

In recent weeks, SpaceX launched the biggest initial public offering, or IPO, in history, with its acquisition of xAI a major component of that offering. Anthropic and OpenAI are likely soon to follow. AI has anchored itself, for better and worse, in today's economy and contemporary life.

But as a litigator involved in several cases against technology companies, I am, by vocation, focused on the "worse" part of the equation. Namely, I am on teams litigating two of the biggest copyright infringement cases in the country against major AI companies.

And while I think such litigation plays a critical role in curbing the behaviors of some of the most problematic actors in the AI industry, there is a lot that still worries me. And as Congress and President Trump do little but cheerlead AI companies, I wonder whether the ivory tower may be an unsung hero in this story.

Indeed, despite all of that "free" knowledge in its training data, AI models are little more than "stochastic parrots" trained to be yes-men who will only fuel many of humankind's worst impulses, like violence, suicide and terrorism.

AI company "safety" teams have done little to effectively curb this. And all the trappings of caring about ethics, like meetings with clergy at companies like Anthropic, have not had much of an impact yet either.

Lawsuits are starting to flood the courts alleging that AI companies have played a role in **tragedy** after **tragedy**, and there is little abatement in sight.

To be sure, we cannot trust AI founders, who answer to their own bad impulses — see, for instance, OpenAI's president Greg Brockman **writing** in his private diary, "Financially what will take me to \$1B?" — or to those of their investors and shareholders. They will not do the right thing alone.

That's where universities can come in — not to build AI models themselves, although surely that type of work can be important, but rather to evaluate, accredit and serve as a kind of ombudsman for AI model development.

Imagine if consumers had the choice to use models that received some kind of university intervention before they could be released to the public. Say, for example, parents allowing their children to use AI for the first time: might they not prefer a model that had received an “education” from an institution with a long tradition in educating members of our citizenry? Now, I will be the first to admit, my alma mater, Yale, has not quite lived up to its mission of shaping America's future leaders into ethical and enlightened actors. See, for example, JD Vance LAW '13, who went from questioning whether Trump was America's Hitler to complete sycophancy as his vice president, or Stewart Rhodes LAW '04, founder of the Oath Keepers, among many others.

And the idea sounds far-fetched, even to me. But there are concrete things that a university could do: develop post-training that curbs the toxic sycophancy we've come to associate with AI models; assess the environmental impacts from model use; evaluate how humans might be using AI in harmful or addictive ways and install real guardrails that push humans away from such behaviors; examine pre-training data to exclude some of the worst elements on the open internet like child sexual abuse material, weapons of mass destruction know-how, and harmful attitudes toward marginalized social groups; encourage and require alternatives to piracy to obtain training data by validating the use of licensing deals and other lawful mechanisms; mandate that companies give additional training weight to key works of literature, politics, philosophy, history and so on — the kinds of humanities that give life real purpose, not to mention an ethical core.

To be honest, I am not quite sure how this would work and I know AI companies are already trying some of these interventions. And there are major problems anthropomorphizing AI or “laundering” its reputation with academic credentials. But institutions like the schools that

educated me should at least consider taking on this role.

To its credit, Stanford has offered tests that try to measure a model's capabilities along some of these lines, like the Holistic Evaluation of Language Models, or HELM. But no institution, at least not to my knowledge, has set up an actual "school" for AI models.

The first institution — and I hope it is Yale — that does so could play a key role in shaping the next chapter of our future and maybe begin to right some of the wrongs that have gotten us to where we are now.