



Proceedings of the Informing Science + Information Technology Education Conference

An Official Publication
of the Informing Science Institute
InformingScience.org

InformingScience.org/Publications

July 26 – 31, 2026

RE-THINKING FAIRNESS IN AI SCREENING: FILTERING BACKGROUND BIAS BEFORE TEACHING THE MODEL [ABSTRACT]

Hsing-Tzu Cathy Lin*	National University of Kaohsiung, Information Management, Kaohsiung, Taiwan (R.O.C.)	cathy@go.nuk.edu.tw
Jung-Yu Lily Wu	National Tsing Hua University, De- partment of foreign languages and literature, Kaohsiung, Taiwan (R.O.C.)	imlilywu@gapp.nthu.edu.tw

* Corresponding author

ABSTRACT

Aim/Purpose	In educational retention and early warning systems, AI models often misinterpret family economic advantage as individual learning potential due to data limitations. This study aims to identify and block hidden socioeconomic signals (e.g., micro-economic stressors) to propose an AI ethics design that reflects diverse values.
Background	Seemingly neutral competency signals (e.g., early academic adaptation metrics, such as first-semester grades) often entail high resource costs and serve as proxies for social class. The White House warns that without calibration, algorithms will systematically exclude disadvantaged students during their academic journey.
Methodology	This research conducts an empirical analysis using the UCI Student Dropout and Academic Success dataset. We identify proxy parameters tied to economic thresholds and simulate a Leveling the Starting Line ethical screening workflow to test its efficacy in mitigating bias.
Contribution	Drawing on socioeconomic theory, this study proposes a series of simulations that feature a resilience weighting for institutional and academic factors.

Accepted by Editor Michael Jones | Received: April 15, 2026 | Revised: May 26, 2026 |

Accepted: June 8, 2026.

Cite as: Lin, H.-T. C., & Wu, J.-Y. L. (2026). Re-thinking fairness in AI screening: Filtering background bias before teaching the model [Abstract]. In M. Jones (Ed.), *Proceedings of InSITE 2026: Informing Science and Information Technology Education Conference, USA*, Article 7. <https://doi.org/10.28945/5831>

(CC BY-NC 4.0) This article is licensed to you under a [Creative Commons Attribution-NonCommercial 4.0 International License](https://creativecommons.org/licenses/by-nc/4.0/). When you copy and redistribute this paper in full or in part, you need to provide proper attribution to it to ensure that others can later locate this work (and to ensure that others do not accuse you of plagiarism). You may (and we encourage you to) adapt, remix, transform, and build upon the material for any non-commercial purposes. This license does not permit you to use this material for commercial purposes.

	This deepens the understanding of algorithmic bias and provides a roadmap for translating ethical principles into technical standards.
Findings	Findings anticipate that without weight adjustments for resource-related features (such as financial distress and enrollment-age barriers), models will reinforce social rigidity. Results show that ethical filtering can effectively reduce evaluation gaps caused by resource inequality, returning talent selection and educational support to its core essence.
Recommendations for Practitioners	We recommend that schools and institutions create a bias checklist before deploying early warning AI. To ensure tech teams understand which economic hardships, such as personal debt or unpaid tuition, might be misleading, and to block or recalibrate them during the design phase.
Recommendations for Researchers	Researchers should shift from static pre-entry bias metrics to dynamic process-fairness audits, investigating how macroeconomic indicators (e.g., family income and parents' occupation) structurally interact with longitudinal student performance data.
Impact on Society	By preventing early warning systems from institutionalizing a punitive poverty penalty, this ethical framework safeguards educational equity and ensures that high-stakes algorithmic decisions enhance socioeconomic mobility rather than reinforce historical disparities.
Future Research	Future research can apply this filtering mechanism to areas such as scholarships, resource allocation, and post-graduation hiring, helping AI continuously learn to ignore new types of structural and macroeconomic bias that change over time, ensuring ethical designs remain current.
Keywords	educational fairness, bias filtering, socioeconomic status (SES), AI ethics, early warning systems

AUTHORS



Cathy S. Lin is an Associate Professor of Department of Information Management at National University of Kaohsiung, Taiwan. She received her Ph.D. in Information Management from National Sun Yat-sen University, Taiwan. Her research focuses on information ethics, artificial intelligence ethics, educational data and analytics, and electronic commerce. Her research work has been published in academic journals, including *Information & Management*, *Journal of Business Ethics*, *Journal of Retailing and Consumer Services*, *Journal of Electronic Commerce Research*, *Online Information Review*, among others.



Lily Jung-Yu Wu, a student in Foreign Languages and Literature at National Tsing Hua University, Taiwan, blends technological adaptability with humanistic inquiry. Her leadership as a Model UN chairperson holding discussions on data ethics and fairness. Extending this commitment to social equity, she is dedicated to multicultural advocacy, particularly for Taiwanese indigenous communities. This passion is shown in her recent research presentation, which advocates for removing ethnic and socioeconomic biases in university admissions to foster inclusivity across both digital and societal realms.