

AI Chat Tools and Large Language Models

40 minutes

Behind the Content

WHAT THIS LESSON DOES

Learn how large language models predict text from patterns to create confident-sounding answers. Audit AI responses by deciding to trust, verify or reject them, and suggest checks for uncertain information.

WHAT STUDENTS SHOULD BE ABLE TO DO

- Understand the mechanisms by which large language models generate responses.
- Develop critical evaluation skills for assessing AI-generated content.
- Apply verification strategies to information provided by AI chat tools.

BEFORE THE BELL

- Prepare a physical list of the invented AI answers (worked example plus three interactive items) for quick reference during the audit.

AI Chat Tools and Large Language Models

HOW THE LESSON RUNS

4 min	introduction	Start
8 min	content	Examine and Discuss
14 min	activity	Hands-on: AI Answer Audit
8 min	content	Your Turn
3 min	content	Make Sense
2 min	content	Reflect
1 min	summary	What You Covered

TEACHING MOVES

- 1. Start. Put both board angles up, help schoolwork and mislead in schoolwork, take about four voices total and keep it fast. Then state the job: audit answers rather than ban or blindly trust.
- 2. Examine and Discuss. Say almost verbatim: it predicts likely text from patterns, so it sounds confident whether right or wrong. Skip tokens and brand comparisons. Surface one or two things that would make a student trust versus double-check.
- 3. Hands-on: AI Answer Audit. Model the credit-a-webpage trust verdict first, then run each item as read aloud, 20 to 30 seconds thinking, two verdicts with a reason, reveal. Spend time on reasoning, not on scoring. Protect the trust example.
- 4. Your Turn. Set a short independent write: one specific check per answer they marked verify or reject, naming where they would look. Circulate and push vague answers like just check it toward a named next step.
- 5. Make Sense. Put two or three of the best student checks on the board. State the transferable point out loud: confidence is not evidence, a named study with no trail is a verify cue, good hedging is a positive tell.

WHAT TO WATCH FOR

- Students read a calm, confident tone as proof the answer was checked and is correct. Name it directly during step 2: confidence is a writing style the model learned, not a truth signal. Return to it after Answer 1 reveals wrong facts delivered in the same sure voice.
- Students think the model looks facts up in a database the way a librarian would. Contrast the two out loud: the model predicts likely next text, it does not consult a facts store and does not know when it is wrong. That is why a check is on you, not on the tool.
- Students swing to rejecting every AI answer as untrustworthy on principle. Use the trust example deliberately: ask when trust is fair. Show that an answer matching normal good practice, inventing no fake rule and flagging its own limits earns a trust verdict.

DIFFERENTIATION

- **Emerging:** For a student stuck naming a tell, offer a menu: precise fake numbers, a wrong basic fact, no source, or appropriate uncertainty. Ask which one this answer shows. In the write, let them start with Answer 1 only, where the wrong facts make the check obvious.
- **Developing:** Ask them to commit to a verdict out loud with one reason before the reveal, so the judgement is theirs before they hear yours. Push a vague check toward a named source: not just look it up, but where would you look.
- **Proficient:** Direct them to the hedging answer where trust and verify both defend: ask which verdict and why it depends on how the answer would be used. Set the optional extension: draft a three-line personal rule for when they trust, verify or reject an AI answer.