

The tells expired. The checklist didn't.

The Economist ran 1.2m words of output from four frontier models against its own prose, against CNN, the *New York Times* and the *Washington Post*, and against sixty years of bestselling fiction. The “how to spot AI writing” checklist still circulating is **half out of date** — and its single most-cited item now points the wrong way.

THE TELL	WHERE THE EVIDENCE ACTUALLY LANDS	STATUS
Em-dashes	Only Claude now exceeds the human rate. ChatGPT uses markedly fewer em-dashes than any other writer in the study , human or machine. Its use spiked in 2025, then fell below the human baseline.	INVERTED
<i>delve, tapestry</i>	The overused vocabulary moved on. These two are no longer characteristic of any model in the study.	EXPIRED
Long, Latinate words eight letters or more	All four models sit well above every human baseline. Gemini and Claude highest . Rare words, scientific register, and nouns made out of verbs.	LIVE
Scientific vocabulary	<i>Parameter, methodology</i> . Every model ranges above <i>The Economist</i> , general news and fiction.	LIVE
“Not X but Y”	ChatGPT and Claude are the outliers . Gemini sits close to human news writing, so the tell identifies a model rather than a machine.	LIVE · UNEVEN
The rule of three	ChatGPT far above everything else , including the other models. The rest cluster just above the human range.	LIVE · UNEVEN
Sparse punctuation	Fewer commas and semicolons, hardly any parentheses , longer sentences. A stronger signal than any banned word, and almost nobody is looking for it.	LIVE · MISSED
No quoted experts	Models do not quote people. On a page of reported prose this is the most structural tell available, and it appears on no checklist.	LIVE · MISSED

THE PART THAT MAKES ALL OF IT TEMPORARY

Every model release moves AI prose closer to human prose. The tells are not being defeated. They are being trained out, release by release, by the same human feedback that made them.

A line graph with a teal line showing a downward trend from 2024 to 2026. A dashed horizontal line is labeled 'MORE AI-LIKE'. The y-axis is labeled '2024' and the x-axis is labeled '2026'.

THREE QUESTIONS BEFORE DETECTION BECOMES A CONTROL

- 01

What does a positive result **authorise**? Name the decision it triggers, not the tool that produced it.
- 02

Who re-validates it **after the next model release**? A named person with a date, not a policy review cycle.
- 03

What happens to the people it flags **wrongly**? Detectors are black boxes. They give no reason for the call.

Every item on that checklist was true when someone wrote it down. Nothing on it has been re-checked since.

Source: *The Economist*, “How to spot AI writing”, 30 July 2026 — 55,940 sentences and 1.2m words of output from ChatGPT-5.6 Terra, Claude Sonnet 5, Gemini 3.5 Flash and Grok 4.5, benchmarked against *The Economist*, an average of the *New York Times*, *Washington Post* and CNN (2018–2022), and NYT bestselling fiction (1950–2022). Models were prompted without web access. Detection vendor Pangram claims 99.98% accuracy; *The Economist* notes detectors are black-box and can return false positives without giving reasons. Status labels above are this sheet’s reading of the published findings, not the newspaper’s framing.

FOLLOW FREDRIK LINDSTROM FOR MORE

Fredrik Lindstrom

AI and cybersecurity governance for boards, CISOs and senior leaders. 25+ years in enterprise security.

A practitioner to guide you.