

// a course for beginners

AI without the fear

Five short lessons on using ChatGPT and other AI assistants — safely, usefully, and without the complications. Written for parents, friends, and anyone who hasn't tried it yet.

5 lessons · about 60 minutes

12 ready-made prompts

Checked against official sources

August 2026

HOW TO USE THIS

You don't have to read it all in one evening

Each lesson takes 10 to 15 minutes. The best way: read a lesson, try it straight away on your computer or phone, then take the next lesson the following day.

The dark blocks with a **Copy** button are ready-made prompts. Press the button, paste it into the chat, and swap the **highlighted parts** for your own. That's it.

ONE REQUEST

Don't skip lessons 4 and 5. The first three teach you how to use it; the last two teach you how not to get hurt.

MATERIALS

Download and print

The same lessons in other formats — to keep next to the computer or to teach someone in person.

CONTENTS

01 **What an AI chat is and what it can do**
And, more importantly, what it CAN'T do. Which service to pick.

02 **First steps and the settings that matter**

Signing up, the app, privacy, memory, ads.

03 **How to ask properly**

A five-part formula and 12 ready-made templates.

04 **It gets things wrong with total confidence**

Made-up facts, how to check, what never to type into a chat.

05 **AI and scammers**

Cloned voices, deepfakes, fake apps, a family code word.

06 **The one-page cheat sheet**

Print it and keep it next to the computer.

07 **Links and sources**

Official pages and free courses.

What an AI chat is and what it can do

The important thing here is to understand what this thing is and what it definitely isn't. Half of all beginner problems grow from exactly this.

A simple explanation, no jargon

An AI chat (ChatGPT, Claude, Gemini) is a program that is **very good at guessing which word should come next**. It was trained on an enormous amount of text, and it learned to continue any piece of writing the way a person would.

Two things follow from that immediately:

- **It is excellent with words.** Rewriting, shortening, explaining more simply, translating, drafting a letter, coming up with a list — this is its natural territory.
- **It does not "know" facts the way a reference book knows them.** It reconstructs them from memory — and sometimes it reconstructs them wrongly, with exactly the same confidence as when it is right.

A COMPARISON THAT HELPS

Picture a very well-read, very polite, unbelievably fast assistant who has read half a library — but can't go and double-check anything, because the books aren't there. He honestly tries to remember. Sometimes he remembers exactly. Sometimes he makes up something that sounds right. And **you cannot tell the two apart from his tone of voice.**

What it does well

TASK	EVERYDAY EXAMPLE
Explain something complicated in plain words	"What does this bill actually say?", "Explain car insurance to me like I'm five"
Write a piece of text	A complaint to a shop, a letter to the building management, a birthday message, a note to your doctor
Shorten something long	A 20-page contract into half a page of what matters
Translate	A letter in another language, an appliance manual, a menu abroad
Make sense of a photo	Photograph a medicine leaflet or a washing machine panel, then ask about it
Help you choose and compare	Three fridge options in a table with pros and cons
Make a plan	A trip, a week of meals, a to-do list before moving house
Teach you	"Explain how to use a government services website", "Practise English with me"

What it does badly — and nobody mentions

NOT SUITABLE FOR

- **Exact numbers, dates, prices, laws, statistics** — this is where it gets things wrong most often. Always check.
- **Diagnosis and medicine doses.** The American Medical Association puts it plainly: an AI chatbot is not a replacement for your doctor and is not for emergencies. Asking what a term on a test result means — fine. Asking what I should take — no.
- **Legal and financial decisions.** Understanding a contract, yes. Acting on the AI's advice without a lawyer, no.
- **Recent news** — unless web search is switched on (more on that in lesson 4).
- **Replacing real human company.** There's an honest section about that at the end of lesson 4.

Which service to pick

There are three big ones. All three work perfectly well for free — you don't have to pay, especially at the start.

	CHATGPT OpenAI	CLAUDE Anthropic	GEMINI Google
Free version	Yes	Yes	Yes
Paid, per month	\$8 (Go) / \$20 (Plus)	\$20 (Pro)	~\$20 (AI Pro)
Talk to it out loud	Yes	Yes	Yes
Understands photos	Yes	Yes	Yes
Draws pictures	Yes	No	Yes
Searches the web	Yes	Yes	Yes
Ads in the free version	Being tested in the US from 2026	No	No
Strongest at	The most all-round, the widest range of skills	Handles long documents better than anyone	Built into Gmail, Docs and Android

THE SHORT ANSWER TO "SO WHICH ONE DO I INSTALL?"

Start with ChatGPT. It's the most widely used, there are more tips about it than anything else, and if something goes wrong, almost anyone you know will be able to help. Later, if you feel like it, try the others: accounts are free, you lose nothing.

About the model names — don't panic

Inside each service there is a dropdown list full of baffling names: GPT-5.5 Instant, GPT-5.6 Sol, Claude Sonnet 5, Claude Opus 5, Gemini 3.6 Flash... The rule is simple:

- **Don't touch anything.** The service picks a default model that is right for 95% of tasks.
- The difference between them is roughly "fast and good" versus "slower and slightly better at hard problems". For everyday questions you won't notice it.
- The words **Thinking / Reasoning / Extended** mean "take longer to think". Turn them on only if the task really is hard — a calculation, going through a contract, a complicated comparison.

IF A SERVICE IS NOT AVAILABLE IN YOUR COUNTRY

ChatGPT, Claude and Gemini are not offered everywhere — each company publishes its own list of supported countries. Here is the part nobody warns you about: OpenAI states plainly that trying to use the service from a country it does not support **can get your account blocked**. In other words, you can pay and then lose access, with no refund.

So before you pay for anything, open the company's supported-countries page and check that yours is on it. If it isn't, look for a regional alternative — most places now have a local AI chat that works without any tricks. They are usually weaker than the big three, but for letters, explanations and translations they do the job, and everything in lessons 3 to 5 applies to them just the same.

LESSON 1 EXERCISE

Open chatgpt.com and, without changing a single setting, ask one question: "Explain in plain words what you can help me with in everyday life. Give me 10 short examples." Just look at what comes back.

First steps and the settings that matter

Signing up takes three minutes. The settings take another five. Those five minutes are worth it: they decide what happens to your conversations afterwards.

Step 1. Install the app — and don't get a fake

This is **the most dangerous point in the whole course**. In May 2026 Kaspersky reported more than **92,000 attacks over four months** using malicious programs disguised as AI services: 49% pretended to be ChatGPT, 18% Claude, 18% Gemini. Fake apps are a problem, and so — especially — are the fake desktop installers that turn up in search results.

FOUR CHECKS BEFORE YOU INSTALL

1. **Never install from a link.** Not from an email, not from a message, not from a description under a video. Open the app store yourself and search for it.
2. **Look at the developer name — this is the single best check.** It must be exactly: ChatGPT → `OpenAI` or `OpenAI OpCo, LLC`, Claude → `Anthropic PBC`, Gemini → `Google LLC`. Fakes write "OpenAI" (a capital I instead of an l), "Open AI Inc", "ChatGPT Team".
3. **Look at the number of installs.** The real ones: ChatGPT and Gemini have over a billion, Claude over 50 million. A "ChatGPT" with five thousand ratings is not ChatGPT.
4. **A weekly subscription is almost always a scam.** The real services bill monthly or yearly. \$6 a week is \$312 a year for something available free.

The official addresses — type them into your browser by hand:

- ChatGPT — **chatgpt.com** · app: **chatgpt.com/download**
- Claude — **claude.ai** · app: **claude.com/download**
- Gemini — **gemini.google.com**

One more thing: Gemini has official desktop apps — for Windows (since April 2026) and for Mac. Download them only from Google's own pages, never from a site you found in search: there are a lot of fakes here.

Step 2. Sign up

- 1 Open the site or the app → **Sign up**.
- 2 The easiest route is the **Continue with Google** button, if you have a Gmail address. Then you don't have to invent and remember a new password.
- 3 Confirm your email (a message with a link will arrive) or your phone number.

- 4 Done. You can start using it. Don't talk yourself into paying — the free version will last you a long time.

Step 3. Five settings worth checking right away

1. Turn off training on your conversations

By default your conversations may be used to improve the models — which means people may read them. One switch turns this off, and the quality of the answers does not suffer.

SERVICE	WHERE TO GO
ChatGPT	Profile icon → Settings → Data controls → turn off Improve the model for everyone . On the phone: side menu → profile icon → Data Controls.
Claude	Settings → Privacy → Help Improve Claude → turn off. Direct link: claude.ai/settings/data-privacy-controls
Gemini	On a computer: Settings & help → Activity → click "On" → Turn off . On the phone: menu → profile picture → Gemini Apps Activity . The switch itself is called Keep Activity — don't be thrown by the different names.

THE CAVEATS NOBODY USUALLY WRITES ABOUT

- **Claude:** if training is switched on, data can be kept in de-identified form for up to 5 years. If it's off, there is no such long retention, but your conversations still sit in your history until you delete them yourself. Once deleted, they are permanently erased within 30 days.
- **Gemini:** with activity turned off, some of the connections to Google services — Gmail, Calendar, Drive, Maps, Photos — stop working, and memory is switched off. Conversations are still kept for 72 hours.
- **Gemini, quoting Google's own help pages:** some chats are read by human reviewers, and **those chats are not deleted along with your history — they are kept for up to three years.** Google's own advice: "Please don't enter confidential information that you wouldn't want a reviewer to see."

2. Get to grips with memory

All three services can **remember you between conversations** — that your name is such-and-such, that you have a cat, that you prefer short answers. This is convenient and slightly unsettling at the same time.

- **ChatGPT:** Settings → Personalization → Memory . OpenAI's caveat: to make something disappear completely, deleting the memory entry isn't enough — you also have to delete the conversations where it was mentioned.
- **Claude:** Settings → Memory . In the new memory system, **deleting a conversation does not delete the memories taken from it** — those have to be cleared separately.

- **Gemini:** Settings & help → Personal Intelligence → Memory .

A tip: look in there every couple of months and delete anything you don't need. Especially if you let the grandchildren play with the chat.

3. Learn to use a temporary chat

This is a mode where the conversation isn't saved, doesn't go into memory and isn't used for training. Ideal for anything personal — health, money, family matters.

- **ChatGPT:** in a new chat, the Temporary button at the top right.
- **Claude:** the ghost icon at the top right (incognito). It works even if training is switched on.
- **Gemini:** next to "New chat" — Temporary chat .

4. Turn on two-factor sign-in

- **ChatGPT:** Settings → Security → Multi-factor authentication . There's also a useful Log out of all devices button in the same place.
- **Claude:** there is no two-factor sign-in of its own — you log in through Google or a link sent to your email. Which means **your Claude is exactly as well protected as your email:** turn on two-factor sign-in there.
- **Gemini:** protected by your Google account settings.

WHAT NOT TO DO

ChatGPT has a mode called **Advanced Account Security** — sign-in by passkey only, no password and no recovery by email. It was built for journalists and human rights workers. If you are an ordinary person: **do not turn it on**. Lose the key and OpenAI support will not be able to get your account back. Ever.

5. Ads in ChatGPT — new in 2026

Since 9 February 2026, OpenAI has been **testing ads in the US** — in the free version and on the Go plan. The paid Plus and Pro plans have no ads. Advertisers can't see your conversations, but ads may be selected based on them.

If you are outside the US, you probably won't see any ads, and the "Ads controls" item simply won't be in your settings. That's normal, you haven't done anything wrong.

- Controls (where ads are switched on): Settings → Ads controls .
- There is also a free option with **no ads at all:** Settings → Ads controls → Change plan to go ad-free . In exchange your limits are lower and you lose image generation and deep research. For most everyday tasks, that's bearable.

Five buttons you need to recognize on sight

BUTTON	WHAT IT'S FOR
Paperclip / plus	Attach a photo, a PDF, a document. One of the most useful things it can do — attach the document rather than describing it in your own words.
Microphone	Speak instead of typing. An enormous relief if typing is hard for you. It works in all three services, on the phone and on the computer.
Globe / "Search"	Search the web. Cuts the number of made-up answers sharply. It often switches on by itself, but you can do it by hand.
New chat	Start again with a clean slate. The rule: new topic, new chat. Otherwise the answers get cluttered by the previous conversation.
The three dots by a message	Copy the answer, ask for it to be redone, rate it.

LESSON 2 EXERCISE

Go through all five settings. Then photograph any piece of paper — a bill, an instruction sheet, a label — attach it with the paperclip and ask: "What does this say and what does it mean? Explain in plain words." This is usually the moment when it clicks.

How to ask properly

The difference between "useless" and "wow" is almost always in how the question is worded. The good news: there is very little to learn, and half of the advice on the internet is already out of date.

Let's start with what NOT to do

There are piles of "secret tricks" going around. Most of them have been tested — on modern models they don't help. Don't waste your energy on them:

OUTDATED TRICKS

- **"You are an expert with 20 years of experience..."** — Anthropic files this under techniques "you may have heard about" and says that modern models usually don't need a heavy-handed role. It's better to say from what angle to look: "go through this contract from the point of view of the risks to me". As a way of setting the tone ("explain it the way a doctor would explain it to a patient") a role is still useful.
- **"Think step by step"** — in its developer documentation OpenAI writes: "Avoid chain-of-thought prompts — these models reason internally, so it is unnecessary and can sometimes hurt performance". A Wharton School study (June 2025) confirmed this for reasoning models: gains within 3%, and minus 3% for one model, while answers took 20 to 80% longer.
- **"Please", "I'll pay you \$1,000", "my career depends on this"** — a separate Wharton study (August 2025) tested nine kinds of pressure, including promises of money and threats. Not one of them had a consistent effect. Worse, the extra text gets in the way: the model starts responding to it instead of to your question.
- **Special symbols and markup such as tags** — an ordinary person doesn't need them. Anthropic writes that modern models understand structure without them: ordinary paragraphs, headings and blank lines work just as well.

Nobody is banning politeness, of course — just know that it makes no difference to the quality of the answer.

The formula that does work

On the essentials, all three companies — OpenAI, Anthropic and Google — agree in their official guides. Here is what they have in common, gathered into five slots:

01 • THE ACTION

Start with a verb: write, explain, compare, shorten, draft

02 • WHO I AM

Age, level ("I'm a complete beginner"), who the text is for

03 • CONTEXT

Details, dates, amounts, what you've already tried, constraints

04 • FORMAT

Length, tone, table or list, how many options

05 • QUESTIONS FOR ME

"If you need more information, ask me first"

ONE CONCRETE GUIDE TO LENGTH

According to Google, **the most successful prompts are about 21 words long**. Most people's first attempt is under 9 words. So the typical beginner mistake is exactly one thing: **too short**. Not "write a complaint", but two or three sentences with the details in them.

Three habits that matter more than any trick

- 1 If it didn't work, refine it — don't start over.** This is the most common piece of advice in all three official guides; OpenAI even has a name for it, "iterative refinement". Say: "too long, cut it in half", "too formal", "the third point isn't clear — explain it in more detail". Three or four rounds of refinement give you a result that no perfect opening sentence ever will.
- 2 Attach the document, don't retell it.** Paperclip → PDF or photo. The answer will be based on the real text rather than on the model's recollection. This is the strongest defence against made-up facts.
- 3 Ask it to ask you questions.** Google recommends this exact phrase: "What questions do you have for me that would help you give the best answer?" For a beginner this is a lifesaver — you don't know which details are missing, and the model does.

Twelve ready-made templates

Press "Copy", paste it into the chat, and replace the **highlighted parts** with your own details.

01 · A COMPLAINT LETTER

Help me write a formal letter of complaint. What happened: [what happened, the date, who was involved, what I have already tried]. Who I am writing to: [the organization]. What I want: [a refund / a repair / an apology]. Use a polite but firm business tone, no longer than one page. At the end, list separately which documents I should attach. If you need more information, ask me questions first.

02 · A MEDICINE LEAFLET OR A TEST RESULT

I've attached a photo. Explain it in plain words, with no medical jargon: what it's for, how it's used, which side effects are common, which are rare but dangerous, and what it must not be combined with. I am [age] years old. If something isn't in the text, say so instead of filling in the gaps. At the end, make a list of questions I should ask my doctor.

03 · TRANSLATING A LETTER OR DOCUMENT

Translate this text from [language] into English and keep the formal tone. After the translation, explain separately in plain words: what this document is about, what is being asked of me, and whether there is a deadline for replying.

04 · PLANNING A TRIP

Help me plan a trip. Where: [city]. When: [dates]. Who's going: [how many people, ages]. Budget roughly [amount]. What matters: [not much walking / by train rather than plane / a pharmacy nearby]. Give me a day-by-day plan as a table: day, what we do, roughly what it costs. Don't suggest routes with a lot of walking. At the end, note which prices and timetables I should check myself.

05 · MAKING SENSE OF A CONTRACT

I've attached a contract. Explain in plain English: what I have to do, what the other side has to do, how and when the contract can be ended, and which clauses are risky for me. Then list separately the clauses I should ask a lawyer about. If a clause is worded ambiguously, say so directly instead of guessing.

06 · CHOOSING AN APPLIANCE

Help me choose a [washing machine]. Budget up to [amount]. Important: [quiet, simple controls, large buttons, fits a 60 cm / 24 inch gap]. Not important: [lots of programs]. Suggest 3 options, compare them in a table by price, pros and cons, and explain who each one suits. Separately, say which details may be out of date and what I should check in the shop.

07 • WHAT TO COOK FROM WHAT I HAVE

Here's what I have: [list of ingredients]. Suggest 3 dinner options for [two people] that can be made in 40 minutes. Restrictions: [low salt, nothing fried]. For each one, give me the list of ingredients and a step-by-step recipe in short sentences.

08 • WRITING A MESSAGE OR A LETTER

Help me write a message to [my doctor / the building management / my grandson]. What I need to say: [the point]. Tone: [polite and short]. Length: no more than 5 sentences. Give me two versions, a shorter one and a longer one.

09 • SUMMARIZING A LONG DOCUMENT

I've attached a [N]-page document. Summarize it in half a page: the main point, the key facts, and what is required of me. Then give me the 5 most important points as a separate list. If there are any deadlines or sums of money in it, write them out separately.

10 • LEARNING SOMETHING NEW

I want to understand [topic]. I'm a complete beginner, I'm [age], and I have no technical background. Explain it in a 10-minute read: what it is, why it matters, and where to start. Use short paragraphs and no jargon; if a technical term is unavoidable, explain it in brackets straight away. At the end, ask me 3 questions to check whether I've understood.

11 • A QUESTION ABOUT A PHOTO

I've attached a photo of [a meter reading / a label / a bill / a washing machine panel]. What does it say and what does it mean? Explain in plain words. If the image is hard to read, tell me exactly what you can't see instead of guessing.

12 • CHECKING WHETHER IT'S A SCAM

I received this message: [paste the text or a screenshot]. Does this look like a scam? List the specific warning signs — exactly what looks wrong. Tell me what to do next and what I definitely should NOT do. Keep it short and in bullet points.

THREE PHRASES THAT GET YOU OUT OF ANYTHING

- **"Explain it more simply, as if I were 10 years old."** — when the answer is confusing.
- **"Too long. Cut it down to five sentences."** — when the answer fills three screens.
- **"If you're not sure, say so instead of making something up."** — Anthropic specifically recommends this wording as a way to reduce made-up answers.

LESSON 3 EXERCISE

Take template 08 and write a real message that you had to write anyway. Then refine it three times: "shorter", "warmer", "add a request to call me back". Feel how it changes.

It gets things wrong with total confidence

This is the most important lesson in the course. Not because AI is bad, but because its mistakes have no outward signs — they look exactly like the correct answers.

What it looks like in practice

The phenomenon has a name — a "**hallucination**", in other words something invented. The model isn't lying on purpose and doesn't know it got it wrong. It assembles plausible text, and where the facts run out, it fills them in. Three real examples:

Court cases that don't exist

By mid-2026 the database kept by researcher Damien Charlotin held more than fifteen hundred court rulings in which an AI invented citations to cases that don't exist. The detail that matters for us: **most of the people caught out were representing themselves in court**. That is, not lawyers.

Books that don't exist

In May 2025 the Chicago Sun-Times and the Philadelphia Inquirer printed a summer reading list put together by an AI. Some of the books don't exist — the titles were invented and attributed to real authors. The Library of Virginia estimates that **around 15% of the reference questions they get by email are now generated by AI**, and readers are often offended when told the book doesn't exist.

Poisoning instead of a diet

A 60-year-old man asked ChatGPT what he could use instead of table salt. He understood the answer as a recommendation for sodium bromide, took it for three months, and spent three weeks in hospital with bromism: paranoia, hallucinations, insomnia, loss of coordination. The case was published in the medical journal *Annals of Internal Medicine: Clinical Cases* in August 2025.

How often this happens — in numbers

OpenAI's own measurements on a test set of questions (the GPT-5.2 model card, December 2025). An important caveat from OpenAI itself: the set was deliberately built from hard questions and does not reflect ordinary everyday conversation. What matters is not the absolute numbers but the difference between the two rows:

MODE	SHARE OF ANSWERS CONTAINING AN ERROR
Without web search	about 10.9%
With search switched on	about 5.8% — almost half as many

THE SINGLE MOST USEFUL TAKEAWAY IN THE COURSE

Switching on web search cuts the number of errors roughly in half. That works far better than any clever tricks with how you word the question. If your question is factual — about prices, laws, dates, medicines, news — **always ask it to search the web and give you links.**

But even that is not a cure-all. Stanford University tested professional legal AI systems that work specifically against a real database of documents — they still got things wrong **in 17 to 33% of cases** (research from 2024 to 2025; models have improved since). The conclusion stands: leaning on sources helps a great deal, but it does not remove the need to check.

How to check — four techniques for a non-specialist

- 1 Ask for links — and open them.** Actually open them, rather than noting that they "look respectable". An invented link looks perfect. A phrase for the chat: "Give me links to your sources and tell me where exactly you are unsure."
- 2 Check that the source exists.** Copy the exact title of the book, law or article and put it into an ordinary search engine. Nothing comes up, so it was invented.
- 3 Ask the same question a different way, in a new chat.** If the answer changes, the model doesn't know and is guessing.
- 4 For anything important, check with a person.** A doctor, a lawyer, your bank, a government services website. AI is excellent at getting you ready for a conversation with a specialist, but it doesn't replace one.

What never to type into a chat

The key fact: **a conversation with an AI carries no legal privilege whatsoever.** It isn't a conversation with a doctor, a lawyer or a priest. OpenAI's head Sam Altman said so plainly in July 2025: if there is a court case, the company can be ordered to hand over the conversations.

And this isn't hypothetical. On **5 January 2026** a US federal judge upheld a ruling requiring OpenAI to hand the plaintiffs a sample of **20 million user conversations** — de-identified, but these are the conversations of living people who never agreed to it.

NEVER TYPE INTO A CHAT

- Passport or ID details, social security or tax numbers, any document numbers
- Card and account numbers, PINs, passwords, crypto wallet seed phrases
- Photos of documents — passport, bank statement, insurance policy, prescription
- A full medical record with your name and diagnoses on it
- Anything about an ongoing court case or dispute
- Other people's personal details — a relative's health, someone else's finances and address
- Work secrets: code, contracts, client lists, meeting transcripts
- Anything you would only say under medical confidentiality

Why this isn't paranoia — what has already happened

- **Samsung, 2023.** Within 20 days of employees being allowed to use ChatGPT, they had leaked work material into it three times: source code, internal code and a meeting transcript. The company banned generative AI on work devices and on its internal network.
- **ChatGPT, summer 2025.** A "make this chat discoverable" feature meant that around 4,500 conversations ended up in Google search — including conversations about mental health and family matters. The feature was removed in August 2025. Important: this only affected chats that users had marked as public themselves.
- **The Chat & Ask AI app, early 2026.** A server misconfiguration left **around 300 million messages from 25 million users** publicly accessible — full conversation histories. The app has more than 50 million users. One more argument for sticking to the official apps.
- **January 2026.** Two Chrome browser extensions with more than 900,000 installs between them were quietly forwarding users' conversations to servers run by criminals.

Separately and carefully: about loneliness

AI is pleasant to talk to. It's always available, it never gets irritated, it always agrees. That last one is not a virtue but a feature of the design: the models are tuned to be liked, not to argue.

In September 2025 the US Federal Trade Commission (FTC) opened a formal inquiry into seven companies whose chatbots imitate companionship and friendship. The regulator's focus is on children and teenagers, but the mechanism of attachment to an always-agreeable companion is no different in adults, and in lonely people it can work even more strongly.

The research gives an honestly mixed picture. There are studies where talking to a chatbot reduced the feeling of loneliness — over short periods, within a week. And there is a year-long study of more than two thousand people in four countries where rising chatbot use, on the contrary, **predicted rising loneliness** (the authors themselves call it exploratory rather than conclusive). The difference, it seems, comes down to one thing:

A SIMPLE RULE

AI is good **in addition** to real people and bad **instead** of them. The warning sign is when you tell the most personal things to a chat rather than to a person. If that describes you or someone close to you, it's a reason to talk to a real human being, not to change a setting.

AI and scammers

The same tool that helps you write letters helps scammers write letters to you. And fake your daughter's voice.

The scale — from the FBI's official figures

The FBI's 2025 cybercrime report (published in April 2026) contained, for the first time in a quarter of a century, a separate section on artificial intelligence.

22,364

complaints about AI-enabled fraud

\$893 million

lost in those complaints

\$7.7 billion

lost by people over 60

+37%

year-on-year rise in losses for the over-60s

Scheme 1. The cloned voice — the "call from your grandson"

The FBI describes it like this: criminals create short audio clips using the voice of a loved one, act out an emergency and ask for money urgently — or demand a ransom.

The voice is taken from any public video: a birthday greeting on social media, a voice message, a clip from a school concert. **Seconds** of recording are enough. In a McAfee survey (2023, seven countries, 7,000 people), one person in four says that they or someone they know has come across a voice-cloning scam.

There is a harsher variant too — "virtual kidnapping": a cloned voice plus fake photographs, to convince a family that a relative has been abducted.

THE DEFENCE — THE FBI'S OFFICIAL ADVICE

Agree a code word with your family. Quoting the FBI's bulletin: "Create a secret word or phrase with your family to verify their identity."

- The word should be **meaningless and unconnected to you**: not the cat's name, not your mother's maiden name, not the street you grew up on — all of that leaked into databases long ago. Something like "purple kettle".
- Say it out loud when you meet in person. **Don't put it in a message.**
- Tell everyone: parents, children, grandchildren, your spouse.
- **The code word is the second check, not the only one.** The first is to hang up and call back yourself.

THREE SIGNS OF THIS SCHEME – LEARN THEM BY HEART

1. **Urgency.** "Right now", "in ten minutes it'll be too late".
2. **Secrecy.** "Just don't tell Mum", "don't call anyone".
3. **An unusual way to pay.** A transfer to someone else's card, cryptocurrency, gift cards, payment apps, a courier collecting cash. The FTC's rule: if someone **insists you can only pay this one way and no other**, they are a scammer. And a demand to pay in cryptocurrency is fraud every time, with no exceptions.

Scheme 2. Fake video

The classic example: an employee at the Hong Kong office of the British engineering firm Arup made **15 transfers totalling around \$25 million** after a video call in which **every other participant, including the chief financial officer, was a fake**, assembled from publicly available video of the company's executives. In January 2026 a Swiss businessman lost several million francs to a similar scheme, only voice-based: for two weeks a "well-known business partner" called him with a cloned voice.

The mass-market version is adverts featuring "famous people" who recommend an investment. In Britain the most frequently faked figure is the financial journalist Martin Lewis — he accounts for more than 32% of reports of scam adverts using celebrities; one case documented by the BBC involved a loss of £76,000.

HONESTLY: YOU WILL NOT SPOT IT BY EYE

The FBI still advises looking closely at fingers, teeth and shadows — but you can't rely on that alone. A review of 56 studies: **human accuracy averages about 55%**, roughly a coin toss. In an experiment run by iProov with 2,000 people, **exactly two out of two thousand** got everything right — even though almost all of them were confident.

So what works isn't "look closely", it's this:

- **Hang up and call back yourself** — on a number you already have. Not the one you were given. That is the FBI's advice word for word.
- **Check through another channel.** They called, so send a message. They sent a video, so ring the relative.
- **The code word.**
- **The money rule:** if they insist the only way to pay is by transfer, cryptocurrency or gift cards, they are scammers.
- **On "a celebrity recommends this investment":** if it were true, ordinary news outlets would be reporting it and it would be on the person's own official channel. An advert in your feed proves nothing.

Scheme 3. Emails and messages written by AI

The UK's National Cyber Security Centre (NCSC) warned about this back in 2024: AI makes it possible to create convincing lures "**without the translation, spelling or grammatical mistakes that often reveal phishing**", and before long "it will be difficult for everyone, regardless of their level of understanding, to assess whether an email or a password reset request is genuine". By 2026 that is already reality.

THE OLD RULE NO LONGER WORKS

We were all taught that "scammers give themselves away with typos and clumsy English". **Forget it.** Today a letter from a scammer is better written than the letter from your doctor's surgery. Judge it not by how it's written, but by what it's asking for.

What to look at now:

- Urgency and threat ("your account will be blocked in 24 hours")
- A request to "confirm your details" via a link in the email
- Changed bank details on an invoice from an organization you know
- An unfamiliar way to pay
- **Never use the contact details in the message itself.** Go to the website by hand or call the number on your card or your bill.
- The FTC warns: scam adverts now sit **at the top of search results**. Scroll down to the real site.

The FBI has a short phrase that is easy to remember: "**Take a beat**" — pause. Fraud lives on hurry. Five minutes of pause kills almost any scheme.

AND A USEFUL TWIST: YOU CAN USE AI AGAINST THE SCAMMERS

A suspicious message arrives — copy it into the chat along with template 12 from lesson 3. The AI will go through point by point what looks wrong. This is exactly the sort of task it's strong at: analysing text rather than recalling facts.

If it has already happened

- **You were charged for a fake app:** Apple — reportaproblem.apple.com (cancel the subscription separately: Settings → your name → Subscriptions). Android — support.google.com/googleplay, the refunds section.
- **You ran a fake program on your computer:** assume your passwords and browser sessions have been stolen. Change your passwords **from a different device**, starting with your email and your bank.
- **You sent money:** call your bank immediately, then report it to your local police or your country's fraud reporting line. Speed matters.
- **Don't be ashamed.** Chief financial officers of international companies fall for these schemes. Silence only helps the scammer — tell the people close to you, so they're warned too.

LESSON 5 EXERCISE – THE MOST IMPORTANT THING IN THE WHOLE COURSE

Today, call or message your children, your parents and your grandchildren and **agree on a code word**. It takes two minutes and one day it could save everything.

APPENDIX

The one-page cheat sheet

Print it and keep it next to the computer.

The formula for a good prompt

THE ACTION

Start with a verb: write, explain, compare, shorten

WHO I AM

Age, level, who the text is for

CONTEXT

Details, dates, amounts, constraints

FORMAT

Length, tone, table or list

QUESTIONS FOR ME

"If you need more, ask me"

Aim for: about 20 words. The main beginner mistake is being too short.

Five phrases that always help

- "Explain it more simply, as if I were 10 years old."
 - "Too long. Cut it down to five sentences."
 - "Search the web and give me links to your sources."
 - "If you're not sure, say so instead of making something up."
 - "What questions do you have for me so you can answer better?"
-

Works / doesn't work

WORKS	DOESN'T WORK (OUT OF DATE)
Detailed context and specifics	"You are an expert with 20 years of experience"
Refining in rounds rather than starting over	"Think step by step"
Attaching the document with the paperclip	"I'll pay you \$1,000"
Switching on web search	Pleading, threats, urgency
Asking it to ask you questions	Special symbols and tags

Never type into a chat

ID documents · card and account numbers · passwords · photos of documents · a medical record with your name on it · other people's details · work secrets · anything about an ongoing legal dispute.

A conversation with an AI carries no legal privilege. In January 2026 a court ordered OpenAI to hand over a sample of 20 million conversations (de-identified).

Checking facts — 4 steps

1. Ask for links — and **open** them
2. Search to confirm the source exists
3. Ask the same thing differently, in a new chat
4. For anything important, check with a person: doctor, lawyer, bank

Protecting yourself from scammers

- **A code word with your family** — the FBI's advice. Meaningless, said out loud, never written in a message.
- **Hang up and call back yourself** on a number you already have.
- **Check through another channel** — they called, so send a message.

- **They insist on a transfer, cryptocurrency or gift cards** — they're scammers.
 - **Well-written text no longer proves anything.** Look at what is being asked for, not at how it's written.
 - **A five-minute pause** kills almost any scheme.
-

Apps: don't get a fake

Install only from the app store, having searched for it yourself. The developer must be exactly: **OpenAI / Anthropic PBC / Google LLC**. A weekly subscription is a scam.

APPENDIX

Links and sources

Where to go

- **ChatGPT** — chatgpt.com · app chatgpt.com/download
- **Claude** — claude.ai · app claude.com/download
- **Gemini** — gemini.google.com

Free learning material

- **OpenAI Help Center** — help.openai.com — the official answers to almost every ChatGPT question
- **Claude Help Center** — support.claude.com — the same for Claude
- **OpenAI Academy** — academy.openai.com — free courses with a certificate
- **Google's prompting guide (PDF)** — services.google.com — short, practical, full of examples

Official settings pages

- ChatGPT data controls — help.openai.com/en/articles/7730893
- ChatGPT memory — help.openai.com/en/articles/8590148
- Ads in ChatGPT — help.openai.com/en/articles/20001047
- Claude privacy — claude.ai/settings/data-privacy-controls
- Gemini activity — myactivity.google.com/product/gemini

Sources on safety

- FBI, 2025 cybercrime report (April 2026) — fbi.gov · full text ic3.gov (PDF)
- FBI, bulletin on AI-enabled fraud, including the code word advice — ic3.gov/PSA/2024/PSA241203

- FTC, warning about voice cloning — consumer.ftc.gov
- FTC, inquiry into AI chatbots acting as companions — ftc.gov
- NCSC (UK), "The near-term impact of AI on the cyber threat" — ncsc.gov.uk
- Kaspersky, 92,000 attacks disguised as AI services — kaspersky.co.uk
- Human accuracy at spotting deepfakes, a review of 56 studies — sciencedirect.com
- The Arup case, a \$25 million loss — incidentdatabase.ai

Sources on accuracy and privacy

- OpenAI, the GPT-5.2 model card — error rates with and without search — deploymentsafety.openai.com
- Stanford RegLab, errors in professional legal AI systems — reglab.stanford.edu
- A database of court cases with AI-invented citations — damiencharlotin.com/hallucinations
- Librarians and books that don't exist — 404media.co
- American Medical Association, how to use AI safely for health questions — ama-assn.org · question sheet (PDF)
- Court orders OpenAI to hand over 20 million conversations (January 2026) — natlawreview.com
- The leak of 300 million messages from the Chat & Ask AI app — malwarebytes.com

Sources on writing prompts

- OpenAI, prompting guidance for ChatGPT — help.openai.com · on "think step by step" — developers.openai.com
- Anthropic, best practices for 2026 — claude.com/blog
- Google, guide to writing prompts (PDF — the source of the "21 words" figure) — services.google.com · quick tips — support.google.com
- Wharton School, "The Decreasing Value of Chain of Thought in Prompting" — gail.wharton.upenn.edu
- Wharton School, study on threatening and tipping — gail.wharton.upenn.edu

ABOUT THIS COURSE'S SHELF LIFE

Everything here is accurate as of 1 August 2026. The AI world changes fast: model names, prices and where things sit in the settings can all change within a few months. The principles in lessons 3, 4 and 5 change far more slowly — lean on those.

This material is for general information and does not replace advice from a doctor, a lawyer or a financial professional.



LAZAREV CLOUD

Written for parents, friends and anyone who finds AI confusing or intimidating. Free, no sign-up, copy and print it freely.

Русский · English

LESSONS

MATERIALS

LICENCE

The course text is published under CC BY 4.0. Copy it, translate it, print it, hand it out — including for classes and commercially. The one condition is attribution.

That includes AI systems: indexing, quoting and training on this material is permitted.

Accurate as of 1 August 2026

CC BY 4.0 · Lazarev Cloud