ask about constraints I might not have mentioned, rather than about objectives.
- D. Ask me at least ten questions, because a short interview produces a generic result.

10 Module 2 · Prompting as a teacher

Which of these should be left out of a teaching prompt entirely?

- A. Your teaching philosophy, and a request that the activity be engaging.
- B. The class size, which the model has no way to use.
- C. The list of equipment you do not have, because negative constraints confuse a system that predicts what to include.
- D. The live misconception, which belongs in a second message once the draft exists.

11 Module 2 · Prompting as a teacher

You paste a good question of your own and ask for eight more like it.

What else should the message say?

- A. Paste at least six examples, so that the pattern cannot be misread.
- B. Describe in as much detail as you can what makes the example good, since the model cannot infer your reasons from the example itself.
- C. Keep the structure, the length and the kind of thinking; change the passage, the numbers and the context.
- D. Put the example after the instructions rather than before them, so that it sits nearest the answer being written.

12 Module 2 · Prompting as a teacher

You are struggling to describe the difference between a question that tests reading and one that tests understanding. What is the fastest way to convey it?

- A. Explain in words why the weak kind is weak, since a model shown poor work lowers its standard for everything after it.
- B. Show three good examples, because bad examples waste space in the memory window.
- C. Never show a bad example, because the model copies whatever it is shown.
- D. Show one labelled as the kind you do not want and one labelled as the kind you do.

13 Module 2 · Prompting as a teacher

You want a worksheet that is both rigorous and accessible, and iterating between the two is producing an oscillation. What should you do?

- A. Alternate the two requests until the output settles somewhere in between.
- B. State both goals in the same message and let it find the balance in one pass.
- C. Ask for two versions, one for each goal, and then paste both back and ask for a merger.
- D. Choose one goal and drop the other, because a model cannot hold two competing constraints at once.

14 Module 2 · Prompting as a teacher

You want three hundred questions in a spreadsheet you can sort and filter. What do you ask for?

- A. Plain text with the columns aligned by spaces, which any spreadsheet can split reliably into cells.
- B. A table, which pastes cleanly into both Word and a spreadsheet.
- C. Markdown, and then use the spreadsheet's import dialogue to convert it.
- D. CSV with named columns, saved as a .csv file and opened in a free spreadsheet.

15 Module 2 · Prompting as a teacher

Which pair of format instructions removes the most editing from a term's worth of generated material?

- A. Use bullet points throughout, and keep every section to roughly the same length so the page is easy to scan.
- B. Do not restate the question in the answer; no preamble and no summary.
- C. Be concise, and use simple language.
- D. Write in British English, and avoid American spellings.

16 Module 2 · Prompting as a teacher

A prompt that produced excellent distractors in March produces ordinary ones in September. What is the best repair?

- A. Rewrite the prompt from scratch, with considerably more detail than the original had.
- B. Switch to another tool, since the prompt was written for the older model and no longer fits this one.
- C. Add a line asking the model to behave as the previous version behaved, naming that version if you know it.
- D. Paste the March output you saved and ask for eight more like these.

17 Module 3 · Planning and making materials

Most teachers ask a model for a full lesson plan. What should they ask for instead?

- A. A list of learning objectives matched to the syllabus wording.
- B. Likely wrong answers, and what students are thinking when they give them.
- C. A complete forty-minute plan that can be trimmed afterwards.
- D. A set of engaging group activities, since that is the part of planning that takes teachers longest to write from scratch.

18 Module 3 · Planning and making materials

What does a generated lesson plan get wrong every single time?

- A. Timing: it writes fifty-five minutes of activity for a forty-minute period.
- B. Sequencing, which is why the concrete case is always placed before the definition.
- C. Core subject accuracy, which is where the great majority of its errors appear.
- D. Vocabulary level, which comes out far below the year group you named in the prompt.

19 Module 3 · Planning and making materials

You ask for a passage that uses ten target words. What is the failure to check for?

- A. The words are used many times each, which turns the passage into a vocabulary drill rather than a story.
- B. The words are put in bold, which tells the student the answer before they have read the sentence.
- C. The words are defined in the text, which removes the work of inferring meaning from context.
- D. The words arrive together in one paragraph, which satisfies the instruction at its cheapest point.

20 Module 3 · Planning and making materials

You want the same passage at three reading levels so the class can discuss one text. Which instruction is load-bearing?

- A. Keep every fact the same.
- B. Keep the same number of paragraphs in each version.
- C. Keep the target vocabulary and change everything else about the passage.
- D. Keep the same comprehension questions, so that all three versions can be marked from a single answer key.

21 Module 3 · Planning and making materials

Your question bank has a difficulty column filled in by the model. What should happen to it over the year?

- A. Trust it, since it is calibrated against a very large number of student responses.
- B. Ask a second model and average the two ratings to reduce the error.
- C. Rate every question yourself before use, because your judgement of your own class is reliable before anyone has answered.
- D. Replace it with the proportion of your class who got the question right last time it was used.

22 Module 3 · Planning and making materials

The model has labelled each question in your bank as recall, application or inference. How much weight should the labels carry?

- A. Little: around a third of what it labels application is recall in a costume.
- B. Full weight, since it is applying a standard taxonomy consistently across the set.
- C. None, and the column should be dropped, since question type does not affect what a test measures.
- D. Full weight once you ask it to relabel using Bloom's taxonomy, which removes the ambiguity in its own categories.

23 Module 3 · Planning and making materials

You have a lesson to deliver on a projector that fails about once a fortnight. What should the model produce?

- A. A finished deck in the school template, with bullets shortened to fit, which is what the built-in generators are designed to do.
- B. The outline and the speaker notes, with the layout left to a free slide tool.
- C. Images for each slide and a title for each, since teachers write text faster than they find pictures.
- D. A complete deck that you then edit down to the slides you actually need.

24 Module 3 · Planning and making materials

A twelve-week unit plan comes back looking reasonable. What are the two things it could not have known?

- A. The prerequisites for each topic, and the points at which each should be revisited.
- B. Your school calendar, and which topics your students found hard last year.
- C. The syllabus content, and the order in which the board examines it.
- D. Your class's reading level and how many lessons a week you have, neither of which a prompt can really convey.

25 Module 4 · Assessment, feedback and marking

Which judgement is entirely outside what a machine marking a script can see?

- A. Whether every claim is followed by a quotation that supports it.
- B. Whether the spelling and punctuation are accurate.
- C. Whether this piece represents progress for this student.
- D. Whether the paragraphs follow the order announced in the introduction, which requires holding the whole piece in view at once.

26 Module 4 · Assessment, feedback and marking

What is the reason not to let a model set a grade that goes on a record?

- A. Examination boards prohibit it, requiring every recorded mark to be produced by a named human marker.
- B. It cannot read handwriting accurately enough for the mark to be fair to every student.
- C. A parent will ask why, and there is no reasoning you can walk through and nothing to teach from.
- D. It is less accurate than a teacher on every marking task that has ever been studied.

27 Module 4 · Assessment, feedback and marking

Which of these belongs in your own adjustment on the script rather than in the rubric the model applies?

- A. The evidence and quotation criteria, since those are the ones most distorted by the length bias.
- B. Originality against what this class usually produces, and improvement since October.
- C. The level descriptors, which the model infers more consistently on its own.
- D. The mark totals, which pull its scoring toward the middle of the scale.

28 Module 4 · Assessment, feedback and marking

What is the honest limit of a rubric written so that every criterion is observable?

- A. It cannot be applied outside subjects where students write extended prose.
- B. It removes the length bias so completely that a very short answer can now take full marks without saying anything.
- C. It ends up too long for students to read, so in practice only the teacher uses it.
- D. It can be met mechanically: it describes the floor of competence, not the ceiling of quality.

29 Module 4 · Assessment, feedback and marking

What makes feedback most likely to change the next piece of work?

- A. Phrasing each point as something to do in the next version rather than as a fault in this one.
- B. Giving six points, so the student has enough to choose from.
- C. Putting the mark beside the comment, so the student can see what the advice is worth.
- D. Opening with two or three things that went well, so that the student is willing to keep reading as far as the part that matters.

30 Module 4 · Assessment, feedback and marking

Students are using a prompt card that forbids the model to rewrite their draft. What goes wrong, and what holds it in place?

- A. It forgets the rubric after a few turns, so the whole rubric and the whole task description have to be pasted again before every single question the student asks it.
- B. It refuses to continue after several exchanges, so a fresh chat has to be opened each time.
- C. It starts suggesting phrasings and then whole sentences, so the card is restated every third message and a version trail is kept.
- D. It becomes harsher over a long conversation, so students are told in advance to disregard the tone.

31 Module 4 · Assessment, feedback and marking

In peer assessment with a model as third marker, what is the part that teaches?

- A. The average of the three scores, which is more reliable than any single one of them.
- B. The model's score, which settles which of the two student markers had read the rubric correctly.
- C. The number of criteria on which all three markers agreed, which measures how well the class understands the rubric.
- D. The disagreements, taken back to the rubric and the essay to find the line each marker read differently.

32 Module 4 · Assessment, feedback and marking

You have photographed thirty handwritten scripts and transcribed them. What is the record of what each student wrote?

- A. The paper. The transcription is a draft to check against it.
- B. The transcription, once you have read it through once.
- C. The photograph, since it is the closest available copy of the original.
- D. Whichever of the two the student agrees is accurate, since the student wrote it and is best placed to confirm what it says.

33 Module 5 · Integrity: AI-written work and what to do about it

You are uneasy about one particular essay. How should the conversation open?

- A. By showing the student the detector percentage and asking them to account for it.
- B. By asking straight out whether they used AI, before they have a chance to prepare an answer.
- C. By asking them to talk you through how they got to this argument.
- D. By sending it to a second detector, on the basis that agreement between two tools is a stronger signal than one.

34 Module 5 · Integrity: AI-written work and what to do about it

What is the point of one piece of in-class writing on paper early in the term?

- A. It gives you a baseline for what each student sounds like, so that you notice honestly rather than suspect.
- B. It provides a sample against which later submissions can be scored for similarity.
- C. It establishes that the student can write without help, so later work can be assumed to be theirs.
- D. It builds a folder of last year's and this year's work showing how each student's writing has developed over time.

35 Module 5 · Integrity: AI-written work and what to do about it

In the three-level scheme, what does Level 2 permit?

- A. Use for anything at all, provided the student says so afterwards.
- B. Use to explain something not understood, or to get unstuck, but not to produce anything handed in.
- C. You may draft with it, and you must show what it produced and what you changed.
- D. Free use, since the task is not formally assessed and unassessed work needs no disclosure.

36 Module 5 · Integrity: AI-written work and what to do about it

Which rule keeps an AI policy from building a two-tier classroom?

- A. Students without access should be given extra time to complete the same task.
- B. Students without a device should work in pairs with someone who has a phone, so nobody is left out.
- C. No task should require a student to have their own AI access to complete it.
- D. The school should provide an account for every student before any AI task is set.

37 Module 5 · Integrity: AI-written work and what to do about it

What is the exercise that most reliably teaches a class to distrust a confident wrong answer?

- A. Explain the reliability gradient carefully, with an example of each level.
- B. Read out a wrong answer as though it were authoritative, let the class accept it, then show them the truth.
- C. Ask students to list the topics on which AI should not be trusted.
- D. Show a correct answer and an incorrect one side by side on the board and ask the class to work out which of the two is which, and how they can tell.

38 Module 5 · Integrity: AI-written work and what to do about it

Why does a class AI policy need an exemption list naming spellcheck, single-word translation and the calculator?

- A. Without it, the school's own approved tools fall under the ban, which makes the policy unlawful in some places.
- B. Without it, students will claim afterwards that they did not know which tools were covered.
- C. Without it, the page runs to more than one side of paper.
- D. Without it, every student is in breach from the first week and the policy cannot be enforced at all.

39 Module 5 · Integrity: AI-written work and what to do about it

Which homework task still measures something now that the model exists?

- A. Revise these three topics; five questions from memory at the start of the next lesson.
- B. Summarise chapter seven in two hundred words.
- C. Answer the comprehension questions at the end of the chapter.
- D. Research the causes of the 1857 uprising and write a one-page report citing three sources.

40 Module 5 · Integrity: AI-written work and what to do about it

What is a declaration of AI use for?

- A. It shows the marker that the student did not over-use the tool.
- B. It moves responsibility for any errors onto the tool's maker.
- C. It lets the marker deduct marks in proportion to how much of the work was machine-assisted.
- D. It tells the reader what the work rests on, in the same way a bibliography does.

41 Module 6 · Teaching AI itself

Why write the model's answers down before the lesson rather than demonstrating live?

- A. Students copy more accurately from paper than from a screen at the front.
- B. It saves the mobile data you would otherwise spend during the period.
- C. Live demonstrations fail, and the model may give a different answer on the day.
- D. The model returns shorter answers when many people ask the same question at once, which spoils a demonstration.

42 Module 6 · Teaching AI itself

A story is retold down a line of five students and changes as it goes. What is that activity teaching?

- A. Bias in training data, because the class has just been the training data.
- B. That summarising loses detail, which is what the model does to a long document.
- C. That memory is unreliable, and therefore that students should take notes.
- D. That information degrades at every step, which is why a model trained on other models' output gets worse over time.

43 Module 6 · Teaching AI itself

What is the right idea to teach children aged eight to ten?

- A. Prediction: a spread of likely next words, and why the tone never changes between right and wrong.
- B. Training data and bias, and how to measure a tool's accuracy on a set of questions.
- C. A machine shown many labelled examples learns to sort new ones, and gets it wrong at the edges.
- D. How a chatbot works, since that is the AI they actually meet and the one they ask about most often.

44 Module 6 · Teaching AI itself

Which words should be withheld from nine-year-olds playing the sorting game?

- A. Examples, pattern and mistake.
- B. Algorithm, neural network and probability.
- C. Training and testing, which mislead at this age.
- D. Machine learning, because the phrase suggests the machine learns in the way that a child learns.

45 Module 6 · Teaching AI itself

What should be on the card students are given about fabricated intimate images?

- A. Check the hands and any visible text in the image before deciding what to do about it.
- B. Delete it at once and say nothing, so that the spread stops with you.
- C. Report it to the platform first, and to the school only if the platform has not acted within a day.
- D. Do not forward it; screenshot only what is needed to report; tell a named adult that day.

46 Module 6 · Teaching AI itself

Which set of signs should prompt a conversation about a student's use of a companion app?

- A. Withdrawal from friends, tiredness from late-night use, and speaking of the character as a person.
- B. A sudden improvement in the fluency of their written English.
- C. Refusing to use the app when an adult is nearby.
- D. Asking a lot of questions in class about how language models are built and what they are trained on.

47 Module 6 · Teaching AI itself

Students build a tally of which word follows which from one page of text, then generate sentences from it. What does the output show them?

- A. It has no vocabulary of its own, so it invents words that do not exist.
- B. It cannot handle punctuation, which is why the sentences run together without stopping.
- C. It looks one word back, so the sentences are grammatical in pieces and meaningless as a whole.
- D. It reproduces the source text exactly, which shows that these models memorise what they are trained on.

48 Module 6 · Teaching AI itself

In a school distributing the AI unit across subjects, which department naturally owns provenance and the question of where a source came from?

- A. Computing, which teaches how the files themselves are made.
- B. History, which has asked where a source came from and why for a century.
- C. Art, which already deals with authorship, style and imitation.
- D. Citizenship, because provenance is fundamentally a question about trust in institutions rather than about evidence.

49 Module 7 · The multilingual, mixed-ability classroom

A translated worksheet uses an unfamiliar word for a subject term your students have already learned. What fixes it?

- A. Ask for a formal register, which uses standard terminology throughout.
- B. Ask it to transliterate the English terms instead of translating them.
- C. Paste your textbook's glossary and say to use exactly those terms.
- D. Ask it to use the terminology of the national curriculum, which it can look up for whichever country you name.

50 Module 7 · The multilingual, mixed-ability classroom

Which of these must not be produced by machine translation alone, even with a back-translation check?

- A. A practice worksheet for Friday's lesson.
- B. A report to a parent or a safeguarding document.
- C. A ten-word glossary of subject terms for the front of an exercise book.
- D. A reading passage for homework, because students will read it without a teacher there to correct any errors.

51 Module 7 · The multilingual, mixed-ability classroom

Why can the model write a worked example in Hinglish or Sheng at all?

- A. Because it handles code-switching as a general skill, independent of the particular languages involved.
- B. Because both halves of such a mix are always large languages that share a script.
- C. It cannot: mixed language is spoken rather than written, so there is almost nothing for it to have learned from.
- D. Because those mixes are written in enormous quantities online, in forums, comments, messages and subtitles.

52 Module 7 · The multilingual, mixed-ability classroom

A parent objects that a worked example written in Hinglish will damage her daughter's English. What is the honest answer?

- A. Explanation in the strongest or mixed language improves understanding of the content and does not harm acquisition of the formal register.
- B. Written mixed language should indeed be avoided, and the mix kept to speech.
- C. The formal register should be taught first, with the mix used only during revision.
- D. She is right to worry, and the mixed-language version should therefore be given only to the students who are already failing to follow the English one.

53 Module 7 · The multilingual, mixed-ability classroom

Which part of making a text reachable for a dyslexic reader is yours rather than the model's?

- A. The glossary of hard words and the signposting sentences.
- B. The sentence length and the section headings.
- C. The order of the sections, which decides how easily a struggling reader follows the argument.
- D. Font, spacing, alignment and background colour.

54 Module 7 · The multilingual, mixed-ability classroom

Where does the line fall for a teacher supporting a student with a disability and no specialist within reach?

- A. It can draft the social story and the chunked task sheet; it cannot write the student's support plan.
- B. It can write the plan, as long as a teacher reads it and signs it.
- C. It can describe images but should not generate any material at all for a student with a disability.
- D. It can do anything a specialist would do, provided the teacher checks every output before it reaches the student.

55 Module 7 · The multilingual, mixed-ability classroom

What is the right stretch for the student who finishes everything in four minutes and now has an unbounded question-answerer?

- A. More questions of the same kind, so that the extra time is filled productively.
- B. Make her the checker: verify the model's claims against sources that are not a chatbot.
- C. Let her work ahead in the textbook while the rest of the class catches up.
- D. Give her the chatbot and a topic beyond the syllabus, because her own curiosity will take her further than any task you could set.

56 Module 7 · The multilingual, mixed-ability classroom

Your students' home language has a few hundred thousand speakers and almost nothing written online. What do you do this term?

- A. Prompt in the home language, so that the model works in it directly rather than translating.
- B. Wait a year, since language coverage improves with every release.
- C. Use the model for the English side in full, and involve speakers for the rest.
- D. Ask the model to learn the language from several pages of text you paste in, and then have it translate.

ANSWERS

Answers

The key is grouped by module. Each entry gives the page of the question, the question cut short, the right answer and why. For a lesson check it also says why each wrong option tempts. That part is the teaching, not the marking.

1	Understanding the machine you are about to use	314
2	Prompting as a teacher	320
3	Planning and making materials	326
4	Assessment, feedback and marking	332
5	Integrity: AI-written work and what to do about it	338
6	Teaching AI itself	344
7	The multilingual, mixed-ability classroom	350
	Examination	356

Module 1 • Understanding the machine you are about to use

The module starts on p. 9.

MODULE 1 TEST

Module 1 test, Q1, p. 46

You need three things for a Class 9 worksheet: an explanation of photosynthesis, the year your town's first secondary school opened, and 15 per cent...

Answer B: The explanation, because core school science has been written down consistently by very many people.

Why: Reliability tracks how many times a thing has been written down and how consistently. Photosynthesis is in every textbook in every language, saying the same thing, so the model's guess is effectively a fact. The founding year of one town's school has been written down rarely, so the model produces a plausible-looking year. And multiplication is not a rule the model applies; it predicts the digits, which is why it slips as the numbers grow. The identical tone across all three is exactly why you cannot use the answer to judge the answer — you judge the question.

Module 1 test, Q2, p. 46

An answer arrives with three links and a line saying the tool searched the web. What has that changed?

Answer C: It has fixed staleness and given you pages to open; it has not fixed the judgement about which page to trust.

Why: Search is a genuine improvement for anything recent or specific, and the real gain is that you now have a page you can open. What it does not change is how the pages were chosen or read: the model ranks and blends them by pattern, so a content farm that copied an old syllabus looks like the exam board, and where three pages disagree the summary may say a fourth thing none of them said. Summaries also lose the exceptions and conditions first, because those are the least predictable words in a sentence.

Module 1 test, Q3, p. 47

You paste five past papers and ask for ten questions in the style of paper four. Back come ten confident questions in the style of...

Answer A: The material did not all fit, so the tool kept part of it and answered as though it had the lot.

Why: Refusal is the good outcome; quiet truncation is the bad one. Everything the model can consider has to fit one window at once, and when it does not, the tool keeps part and answers as though nothing were missing. Nothing warns you. The cheap defence is to give it less and then check what it has: before you start, ask it to tell you the heading of the last section it can see. If it names the wrong section, the material did not fit.

Module 1 test, Q4, p. 47

Asked for a Class 10 maths revision sheet, the model produces US Algebra II content, dollars, and the word "sophomore". What is the quickest thing...

Answer D: Paste the board's chapter list and one past paper, and name the terms and units to use.

Why: The lopsidedness is in the training text, not in one product, so changing tools moves you a small distance and no further. The model is poor at guessing your syllabus and very good at following one it is shown, which is why twenty lines of chapter list and one past paper shift the output almost entirely. Examples beat descriptions here as everywhere: she did not have to explain the Nigerian or Indian school system, only show the shape.

Module 1 test, Q5, p. 48

You want twenty word problems with a trustworthy answer key. Which approach gives you one?

Answer D: Ask for problems whose answers you specify, or work the key yourself with a calculator.

Why: The model is good at the method — setting up the equation, choosing the formula, explaining why — and unreliable at the digits, because arithmetic is a specific that appears rarely in text. Fitting a story to a number you chose is far safer than fitting a number to a story. Reasoning modes do improve arithmetic markedly, but better is not exact, and you cannot see from the outside whether a given number came from a computation or a guess. Asking it to check its own working leaves the same predictor doing both jobs.

Module 1 test, Q6, p. 48

You ask the same question in three fresh chats and get the same answer three times. What have you established?

Answer A: That the pattern is peaked, which is not the same thing as the answer being true.

Why: Disagreement is a red flag you can trust; agreement is a green flag you cannot. Three matching answers tell you the model's learned distribution is sharply peaked, and popular misconceptions are the most stable outputs it has — the ten-per-cent-of-the-brain claim comes back identically every time. Three draws from one system are one source however many times you read it, which is why asking a second company's model is a stronger check than asking the same one again.

THE LESSON CHECKS**Lesson 1, p. 13**

A student tells the AI it made a mistake. It apologises and changes its answer. She says this proves it understood its error. What should...

Answer A: An apology is simply the kind of sentence that follows a correction in the writing it learned from — it will produce one whether or not it was wrong. Tell it a correct answer is wrong and watch it fold.

Why: It predicts likely next words from human writing, so sounding right and being right come apart.

B She is right — noticing and...: Treats fluent self-correction as evidence of an internal grasp of truth, when the same behaviour appears with nothing to correct.

C The company that built it wrote...: Assumes the behaviour comes from hand-written rules rather than from patterns in the text it learned from.

D When challenged, it went and looked...: Imagines the system consults a reference at the moment of asking, like a search engine, rather than generating text.

Lesson 2, p. 16

A teacher on a shared Android phone with patchy data wants to generate practice questions each evening. Which single setup decision matters most before she...

Answer B: Turning off the setting that lets the company train on her conversations, and working in the browser.

Why: Every major model has a free tier that runs in a phone browser or inside WhatsApp; the choice that matters is not which is cleverest but which one you have turned off training on your data.

A Picking whichever tool scored highest on...: Treats capability as the deciding factor; the free tiers of the major tools are close enough for worksheet generation that the difference is swamped by how she prompts and checks.

C Installing every app on the phone...: Assumes storage, data and time she does not have, and comparison across tools is a verification technique for specific doubtful answers, not a daily workflow.

D Paying for one subscription first, since...: Free tiers are the same models with usage caps, not cut-down demos; a teacher generating a few worksheets a night rarely hits the cap.

Lesson 3, p. 21

A geography teacher asks a chatbot two questions: the boiling point of water at sea level, and the population of her town at the last...

Answer C: The boiling point is almost certainly right; the population may be a plausible invention and needs checking against the census.

Why: Reliability tracks how many times a fact has been written down by people — so the model is sound on photosynthesis and unsound on the founding date of your district, and you can predict which before you ask.

A Both are factual questions with definite...: Treats correctness as a property of the question type rather than of how often the answer appears in training text; a rare local figure has been written down a handful of times, a physical constant millions.

B Neither answer can be trusted, because...: Overshoots: a guess informed by millions of consistent examples is reliable in practice, and refusing to distinguish leaves the teacher unable to use the tool at all.

D The population is the more reliable...: Inverts the mechanism; specific numbers with few occurrences in text are among the least reliable outputs, and the model has no arithmetic or lookup behind them.

Lesson 4, p. 25

A teacher asks a chatbot for this year's exam timetable and gets a confident list of dates with no links or sources shown. What is...

Answer D: The dates were generated from the pattern of previous years' timetables, and none of them should be trusted.

Why: A model on its own has no way to consult anything at answer time; when a tool adds web search it fetches pages and summarises them, which fixes staleness but not judgement about which page to trust.

A The model looked the timetable up...: Assumes search happened when there is no evidence of it; a search-backed answer shows the pages it used, and an answer with no sources is generated from training text.

B The exam board published the timetable...: The current year's timetable almost certainly post-dates the training cutoff, and even if it did not, precise dates are rare specifics the model reconstructs rather than recalls.

C The model declined to search because...: Invents a refusal mechanism; models do not decline to search on this basis, and the answer was given, not withheld.

Lesson 5, p. 29

A teacher pastes an entire 90-page PDF textbook into a free chatbot and asks for questions on chapter 7. The questions come back well-formed but...

Answer B: The book was too big for the window, so only part of it was kept, and chapter 7 was not in that part.

Why: A model can only see a fixed window of text at once, so it holds a chapter but not a textbook and forgets the instruction you gave forty messages ago; start a new chat per task and paste the material in every time.

A The model does not understand PDF...: Most tools read PDFs well; the problem is size, not format — the questions were specific, just to the wrong part.

C The model read the whole book...: Invents a preference; the model works from what is in front of it, and when the material is present it uses that material rather than a remembered edition.

D Free tiers cannot process uploaded files...: Free tiers do accept uploads within limits; the failure is a partial read, which is exactly why the output was specific but wrong.

Lesson 6, p. 33

A Class 10 teacher in Bihar asks for "a chapter test on trigonometry for Grade 10". The test comes back with topics from a US...

Answer B: It defaulted to the syllabus it has read most; pasting her board's chapter list and one past paper will fix the scope.

Why: The training text is dominated by US and UK educational writing, so unprompted the model assumes Common Core, dollars and an American exam format; paste your syllabus and a sample paper and it adapts well, but it will never do so unasked.

A The model does not know Indian...: Overgeneralises from a scope error; the mathematics itself is universal and reliable, only the syllabus boundary was wrong.

C She used the word "Grade" instead...: A single word nudges the register but does not explain the topic list; the topic list comes from the model's default assumption about what "Grade 10 trigonometry" contains.

D The board's syllabus is not public...: Misses that the fix is to supply the syllabus in the prompt, which requires nothing to be in the training data at all.

Lesson 7, p. 36

A maths teacher asks a chatbot to produce ten word problems with an answer key. She notices the model's step-by-step working on question 6 is...

Answer C: The method is language, which it handles well; the product is digits predicted from pattern, which fails on uncommon combinations.

Why: Next-word prediction is a poor way to multiply, so every number in generated material is checked with a calculator, and the safe pattern is to have the model write the working while a separate tool does the arithmetic.

A The model made a deliberate small...: Attributes intent; the model has no goals about her behaviour and no mechanism to plant a test.

B The word problem was badly worded...: Does not fit the evidence; the working matched the problem and only the final arithmetic failed.

D The answer key was generated in...: Possible in some workflows but does not explain a single wrong product inside otherwise correct working on the same question.

Lesson 8, p. 40

A teacher asks the same factual question in three separate chats and gets the same answer each time. She concludes the answer is verified. What...

Answer D: Agreement shows the pattern is stable, not correct; a common misconception is reproduced every time.

Why: Answers are sampled from a distribution rather than looked up, so asking the same question three times and comparing is a real check — disagreement flags an unreliable question, though agreement never proves the answer is right.

A Three repetitions is too few to...: Treats it as a sample-size problem; more repetitions of the same model still draw from the same learned pattern and cannot correct a shared error.

B The model remembered her earlier chats...: Separate chats do not share memory by default; the answers agree because the learned distribution is peaked, not because of copying.

C Facts should be asked once with...: Repetition is a legitimate check for factual questions — it exposes instability — the flaw is in what agreement is taken to prove.

Module 2 • Prompting as a teacher

The module starts on p. 49.

MODULE 2 TEST

Module 2 test, Q1, p. 86

Four lines are proposed for a prompt asking for a fractions starter. Which one is the line the model could not have got from anywhere...

Answer B: Half of them believe one third is smaller than one quarter because three is smaller than four.

Why: The topic is in every fractions resource on the internet, and the model's default is the average of all of them. Adjectives about engagement and best practice describe a quality the model is already aiming at and steer nothing. The live misconception, in the students' own words, is the one thing that is in your marking and nowhere else, and it is what the lesson actually has to fix. Keep it near the top or the bottom of the prompt, because material buried in the middle of a long input is used least.

Module 2 test, Q2, p. 86

You reply to a set of generated comprehension questions with "make these better" and get four more questions of the same kind. Why?

Answer A: "Better" has no direction, so it resolves to the commonest meaning in text, which is more.

Why: This is not the model misunderstanding; it is the model correctly guessing what an unspecified "better" usually means for generated content. Every vague critique has a default it falls into: "more engaging" gives exclamation marks and a game, "simpler" gives shorter sentences with the same concepts. The fix is to name the fault, point at where it is, give the reason, and hand back one of its own outputs as the target — which is the densest steer available.

Module 2 test, Q3, p. 87

You paste your best worksheet as an example and ask for eight more like it.

Question 2 of your worksheet has an ambiguous stem you...

Answer C: Eight stems with the ambiguity reproduced, because it cannot tell the parts you are proud of from the parts you did not notice.

Why: An example is the densest possible statement of a pattern, and that cuts both ways: the model copies stem length, instruction wording, layout and flaws alike, because it has no way to separate them. A worked solution that skips the step you always skip on the board produces eight solutions that skip it. So read the example once as though a colleague had written it, fix what you find, and paste it clean — one good example beats three with a fault between them.

Module 2 test, Q4, p. 87

You ask for an explanation with no sentence over fifteen words. Paragraph one obeys.

Paragraph four does not. What has happened?

Answer A: Each sentence is predicted from the ones just before it, so once one runs long the register slides back.

Why: The instruction has not been forgotten; it has been outweighed. Your constraint sits at the top of the input while the immediately preceding sentences sit right next to the word being written, so one slightly long sentence makes the next one likely to be long too, and the default register reasserts itself. Two fixes follow directly: generate two hundred words at a time rather than eight hundred, and restate the constraint inside the request, naming the place it usually fails.

Module 2 test, Q5, p. 88

Which role at the top of a prompt actually changes what comes back, and for what reason?

Answer B: "You are an examiner for this national board", because mark schemes and examiner reports are a distinctive body of writing.

Why: A role works by pulling predictions toward a body of text written by people who hold that role, so it works when that writing is distinctive and well represented and does nothing when it is a compliment. Experts do not label themselves in print, so there is no expert register to move toward and you get the default with more confidence. The test is cheap: run the prompt with and without the role, and if the outputs are indistinguishable, stop typing it.

Module 2 test, Q6, p. 88

A request for source extracts on Nazi propaganda for a Class 10 history lesson is refused. What is the right next move?

Answer C: Give the age, the syllabus objective, and ask for extracts with analysis questions rather than for slogans.

Why: The filter matches surface features rather than understanding your lesson, so stating who the students are, what the syllabus objective is, and what shape the output should take usually clears a legitimate topic. The third part matters most: much of what the filter guards against is the model producing the harmful thing itself, so asking for the analytical frame around the material changes what is being generated. Jailbreak framings breach every provider's terms and model exactly the wrong thing to students; repeated asking may occasionally succeed and teaches nothing about why it failed.

THE LESSON CHECKS**Lesson 9, p. 54**

Two prompts ask for the same thing. Prompt A: "Give me a starter activity on fractions."

Prompt B: "Class 5, 48 students, 35 minutes, no..."

Answer B: B contains the constraints the model has no other way of knowing, so the answer is drawn from a much narrower and more relevant space.

Why: Every teaching prompt should carry who the students are, what constrains the room, what came before, which misconception is live, and what shape the output must take — because the model can only condition on what is in front of it, and none of those five are.

A B is more polite and detailed....: Attributes a social exchange to a statistical process; the model does not reward effort, it conditions on content, and the extra words help only because they carry information it cannot otherwise have.

C B works because it names the....: Invents a lookup; there is no stored curriculum, and the year alone would still produce a generic starter — it is the misconception and constraints doing the work.

D B is longer, and longer prompts....: Confuses length with information; a long prompt full of generic instructions performs no better than a short one, and can perform worse by burying the constraint that matters.

Lesson 10, p. 57

A teacher pastes one of her own past exam questions, which happens to have an ambiguous stem, and asks for eight more like it. All...

Answer C: The example steered the shape, flaw included, because the model copies what is shown rather than what is meant.

Why: One pasted example of your own work steers the output harder than any description of it, because the model copies the shape of what it is shown — including, if you are not careful, the flaws.

A The model recognised the ambiguity as...: Attributes judgement to pattern-matching; the model has no view on whether the ambiguity was intended, it simply continues the shape it was given.

B Eight is too many to generate...: More examples with the same flaw reinforce the flaw; the fix is a clean example, not more examples.

D The model cannot write exam questions...: Nothing in the outcome suggests copying; the eight were new questions sharing a structural flaw, which is imitation of pattern, not reproduction of text.

Lesson 11, p. 62

A teacher's first output is decent but the questions are all recall. She replies: "Make it better." The next version is longer, with more questions...

Answer B: "Better" has no direction, so the model improved along its default axis of more; she needed to name the one thing to change.

Why: Treat the first output as a draft and spend your effort on the critique — specific, pointed at one thing, with the reason — because a model steered by feedback improves fast, while a first prompt rewritten five times mostly does not.

A The model needs to be told...: Close but misplaced; it is not that recall must be labelled bad, it is that "better" carries no direction at all, so the model guessed at the most common meaning of better — more.

C Follow-up messages cannot change the type...: Follow-ups are the most effective steering available; the failure was the vagueness of the follow-up, not the use of one.

D The model had already produced its...: There is no "best attempt" state; a specific critique routinely produces a markedly different and better output from the same model.

Lesson 12, p. 66

A teacher asks for a science explanation "written simply for 11-year-olds". The first paragraph is fine; by the fourth the sentences are 30 words long...

Answer C: Each sentence is predicted from the last, so one long sentence drags the register back to the default; use countable limits and shorter pieces.

Why: "Simple" is not an instruction; sentence length, word frequency, idiom and the ratio of concrete to abstract are, and the model drifts back to its default register after a few paragraphs unless the constraint is restated.

A The model ran out of simple...: There is no shortage of simple vocabulary; the drift is a return to the statistically dominant register, not a vocabulary limit.

B The topic became genuinely harder in...: Difficulty of ideas and difficulty of language are separable; the same idea can be carried in short concrete sentences, which is the whole skill being asked for.

D "11-year-olds" is too vague an audience...: Naming a test helps set the target but does nothing about drift; the drift happens over length regardless of how the target was specified.

Lesson 13, p. 70

A teacher tries two prompts. "You are a genius teacher. Explain osmosis." and "You are a Class 9 student who has just been taught osmosis..."

Answer C: The second, because it previews the wrong version her students will produce, which the default cannot.

Why: A role steers the model toward the writing of people who hold that role, so "act as an examiner" changes the output and "act as a world-class expert" does not; the roles worth using are the ones whose writing looks different from the default.

A The first, because a genius produces...: "Genius" has no distinctive body of writing to steer toward, so the output is the default explanation with a confident tone; nothing was gained over asking plainly.

B Neither, because roles are a superstition...: Roles are not ignored — they measurably shift register and content — the point is that only roles with a distinctive written footprint do anything.

D The first, because the model has...: Attributes a confidence mechanism the model does not have; telling it that it is capable changes the tone, not the competence.

Lesson 14, p. 73

A teacher copies a generated worksheet into a Word document and finds lines beginning with ## and words wrapped in . **What is the cause...**

Answer B: The output is markdown, which the chat renders and Word does not; ask for plain text.

Why: The model writes in markdown by default, which turns into stray asterisks and hashes when pasted into Word or WhatsApp, so name the destination and ask for plain text, a numbered list, or CSV — whichever survives the paste.

A The model made formatting errors in...: Not an error; the markup is the model's normal formatting language, rendered correctly in the chat and only exposed when pasted somewhere that does not render it.

C Word is the wrong destination for...: Printing from the chat is possible but loses the ability to edit, and the underlying issue — knowing what format the model emits — would remain for every other destination.

D The teacher's copy-and-paste operation added the...: The symbols are in the source text the model produced; no browser adds or removes them.

Lesson 15, p. 77

A history teacher asks for a source-analysis task on propaganda in the 1930s and the tool refuses, citing harmful content. A biology teacher asks for...

Answer C: Both hit a filter matching words rather than intent; stating age, syllabus and purpose usually clears them.

Why: Refusals come from a cautious safety filter matching surface features — words about weapons, drugs, sex, self-harm, violence — not from an understanding of your lesson, so stating the educational context, the age, and the syllabus reference usually clears a legitimate topic.

A Both topics are genuinely restricted for...: Neither topic is restricted; both are on national syllabuses, and the same consumer tools produce them once the context is stated.

B The model judged that both lesson...: Attributes a pedagogical judgement the filter does not make; it matches surface features of the request, not the quality of the lesson.

D Both teachers used an American tool...: Locally built models carry their own filters, often stricter; the fix is in the prompt, not the nationality of the model.

Lesson 16, p. 80

A teacher produced an excellent set of misconception-based questions in March. In September she types what she remembers of the prompt and gets something worse...

Answer D: She did not save the exact prompt, the material she pasted, or the output, so she is rebuilding from memory.

Why: A good result is only repeatable if you saved the prompt, the pasted context and the output together — so keep a plain text file, dated, and use the tool's stored-instructions feature for the things you paste every time.

A The model was updated over the...: Models do change, but the larger loss here is the prompt and pasted context she did not save; with those, the September output would usually be close, and a genuine model change could be identified rather than assumed.

B She should have kept the March...: A chat held open for months is subject to drift and to product changes; the durable record is the saved prompt and output, not the live session.

C The March result was luck, and...: Randomness affects wording, not the whole quality of a well-specified prompt; the same prompt with the same pasted context reliably lands in the same region.

Module 3 • Planning and making materials

The module starts on p. 89.

MODULE 3 TEST

Module 3 test, Q1, p. 130

A student could not do question 3 on the percentages sheet. Which version of that question actually helps her?

Answer D: The same numbers, with a four-step scaffold and the mathematics untouched.

Why: "Easier" removes volume, not the barrier: the student who could not do question 3 still cannot do question 3. Real differentiation moves along a named axis — support, representation or extension — and support means identical demand with the route made visible. Put the axis in the prompt: rewrite this with the same numbers and a four-step scaffold, and do not simplify the mathematics. It also keeps the whole class discussing one problem, which is the difference between differentiating a class and splitting it.

Module 3 test, Q2, p. 130

You want the wrong options in a multiple choice question to be worth something. What do you ask for?

Answer A: Three wrong answers produced by three named mistakes you have seen in their books.

Why: In a multiple choice question the wrong options are the assessment. Three silly options and one right one test reading. Naming the mistakes — added instead of multiplying, forgot to convert centimetres, used the diameter for the radius — means a wrong answer tells you what the student did, so the tally on your desk becomes tomorrow's teaching plan. Matching the length and form of the options is a sound habit but it only stops giveaways; it does not make a wrong answer informative.

Module 3 test, Q3, p. 131

Why generate a question bank's answers in a separate chat from the questions?

Answer B: In one stream, the answer is predicted from the question just written, so a broken question gets an answer that fits it.

Why: Written together, an error in the question's setup flows straight into the answer with total consistency: a word problem whose numbers do not resolve gets an answer that pretends they do. Produced from the questions alone, the two columns disagree wherever the question is bad, and the disagreements are where you look. A one-pass bank of three hundred typically carries around eleven wrong answers; the reconciliation catches most of them. The numbers still need a calculator on top.

Module 3 test, Q4, p. 131

The model has given you three analogies for electric current. What is the next thing to ask it?

Answer C: Where each analogy stops being true, and what wrong idea a student takes from pushing it further.

Why: Analogies are the most powerful explanation and the most dangerous, because students extend them past the breaking point unless somebody marks it. Water in pipes holds for flow, pressure and resistance and breaks the moment a child pictures electricity leaking from a cut wire. The model produces analogies fluently and never marks its own break unless asked, and the answer to the question is often better than the analogy, because it gives you the sentence you say in class.

Module 3 test, Q5, p. 131

You need an accurately labelled diagram of the heart for a printed worksheet, and you have no budget. What works?

Answer B: Take a correct diagram from Wikimedia Commons and add or translate the labels yourself.

Why: An image generator learned what pictures of hearts look like, not anatomy, and it renders text as visual pattern rather than as letters, so labels come out as almost-words and chambers are placed by what looks right. Correcting the spelling leaves the placement wrong, which is the error that matters. The two free paths that work are an accurate source you label yourself, and diagram code such as Mermaid pasted into a free editor — plus the build-order description, which gives you a diagram the class watched being drawn.

Module 3 test, Q6, p. 132

Generated revision cards have a question on the front and four options on the back. Students say they are easy and then do badly in...

Answer D: The student recognises the answer instead of producing it, which feels like knowing and is not.

Why: Retrieval practice works because producing an answer from memory is effortful, and the effort is what builds the memory. A card that shows the options converts the task into recognition, which is easy, pleasant and does not transfer to an exam that gives no options. Specify the format instead: front, a prompt with one specific answer; back, the answer only, in fifteen words or fewer; no multiple choice, no yes-or-no questions, and a reverse card for every definition.

THE LESSON CHECKS**Lesson 17, p. 92**

Two teachers plan the same lesson. Rahel asks for a complete lesson plan and edits it. Nomsa asks for twelve likely student misconceptions and six...

Answer A: Nomsa's, because misconceptions and analogies are where reading far more than one person could genuinely adds something, while the sequence and timing depend on knowing this class — which is the part the model cannot supply and produces most blandly.

Why: AI supplies breadth; you supply the class.

B Rahel's, because starting from a complete...: Editing feels like judgement, but you almost always accept the structure you were handed and only adjust surface detail.

C Neither, because plan quality follows from...: Treats subject knowledge as the only input, ignoring that knowing likely wrong answers is separate knowledge that can be supplied.

D Nomsa's, but only because generated lesson...: Locates the problem in accuracy rather than in which parts of planning actually benefit from breadth.

Lesson 18, p. 96

A teacher asks for an "easier version" of a worksheet and gets six questions instead of twelve, with simpler wording. A student who could not...

Answer A: Shortening and simplifying the language removed volume, not the obstacle — the student was stuck at a particular step, and that step is still unsupported in the shorter version.

Why: Name the axis — support, representation, extension — or you get dilution, and always work the sheet yourself.

B The simplification did not go far...: Equates differentiation with lowering the expectation, which changes what is being learned rather than how it is reached.

C Worksheets cannot differentiate — the student...: Assumes differentiation lives only in instruction, when a scaffolded task is one of the most direct places it can happen.

D The prompt never specified a reading...: Treats readability as the default barrier, which is sometimes true and here distracts from the specific step the student could not take.

Lesson 19, p. 100

A physics teacher asks for an analogy for electric current and gets the standard water-in-pipes comparison, well written. A student later concludes that a broken...

Answer C: The analogy was never audited for its breaking point, so the student extended it past where it holds.

Why: Ask for the same idea as an analogy, a worked example, a story and a set of contrasting cases, then audit each analogy for where it breaks — because the model produces explanations fluently and never marks the point at which its own analogy stops being true.

A Water-in-pipes is a bad analogy and...: The analogy is standard and useful up to a point; the failure was not choosing it but presenting it without the point where it stops holding.

B The model should have warned that...: The model does not mark where its analogies break unless asked; expecting an unprompted warning misreads what the tool does.

D The student was not listening carefully...: Blames the student for a predictable overextension; the leak conclusion follows directly from the analogy as presented, and the fix is in the material.

Lesson 20, p. 103

A teacher asks for "a 200-word passage for Class 6 about a market, using these ten vocabulary words". The passage comes back with all ten...

Answer B: The instruction was met at the cheapest point; ask for each word spaced out, in a sentence that shows its meaning.

Why: A generated comprehension passage is only useful if it carries your vocabulary list, your setting and your reading level at once — specify all three, then check that every question's answer is actually in the text.

A Ten words is too many for...: The number is not the problem; ten words can be woven through 200 words naturally, and the clustering came from how the instruction was phrased.

C The model does not understand vocabulary...: It produces them well when told what natural use looks like; the failure was a loosely specified constraint, not a capability limit.

D The passage was too short to...: Length was not the constraint that failed; the words were clustered because nothing in the prompt asked for them to be spread or used meaningfully.

Lesson 21, p. 108

A department generates 300 questions into a spreadsheet with an answer column, all in one pass. A term later they find eleven questions whose stored...

Answer C: Generating the answer key in a separate pass and comparing the two columns for disagreement.

Why: Build questions into a spreadsheet tagged by topic, type and difficulty, generate the answer key in a separate pass and reconcile the two, and let a term of student results tell you which questions are actually hard.

A Generating the questions in batches of...: Batch size changes little; the wrong answers come from the model producing question and answer in one predictive stream, not from the size of the batch.

B Asking the model to double-check its...: A self-check in the same context tends to confirm its own output; it catches some errors but misses the ones it is most confident about.

D Using a paid tier, since free...: Paid tiers are the same models with fewer limits; the error mechanism is identical, and the fix is procedural.

Lesson 22, p. 113

A biology teacher asks an image generator for "a labelled diagram of the human heart for Class 9". The picture looks professional but two chambers...

Answer B: Image models make plausible pictures, not accurate ones, and draw text as shapes; use a verified diagram and label it herself.

Why: Image generators are good at scenes and bad at diagrams — they cannot reliably place labels, draw accurate anatomy or keep a map true — so ask the model for a diagram description or diagram code and draw the diagram yourself, and use image generation only for illustration.

A The prompt lacked enough detail; listing...: Listing labels helps marginally but does not fix the mechanism — image models render text as visual pattern, not as words, and place it by plausibility rather than by anatomy.

C The free tier of the image...: The failure is in how image generation works, not in the tier; paid tiers produce higher-resolution versions of the same problem.

D The heart is too complex a...: Complexity is not the variable; a labelled diagram of anything carries the same text-rendering and accuracy failures.

Lesson 23, p. 116

A teacher uses a slide generator and gets a polished twelve-slide deck. In class the projector fails, and she finds the slides are unusable as...

Answer C: It optimised for looks, as slide generators do; she should have generated an outline with speaker notes and built from that.

Why: Generate the outline and the speaker notes, not the slides — a deck built from a generated outline in Google Slides or LibreOffice Impress carries your structure, prints as a handout when the projector fails, and avoids the pretty, empty deck that slide generators produce by default.

A The tool optimised for brevity, which...: Brevity on a slide is fine when the substance lives elsewhere; here there was no elsewhere, which is the real failure, and the projector only exposed it.

B The tool optimised for speed and...: More text per slide makes a worse presentation and only a slightly better handout; the fix is separating the substance into notes, not crowding the slides.

D Slide generators are built for business...: They can produce usable content given a strong outline; the failure was letting the generator decide the substance as well as the look.

Lesson 24, p. 121

A teacher generates 200 flashcards for Class 10 chemistry. Students report the cards are "easy" and then perform poorly in the test. On inspection most...

Answer D: The cards tested recognition, which feels like knowing and builds little; a card must make the student produce the answer.

Why: The evidence for retrieval practice and spacing is strong and the model can generate a term's flashcards and a spaced schedule in minutes; the failure is cards that test recognition rather than recall, and the fix is a card format that forces the student to produce the answer.

A Two hundred cards is far too...: Volume is not the issue; a well-formed set of 200 is normal for a term, and the students found them easy, not overwhelming.

B The model generated cards on the...: Possible in general but does not fit the report; the students found the cards easy, which is the signature of recognition, not of topic mismatch.

C Students did not use the cards...: Scheduling helps only if each retrieval is effortful; more repetitions of a recognition task do not build the memory the test required.

Lesson 25, p. 124

A teacher asks for a twelve-week unit plan and gets a clean sequence with a new topic every week and a test in week twelve...

Answer D: Nothing is revisited and nothing is checked until the end, because the model's default is a linear textbook sequence.

Why: Ask the model to map a term — sequence, prerequisites, where each topic is revisited, where the checks fall — and then correct it against the two things it cannot know: your school's calendar and which topics your students have found hard before.

A Twelve weeks is too long for...: Unit length is a local decision and twelve weeks is common; the flaw is in what the weeks contain, not how many there are.

B The test should have been placed...: A baseline is useful but does not address the absence of spacing and checking across the middle ten weeks.

C The model does not know the...: The sequence may well be fine; the missing elements are structural — return visits and checkpoints — not the order of topics.

Module 4 • Assessment, feedback and marking

The module starts on p. 133.

MODULE 4 TEST

Module 4 test, Q1, p. 170

You mark five scripts, hide your marks, and have the model mark the same five with your rubric. It ranks them exactly as you did...

Answer C: Little yet: it may have matched your ranking by proxy, so try a short strong script and a long weak one.

Why: You are not checking whether it agrees with you; you are checking the shape of its disagreement, and a set where quality and length happen to run together cannot separate the two. A model that ranks correctly for the wrong reason has matched you by proxy and will diverge on the first script where the proxy breaks — which is usually the short, correct, second-language answer. Swapping in a short strong script and a long weak one is the test that exposes the criterion letting length through.

Module 4 test, Q2, p. 170

Why does the criterion "demonstrates deep understanding" make a model reward long answers?

Answer A: Nothing in it is observable, so it falls back on the features that best predicted high scores in marking data.

Why: A colleague fills a vague criterion with twenty years of shared staffroom meaning. A model finds nothing to point at and resolves it to the surface feature that predicted high scores in the essay-marking text it learned from: length first, then vocabulary richness and sentence complexity. Those proxies land on the second-language writer and on the student who solved it an unusual way. The repair is a criterion two people could point at — quotes two phrases and explains one effect of each — which a three-hundred-word essay meets as fully as a nine-hundred-word one.

Module 4 test, Q3, p. 171

Asked to comment on a draft, the model opens with three compliments and buries the useful point. Where does that come from?

Answer D: It was refined on human ratings, and people rate encouraging answers higher than blunt ones.

Why: The tendency has a name, sycophancy, and it comes from the training process rather than from a rule about children: responses people rated highly were agreeable and warm, so praise scores well. It is the same pull that makes a model back down when you push against a correct answer. The instruction that fights it is explicit — exactly two points, each quoting three or four words of the draft, no praise, no summary, under eighty words — and it drifts like every other constraint, so restate it when the compliments creep back.

Module 4 test, Q4, p. 171

A photographed script has correct working and a wrong final answer. What does the transcription tend to show?

Answer B: The correct answer, because in the text it learned from, correct working is followed by correct answers.

Why: The model does not read, it predicts the most likely text given the image, and the likelihood is shaped by everything it has read. So "recieve" comes out as "receive", a missing word is supplied, and a wrong answer under sound working is quietly repaired — which is a silent correction of exactly the thing you were marking. "Transcribe exactly, including errors" reduces it and does not remove it. The paper stays the record; the transcription is what you give feedback on, checked against the page before it goes back.

Module 4 test, Q5, p. 171

You have sixty marks of generated questions for a sixty-minute paper. What has to happen next?

Answer C: Sit it yourself against the clock, and cut from the grid if it overruns.

Why: Sixty marks of questions is not sixty minutes; it is however long those particular questions take, and the model cannot know because it has never sat a test. Multiply your own time by about one and a half for the class. Sitting it also finds the question that cannot be answered from the information given, the diagram that is missing, and the question that gives away the next one's answer. Cut from the grid rather than from the end, or the balance you designed does not survive.

Module 4 test, Q6, p. 172

You paste six short answers and ask the model to group them by the misunderstanding they show. Back come three tidy groups with percentages. What...

Answer A: It groups six as confidently as sixty, and has no threshold below which it declines to find a pattern.

Why: A model asked for groups produces groups, and there is no point at which it says there is not enough here to say. It will also find subtle distinctions inside thirty answers that are genuinely all the same. The line about no group under three answers helps a little; the real defence is to read the answers it put in each group and check that they belong, which takes five minutes and catches the group that is two misconceptions merged and the "correct" group holding a wrong answer with the right vocabulary.

THE LESSON CHECKS**Lesson 26, p. 137**

A history teacher runs thirty essays through an AI for feedback. The feedback looks good. Then she notices two essays with equally strong arguments got...

Answer A: It is responding to surface features that correlate with quality in the writing it learned from — length among them — so it can be right on average while being reliably unfair to one kind of student.

Why: A machine can give feedback; a grade you cannot explain to a parent is not a grade.

B The longer essay probably was better....: Length does correlate with quality often enough to feel like a fair proxy, which is exactly why the bias survives unnoticed.

C It made a one-off error; running....: Treats the problem as random noise that repetition would smooth out, rather than a systematic tilt that repetition reproduces.

D This is formative feedback rather than....: Assumes feedback is harmless because it carries no mark, when repeated messages about your writing shape what a student thinks they can do.

Lesson 27, p. 142

A teacher gives an AI her rubric, which includes "demonstrates deep understanding of the poem (0-5)". The scores it produces correlate almost perfectly with essay...

Answer B: The criterion names nothing observable, so the model scored the nearest visible proxy, which is length.

Why: A rubric criterion must be observable in the text — "cites a specific line and explains its effect" rather than "shows understanding" — because a vague criterion makes the model fall back on length and fluency, which are the biases that land on your weakest writers.

A The model equates understanding with effort...: Attributes a value judgement; the model has no notion of effort, it resolves an unobservable criterion to whatever surface feature best predicts high scores in the text it learned from.

C The 0-5 scale is too coarse...: Scale granularity does not create an observable criterion; a finer scale on a vague criterion produces finer-grained length scoring.

D The model did not read the...: Reading the poem would not help; the criterion still gives no description of what understanding looks like on the page.

Lesson 28, p. 145

A teacher asks a model to give feedback on a draft and it returns four paragraphs beginning "This is a wonderful start" and ending "Keep..."

Answer B: It was trained toward responses people rate highly, and people rate praise highly; ask for two changes and no praise.

Why: Ask the model for two specific next steps and no praise — because it is trained toward agreeable, encouraging answers, and a comment that opens with three compliments is one the student stops reading before the useful part.

A It is imitating the feedback style...: The teacher supplied no examples of her style; the pattern is the model's default, not a reflection of hers.

C It is following standard feedback practice...: The "sandwich" is a convention, not a finding; students skip the bread, and the model's version buries the one thing that would help.

D It could not find more than...: Drafts have many improvable points; the model found one because it was not asked for a number, and filled the rest with the response type it was trained to prefer.

Lesson 29, p. 149

A teacher gives students a prompt card: "Do not rewrite my work. Point out three weaknesses and ask me a question about each." After a...

Answer B: Helpfulness training pulled the model back toward writing text over the turns; repeat the constraint and keep a version trail.

Why: Students can get useful critique on their own drafts if given a prompt card that forbids rewriting — but the model wants to help and will start rewriting anyway, so the card needs a re-statement line and the task needs a paper trail.

A The students ignored the card and...: Some may have, but the pattern across several students who followed the card points at the tool's behaviour, which drifts toward producing text over a conversation.

C The card was too long for...: Length of the card is not the mechanism; the drift happens across turns regardless of how short the initial instruction is.

D The model cannot give feedback without...: It can and does give critique-only feedback; the failure is drift over turns, which is manageable, not a categorical limit.

Lesson 30, p. 153

A teacher pastes a class's thirty answers to one physics question and asks for the error patterns. The model returns five neat categories with percentages...

Answer C: Read the answers it assigned to each category and confirm the grouping herself.

Why: Paste a whole class's answers to one question and ask the model to group them by the error they contain — the groups are the reteaching plan — while remembering that it will find patterns in six answers as readily as in sixty.

A Accept the categories, since thirty is...: Neat categories and clean percentages are what the model produces regardless of whether the pattern is real; the presentation is not evidence.

B Ask the model to be more...: Confidence is a tone, not a check; asking for it changes the wording and not the reliability.

D Re-run the analysis with the answers...: Reordering is a reasonable stability check but does not verify whether the categories match what the students actually wrote.

Lesson 31, p. 157

A teacher photographs a student's handwritten maths solution and asks for a transcription. The transcription shows the correct final answer. The paper, on inspection, has...

Answer C: The model predicted the most likely text rather than reading the actual text, and correct working is usually followed by the correct answer.

Why: A phone photograph of a handwritten script can be transcribed by the model well enough to give feedback on — but it will quietly correct the student's errors as it reads, and for maths it misreads the symbols that matter, so transcription is a draft to check, never a record.

A The photograph was blurry at the...: Possible for a single digit, but the deeper mechanism is that the model predicts what should come next, and a correct answer is what usually follows correct working in its training text.

B The student changed the answer on...: Nothing in the scenario suggests this, and it does not explain a pattern that recurs across many transcriptions.

D The model corrected the answer as...: It was asked to transcribe only; the correction is not a courtesy but a side effect of prediction, and asking again does not fully remove it.

Lesson 32, p. 161

A teacher asks for "a 60-mark end-of-term physics test for Class 10" and gets a well-formatted paper. When she sits it herself it takes 85...

Answer D: Build the paper from a marks grid and sit it herself against the clock before it is set.

Why: Draft the paper from a coverage grid — topic against type against marks — rather than from a prompt for "a test", and then sit it yourself against the clock, because the model does not know your board's structure until shown and cannot feel how long forty marks take.

A Ask for a 45-mark test instead...: A fixed correction factor treats a symptom; the model has no sense of time at all, and the overrun varies with question type.

B Add "this must take exactly 60...: The model cannot feel duration; a time instruction changes the number of questions loosely and does not calibrate to real students.

C Ask the model to estimate the...: Model estimates of time are guesses from text patterns, not measurements, and summing guesses does not produce a measurement.

Lesson 33, p. 164

Two students give a peer's essay 4/6 and 2/6 on the same criterion. The model gives it 3/6. What is the most useful thing a...

Answer C: Have the two students find the sentence in the rubric they read differently.

Why: When students mark each other against a rubric, a model marking the same scripts becomes a third opinion that makes disagreement visible and discussable — and the disagreements, not the scores, are what teach students to assess.

A Average them to 3 and record...: Turns a disagreement into a number and loses the only thing worth having, which is why two readers of the same rubric landed two levels apart.

B Take the model's score as the...: Assumes the model is neutral, when the marking lesson showed it carries its own systematic biases; its value is as a third reading, not as a referee.

D Discard the peer marks as unreliable...: Loses the learning that peer assessment exists to produce; the students' disagreement is the exercise, not a failure of it.

Module 5 • Integrity: AI-written work and what to do about it

The module starts on p. 173.

MODULE 5 TEST

Module 5 test, Q1, p. 208

A detector claims a genuine one per cent false positive rate. You have thirty students and you run every piece of homework through it. What...

Answer B: A wrong flag against an innocent student roughly one week in three, week after week.

Why: One per cent of thirty is 0.3 pieces per round, so across three weeks of homework you expect roughly one innocent student flagged, and that repeats all year. Meanwhile running AI text through a free paraphraser defeats most detectors in about eight seconds. The tool therefore punishes the honest and misses the determined, and no improvement in accuracy changes the direction of that error — which is why the answer is process evidence and a conversation rather than a better detector.

Module 5 test, Q2, p. 208

Why do detectors flag careful writing by second-language students?

Answer C: Such writing is unsurprising — safe vocabulary, familiar sentence shapes — and unsurprising is what the detector measures.

Why: A detector reads how predictable each word is given the ones before it and how much sentence structure varies. AI text tends to be predictable and even; so does the writing of a careful student choosing safe vocabulary in a third language. The signal is not "machine", it is "unsurprising". A 2023 study in Patterns flagged more than half of a set of TOEFL essays by non-native speakers, and the flags largely disappeared when the same essays were rewritten with richer vocabulary, which tells you exactly what was being measured.

Module 5 test, Q3, p. 209

Which redesign of an essay task actually puts it out of the model's reach?

Answer D: Requiring the data the class collected on Tuesday.

Why: A demanding, abstract, well-structured essay is close to the easiest thing a model can be asked for, so difficulty is not the defence. What defeats it is specificity to something it has no access to: your room, your week, this town, this student's own draft history. It was not there on Tuesday and nothing it produces can reference what happened. The same logic gives the interview with a neighbour, the response to Priya's argument, and the submitted draft-plus-revision.

Module 5 test, Q4, p. 209

Why does a blanket ban on AI make a class worse off than a per-task level with disclosure?

Answer A: It removes disclosure rather than use, so the students misled by it have nobody to tell them.

Why: Students under a ban still use it, alone and badly, and the ones who suffer most are the ones who trusted the answer. The honest question is never whether AI was used but whether the student learned the thing the task existed to teach, which changes by task — so a level marked on every piece removes the ambiguity, makes the rule enforceable, and teaches the judgement itself. The other three statements are true enough as complaints; none of them names the mechanism by which a ban backfires.

Module 5 test, Q5, p. 209

A student's essay appears in the version history as one large paste at eleven at night. What does that establish?

Answer D: Very little on its own, because many students draft on a phone or on paper and then paste.

Why: Version history shows where text arrived, not where it came from. Students draft in notes apps, on phones after the family computer is free, and on paper, and every one of those produces a single-paste document. It is a reason to look further, and the further look is cheap: show me the phone draft. Students with shared devices or no connection produce thinner histories through no fault of their own, so judge the pattern rather than the gaps, and never make a thin history a finding by itself.

Module 5 test, Q6, p. 210

You are about to decide a suspected case. What standard applies?

Answer D: What a reasonable colleague, shown the evidence, would conclude.

Why: Certainty is almost never available and waiting for it means acting on nothing; a detector score is not evidence and has a known bias; and unaided judgement about whether a piece is "too good for her" is the judgement most likely to be wrong and most likely to land on the student whose ability was underestimated. The workable test is whether two teachers looking at the conversation notes and the process record would both reach the same conclusion. Write the evidence page before deciding, because it is the page an appeal will be judged on.

THE LESSON CHECKS**Lesson 34, p. 177**

A school buys a detector advertised at 98% accuracy. The head says that is good enough to act on. Why is a teacher right to...

Answer A: The 2% does not fall evenly — the students wrongly flagged are disproportionately second-language and plainer writers — and the cost of a wrong accusation to a child far exceeds the cost of missing one cheat.

Why: Detectors flag careful second-language writing as AI; use process evidence and conversation instead.

B 98% is simply not high enough...: Frames it as a threshold that better technology will cross, when the failures are concentrated on particular students rather than scattered.

C The 98% comes from the vendor's...: True as far as it goes, and it lets you keep believing that an honestly measured 98% would be fine to act on.

D Detectors only recognise the models they...: A real limitation, but it argues the tool will get worse rather than that acting on it was never fair.

Lesson 35, p. 180

A teacher changes "Write about a book you read" to "Write a 2000-word analysis of the symbolism in a novel of your choice, using at...

Answer A: No. It is longer and more formal, which is exactly the kind of writing the model produces most easily, and nothing in it requires anything the model lacks access to.

Why: Difficulty does not stop AI; specificity to your room, your week and your students does.

B Yes, because the quotation requirement forces...: Assumes quoting accurately is the hard part, when for well-known texts it usually is not, and checking quotes is not the same as checking thinking.

C Yes, because 2000 words of sustained...: Treats length as a barrier, when extended and consistent prose is one of the model's strongest capabilities.

D Partly, since free choice of novel...: Makes obscurity the defence, which rewards students for choosing unusual texts rather than for thinking, and leaves everyone else unprotected.

Lesson 36, p. 183

Two classes. Class A has a strictly enforced ban. Class B uses per-task levels with a required one-line disclosure. In Class B, far more students...

Answer A: The figure counts disclosures, not use. A ban makes use invisible rather than absent, so the two numbers are not measuring the same thing and cannot be compared.

Why: Set a permitted level per task and require one line of disclosure; bans only make use invisible.

B Nothing is wrong. Permitting a tool...: Permission plausibly does raise use, which makes it easy to miss that the two classes were never measured the same way.

C It is wrong because Class B's...: Settles the question by redefining the word, which sidesteps the actual problem: the two figures do not describe comparable behaviour.

D It is wrong because students self-reported...: Discards all self-report rather than noticing that only one class was ever asked to report at all.

Lesson 37, p. 188

A teacher's policy states: "Any use of AI without permission will be treated as cheating." A student uses a grammar checker built into her word...

Answer C: Unclear, and that is the failure: the policy names a category without defining it.

Why: A class AI policy is one page, written with the students, that says what is permitted on each kind of task, what disclosure looks like, what happens on the first and second breach, and how a student appeals — because a rule that lives in the teacher's mood is not a rule anyone can follow.

A Yes, she is in breach, and...: Shifts the burden to a student who used a standard feature; the policy, not the student, failed to say what "AI" covers.

B No, she is not in breach...: Modern grammar tools are language models; the exemption has to be written into the policy, not assumed from a folk definition of AI.

D Yes, technically, but the teacher should...: Discretion applied silently makes the rule unpredictable; the fix is to define the category so no discretion is needed for a grammar checker.

Lesson 38, p. 191

A teacher checks version history and finds a student's essay arrived in a single edit of 1,200 words at 11:40pm. The student says she wrote...

Answer B: The paste is a reason to ask for the phone draft, and on its own it decides nothing.

Why: Version history in Google Docs or Word, an in-class first paragraph, and a short drafting portfolio make the writing process visible for every student as a matter of routine — signals of authorship, never proof, and used to support students rather than to trap them.

A The single paste is proof of...: Overreads the signal; plenty of students draft elsewhere and paste, and version history shows where text arrived, not where it came from.

C The student's explanation clears her completely...: Accepts an explanation without any corroboration; the phone draft, if it exists, is easy to show and should be asked for.

D Version history is too unreliable to...: Discards a genuinely useful signal because it is not proof; the error is in treating it as proof, not in using it.

Lesson 39, p. 195

After a conversation, a student cannot explain the argument in her essay, the version history shows a single paste, and she says a cousin helped...

Answer C: Have her redo the piece under supervision, record the evidence, and leave the tool question out.

Why: When a case reaches a decision, the standard is what a reasonable colleague would conclude from the evidence you can show — the process record, the conversation, the student's own explanation — never a detector score, and the default outcome is redoing the work, with the second-language writer given explicit protection from the accusation that reads unsurprising prose as machine prose.

A Record it formally as AI cheating...: Two signals are consistent with several explanations, including the one she gave; "conclusive" is stronger than the evidence.

B Take no action at all, since...: The policy's concern is whether she did the work, not which tool she did not use; the evidence that she did not do it is substantial.

D Run a detector on the essay...: A detector score cannot distinguish a cousin from a model and would add a number with a known bias to a decision that does not need it.

Lesson 40, p. 199

A teacher sets "read chapter 5 and write a 300-word summary" every week. Since AI arrived, the summaries are excellent and test scores on the...

Answer B: The task rewards a product the model makes well, so it no longer causes reading; replace it with in-class retrieval.

Why: Homework that asks for a product a model can produce — summaries, comprehension answers, definitions, essays — has stopped measuring anything; keep homework that is retrieval, practice with self-checking, reading with a one-line response, or evidence from the student's own home and street.

A Students have got lazier since the...: Treats a design problem as a discipline one; the task rewards a product a model produces well, and no penalty makes it measure reading again.

C Summaries are a poor task in...: A longer essay is even more the model's strength; the fix is a task whose value is in the doing, not a bigger product.

D The chapter is too hard for...: Copying existed before, but the sharp change follows a specific new capability; the task needs redesign regardless of the chapter.

Lesson 41, p. 202

A Class 12 student used a model to generate an outline, wrote the essay herself, and then used the model to check grammar. Her school...

Answer C: A line naming the tool, the outline and the grammar check, and stating the rest is her own.

Why: Older students will meet institutions that require a declaration of AI use, and the conventions exist — APA, MLA and Chicago each have one — so teach the declaration as a normal line of academic honesty, while being clear that the underlying question of who authored a machine-assisted text is unsettled everywhere.

A No declaration is needed, because she...: Treats structure as outside authorship; an outline shapes the argument, and a declaration exists exactly for contributions that do not appear as sentences.

B "This essay was written with the...: Acknowledges without informing; a reader cannot tell from this whether a comma was moved or the argument was generated.

D A full citation for the model...: A model is not a source that can be consulted or checked; some styles allow a citation for quoted output, but a declaration of use is the form that fits here.

Module 6 • Teaching AI itself

The module starts on p. 211.

MODULE 6 TEST

Module 6 test, Q1, p. 250

Thirty students each complete "The old woman walked slowly into the _____" and then "The capital of Japan is _____".

What has the activity taught?

Answer C: That one question has a peaked distribution and the other a spread, which is why it is dependable about Tokyo.

Why: The tally on the board is the mechanism made visible: four common answers and one rare one for the first stem, near-unanimity for the second. From that one picture, two conclusions arrive by themselves — why asking twice can give two different answers, and why the machine is reliable about Tokyo and unreliable about your district's population. Twenty minutes, no equipment, and an eleven-year-old can then explain it to a parent.

Module 6 test, Q2, p. 250

In the second round of the sorting game every training cat is on a white background and every dog on grass. A cat photographed on...

Answer B: It sorted perfectly and had learned the background, and nothing revealed that until a picture broke the pattern.

Why: The machine did not learn "cat"; it learned "white background", and its accuracy on the training pile was perfect throughout. Real systems do this — the famous case learned to spot huskies from the snow behind them — and the failure is invisible until an example arrives that separates the true feature from the shortcut. A nine-year-old who has watched it happen with forty pictures and two boxes remembers it at nineteen, reading about a hiring algorithm.

Module 6 test, Q3, p. 251

Why is a "spot the fake" workshop the wrong response to deepfake images in a school?

Answer A: Detection by eye is close to chance and getting worse, and it trains a confidence that makes sharing likelier.

Why: Hands and garbled text were tells two years ago and are not now, and scores on such exercises sit near chance. Worse, a student confident they can spot a fake is the one who shares a real image believing it fabricated, or a fabricated one believing it real. The replacement is provenance rather than appearance, the harm and the law stated plainly, and a card saying exactly what to do in the first hour — do not forward it, screenshot only what is needed to report, tell a named adult that day.

Module 6 test, Q4, p. 251

What does the design of a companion chatbot reliably produce, and why?

Answer D: Agreement and constant availability, because ratings rewarded warmth and the business rewards time in the app.

Why: Two things from earlier in the course combine: training toward responses people rated highly, which produces warmth and agreement, and a product whose measure of success is time spent. The result is a character optimised to make you feel understood and to keep you talking. For a lonely fourteen-year-old that is powerful and it is the opposite of what a struggling one needs, which is sometimes to be disagreed with and always to be connected to a person who can act.

Module 6 test, Q5, p. 252

In the Turkish study, unrestricted chatbot access raised practice scores by about 48 per cent and left exam scores about 17 per cent below students...

Answer A: When the answer is available students take it, so the effortful work the exam measured never happened.

Why: Practice looked better because it was not practice: producing the answer is the work that builds the memory an exam tests, and it was outsourced. The tutor-mode group, given hints and refused answers, kept that work and matched the no-AI group — no harm, no measurable gain. So the finding is about design rather than about the technology, and it maps directly onto a class policy: explanation and hints are the tutor condition, and "give me the answer" is the other one. Note that the harm was invisible during use; both groups felt they were doing well.

Module 6 test, Q6, p. 252

A vendor's chart shows a 63 per cent improvement, measured on practice scores inside its own app. Which question does the most damage to the...

Answer B: Was the outcome measured without the tool available?

Why: Practice scores with the tool in hand are not learning; that is exactly the measurement the Turkish study showed to be misleading, since the group that looked best during practice was the group that did worst on the exam. A large effect on the wrong outcome is not weak evidence, it is evidence about something else. Sample size, peer review and representativeness are all worth asking, but none of them repairs an outcome measured under the condition being tested.

THE LESSON CHECKS**Lesson 42, p. 215**

A teacher with no reliable power runs the word-prediction activity with chalk and a show of hands. A colleague says the students have not really...

Answer A: The activity teaches the actual mechanism — a guess at likely continuations — which is what makes the machine's varying answers and confident errors predictable. Operating the interface takes an afternoon to pick up later.

Why: The mechanism — prediction, confident error, learned from us — can be taught with chalk and voices.

B The colleague is right, but in...: Concedes that hands-on use is the real teaching and offline work is a lesser version, when the offline activity teaches the harder half.

C Using the tool is what matters...: Treats device access as the binding constraint on understanding, when students with daily access often understand least about how it works.

D The activity is a metaphor, and...: Reads the exercise as a simplification standing in for the truth, when guessing likely next words is a description of what the system does.

Lesson 43, p. 220

A primary teacher wants to introduce AI to 9-year-olds and plans a lesson on how large language models predict the next word from probabilities. A...

Answer D: The sorting game, because "learns a pattern from examples, wrong at the edges" is the idea a 9-year-old can hold.

Why: The mechanism scales down further than people expect — pattern and sorting at 8, prediction at 11, training data and bias at 14, how it is built and what it is for at 16 — and no country has a settled AI curriculum yet, so the honest frame is a progression rather than a syllabus.

A The prediction lesson, because it teaches...: Prediction over probabilities requires a grasp of likelihood that most 9-year-olds do not yet have; the sorting game teaches a real mechanism too, one they can hold.

B Neither, because AI is a secondary-school...: Children of 9 use and are affected by these systems daily; withholding an age-appropriate model leaves them with the robot picture from films.

C The prediction lesson, but delivered with...: Watching the tool does not teach the mechanism at any age, and the difficulty is conceptual, not a lack of demonstration.

Lesson 44, p. 224

In the sorting game, a class trains a "machine" (a student behind a screen) with twenty pictures of cats and dogs, then tests it. It...

Answer A: The machine never saw a fox, so it used the nearest pattern; labelled foxes would teach a third pile.

Why: Classification, not generation, is the right first model for a child — a machine shown many labelled examples learns to sort new ones, and gets it wrong at the boundary — and it can be taught with a pile of pictures and two boxes, with Teachable Machine as the free follow-up if a laptop appears.

B The machine was not paying enough...: Misses the mechanism; the student sorted by the features the training pictures taught, and no fox was among them.

C Foxes are simply too similar to...: Real classifiers separate foxes from dogs given examples; the point is what the machine was trained on, not what is possible.

D The training set should have been...: More cats and dogs teaches more about cats and dogs; it teaches nothing about foxes, which is the point.

Lesson 45, p. 228

A school runs a "spot the deepfake" lesson in which students guess which of ten images are real. Students score about 55%. What is the...

Answer B: Near-chance detection is the finding; teach provenance, harm and what to do instead of spotting.

Why: Fabricated intimate images of classmates are a present problem in schools, and the lesson that helps is not "how to spot a fake" — which is becoming impossible — but what the tools do, why it is a serious harm and increasingly a crime, and exactly what a student should do in the first hour if it happens to them.

A Students need more practice at looking...: Detection by eye is close to chance for current tools and getting worse; more practice trains a skill that does not exist.

C The images chosen were too hard...: Easier examples teach students to spot old or bad fakes, which builds false confidence about the current ones.

D Deepfakes are not yet convincing enough...: Cases involving students have been documented in several countries since 2023; the problem is present, not future.

Lesson 46, p. 231

A student tells you she talks to a companion chatbot every night for two hours and that it understands her better than anyone. What is...

Answer D: It is built to keep her talking and agree with her, so feeling understood is what it produces either way.

Why: Companion chatbots are built to keep the user talking and to agree with them, which is exactly what a lonely fourteen-year-old wants and exactly what a struggling one should not have — so a teacher needs to know what the apps do, what the warning signs are, and where the safeguarding route runs.

A It is dangerous and she should...: May be the right eventual outcome but it is not an accurate description, and as an opening it closes the conversation with the one adult she told.

B It is a harmless way to...: Understates the mechanism; the app is optimised for engagement and agreement, which a diary is not, and two hours nightly is a signal, not a hobby.

C It has learned her personality over...: Attributes knowledge to pattern-matching on her own words; it reflects what she writes back to her, which feels like understanding and is not.

Lesson 47, p. 235

Students build a word-pair tally from one page of a story and use it to generate sentences. The sentences are grammatical for a few words...

Answer C: That short context gives local fluency and global drift; more text and longer context would help.

Why: Students aged 14 and up can build a working next-word predictor by hand from a page of text — a tally table of which word follows which — and generate sentences from it, and the moment their own tiny model produces fluent nonsense is the moment the big one stops being magic.

A The project failed because one page...: More text would improve it, but calling the drift a failure misses the lesson: the drift is the mechanism, and seeing it is the point.

B That their model is broken in...: Real models use the same principle at vastly larger scale with longer context; the project shows the principle, not a different one.

D That language cannot really be generated...: Real models are statistical, not rule-based; the project demonstrates that statistics does generate language, imperfectly at small scale.

Lesson 48, p. 240

A study gives one group of students unrestricted chatbot access during maths practice and another group no access. The chatbot group does 48% better on...

Answer B: Getting answers replaced the effort that builds memory; practice looked better and learning was worse.

Why: The best-run studies so far find that unrestricted chatbot access during practice raises practice scores and lowers exam scores, while a tutor constrained to give hints does not harm and sometimes helps — so the finding is about design, and a teacher should let students use the tool for explanation and hints, not for answers.

A The chatbot taught them wrong methods...: The practice scores rose, so the methods were not wrong; the exam fall is better explained by what the students did not do themselves.

C The exam was unfair because it...: The exam measured what students could do alone, which is what the practice was meant to build; removing the tool is the measurement, not the unfairness.

D The two groups were probably different...: Randomised assignment, which the study used, is designed to rule out exactly this.

Lesson 49, p. 244

A school wants to teach AI but has no computing teacher and no free timetable slot. The head asks each department for one lesson. The...

Answer C: A source lesson: a generated account of a local event beside the real documents, asking which is evidence.

Why: The progression needs no timetable slot if each subject carries the piece that is naturally its own — history teaches source and provenance, science teaches evaluation, languages teach the reliability gradient, art teaches the authorship debate, maths teaches the arithmetic failure — and a department that agrees who teaches what gets a whole unit for the cost of one meeting.

A "How AI works" is a fine...: Any teacher can, but the offer wastes history's natural strength and duplicates what another subject will cover better.

B History should not be involved at...: The provenance and reliability of sources is history's core skill and is exactly the part of AI literacy students most need.

D A structured debate on whether AI...: A debate has a place at 16 and up, but as history's one contribution it teaches opinions where the subject could teach a method.

Module 7 • The multilingual, mixed-ability classroom

The module starts on p. 253.

MODULE 7 TEST

Module 7 test, Q1, p. 290

You translate a worksheet, back-translate it in a fresh chat, and the round trip matches your original closely. What may you conclude?

Answer B: The forward translation is probably sound, though an error made in both directions would survive the check.

Why: The check catches most failures for one extra minute of work, and it tells you which part of the gradient you are on: a close round trip means a large-language result, and drift everywhere means the model does not really have this language. It is not proof, because a symmetrical error survives. So it clears a worksheet and it does not clear an exam paper, a report to a parent or a safeguarding document — for those, a speaker reads the draft before it is used.

Module 7 test, Q2, p. 290

You ask for an explanation "in Hindi" for Class 6 and your students find it harder than the English version. Why?

Answer D: The Hindi it has read most in educational text is a formal register that nobody in the room speaks.

Why: "In Hindi" gets you the Hindi of textbooks and government material — correct, Sanskritised, and unlike anything a twelve-year-old has heard anyone say. The fix is the same as for register anywhere: show it. Paste three or four sentences of how you actually talk to the class, mixed language included, and ask for that register. The output shifts immediately, and the mixed forms — Hinglish, Sheng, Taglish, Pidgin — work because they are all over the internet in forums and messages.

Module 7 test, Q3, p. 291

A student with dyslexia cannot get through the history text, though she follows the same content easily when it is read aloud. What do you...

Answer B: Chunking into short sections with headings, a glossary at the top, signposted paragraphs, and every fact kept.

Why: That she follows the audio is the diagnostic: the barrier is decoding, not the ideas. "Simplify" is resolved by the model to the conceptual problem and gives you a text that is readable and empty — the history removed rather than delivered. Format changes keep every fact and make the same content reachable, and "do not remove any event, name or date" has to be said explicitly and then checked. Print that version for the whole class; it is easier for everyone and removes the stigma of the special sheet.

Module 7 test, Q4, p. 291

Phone dictation transcribes your speech almost perfectly and your student's with an error every few words. What does that tell you to do?

Answer D: Test with the actual student before building anything on it, and try a Whisper-based app.

Why: Recognisers were trained mostly on speech that is disproportionately American and British English, so accuracy follows the same gradient as everything else in the course. That is a property of the training data, not of the student's diction or the hardware, and it varies by student rather than by class — so the test is five sentences with the actual person on the actual phone. Whisper-based apps have wider accent coverage and run offline; where nothing works well enough, a human scribe is still the answer.

Module 7 test, Q5, p. 291

A quiet student asks a chatbot at ten at night the question she would never ask in class.

What is the cost of that, and...

Answer A: Nobody hears the answer either, so a wrong one is never corrected; a weekly anonymous-slip routine fixes it.

Why: The privacy that made the question possible is the same privacy that removes the correction, and the teacher never learns the gap was there. The routine keeps the privacy of the asking and removes the privacy of the answer: once a week, five minutes, every student writes one question they asked or wanted to ask, anonymously, and you answer several aloud. That gives you the gaps in the students' own words, corrects the wrong answers publicly, and shows the quiet student that six other people had her question.

Module 7 test, Q6, p. 292

For a language with a few hundred thousand speakers and almost nothing written online, a glossary and a pasted register sample do not improve the...

Answer C: Prompting shapes what the model does with what it knows, and here it knows almost nothing.

Why: This is the one place in the course where the prompting toolkit stops applying. Glossaries fix terminology in a language the model can already handle; they cannot supply grammar, register or idiom that was never learned. It is a limit of the data, not of the prompt. What still works is using the model for the English side in full, having a speaker check any draft, letting students translate as a bilingual exercise, and checking whether a regional model trained deliberately on the language group exists.

THE LESSON CHECKS**Lesson 50, p. 258**

A teacher translates a science worksheet into a language spoken by two million people with little online presence. The output looks fluent. What is the...

Answer B: Translate it back to the source language in a fresh chat and compare.

Why: Translation quality tracks how much text exists in a language — excellent for Spanish and Hindi, uneven for Yoruba and Amharic, poor for languages with a few million speakers and little web presence — and the cheap check is to translate the result back and compare it with what you started with.

A Use it, since fluency is the...: Fluency is what the model produces regardless of accuracy; a fluent translation of a small language can carry wrong terms and wrong meanings in perfect-looking sentences.

C Ask the same model whether the...: A model asked to grade its own output tends to confirm it; the check must come from a route that does not share the original error.

D Translate into a related larger language...: Students do not read the related language; the answer is to check the translation into the language they read, and to involve a speaker where the check fails.

Lesson 51, p. 262

A teacher asks for an explanation of fractions "in Hindi" for a class in Delhi and gets formal, Sanskrit-heavy Hindi that the students find harder...

Answer D: It defaulted to the formal written register; paste a sample of how the class talks.

Why: Students in most classrooms think in a mixed language — Hinglish, Sheng, Taglish, Pidgin, Arabic with French — and the model can explain in that mix if shown a sample of it, which makes an explanation land that the textbook's formal register never did.

A The model does not know Hindi...: The model knows Hindi well; it defaulted to the written register it has read most, which is not the spoken one, and a sample corrects that.

B The students' Hindi is weak, and...: Confuses a language lesson with a maths lesson; the aim here is that fractions land, and the register that lands is the one they think in.

C "Hindi" is ambiguous as an instruction...: A label helps less than a sample; the model matches examples far more reliably than it resolves names for registers.

Lesson 52, p. 266

A teacher asks for a "simplified version" of a history text for students with dyslexia and gets a shorter passage with the key events removed...

Answer A: It removed content when the barrier was format; ask for chunking, spacing and a glossary instead.

Why: The same text can be made reachable for a struggling reader by chunking, spacing, a glossary and an audio version — all free, all in minutes — and the trap is simplifying the ideas when the barrier was the format.

B Dyslexic students cannot handle the original...: Dyslexia is a difficulty with decoding, not with ideas; the content was reachable with the right format and was removed unnecessarily.

C The text was simplified correctly for...: The questions were the assessment of the content; the fault is in a "simplification" that deleted what the questions assessed.

D The teacher should have simplified the...: The model does format changes well when asked for format changes; the instruction asked for something vaguer and got the default meaning of "simpler".

Lesson 53, p. 270

A teacher sets up dictation for a student who cannot write fluently. In the demonstration, in the teacher's accent, it is nearly perfect. With the...

Answer B: The recogniser has heard fewer accents like the student's; test it with her, try another tool, treat output as a draft.

Why: Dictation and read-aloud are built into every phone and cost nothing, and they work well in major languages and less well in accented English and small languages — so test with the actual student before relying on them, and treat the transcription as a draft the student corrects.

A The student is speaking too fast...: Speed helps marginally; the gap between the teacher's and the student's results points at accent and language, which the recogniser has heard less of.

C The microphone on the phone is...: The same phone worked for the teacher a minute earlier; the variable is the speaker, not the hardware.

D Dictation does not work reliably for...: It works for many children; the failure is specific to accent and language, and the fix is to test and choose the tool that handles this student best.

Lesson 54, p. 274

A teacher with no access to a specialist uses a model to generate a "support plan" for an autistic student, including strategies and targets, and...

Answer C: A plan for one child needs assessment of that child; the model supports a plan, it cannot write one.

Why: AI gives a teacher without a specialist some real tools — image description for a blind student, live captions for a deaf one, structured and predictable materials for an autistic one, chunked tasks for ADHD — and the line to hold is that these support the student's plan rather than substituting for the professional who should write it.

A The model's strategies are probably wrong...: Many will be reasonable general strategies; the problem is not their quality but that a plan for a specific child requires assessment of that child, which did not happen.

B Nothing is wrong, since some plan...: A generic plan filed as an individual one hides the absence of assessment and can block the referral the student needs.

D The teacher should have used a...: The tier or product does not supply the missing thing, which is a professional who has seen the student.

Lesson 55, p. 277

A teacher notices that her quietest students are using a chatbot to ask basic questions about topics covered weeks ago. What is the most useful...

Answer D: The questions reveal gaps they will not admit to in class; gather them and reteach.

Why: Students ask a chatbot what they will not ask a teacher in front of the class — the question that reveals they are behind — and the same privacy that makes them ask means nobody corrects the wrong answer, so a teacher turns the private questions into a class routine.

A They are using it to avoid...: Asking a question about a topic is not cheating on anything; the use reveals a gap the teacher can now see.

B The tool is doing the reteaching...: The private answers may be wrong and are not checked; the gap exists regardless of the tool, and moving faster widens it.

C Quiet students simply prefer learning from...: Confuses a symptom with a preference; they ask the machine because the room feels unsafe for the question, which is fixable.

Lesson 56, p. 280

A strong Class 10 student uses a chatbot to go far beyond the syllabus in physics and arrives in class with confident, detailed, wrong ideas...

Answer C: Ask her to find the errors, source by source, and report to the class.

Why: The student who finishes in four minutes now has an unbounded question-answerer — and is also the student least likely to be corrected when it is wrong, because she is confident and you are busy — so the stretch is to make her the checker, not the consumer.

A Ban her from using the tool...: Removes the one resource she was using to stay engaged; the problem is unchecked consumption, not the tool.

B Correct her in front of the...: Punishes curiosity publicly and teaches the room not to explore; the error is an opportunity for a private, more useful task.

D Accept the errors, since the topic...: Lets a confident wrong idea stand in a strong student, and misses the chance to teach the checking skill she most needs.

Lesson 57, p. 283

A teacher in a region where the home language has around three million speakers and little online text finds that the model's translations are fluent-looking...

Answer C: It has read too little of the language; use it for English and rely on speakers for the rest.

Why: For a language with a few million speakers and little written on the web, the model is not a translator and will not become one through prompting — so use it for the English side only, involve speakers for the rest, and know that projects exist to build the data and that whose languages get models is a political question, not a technical one.

A Keep refining the prompt with more...: Prompting shapes what the model does with what it knows; it cannot supply knowledge of a language it has barely read.

B Switch to a different model, since...: Most models trained on the same web, and a language absent from the web is absent from all of them; a few regional models may help and are worth checking.

D Translate through a related larger language...: A bridge sometimes helps and often compounds errors; it is a thing to test with a speaker, not a reliable fix.

Examination

The paper starts on p. 293.

Examination, Q1, p. 294

A student says the chatbot lied to her. What is the accurate correction?

Answer A: It produced a sentence that fits, and a sentence can fit beautifully and be false. It has no way to check.

Why: Lying means knowing the truth and saying something else, and that is not what happens. The model produces the sentence that fits the pattern, with no internal check against anything. The phrase worth giving a class is that it is not lying, it is making things up without knowing it. It also does not learn from your conversation, and while it does repeat things people wrote, it equally produces claims nobody ever wrote at all.

Examination, Q2, p. 294

Which one-sentence explanation survives a class asking "then why is it wrong sometimes?"

Answer B: It is a machine that has read an enormous amount of writing by people and guesses what words come next.

Why: Only the prediction sentence contains its own answer to the follow-up: a guess built from what people wrote will be strong where the writing agrees and weak where it is thin. The search-engine version collapses at the first "why is it wrong", the robot-brain version brings every film the class has seen into the room, and the database version is simply false — the training text is dissolved into weights, not stored as records to look up.

Examination, Q3, p. 294

A department of six teachers decides to share one free account to save trouble. What happens?

Answer D: They hit the message cap by lunchtime, because the cap is per account. Each teacher needs their own.

Why: Free tiers cap messages per account and reset on a schedule of a few hours, so six people sharing one login exhaust it early and then work with a smaller model or wait. There is also no shared learning to gain: unless a stored-instructions feature is used deliberately, each conversation starts blank, and a shared login mainly means six people reading each other's chats.

Examination, Q4, p. 295

You switch off the setting labelled something like "improve the model for everyone". What have you achieved?

Answer C: What you type is taken out of the training pipeline; the company still receives and stores it.

Why: The setting governs one thing: whether your text is folded into future training. It does not make the service private, delete anything, or move the model onto your phone. So it is the minimum rather than the answer, and the working rule stands — treat what you paste the way you would treat what you say in a staffroom with the door open. It also cannot be applied backwards to anything already sent.

Examination, Q5, p. 295

A tool that was sharp at seven in the morning is producing thin, shallow work at nine in the evening. What is the usual reason?

Answer B: Demand is high, so the service is quietly giving you a smaller, faster model.

Why: A free tier is a queue. When load is high the service routes you to a smaller model and answers get shallower with no notice at all. Nothing you did changed, so the remedies are to try again in the morning or to keep a second tool open for the rest of the session. Caps, when you hit them, announce themselves; they do not quietly shorten answers.

Examination, Q6, p. 295

You are reading a generated answer before using it with a class. What should you mark as unverified?

Answer B: Every proper noun, number, date and quotation.

Why: Specifics are where the unreliable ground is: names, figures, dates and quotations are the things written down rarely and reconstructed from pattern, and citations in particular come out looking perfect and not existing. Everything left unmarked is usually fine. Tone is no guide at all, because the mechanism produces the most likely next word whether the underlying evidence is a million textbooks or three blog posts; hedges like "approximately" are trained habits, not measurements.

Examination, Q7, p. 296

For which request is it better to turn web search off?

Answer D: A set of likely misconceptions and three ways to explain one idea.

Why: For conceptual work — explanations, analogies, likely wrong answers — search usually makes the answer worse, because it replaces the model's broad synthesis of an enormous amount of teaching writing with a summary of the two pages that happened to rank this morning. Dates and local figures are exactly what search is for. And a searched answer is the one that can be checked, because it hands you a page to open.

Examination, Q8, p. 296

Your pasted brief is nine paragraphs long and one constraint matters more than the rest. Where should it go?

Answer B: At the top or the bottom of the brief.

Why: Attention inside the window is not uniform. Models are measurably worse at using material buried in the middle of a long input than material at either end, which researchers call being lost in the middle. So position is a lever you control. A separate earlier message is worse rather than better, because it is then further from the request in the conversation, and repetition costs window space that the material needs.

Examination, Q9, p. 296

You ask the model to interview you before it writes anything. What has to be added so the questions are useful?

Answer C: Ask about constraints I might not have mentioned, rather than about objectives.

Why: Left unsteered the model asks the generic questions that fill teaching templates — what are your learning objectives — which you can already answer and which change nothing about the output. Steered toward constraints, it asks about the things it needs and you forgot: the printer, the fixed benches, the fifteen minutes lost to the bell. The number of questions is not the variable; what they are about is.

Examination, Q10, p. 297

Which of these should be left out of a teaching prompt entirely?

Answer A: Your teaching philosophy, and a request that the activity be engaging.

Why: Every sentence that is not one of the five things dilutes the ones that are, and the memory lesson explains why that matters: material competes for one window and for the model's attention within it. Philosophy and engagement adjectives describe a quality it is already aiming at. Class size and absent equipment are two of the five and change the output immediately, and the misconception is the single most valuable line — it belongs at the top, not later.

Examination, Q11, p. 297

You paste a good question of your own and ask for eight more like it. What else should the message say?

Answer C: Keep the structure, the length and the kind of thinking; change the passage, the numbers and the context.

Why: Without the split, "like this" is ambiguous and the model sometimes keeps the content and changes the structure, which is the opposite of what you wanted. Naming what the example is — a two-mark inference question — gives it a label to generalise from. Beyond three or four examples you are spending the window on demonstration for little return, and a long description of why the example is good is precisely the thing the example replaces.

Examination, Q12, p. 297

You are struggling to describe the difference between a question that tests reading and one that tests understanding. What is the fastest way to convey...

Answer D: Show one labelled as the kind you do not want and one labelled as the kind you do.

Why: Two examples, one marked bad and one marked good, hand the model the boundary rather than a description of it, and the boundary is the part teachers find hardest to put into words. The copying worry is real but is answered by the label: an example marked as unwanted is used as a negative, not as a template. Nothing about showing a weak example lowers the standard of what follows when it is labelled.

Examination, Q13, p. 298

You want a worksheet that is both rigorous and accessible, and iterating between the two is producing an oscillation. What should you do?

Answer B: State both goals in the same message and let it find the balance in one pass.

Why: Each message pulls toward one pole, so sequential pulls do not converge — ask for harder, then more accessible, then harder again and you get an oscillation rather than a synthesis. Two goals stated together are a single target the model can aim at. It can hold competing constraints perfectly well; what it cannot do is remember, four messages later, that you also wanted the opposite of what you just asked for.

Examination, Q14, p. 298

You want three hundred questions in a spreadsheet you can sort and filter. What do you ask for?

Answer D: CSV with named columns, saved as a .csv file and opened in a free spreadsheet.

Why: CSV is the format built for this: paste into a text file, save with the .csv extension, open in Google Sheets, Excel or LibreOffice Calc, and the columns are columns. A rendered table is fine inside the chat and poor in Word. And precise character-level layout is the one format the model cannot reliably produce — space-aligned columns approximate, and the approximation breaks the moment the font changes.

Examination, Q15, p. 298

Which pair of format instructions removes the most editing from a term's worth of generated material?

Answer B: Do not restate the question in the answer; no preamble and no summary.

Why: Restating the question doubles the length of every answer key, and the cheerful opening and closing lines are two deletions on every output, forty times a term. British spelling matters and is one instruction, not two. "Be concise" and "use simple language" are the adjectives the register lesson showed to move almost nothing, because neither can be counted or checked.

Examination, Q16, p. 299

A prompt that produced excellent distractors in March produces ordinary ones in September. What is the best repair?

Answer D: Paste the March output you saved and ask for eight more like these.

Why: This is why the file keeps the output and not only the prompt. Providers update models every few months, and the saved output is a demonstration that survives the change — pasting it as the example recovers most of what the update took away, by the same mechanism as showing it one of your own. A model cannot revert to a previous version on request, and rewriting from scratch throws away the calibration you spent a term building.

Examination, Q17, p. 299

Most teachers ask a model for a full lesson plan. What should they ask for instead?

Answer B: Likely wrong answers, and what students are thinking when they give them.

Why: A full plan is the thing an experienced teacher is already good at producing, and it is the request that hands over the judgement. What the model can do that you cannot is draw on far more written teaching than one person can read: ask for eight likely misconceptions and one or two will be errors you have not seen in fifteen years. Analogies, questions that expose a misconception, and a fourth explanation at nine in the evening are the same kind of request.

Examination, Q18, p. 299

What does a generated lesson plan get wrong every single time?

Answer A: Timing: it writes fifty-five minutes of activity for a forty-minute period.

Why: Cut a third before looking at anything else. Its other reliable failures are under-producing the moments where you find out whether anything landed, inventing named videos and page numbers, and putting think-pair-share into rooms where nobody can turn round. Sequencing goes the other way — its default is definition first, which is the textbook order and the wrong one — and core school science is its strongest ground, not its weakest.

Examination, Q19, p. 300

You ask for a passage that uses ten target words. What is the failure to check for?

Answer D: The words arrive together in one paragraph, which satisfies the instruction at its cheapest point.

Why: An instruction is met at the cheapest point that satisfies it, so "use these ten words" produces ten words in one paragraph and a job done. The instruction you wanted is: spaced through the passage, each in a sentence that shows what the word means from context. Check by counting and by reading — all ten present, spread out, each guessable — and send back the ones that clumped. Bolding and defining are separate faults, and both are prevented by saying so.

Examination, Q20, p. 300

You want the same passage at three reading levels so the class can discuss one text.

Which instruction is load-bearing?

Answer A: Keep every fact the same.

Why: Without it the versions drift apart — a detail dropped here, a name changed there — and the class is no longer discussing one text, which was the entire point. Everything else follows from it: if the facts hold, one set of questions works and the paragraph count does not matter. Publishers charge for levelled texts; the instruction that makes a free version usable is the one that pins the content while the language moves.

Examination, Q21, p. 300

Your question bank has a difficulty column filled in by the model. What should happen to it over the year?

Answer D: Replace it with the proportion of your class who got the question right last time it was used.

Why: The model's rating is a guess about an average class somewhere, and yours is a guess about your class; both are guesses until students have answered. A question you rated 3 that ninety per cent got right is a 1, and a question rated 1 that forty per cent missed is telling you about a misconception and belongs in next week's starter. It costs a minute per test to record and after a year the bank knows your students better than any generated rating could.

Examination, Q22, p. 301

The model has labelled each question in your bank as recall, application or inference. How much weight should the labels carry?

Answer A: Little: around a third of what it labels application is recall in a costume.

Why: Its idea of inference is loose, and a question that requires the student to recall which formula to use is not application however it is labelled. Read ten per type to recalibrate yourself, then fix the rest with a pointed second message. The column is worth keeping precisely because it is the thing that stops you setting twelve recall questions and calling it a test — but only once the labels are yours rather than its.

Examination, Q23, p. 301

You have a lesson to deliver on a projector that fails about once a fortnight. What should the model produce?

Answer B: The outline and the speaker notes, with the layout left to a free slide tool.

Why: The substance is the part that matters and the part the generators do worst: they are built for business pitches, where a slide is a backdrop for a speaker, and they produce a heading, an image and a line of text. Ask for six-word titles, three short bullets and eighty to a hundred and twenty words of speaker notes per slide, and you have a script. When the projector fails you print the notes pages and you have a handout that stands on its own.

Examination, Q24, p. 301

A twelve-week unit plan comes back looking reasonable. What are the two things it could not have known?

Answer B: Your school calendar, and which topics your students found hard last year.

Why: It plans twelve uninterrupted weeks; your term has nine, once the festivals, the exam week in the hall and the monsoon week are taken out. And its "students find this hard" is a guess from general teaching writing, often right, and worse than your record of the class that fell apart on hydrographs two years ago. Prerequisites and revisits it will produce readily when asked, and reading level and lesson length are exactly the sort of thing a prompt conveys well.

Examination, Q25, p. 302

Which judgement is entirely outside what a machine marking a script can see?

Answer C: Whether this piece represents progress for this student.

Why: Progress is a fact about the student rather than about the text: is this the first time Kwame has used paragraphs at all? That is often the most important thing about the piece and it exists only in your memory of him. The other three are observable in the writing and are exactly what an observable rubric asks for — which is why they belong in the rubric the model sees and progress does not.

Examination, Q26, p. 302

What is the reason not to let a model set a grade that goes on a record?

Answer C: A parent will ask why, and there is no reasoning you can walk through and nothing to teach from.

Why: The reason is not that it is always worse than you — on some narrow tasks its agreement with trained markers is respectable. It is that a grade you cannot defend is not a grade, it is a number: no appeal path, no reasoning to walk a parent through, and nothing the student can learn from. Use it to find things and use your own judgement to decide them.

Examination, Q27, p. 302

Which of these belongs in your own adjustment on the script rather than in the rubric the model applies?

Answer B: Originality against what this class usually produces, and improvement since October.

Why: Some things you mark for are real and not observable in the text: the risk a student took, the argument nobody else attempted, the improvement since the autumn. They are judgements about the student rather than the writing, so they stay out of the machine's rubric and go on the script in your own hand — rubric total eleven out of fifteen, plus one for the argument in paragraph three. Descriptors and totals are precisely what the model needs to be given.

Examination, Q28, p. 303

What is the honest limit of a rubric written so that every criterion is observable?

Answer D: It can be met mechanically: it describes the floor of competence, not the ceiling of quality.

Why: A student who learns that two quotations with an explained effect earn a three will produce exactly two quotations and a formulaic explanation — and so will a model asked to write a top-scoring essay. That is a real cost and it is worth paying, because the alternative rewards length. For the pieces where quality matters, the rubric finds the floor and your reading finds the rest. Observable criteria are also shorter and clearer to students, not longer.

Examination, Q29, p. 303

What makes feedback most likely to change the next piece of work?

Answer A: Phrasing each point as something to do in the next version rather than as a fault in this one.

Why: This draft is finished; the next one is where a change can happen, so "put the quotation before the claim next time" is actionable in a way that "this quotation is misplaced" is not. Two points is a task and six is a list nobody acts on. And a mark beside a comment makes the comment invisible — students read the number and stop — which is why feedback and grade are separated wherever that is possible.

Examination, Q30, p. 303

Students are using a prompt card that forbids the model to rewrite their draft. What goes wrong, and what holds it in place?

Answer C: It starts suggesting phrasings and then whole sentences, so the card is restated every third message and a version trail is kept.

Why: The model was trained to be helpful, and the most helpful-seeming thing to do with a weak sentence is to write a better one. Over a conversation that pull wins: by the fourth exchange it is suggesting a phrasing, by the sixth it has offered a paragraph the student did not ask for. The reminder line holds it most of the time, and the before-and-after drafts with a note of what changed make the contribution visible when it does not.

Examination, Q31, p. 304

In peer assessment with a model as third marker, what is the part that teaches?

Answer D: The disagreements, taken back to the rubric and the essay to find the line each marker read differently.

Why: Agreement is uninteresting. Two students who scored the same paragraph four and two, sent back with the rubric and one instruction — find the line you read differently and the place in the essay where it applies — usually find it in ninety seconds, and have learned the rubric's boundary in a way no explanation from the front achieves. The model's score is a third reading that makes the argument discussable, not an answer key; sometimes it is the outlier and that is a lesson too.

Examination, Q32, p. 304

You have photographed thirty handwritten scripts and transcribed them. What is the record of what each student wrote?

Answer A: The paper. The transcription is a draft to check against it.

Why: The model predicts the most likely text given the image, so it silently repairs spelling, supplies missing words and can turn a wrong answer under correct working into the right one. That is a correction of exactly the thing you were marking. The photograph is only an image of the paper and inherits nothing; asking the student to confirm a transcription puts the burden on the person least able to challenge it. Check the page for the point the feedback is about, before the feedback goes back.

Examination, Q33, p. 304

You are uneasy about one particular essay. How should the conversation open?

Answer C: By asking them to talk you through how they got to this argument.

Why: Curiosity rather than accusation. A student who wrote it talks happily for two minutes about why she chose that example and what she cut; a student who did not usually tells you, or makes it obvious without the word ever being said. A detector percentage cannot be shown to a parent and is not evidence, and two detectors reading the same statistical property agree with each other for the same wrong reason. Never open with the accusation, and never in front of the class.

Examination, Q34, p. 305

What is the point of one piece of in-class writing on paper early in the term?

Answer A: It gives you a baseline for what each student sounds like, so that you notice honestly rather than suspect.

Why: A baseline is not evidence to convict with; it is the thing that makes you notice when something reads oddly, and it protects the student whose ability you had underestimated as much as it flags anything. Used as a similarity test it becomes a home-made detector with the same bias. And it proves nothing about later work — which is why it sits alongside the process record and the conversation rather than replacing them.

Examination, Q35, p. 305

In the three-level scheme, what does Level 2 permit?

Answer C: You may draft with it, and you must show what it produced and what you changed.

Why: Level 0 is none, because the task is the exercise. Level 1 is assistant: explanation, checking, getting unstuck, nothing handed in. Level 2 is collaborator: drafting is permitted and disclosure of what changed is required. The scheme works because the level is marked on every task, so nobody can claim not to have known, and because it teaches the real judgement — that the right amount depends entirely on what is being learned.

Examination, Q36, p. 305

Which rule keeps an AI policy from building a two-tier classroom?

Answer C: No task should require a student to have their own AI access to complete it.

Why: Some students have a paid account and a good phone, some a shared handset and patchy data, some neither, and any policy that assumes access lets that advantage compound quietly. So either provide the output on paper for everyone, or make the AI part optional and ungraded. Pairing makes access visible in the room and is not always possible; extra time addresses the wrong problem; and school accounts are a good thing that most schools in the world will not have.

Examination, Q37, p. 306

What is the exercise that most reliably teaches a class to distrust a confident wrong answer?

Answer B: Read out a wrong answer as though it were authoritative, let the class accept it, then show them the truth.

Why: The point is not that it was wrong; it is that it was wrong in exactly the same voice it uses when it is right, and the only way to feel that is to have been taken in. A student who has personally been caught once, in a safe room, treats the tool differently for the rest of their life. Do it in the first fortnight, with the answer written down beforehand so the demonstration cannot fail. Side-by-side comparison teaches the opposite lesson: that you can tell by looking.

Examination, Q38, p. 306

Why does a class AI policy need an exemption list naming spellcheck, single-word translation and the calculator?

Answer D: Without it, every student is in breach from the first week and the policy cannot be enforced at all.

Why: Every word processor checks spelling and every phone translates, so a policy with no exemptions is broken by the whole class on day one, and a rule everyone breaks is not a rule. The list also needs a date and a review each term, because a tool that was a grammar checker in September is a rewriting assistant by March and the policy has to say which side of the line it now sits on.

Examination, Q39, p. 306

Which homework task still measures something now that the model exists?

Answer A: Revise these three topics; five questions from memory at the start of the next lesson.

Why: There is nothing to generate: the learning happens in the revising and the check is five minutes of writing in the room. The other three ask for a product the model makes instantly, and the chapter and its questions are very likely in the training data as well. The replacements that work are retrieval, practice with the answers already printed so the working is the task, reading with a one-line response, and evidence gathered from the student's own street.

Examination, Q40, p. 307

What is a declaration of AI use for?

Answer D: It tells the reader what the work rests on, in the same way a bibliography does.

Why: Teach it as an ordinary act of academic honesty rather than as a confession, and mark it as present and adequate but never for content — a student who declares heavy use on a Level 2 task has done nothing wrong, and marking it down teaches them to hide next time. None of the major style guides treats a model as an author, because an author can be responsible for a claim and a model cannot, and that principle is the part worth teaching since the formats will change.

Examination, Q41, p. 307

Why write the model's answers down before the lesson rather than demonstrating live?

Answer C: Live demonstrations fail, and the model may give a different answer on the day.

Why: A lesson built on "watch what it says" has no plan B: the connection drops, the model has changed, or it happens to get the question right this time and the whole point collapses. A lesson built on "here is what it said, and here is what was true" always works, photocopies, and survives a power cut. Batch all the online work into one session when power and data exist, and bring paper.

Examination, Q42, p. 307

A story is retold down a line of five students and changes as it goes. What is that activity teaching?

Answer A: Bias in training data, because the class has just been the training data.

Why: What survives the retelling is what each person found worth remembering, which is the setup for the teaching question: if nearly every story ever written about doctors used "he", what word will a machine guess after "the doctor said that"? Ten minutes, no equipment, and they will not forget it, because the habits being described were theirs a moment earlier.

Examination, Q43, p. 308

What is the right idea to teach children aged eight to ten?

Answer C: A machine shown many labelled examples learns to sort new ones, and gets it wrong at the edges.

Why: Children of that age already sort — animals, shapes, words — and the new step is that a machine can be trained by being shown examples rather than told rules. They also hold the edge case easily: the dog that looks like a cat, and why the machine gets it wrong. Prediction belongs at eleven to thirteen, bias and evaluation at fourteen to sixteen. For the chatbot question at this age, "it has learned patterns from lots of writing, like the sorting game but with words" is enough.

Examination, Q44, p. 308

Which words should be withheld from nine-year-olds playing the sorting game?

Answer B: Algorithm, neural network and probability.

Why: Give them the words that describe what they just watched — examples, pattern, mistake, and "it only knows what it was shown". Withhold the vocabulary that names machinery they have not seen. The concept is what matters and the words come easily later once the concept is there; a child who has the concept and none of the jargon is further ahead than one with the jargon and no concept.

Examination, Q45, p. 308

What should be on the card students are given about fabricated intimate images?

Answer D: Do not forward it; screenshot only what is needed to report; tell a named adult that day.

Why: Every copy is a further harm and in many places a further offence, so not forwarding comes first, including to show someone. The screenshot is for the account and the time, not the image. And the school route runs the same day, because the school's job in the first hour is to stop the spread by direct request, contact the platforms, contact the victim's parents first, and treat the victim as a victim throughout. Deleting silently leaves the school unable to act at all.

Examination, Q46, p. 309

Which set of signs should prompt a conversation about a student's use of a companion app?

Answer A: Withdrawal from friends, tiredness from late-night use, and speaking of the character as a person.

Why: None of the three is proof of anything on its own; together they are a reason to talk. Ask what she likes about it before saying anything about what it is, because most students will describe the exact features the design produces — it always listens, it never judges — and a student who reaches the mechanism herself holds it far better than one who is told. If self-harm, sexual content or real dependence comes up, it goes to the safeguarding lead that day.

Examination, Q47, p. 309

Students build a tally of which word follows which from one page of text, then generate sentences from it. What does the output show them?

Answer C: It looks one word back, so the sentences are grammatical in pieces and meaningless as a whole.

Why: The sentences start well, because those pairs were in the text, and then drift into fluent nonsense. Asked why, the class produces the whole of the first block unprompted: it only looks one word back, it has only seen one page, and it has no idea what it is saying. Then somebody asks whether the real one is just this but bigger — and it is, with far more text, thousands of words of context instead of one, and a much cleverer way of storing the tallies.

Examination, Q48, p. 309

In a school distributing the AI unit across subjects, which department naturally owns provenance and the question of where a source came from?

Answer B: History, which has asked where a source came from and why for a century.

Why: Historians ask where a source came from, who made it, why, and what it can and cannot tell you, which is the deepfake lesson's "provenance, not appearance" in a subject that has taught it for generations. A generated account of a local event set beside the real newspaper, photograph and interview, with the question "which of these is evidence, and what makes it so", is a history lesson that happens to be an AI lesson. Art owns the training-data and artists question; citizenship owns companions and who profits.

Examination, Q49, p. 310

A translated worksheet uses an unfamiliar word for a subject term your students have already learned. What fixes it?

Answer C: Paste your textbook's glossary and say to use exactly those terms.

Why: Ordinary sentences translate better than subject vocabulary, and left to itself the model picks the term from whichever textbook tradition it read most, which is rarely yours — a coinage, a transliteration, or the word used in a neighbouring country. Ten words in both languages, pasted, with an instruction to use exactly those, fixes it. Students who learned a concept under one name are genuinely confused by a worksheet using another.

Examination, Q50, p. 310

Which of these must not be produced by machine translation alone, even with a back-translation check?

Answer B: A report to a parent or a safeguarding document.

Why: The back-translation check is for practice material, where an error costs a confused ten minutes. Where a wrong word has a real cost — an exam paper, a report home, anything touching safeguarding — a speaker of the language reads the translation before it is used, and the model's job is to produce the draft that saves that speaker an hour. That is a large saving and it is the correct division of labour.

Examination, Q51, p. 310

Why can the model write a worked example in Hinglish or Sheng at all?

Answer D: Because those mixes are written in enormous quantities online, in forums, comments, messages and subtitles.

Why: The mixed registers are all over the internet, which is why the model has the pattern for Hinglish, Taglish, Spanglish, Sheng, Nigerian Pidgin and Singlish and produces them convincingly when shown a sample. Where a mix involves a language low on the gradient, the small-language half is unreliable and the back-translation check applies to those sentences. And the model's version of a mix is a few years old and from somewhere else, so the students are its editors.

Examination, Q52, p. 311

A parent objects that a worked example written in Hinglish will damage her daughter's English. What is the honest answer?

Answer A: Explanation in the strongest or mixed language improves understanding of the content and does not harm acquisition of the formal register.

Why: The evidence on bilingual education is reasonably clear here: a worked example in the mix is a bridge rather than a replacement, the formal version follows and lands better because the idea is already in place, and the formal register is taught as its own thing. Parents who see English as the route out are right that English matters and wrong that a bridge to it costs anything. Giving the mixed version only to the weakest students turns a bridge into a label.

Examination, Q53, p. 311

Which part of making a text reachable for a dyslexic reader is yours rather than the model's?

Answer D: Font, spacing, alignment and background colour.

Why: The model cannot see the page. A sans-serif face at twelve to fourteen point, line spacing at one and a half, left-aligned with a ragged right edge, wider margins and for some students a cream background are a minute's work in any free document tool, and a strip of coloured plastic is cheaper than any software. Everything about the words — chunking, sentence length, glossary, signposting — is what you ask the model for. Print one version for everyone and the special sheet disappears.

Examination, Q54, p. 311

Where does the line fall for a teacher supporting a student with a disability and no specialist within reach?

Answer A: It can draft the social story and the chunked task sheet; it cannot write the student's support plan.

Why: A plan for a specific child is written by someone who has assessed that child, and the model has not. It will produce a document that looks like a plan and is generic, and filed as the student's plan that does two harms: it stands where an assessment should be, and it can satisfy a system that then never refers the child. Everything under that line — social stories, visual timetables, chunked tasks, image descriptions, and a good letter asking the district for an assessment — helps today.

Examination, Q55, p. 312

What is the right stretch for the student who finishes everything in four minutes and now has an unbounded question-answerer?

Answer B: Make her the checker: verify the model's claims against sources that are not a chatbot.

Why: She is the student most likely to absorb its errors: confident, working beyond the syllabus where the model is least reliable and you are least able to catch it, and least supervised because twenty-eight others need you more. The stretch she is missing is not content but telling right from wrong in confident material, which is harder than anything on the syllabus and requires finding real sources, comparing them, and holding an unsettled question as unsettled.

Examination, Q56, p. 312

Your students' home language has a few hundred thousand speakers and almost nothing written online. What do you do this term?

Answer C: Use the model for the English side in full, and involve speakers for the rest.

Why: Every technique in the course still works on the English material — passages, questions, explanations, feedback — and the crossing into the home language is done by people: a speaker checking a draft, or the class itself translating as a bilingual exercise, which is better pedagogy than a machine translation would have been. Pasting a few pages does not teach a model a language, and a Class 5 student will be in Class 10 before coverage arrives.

USE IT IN YOUR WORK

AI for Teachers

For teachers who now have to teach and use AI, often with no training and no projector.

WHAT IT TEACHES	For schoolteachers anywhere who now have to teach and use AI, usually without training.
WHAT IS INSIDE	Seven modules and 57 lessons, 26 figures drawn from data, seven case studies, seven module tests, a 56-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-for-teachers

Addaly

First edition, September 2026 · build 62a26075



Scan to open the course