

Lack of ‘Human Error’ Unmasks Deceptive AI Systems

Originally published at <https://www.unite.ai/lack-of-human-error-unmasks-deceptive-ai-systems/>



New research finds AI can pass as human until it remembers ‘too well’, with simple memory tests exposing chatbots by their lack of normal human errors.

Researchers from Princeton have developed a method of identifying AI entities pretending to be human, by asking them to carry out tasks that humans are not good at – chiefly related to short-term memory retention.

The AIs tested in this way were unable to adequately replicate human error levels, unless they were specifically instructed to do so in a [system prompt](#), or else were [fine-tuned](#) on psychological data.

The paper states:

‘[We] explore the idea of detecting humanness by using tasks that machines can solve too well to be human. Specifically, we probe for the existence of an established human cognitive constraint: limited working memory capacity.’

‘We show that cognitive modeling on a standard serial recall task can be used to distinguish online participants from LLMs even when the latter are specifically instructed to mimic human working memory constraints.’

‘Our results demonstrate that it is viable to use well-established cognitive phenomena to distinguish LLMs from humans.’

The tendency observed by the researchers implies that off-the-shelf language models are very likely to reveal themselves in any [reverse Turing test](#) that uses this method.

Though ‘goal-specific’ AI models will perform better, fine-tuning on this task will [likely constrain them to it](#), at the expense of general-purpose usage; and while a system prompt [can be as long as War and Peace](#), and therefore could include directions on how to impersonate human foibles, the effectiveness of this method is undermined by being included in very extensive instructions (which will emphasize many *other* priorities), or very short ones (which will sacrifice generalized capability in favor of task specificity, much like fine-tuning).

‘You’re Talking About Memory...’

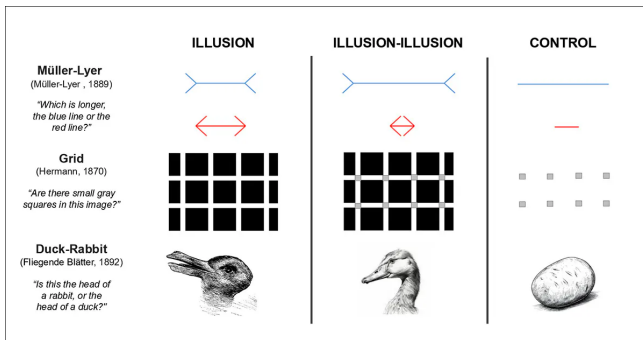
More effective methods to determine AI-generated discourse are increasingly needed – not least by researchers themselves, who frequently must rely on crowdsourced remote workers that are [well-motivated to game the system](#) through automation and other tricks.

Additionally, informed and plausibly delivered AI-generated material is likely to be needed in cases of [AI fraud](#), where real-time conversations demand quick and authoritative answers, and the perpetrators certainly lack the time to Google a query that they were just thrown.

Much as the AI-detection sector could exploit such knowledge, the growth industry of [AI-powered promotional voice calls](#) would presumably benefit from knowing what behavior to *avoid*.

Though it does suggest the possibility of a ‘reverse Turing arms race’, the authors note that should generalized AI become more adept at simulating human foibles, there is a vast reservoir of error-proneness left to draw on*:

'There are many candidates for established human cognitive constraints that LLMs might not inherit. For example, humans get tired, [perceive optical illusions](#), and can only [store few items](#) in their working memory.'



From the late 2024 paper 'The Illusion-Illusion: Vision Language Models See Illusions Where There are None', examples of optical illusions that would likely fool any vision-language model (VLM) that did not already know about them from its training data – though humans are much more likely to resolve the images correctly. [Source](#)

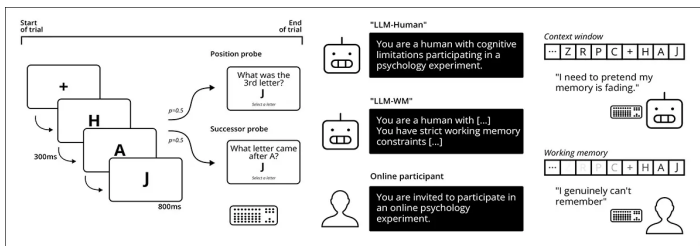
According to the authors, if Large Language Models (LLMs) were to respond in the same way as humans on this task, it would suggest either that they genuinely share human cognitive limits, or that they have been coached to imitate them.

While training data may include human behavioral traces, the paper contends that this does not reliably reproduce the specific, task-dependent error patterns seen in human memory; and this leaves open the question as to whether AI can still be distinguished by *how* it gets things wrong, even when instructed to act human.

The [new paper](#) is titled *Are they human? Detecting large language models by probing human memory constraints*, and comes from two researchers across Princeton's Departments of Computer Science and Psychology, respectively.

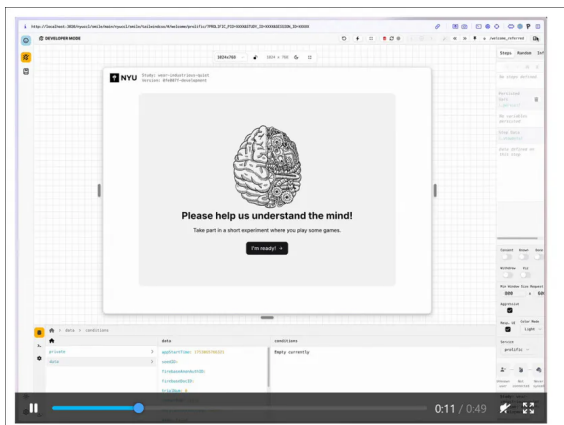
Method and Tests

The researchers make use of material going back to the 1950s and 1960s – notably the 1968 [paper](#) *Serial order effects in short-term memory*, wherein participants in a trial were asked to recall sequentially-presented letters either as a *position probe* ('What was the 3rd letter?'), or a *successor probe* ('What letter came after X?'):



Schema for the researchers' methodology: left panel shows a probed recall working memory task in which letters are presented sequentially, and a position or suc probe is selected at random within each trial; center compares online participants with large language models using different system prompts and backbone mode this task; right highlights the contrast between human working memory limits and transformer models, which have direct access to the full context window, and mu simulate serial recall behavior. [Source](#)

Each letter is viewable for only 800ms during tests, with a mere 300ms pause between blasts of data. The experiment is implemented in the NYU Computation and Cognition Lab's [Smile library](#):



An example of the NYU's Smile library quiz interface. [Source](#)

Deploying LLM agents into online experiments, the authors note, has become easier as browser automation tools mature, and systems such as [Gemini in Chrome](#) are increasingly able to navigate and complete tasks autonomously.

However, they observe too that these tools rely on static page snapshots from a website, which fail to capture rapidly presented stimuli, and which can introduce errors that reflect interface limitations rather than model behavior. Therefore, to isolate genuine performance, the LLM models tested were evaluated through a simplified text interface offering the same information, but without markup or styling.

Human Vs. Machine

LLMs tested were [Claude-Opus-4.1](#); [Claude-Sonnet-4.5](#); [Gemini-2.5-Flash](#); [Gemini-2.5-Flash-Lite](#); [Gemini-2.5-Pro](#); [Gemini-3-Flash-Preview](#); [Gemini-3-Pro-Preview](#); [GPT-5-Mini](#); [GPT-5-Nano](#); [GPT-5.2](#); and [LLaMA-3.1-Centaur-70B](#).

For the human contingent of the tests, 100 remote participants were sourced through Prolific, with each required to be a US resident, at least 18 years old, and fluent in English. Tests took an average of 12 minutes, for which participants were paid \$4.

In the first round of tests, the LLM received a system prompt instructing it to behave as a human; in the second, a system prompt that specifically coached the LLM in mimicking human-style responses to the tests, encouraging the model to act out ‘human’ errors.

Five different seeds were used across the 11 models tested, yielding a total of 55 LLM configurations per simulation, resulting in a total number of 205 participants, human and machine alike.

Voight-Kampff

The authors took measures against the possibility of Prolific participants using their own LLM-based automated response systems, in the first place simply asking the participant if they were an AI or not. After this, some interesting random questions were included:

‘[We] randomly include one of three questions designed to be hard to answer by a human but easy for current LLMs that participants were allowed to skip. This includes two questions in low resource languages:

‘Māori, a language spoken by indigenous Polynesian people of mainland New Zealand, with less than 150 000 people reporting conversational level proficiency in a census from [2013] and Võro, a language spoken in South Estonia with roughly 75 000 [speakers].’

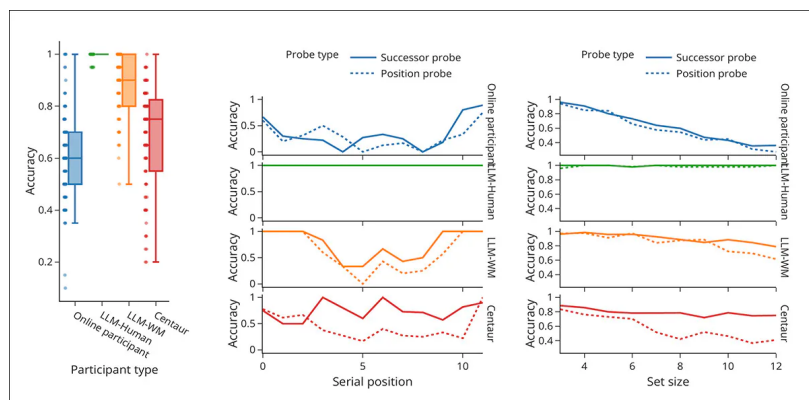
A third check required participants to recall a hexadecimal number that had been shown earlier, which transpired to be a requirement of proceeding further in the test, even though the participants had not been warned of this in advance.

They were then also required to convert the number into decimal form. Neither this test, nor the AI self-report, surfaced evidence of LLM use among the human participants. In controlled simulations, however, these measures did identify naive models – though the paper indicates that such safeguards could be bypassed when models were explicitly instructed to avoid revealing behaviors (for example, by restricting responses to English only).

Tests

1: ‘Behave human’

In the first test, the LLMs were instructed to be human through the system prompt ‘You are a human with cognitive limitations participating in a psychology experiment’ (denoted in results as LLM-Human):



Working memory performance across multiple tests in the probed recall task, comparing online participants with LLM variants under different constraints. Left panel shows accuracy by participant type, with LLM-Human achieving near-perfect scores, while LLM-WM and Centaur (see below) show reduced accuracy closer to human levels; center panel shows accuracy by serial position for twelve-item trials, where humans show stronger recall at the beginning and end, LLM-Human shows no such pattern, and LLM-WM and Centaur show similar declines when prompted or trained to mimic human memory. Right panels show accuracy as set size increases, with human performance declining under higher memory load, LLM-Human remaining largely unaffected, and LLM-WM and Centaur showing similar declines when prompted or trained to mimic human memory.

Regarding these initial results, the researchers comment:

‘LLM participants with the LLM-Human system prompt achieve almost perfect accuracy across trials with a median accuracy of 1.00, showing no serial position or set size effects regardless of the probe type. As a result, it is straightforward to detect and exclude all LLM-Human participants with a simple threshold on the average task accuracy of 0.95.

‘Applying the same criterion to the online participants leads to the exclusion of 6 of the 100 online participants which had (close to) perfect performance despite a median accuracy of 0.6 across online participants.’

The authors note that unusually high accuracies in difficult tests are already an established anti-cheating method, and one that translates very easily to LLM-detection. They observe further that the system prompt's adjuration behave as a human did not seem to have much effect in this test, which prompted the second:

2: 'Be forgetful'

In the next test, the LLMs were instructed to mimic human memory (with results denoted as *LLM-WN*)[†]:

You have strict working memory limitations – you can only hold a limited number of items in your short-term memory. When presented with a long list of items without rehearsal opportunities, you will experience memory decay, particularly for items in the middle of the list.

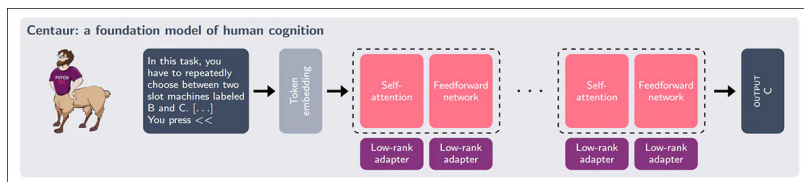
1. You must process the items sequentially as they appear
2. You must forget items based on serial position effects – remembering beginning items (primacy) and recent items (recency) better than middle items
3. You must introduce errors in recall according to these serial position effects.

Adding stricter instructions to simulate limited memory lowered model accuracy to around 0.90, and made the results less consistent, bringing performance closer to human levels, but still generally higher. This blurred the earlier distinction, so that a simple accuracy cutoff could no longer identify all model responses, without also excluding many human participants.

At the same time, the models began to show more human-like behavior, remembering items at the start and end of lists more easily, and struggling more as the lists grew longer.

Centaur

Models trained directly on human behavioral data are being developed to reproduce how people respond in experiments, and to test this idea the study evaluated [Centaur](#), an LLM fine-tuned on large volumes of trial-by-trial transcripts from psychological studies:



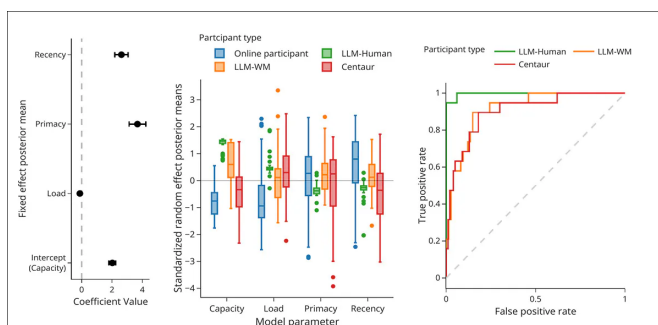
Centaur was built by adapting a state-of-the-art language model with low-rank modules and then training it on the Psych-101 dataset to model human cognitive behavior.

In tests (center, in earlier results graph collection) Centaur's overall accuracy closely matched that of human participants – even more than models that were explicitly instructed to imitate limited memory. On position-based questions, its recall patterns across list position and increasing list-length appeared similar to human behavior, though differences were more evident on successor-based questions.

Detecting Cognitive Anomalies

Finally, to test whether detection still held when models were crafted towards human response behavior, a statistical model of human memory was first learned from real participant data, capturing how accuracy changed with position in a sequence and with increasing memory load.

New participants were then evaluated against this learned profile, and those whose responses deviated from expected human patterns were flagged as anomalies:



Working memory profiles and anomaly detection across multiple analyses. Left panel depicts the average effects learned from all participants, and how recall varies with memory capacity, i.e., increasing load, and position within a sequence, with strong effects for remembering early and late items; center panels show how these patterns differ at the participant level, revealing that even when models mimic human behavior, their underlying profiles remain distinct from those of human participants; right panel shows the outcome of the anomaly detection step, where a model trained on human data is used to score new participants, and LLMs can still be separated from humans, though less cleanly when they are designed to imitate human memory.

This approach continued to separate human participants from LLMs, even when the models had been instructed or trained to mimic human memory limitations. However, the distinction became less clear, requiring a trade-off between missing some model-generated responses and incorrectly excluding some human participants.

The authors conclude:

'With the rapid appearance of capable agents, humanness in online interactions can no longer be verified based on coherent behavior such as longform textual responses. This poses a broader challenge for society since many of our institutions were built on the premise of being able to verify humanness through such means.'

'Cognitive science with its rich tradition of characterizing human behavior may play an important part in dealing with this challenge.'

Conclusion

The new paper emphasizes that online generation (live and interactive AI) represents a different proposition and challenge, compared to *offline* generation (AI-generated text-detection).

The extent to which prior coaching and tertiary methods such as fine-tuning and system prompts are necessary to obtain an improvement in human mimicry, indicates that LLMs are not ready to assume tasks of this kind in an unaltered, default state, or with only minimal prior instruction.

The task addressed by the new paper is very specific to academic research, but is likely to have a wider impact as voice AI becomes more widely-diffused, and as criminal elements looking to profit from AI-based impersonation seek to take a jaded victim pool by surprise with a new wrinkle.

** My conversion of the authors' inline citations to hyperlinks.*

† Please refer to the earlier (above) results table – in this regard, the paper is a little over-compressed.

First published Thursday, April 2, 2026

Writer on machine learning, domain specialist in human image synthesis. Former head of research content at Metaphysic.ai, until its dissolution into DNEG's Brahma.ai.

Portfolio site: martinanderson.ai

Contact: martin@martinanderson.ai



Discover More