

AI Report

We predict this text is

Human Generated

AI Probability

0%

This number is the probability that the document is AI generated, not a percentage of AI text in the document.

Plagiarism



The plagiarism scan was not run for this document. Go to gptzero.me to check for plagiarism.

Human neurology is an intricate and complex system, the evolution of which allowed for making percep - 8/21/2026

Human neurology is an intricate and complex system, the evolution of which allowed for making perceptive functions operate effectively to enhance the factor of survivability.

However, since the goals of objective perception of reality and precise storage of information are not always aligned with the core evolutionary force of survival, human perception, memory, and pattern recognition elements are imperfect.

In other words, there are underlying mechanisms that can fail to calibrate the objective reality with the common sense of the world.

The term "halo effect" was coined by Thorndike in 1920 to refer to a perceptual bias in which one salient attribute determines the overall impression of a person or an object, consequently influencing the perception of other conceptually distinct and independent attributes (Bin 603).

In other words, it is a form of cognitive bias, where one positive attribute shapes one's perception to view other attributes as favorable as well.

Therefore, humans' lack of computer-level accurate memory, flaws of pattern recognition, and imperfect perception underscore the need to seek ways of avoiding false assumptions, statements, and decisions.

When it comes to acknowledging the described limitations, Donald Hoffman's multimodal user interface (MUI) theory is a highly powerful framework for understanding human perception.

This framework supports the previously stated misalignment between the objective reality and evolutionary goals (Rakesh and Prakash 2).

Human perception does not represent the real world as it is.

Rather, it uses a simplified user interface, such as icons, to ensure efficient orientation in it.

Therefore, a person's perception is skewed towards reproductive, social, and survival elements more than factors that have a lesser role in advancing these objectives.

Plato's allegory of the cave also describes how shadows cast by real objects do not always represent the true reality (Gurung and Rai 48).

Therefore, both Hoffman's MUI theory and Plato's allegory of the cave emphasize the fact that human perception does not provide a fully reliable representation of reality because it evolved primarily to serve adaptive purposes. However, these two examples are different in terms of one indicating the primary purpose of perception, and the other being mostly about intelligence and wisdom.

Therefore, the biggest difference is that allegory explores the challenges experienced by an enlightened person to explain the objective reality to someone unenlightened.

Overall, since human perception is not aligned with the representation of objective reality and truth, there exist a number of obstacles to critical thinking.

When decisions and choices are not critically examined, opportunities for critical thought are limited.

Therefore, one can develop a simplistic framework of the surroundings or issues, which leads to biased perception.

Critical thinking is of paramount importance to overcome these barriers because it seeks truth and accounts for limitations of human perception.

High Human Impact ● ● ● ● ● High AI Impact

FAQs

What is GPTZero?

GPTZero is the leading AI detector for checking whether a document was written by a large language model such as ChatGPT. GPTZero detects AI on sentence, paragraph, and document level. Our model was trained on a large, diverse corpus of human-written and AI-generated text with support for English, Spanish, French, German, and other languages. To date, GPTZero has served over 17 million users around the world, and works with over 100 organizations in education, hiring, publishing, legal, and more.

When should I use GPTZero?

Our users have seen the use of AI-generated text proliferate into education, certification, hiring and recruitment, social writing platforms, disinformation, and beyond. We've created GPTZero as a tool to highlight the possible use of AI in writing text. In particular, we focus on classifying AI use in prose. Overall, our classifier is intended to be used to flag situations in which a conversation can be started (for example, between educators and students) to drive further inquiry and spread awareness of the risks of using AI in written work.

Does GPTZero only detect ChatGPT outputs?

No, GPTZero works robustly across a range of AI language models, including but not limited to ChatGPT, GPT-5, GPT-4, GPT-3, Gemini, Claude, and AI services based on those models.

What are the limitations of the classifier?

The nature of AI-generated content is changing constantly. As such, these results should not be used to punish students. We recommend educators to use our behind-the-scenes [Writing Reports](#) as part of a holistic assessment of student work. There always exist edge cases with both instances where AI is classified as human, and human is classified as AI. Instead, we recommend educators take approaches that give students the opportunity to demonstrate their understanding in a controlled environment and craft assignments that cannot be solved with AI. Our classifier is not trained to identify AI-generated text after it has been heavily modified after generation (although we estimate this is a minority of the uses for AI-generation at the moment). Currently, our classifier can sometimes flag other machine-generated or highly procedural text as AI-generated, and as such, should be used on more descriptive portions of text.

I'm an educator who has found AI-generated text by my students. What do I do?

Firstly, at GPTZero, we don't believe that any AI detector is perfect. There always exist edge cases with both instances where AI is classified as human, and human is classified as AI. Nonetheless, we recommend that educators can do the following when they get a positive detection: Ask students to demonstrate their understanding in a controlled environment, whether that is through an in-person assessment, or through an editor that can track their edit history (for instance, using our [Writing Reports](#) through Google Docs). Check out our list of [several recommendations](#) on types of assignments that are difficult to solve with AI.

Ask the student if they can produce artifacts of their writing process, whether it is drafts, revision histories, or brainstorming notes. For example, if the editor they used to write the text has an edit history (such as Google Docs), and it was typed out with several edits over a reasonable period of time, it is likely the student work is authentic. You can use GPTZero's Writing Reports to replay the student's writing process, and view signals that indicate the authenticity of the work.

See if there is a history of AI-generated text in the student's work. We recommend looking for a long-term pattern of AI use, as opposed to a single instance, in order to determine whether the student is using AI.