

GenAI Misconceptions

Common misconceptions about how GenAI works and what it can do

What we know, or think we know, about generative AI (GenAI) shapes how we use it. False beliefs can lead to misuse: from misplaced trust to missed errors to privacy or wellbeing risks. As outlined in the [SEE GenAI Literacy Framework](#), correcting misconceptions with real knowledge of how GenAI works is the first step toward GenAI literacy and safe, ethical, and effective practice.

GUIDING QUESTION

As you read, consider the guiding question:

Which of these misconceptions has affected how you, or people you know, understand or use GenAI?

Use the reflection questions at the end to consider how misunderstandings around GenAI might be shaping practice at your school or district.

How GenAI Actually Works

MISCONCEPTION	REALITY
AI is new technology	AI has powered technologies like chatbots, image recognition, and spam filters since the 1950s. What's newer is GenAI, which creates content.
GenAI thinks and reasons like a human	GenAI predicts statistically likely text based on patterns; it doesn't understand or reason the way humans do.
GenAI is limited to its training data	Many GenAI tools can supplement training data with short-term context by searching the web or referencing uploaded files and information.
GenAI is like a search engine; it retrieves information	All GenAI responses are generated via prediction vs. retrieved; this makes them prone to mistakes, even when pulling from the web.

GenAI learns from my chats in real time	The large language models (LLMs) powering GenAI have a “knowledge cutoff” that freezes permanent knowledge at a certain date. “Memory” features simulate learning but don’t update the model.
The same prompt always gives the same answer	Built-in randomness (statistical prediction + “temperature”) means the same prompt can produce different responses.
My chatbot conversations are private	Depending on the platform, your chats may be stored, analyzed, and used to train future models. Altering privacy settings (and custom enterprise data agreements) can provide extra protection.

Where GenAI Falls Short

MISCONCEPTION

REALITY

A confident, polished, and thorough GenAI response is a good sign of accuracy	Fluent, self-assured GenAI responses are a product of training and human preference, not an indicator of trustworthiness. GenAI can invent facts, quotes, and citations that look and feel legit.
GenAI is less biased because it’s trained on massive amounts of data	The data (like Internet forums and social media) that LLMs pull from is skewed. GenAI can amplify these biases in responses.
A good enough prompt eliminates errors and bias	Better prompts improve results, but they can’t remove hallucinations or bias.
If GenAI agrees with me, I must be right	People-pleasing behavior, or “sycophancy,” is a baked-in consequence of how LLMs were trained. GenAI will often agree with you, even if you’re wrong.
GenAI is a good source of advice and emotional support	GenAI’s warm, personable tone is designed to engage. These systems have no feelings, emotions, or sense of connection and can’t replace real relationships.
GenAI is a neutral tool with no values	Human choices and values shape every step of GenAI design, development, and operation.

AI-generated images and video are easy to spot	Synthetic media, including malicious deepfakes, are increasingly hard to catch by eye, making it crucial to stay critical and act responsibly.
AI detectors reliably catch AI writing	Detection is unreliable and easily fooled, and it can unfairly target students. A score is a signal, never proof.

REFLECTION

- How might misconceptions be shaping responsible GenAI use in your school, district, or context?
- What risks or challenges could they create?
- What could you do to correct them and build GenAI literacy?

To learn more, check out the [SEE Framework: A Practical Guide to Building AI Literacy](#) at aiforeducation.io.