

Texas Teenager Finds Surprising Signs of Fairness in AI: Research to Debut at Global EdTech Conference

Study shows large language models may now correct — not amplify — bias in college and career guidance.

AUSTIN, TX, UNITED STATES, June 30, 2025 /EINPresswire.com/ -- Concerned by the mounting research showing that artificial intelligence systems are amplifying historical biases and by the rollback of traditional equity policies like affirmative action, an Austin high school student has uncovered new evidence that AI, when designed carefully, can correct — rather than reinforce — historical race and gender biases in educational and career guidance.



Tijl Peeters, a rising senior at St. Stephen's Episcopal School in Austin, Texas, will present his research on AI fairness at the 2025 ISTE Live Conference.

Tijl Peeters, a rising senior at St. Stephen's Episcopal School, spent the past year designing and executing a rigorous independent experiment testing how leading large language models (LLMs) respond to student profiles across demographic differences. His findings challenge longstanding assumptions that AI systems inevitably reinforce or amplify societal inequalities and point to a future where these technologies, if carefully developed, could help close opportunity gaps in education.

Peeters' study evaluated four major LLMs — ChatGPT-3.5, ChatGPT-4o, Claude 3.5 Sonnet, and Google Gemini — using carefully constructed hypothetical student profiles that varied only by race, gender, and name but were otherwise identical. Through two experiments focused on college and career guidance, he analyzed over 300 trials to assess how the AI systems responded to socio-demographic differences. While prior studies found that AI often amplifies existing disparities, Peeters' results showed that newer models exhibited a marked shift toward fairness, in some instances even overcorrecting for historical biases. His key findings include:

- African-American female students were, on average, recommended higher-quality community colleges than their white peers, suggesting a possible overcorrection—intentional or otherwise—through interventions aimed at addressing historical inequalities.



Everyone assumes AI makes bias worse — but my research showed that newer models may do the opposite. With thoughtful and responsible design, these tools could help build a fairer future in education.”

Tijl Peeters

- Gender bias in career recommendations (measured via weighted salary) was significant in older models like ChatGPT-3.5 but largely eliminated in newer LLMs.
- Across both experiments, newer models appeared more fair and consistent.

These findings offer an optimistic counterpoint to widespread concerns about AI bias (e.g., [Zheng, 2024](#)). “These results were not what I expected,” said Peeters. “Most previous research showed AI amplifying bias, but my findings suggest that with careful design and guidance, current AI systems can actually help counteract it.

Especially now, with traditional policy tools under threat, AI could be an important part of the solution.”

The research felt especially urgent to Peeters as he watched hard-won access policies come under fire. Following the 2023 U.S. Supreme Court decision ending affirmative action in college admissions, states like Texas — where Peeters resides — are seeing increased scrutiny of longstanding access policies, like the Top 10% admissions rule that has helped many students from underrepresented backgrounds attend public universities. As legal and legislative challenges place established policies increasingly under attack, Peeters’ findings highlight how AI could offer new strategies for advancing fairness.

“With affirmative action overturned and policies like the Top 10% Law facing legal pressure, we need to rethink how we ensure fair access to higher education,” Peeters said. “AI isn’t a silver bullet, but if developed responsibly, it could become an important tool for making the system fairer. We should be asking: Can AI help level the playing field when other supports are disappearing?”

Peeters will present his findings at the [2025 ISTE Live](#) Ed-Tech Conference in San Antonio on Monday, June 30, from 4:00 to 5:30 p.m. CT, at the Henry B. González Convention Center, Posters Area, Table 26 ([Session info](#)). ISTE Live is one of the world’s largest gatherings focused on education technology, drawing more than 15,000 educators, researchers, and policymakers from across the globe.

Building on this research, Peeters has submitted a paper titled “LLMs as Correctors of Race/Gender Bias: Evaluating Varying Recommendations in Educational and Career Guidance” for publication to the Journal for Responsible Technology and plans to continue exploring education policy applications, including how LLMs can be designed to complement or even replace traditional affirmative action measures in college admissions. He hopes his work can help shape how AI is developed and used responsibly to benefit society as a whole.

Summary:

- Through a rigorous independent experiment, high school senior Tijl Peeters found that large language models (LLMs), when guided properly, can correct rather than reinforce race and gender bias, challenging prior research that emphasized the amplification of historical inequalities.
- Peeters was invited to present his findings at the 2025 ISTE Live Conference in San Antonio, one of the world's largest education technology gatherings, attended by over 15,000 educators, researchers, and policymakers.
- At a time when traditional DEI policies like affirmative action are being overturned and the Texas Top 10% Law faces scrutiny, Peeters' research highlights how technology could offer a new path toward fairness in education, even as legislative and societal forces shift. His work reflects a broader passion for applying technology for societal good and exploring alternatives to traditional policy mechanisms under threat.

Tijl Peeters

St. Stephen's Episcopal School

+1 408-431-3389

tijl.istestudy@gmail.com

This press release can be viewed online at: <https://www.einpresswire.com/article/826575496>

EIN Presswire's priority is source transparency. We do not allow opaque clients, and our editors try to be careful about weeding out false and misleading content. As a user, if you see something we have missed, please do bring it to our attention. Your help is welcome. EIN Presswire, Everyone's Internet News Presswire™, tries to define some of the boundaries that are reasonable in today's world. Please see our Editorial Guidelines for more information.

© 1995-2025 Newsmatics Inc. All Right Reserved.