

Can a Preventative Social Media UI Break "Fake News?"

Beam, Ryan (School: Scotts Valley High School)

In the wake of the 2016 Presidential Election, Social Media giants such as Facebook and Twitter have implemented a variety of changes to their respective sites' user interfaces, in an effort to crack down on the spread of "fake news" via their platforms. Both have claimed that preventative in-product features are enough to keep deliberately misleading information out of users' feeds, but their failure to release any accompanying data, combined with worrying observations from lawmakers and academics alike, has left many wondering whether there really are any feasible methods for preventing the spread of deliberately misleading information on social media. The project examined how people (examining a mock social media newsfeed) interact with "fake news" Three User Interfaces were tested, one which served as a control, one which implemented Facebook's "Warning Icon" feature, and one which implemented Twitter's "Sensitive Content" feature, which requires users to acknowledge a disclaimer before viewing potentially misleading information. The hypothesis predicts that UI 1 will see rampant sharing of "fake news", UI 2 will significantly cut down on the sharing of "fake news" among the unaffiliated, but fail to overcome confirmation bias, and UI 3 will most effectively overcome all obstacles, keeping a person from sharing potentially misleading information whether it aligns with their ideology or not. The project confirms the hypothesis, although leads to an unexpected conclusion: The "Warning Icon" feature on UI 2, while failing to breach so-called "echo chambers" like UI 3, consistently kept misleading information from going viral as it did on UI 1, and did so just as effectively as UI 3.

Awards Won:

American Psychological Association: Certificate of Honorable Mention