

Industrializing Humanization

Originally published Thursday 17th September 2026 at <https://www.unite.ai/industrializing-humanization/>



The gig postings and service offerings at casual work outlets such as Upwork and Fiverr heavily feature the word *'humanize'* these days, with a range of jobs and providers proposing various levels of obfuscation of the fact that an article was written, or partly written, or co-written, with AI.



The *'Humanize'* trend across real posts at Upwork.

The demand for human and automated services that can rewrite AI-'polluted' output so that it seems a human sat down and sweated for a few hours, suggests that a universal consensus has been reached that the 'human touch' is needed for connection and outreach – and that the sheer volume of content in the AI age [will make that connection essential again](#).

I would personally like to think so, but it isn't necessarily so – not least since Google's own guidelines on the matter [treat of the quality of the text](#), rather than the means by which it was created.

The truth is, we don't currently know exactly *how much or what kind of* AI we are willing to tolerate, or the extent to which we can rationalize the benefits that billions of us get from AI every day against a feeling of 'complicity' and moral seduction in regard to the technology's wider implications.

The aforementioned job listings reflect this too; there, creators and companies are slaving for automation, presumably because they want to publish at scale, which is exactly the practice that Google promises to specifically punish in its ranking algorithm. At the same time, they are seeking both automated and human-based methods of 'industrializing' humanization.

It reminds me somewhat of [one of the best jokes](#) in Steven Wright's neurodivergent 1980s comedy routines:

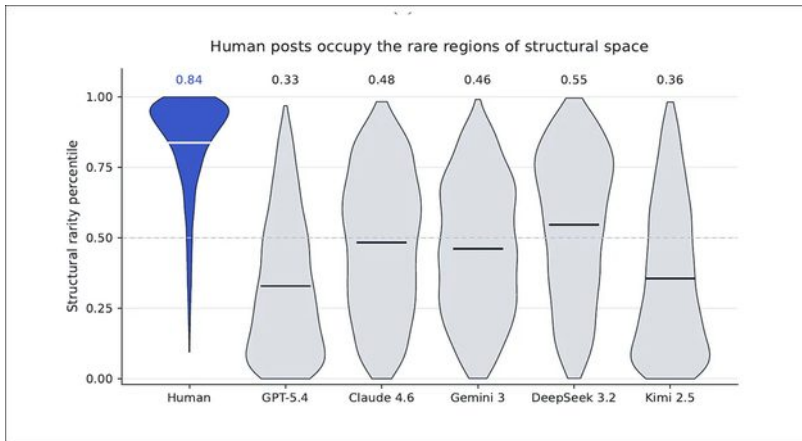
'In my house, I have a microwave fireplace. I can lay down in front of the fire for the evening in just 8 minutes.'

The Shape of Slop to Come

As usual, these thoughts are inspired by my daily trawl through Arxiv, where today I especially lighted on a [paper](#) by Jochen Madler at the marketing automation company Sitefire, titled *SlopShape: Identifying AI-Generated Commercial Web Content*.

I normally discount papers with sole authors, on the 'Victor Frankenstein' principle, but the premise of this one is Spartan, and the testing methodology robust.

The study is essentially saying that when an AI-generated post is completely rewritten, even by humans, it retains a distinct architectural signature that resides at a level much higher than word-level, comprising a characteristic 'shape' that can't be eradicated by any amount of rephrasing or surface-level amendment of the text:



Comparison of structural patterns in human and AI-written posts. Human writing tends to use rarer combinations of structures, while all five AI models favor more common structures. [Source](#)

Obviously one can get around this issue (if you consider it an issue) by actually planning the course of the writing yourself, and either using no AI at all, or using it only on the separate components that you planned yourself.

However, as the rabid ads at Upwork attest, this is not what's wanted; one can infer from the tone of such postings that article content should preferably not only be structured by AI, but *ideated* by AI – that is to say, that an LLM would suggest suitable topics aligned with a company or entity's overarching goals.

After this, it's fully expected that the AI will produce various progressive drafts of the content, with the annoying meatware component kicking in only for a closing and casual gloss of humanity (though ideally, this process would eventually evolve into a Python workflow).

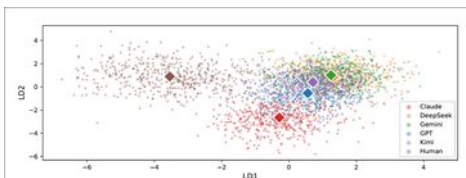
Striking Features

The author of the new paper obtained these insights by collecting a few thousand company blog posts immediately predating the release and diffusion of ChatGPT (2020-2022) and inferring prompts from them.

That is to say, he used [Gemini 3 Flash](#) to read the real, human-written papers and reverse-engineer a prompt designed to produce a similar work across some of the leading LLMs, namely [GPT-5.4](#), [Kimi 2.5](#), [DeepSeek V3.2](#), [Claude 4.6](#), and [Gemini 3](#).

These prompts were then used to create five AI versions of each of the real human posts, from which [features](#) were extracted which would eventually describe the LLM-specific 'shape' of content across the five models; and Madler notes that 'AI posts share a tidy, self-announcing shape'.

The entire study is a re-run of the [StoryScope project](#) from COLM 2026 this August, a collaboration between the University of Maryland and Google DeepMind, which applied essentially the same methodology to fiction, finding that AI-generated stories converge on a narrower and more predictable range of narrative structures than human writing:



From the StoryScope study, a scatter-plot comparison of narrative patterns in human and AI-generated fiction.

Human-written stories occupy a distinct region, while stories from five leading AI models cluster more closely together, indicating greater similarity in their underlying narrative choices. [Source](#)

Madler decided to re-run StoryScope on company blogs because of the extensive financial impetus to automate these, compared to the more eccentric field of fiction-writing, and the great extent to which they are already automated (50% automation of online content by May of 2026, according to [a Graphite study](#) cited in Madler's work).

Perhaps the most interesting analysis is the picking apart of features extracted from the AI-produced material based on the pre-2022 human-written posts:

'The values most commonly produced by AI models assemble into what we call the tidy, self-announcing blog post: it promises the payoff already in the title, states its thesis and announces its structure before the first section, speaks in an editorial-explainer voice, and closes with a stage that summarizes or restates the thesis.'

'In a typical AI post, the payoff is promised in the title: "How to Cut Onboarding Time in Half." The thesis is stated and the flow announced before the first section: "In this post, we'll cover why onboarding stalls, three fixes that work, and how to measure the difference."

'And the close restates the thesis: "In short, structured onboarding saves time." The human-leaning values describe posts without those signposts: no announced section flow, no escalation of stakes, and no closing section that restates the thesis.'

A Faulty Premise

Incidentally, most of the tropes and cliches present in the described template/s were derived from [SEO-obsessed strictures](#) of the pre-LLM marketing scene. Where do we think that AI learned these habits?

Templated marketing copy was [always hated](#), long before AI, and always produced at scale, [long before machines could do it well](#); and because of the requirement to conform to SEO practices, some horrific conventions were established, such as [keyword placement](#), and [crafted design of links](#).

In the latter case, AI may therefore have learned the wrong lesson, since the text-only corpora on which it was trained stripped out so many of the SEO-facing tactics that were originally included in the source material, such as strategic diagrams, visual explainers; and most especially hyperlinks.

Entire sections of an article may have been crafted to frame a hyperlink or a YouTube embed that a client wanted to push; but by the time the data was trained into an LLM, the embed or the link was stripped out, and the text designed to highlight it shorn of its original context.

Therefore, arguably, the great error in designing training methodologies around such material, was treating web content as if it was 19thC literature (entirely self-reliant) instead of multimedia (significantly affected by non-text factors such as video, hyperlinks and illustrations).

The other great error was training on material [that people already hated](#), and which was tormented into its unpopular shape by rules that [probably aren't valid any longer](#).

Attention, Humans

Nonetheless, if you can believe the Upwork trends, there's a significant demand for more of the same, produced at scale by LLMs, and 'humanized' through the least idiosyncratic process available, at the highest speed, and at the lowest cost available.

All this, to produce the effect of *attention*: the idea that someone like you (i.e., in possession of an endocrine system, at the very least) sat down for a few hours and thought about what they wanted to say to you.

The perception of receiving human attention is meaningful to us [even when it is not helpful](#); therefore it's certainly worth industrializing.

I can't help but feel that Madler's paper is in some way likely to make the situation worse, since it identifies characteristic higher-dimension traits that can now, presumably, be addressed with *more* AI, so that the default signatures will assume a more human shape as the models evolve.

First published Thursday, September 17, 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