

Policing and Generative AI is Immoral

Douglas C. Youvan

March 20, 2023

The most dangerous and subversive aspect of AI that I've seen so far is the inclusion of flagged social media posts in the AI training set which ChatGPT actually acknowledges. It makes sense that flagging posts as community standards violations was originally done by humans at places like Facebook. Those people would have been largely atheist and WOKE. At some point, FB's policing standards were handed over to their AI that was trained in these "morals", outside of any law or concept of freedom of speech. Now, we see the same morals in the FB AI police and the generative capacities of ChatGPT. The AI art program Dall.E also has these standards. Microsoft Power Apps (e.g., AI-assisted MS Word) are being released and this time, and they too will have such morals. Hence, our searches, writing, and art will be set by the original judgment calls of humans at tech companies – now set in Silicon, not stone. Expect this logic in social credit points, too.