



4th Interdisciplinary Research Conference Proceedings

“TECHNOLOGY AND HUMANITY: EXPLORING THE IMPACT OF DATA-DRIVEN TECHNOLOGIES ON HUMAN DEVELOPMENT”



Madison, NJ / United States

October 10, 2025

Edited by

Khaled Mohammed Saifuddin, Ph.D. [k.saifuddin@siu.edu]

Preface

The 4th Society of North American Scholars (SNAS) Interdisciplinary Research Conference Proceedings highlights the high-quality research and intellectual contributions of our attendees. The SNAS Board initiated this publication as a testament to the organization's commitment to academic excellence and as a platform to showcase the valuable work presented at our conference.

The 4th Annual Interdisciplinary Research Conference, held on October 10, 2025, at Fairleigh Dickson University in Madison, New Jersey focused on *Technology and Humanity: Exploring the Impact of Data-driven Technologies on Human Development*. This theme is timely and critical as AI technologies continue shaping various sectors, notably higher education. The conference provided an essential forum for scholars, educators, administrators, and technology experts to examine how AI transforms teaching, learning, research, and institutional operations.

The Organizing Committee invited participants to submit 3-to 5-page summaries of their presentations to establish a formal conference record. These summaries underwent an editorial review. Each contribution reflects the growing interaction between technology and humanity and the innovative approaches being developed to harness this intersection.

For permission requests or access to full papers, [please contact the respective authors via the email addresses provided](#).

Cite this volume (APA style):

Author, A. A. (2025). Title of paper. In K. M. Saifuddin (Eds.), Proceedings of the 4th Interdisciplinary Research Conference (pp. xx-xx). Society of North American Scholars.

TABLE OF CONTENTS

1. **AI and Ethics in Healthcare: Predicting Cancer with EHR Data and Addressing Equity in Model Development**
Pages 5–11
2. **AI, Linguistic Justice, And Health Equity: Bridging Social Determinants Of Education Gaps in Diverse Communities**
Pages 12-18
3. **Designing AI-Supported Interventions to Strengthen Pre-service Mathematics Teachers' Content Knowledge**
Pages 19–23
4. **A Machine Learning Approach to Predicting Environmental Performance and Carbon Emissions: An Examination of the Major U.S. Pharmaceutical Firms**
Pages 24-30
5. **AI-Driven Network Security Monitoring: Leveraging Open-Source Tools and Affordable AI Alternatives**
Pages 31-39
6. **Enhancing the Effective Use of Artificial Intelligence in Higher Education: A Case Study with Practices and Recommendations**
Pages 40-48
7. **Exploring AI for Teaching and Learning: A Practical Instructor's Perspective**
Pages 49-55
8. **Exploring The Barriers and Enablers of Ai-Driven Energy Efficiency in SMEs**
Pages 56-67

9. **Revealing the Unknown: LLM-based Forensics Triage Agent for Android Mobile Forensics**
Pages 68-72
10. **Strategy Writing Evaluation with Large Language Models in Introductory Physics**
Pages 73-77
11. **Understanding Student and Faculty Perceptions of ChatGPT and the Acceptance of Large Language Models in Higher Education**
Pages 78-84
12. **Emotional Harms of Artificial Intelligence on Human Psychology and Relations**
Pages 85-94
13. **Predictive Performance of LSTM vs GRU on Google Stock Prices**
Pages 95-108
14. **The Relative Importance of Healthcare in Individuals' Perception of Quality of Life: A Fuzzy Pairwise Comparison Assessment**
Pages 108-115

AI AND ETHICS IN HEALTHCARE: PREDICTING CANCER WITH EHR DATA AND ADDRESSING EQUITY IN MODEL DEVELOPMENT

Zeynep Akcay Ozkan, Ph.D. (Queensborough Community College, City University of New York)- zakcayozkan@qcc.cuny.edu

Abstract: *Cancer is one of the leading causes of mortality worldwide, and early detection remains critical for improving patient outcomes. Advances in artificial intelligence (AI) and the availability of large-scale electronic health record (EHR) datasets offer powerful new opportunities for predictive modeling. However, the rapid adoption of complex AI systems, particularly large language models (LLMs), raises important ethical questions related to equity, transparency, and accessibility in healthcare.*

Using the All of Us Research Program, we developed machine learning and deep learning models to predict pancreatic cancer occurrence from longitudinal EHR data. Our models demonstrate that relatively lightweight approaches can yield meaningful predictive performance, while remaining more transparent and reproducible than computationally intensive LLMs. We provide a comparative discussion of model accuracy, interpretability, and feasibility, highlighting challenges associated with advanced models, such as the non-disclosure of parameters due to privacy concerns and the high resource requirements that limit widespread use.

This work underscores the ethical and practical challenges of integrating AI into healthcare. Disparities in institutional capacity mean that only well-funded centers can deploy state-of-the-art LLMs, potentially widening existing healthcare inequities. By contrast, accessible and transparent models can promote broader adoption and trust in AI-assisted care. We argue that responsible deployment of AI in healthcare must balance innovation with fairness, equity, and patient trust to ensure that advancements in diagnostics and personalized medicine benefit all populations.

Introduction

Pancreatic cancer has the lowest 5-year survival rate among major cancers, making early detection crucial for improving outcomes [1], [2]. Timely identification of high-risk patients could lead to earlier interventions, yet existing clinical screening approaches remain limited. Electronic Health Records (EHRs) offer a rich, longitudinal view of patients' health, creating new opportunities to leverage artificial intelligence (AI) for predictive modeling.

A growing body of research has demonstrated associations between certain clinical features, such as elevated HbA1c levels, increased pancreatic cancer risk [3], [4]. Developing predictive models capable of identifying these patterns early could transform clinical care. However, the adoption of AI in healthcare raises key concerns related to equity, data accessibility, and model reproducibility.

This study explores the use of deep learning methods to predict pancreatic cancer onset using EHR data, with particular emphasis on the challenges of model transparency, privacy, and equitable access to AI technologies.

Literature & Motivation

Machine learning has shown significant promise in clinical prediction [5], but most classical models (e.g., logistic regression, k-nearest neighbors) operate on static cross-sectional data. In contrast, modern healthcare data are longitudinal, capturing health trajectories over time through lab results, diagnoses, vital signs, and clinical notes.

Neural network architectures, such as convolutional neural networks (CNN), recurrent neural networks (RNN), and large language model (LLM)-based transformers, enable modeling of these temporal dynamics. Prior studies [6], [7] have demonstrated that previous visit sequences can be used to effectively predict disease occurrence in the next visit. Similarly, sentiment analysis in natural language processing provides a useful conceptual parallel, where trends and contextual signals are modeled over time rather than through isolated points. Despite these advances, barriers remain: lack of open model parameters, limited access to large

and diverse datasets, and unequal computational resources restrict broad participation in AI innovation.

Methods

We used data from the All of Us Research Program [8], a national dataset designed to include diverse and longitudinal EHR data from across the United States. The dataset contains: patient demographics (age, gender, race, ethnicity), medical history (family history, conditions), medications and prescriptions, lab and diagnostic test results, vital signs and measurements, lifestyle surveys (e.g., smoking, alcohol use) and unstructured clinical notes.

Patients included in the study were required to have at least three recorded medical conditions. Individuals who were 18 years old or younger, as well as those 89 years old or older, were excluded from the cohort. A total of 686 pancreatic cancer cases were matched to 686 control patients based on age, gender, and race to ensure comparability between groups.

The modeling approach in this study was informed by the work of Rasmy et al. [9], and their algorithms were applied to the pancreatic cancer cohort. Specifically, several recurrent neural network architectures were implemented to capture temporal patterns in longitudinal EHR data. These included a Vanilla RNN cell with a tanh activation function, Long Short-Term Memory (LSTM) networks, Gated Recurrent Units (GRU), and the RETAIN (Reverse Time Attention) model [10]. The RETAIN model was further explored with standard, bidirectional, and dilated connections to enhance its ability to capture complex temporal dependencies.

Null ICD codes, duplicate diagnostic entries and all records occurring after the pancreatic cancer diagnosis (ICD code C25) were excluded to ensure that only pre-diagnosis data were used for model training and prediction. In total, the dataset included 220,143 medical condition records for the pancreatic cancer cases and 232,205 condition records for the control group.

Results / Preliminary Findings

The RETAIN model with LSTM cell type achieved the best performance among the architectures tested with an AUROC of 0.828. This aligns with prior findings that

attention-based models are well suited for healthcare prediction tasks because they provide a degree of transparency in identifying key visits and codes.

To address data imbalance, an equal number of pancreatic cancer and control patients were included, which reduced the overall training set and may have limited model performance. The model also did not exclude patient records from the last three or six months prior to diagnosis and relied primarily on diagnostic codes rather than incorporating lab results or vital signs. Future work could apply alternative methods to handle class imbalance and expand the feature set, which may further enhance predictive accuracy and generalizability.

Overall, these findings illustrate that predictive modeling of pancreatic cancer using longitudinal EHR data is promising despite the inherent challenges of data limitations, class imbalance, and modeling complex temporal relationships.

Discussion & Implications

This work highlights several key challenges in applying advanced AI methods to healthcare data. Although our initial goal was to use existing pre-trained models from prior studies and build upon them, we were unable to obtain their trained parameters. Access to these model parameters was not possible, which limited reproducibility and model comparison. We also planned to experiment with transformer-based architectures [6]; however, these models required specific system configurations that could not be implemented within the All of Us Researcher Workbench environment. The All of Us support teams were extremely helpful throughout the research process and provided clear guidance, but they confirmed that such system-level changes were not feasible due to infrastructure constraints.

More broadly, this experience reflects the wider barriers to reproducibility, transparency, and equitable access in healthcare AI. Strict privacy regulations such as HIPAA and GDPR, while essential, restrict large-scale data sharing across institutions. As a result, health data remain fragmented across institutional silos, preventing the development of globally shared, high-performing medical AI models. Furthermore, because many advanced architectures and pre-trained parameters are closed-source, only large and well-resourced institutions can replicate or extend

state-of-the-art methods. These realities underscore that while predictive modeling with longitudinal EHRs holds great promise, realizing that potential requires balancing innovation with reproducibility, accessibility, and ethical stewardship of patient data.

Conclusion

Predicting pancreatic cancer using longitudinal EHR data is promising but complex. Sophisticated models like RETAIN and transformer-based architectures can detect subtle trends and temporal patterns, offering opportunities for earlier intervention. Yet, without addressing barriers of data privacy, computational accessibility, and model transparency, AI systems risk reinforcing existing inequities in healthcare. The path forward requires balancing accuracy, interpretability, and equity through strategies such as federated learning, which trains models across institutions without sharing patient data, privacy-preserving techniques, and investment in computational infrastructure for smaller institutions.

References

- [1] National Health Service England, "Cancer Survival in England, cancers diagnosed 2016 to 2020, followed up to 2021," 2023.
- [2] R. L. Siegel, T. B. Kratzer, A. N. Giaquinto, H. Sung, and A. Jemal, "Cancer statistics, 2025," *CA. Cancer J. Clin.*, vol. 75, no. 1, pp. 10–45, Jan. 2025, doi: 10.3322/caac.21871.
- [3] D. McDonnell *et al.*, "Elevated Glycated Haemoglobin (HbA1c) Is Associated with an Increased Risk of Pancreatic Ductal Adenocarcinoma: A UK Biobank Cohort Study," *Cancers (Basel)*, vol. 15, no. 16, p. 4078, Aug. 2023, doi: 10.3390/cancers15164078.
- [4] S. Grewe *et al.*, "Elevated HbA1c Levels Are Associated with a Risk of Pancreatic Cancer: A Case–Control Study," *J. Clin. Med.*, vol. 13, no. 18, p. 5584, Sep. 2024, doi: 10.3390/jcm13185584.
- [5] M. A. Ahmed, A. AbdelMoety, and A. M. A. Soliman, "Predicting cancer risk using machine learning on lifestyle and genetic data," *Sci. Rep.*, vol. 15, no. 1, p. 30458, Aug. 2025, doi: 10.1038/s41598-025-15656-8.
- [6] Z. Yang, A. Mitra, W. Liu, D. Berlowitz, and H. Yu, "TransformEHR: transformer-based encoder-decoder generative model to enhance prediction of disease outcomes using electronic health records," *Nat. Commun.*, vol. 14, no. 1, p. 7857, Nov. 2023, doi: 10.1038/s41467-023-43715-z.
- [7] D. Placido *et al.*, "A deep learning algorithm to predict risk of pancreatic cancer from disease trajectories," *Nat. Med.*, vol. 29, no. 5, pp. 1113–1122, May 2023, doi: 10.1038/s41591-023-02332-5.
- [8] National Institutes of Health, "All of Us Research Program." <https://allofus.nih.gov/>.
- [9] L. Rasmy *et al.*, "Simple Recurrent Neural Networks is all we need for clinical events predictions using EHR data," 2021. [Online]. Available: <https://doi.org/10.48550/arXiv.2110.00998>.

- [10] E. Choi, M. Bahadori, J. Kulas, A. Schuetz, W. Stewart, and J. Sun, "RETAIN: An Interpretable Predictive Model for Healthcare using Reverse Time Attention Mechanism," 2017.

AI, LINGUISTIC JUSTICE, AND HEALTH EQUITY: BRIDGING SOCIAL DETERMINANTS OF EDUCATION GAPS IN DIVERSE COMMUNITIES

Madjiguene Fall, Ph.D. (Kean University) – madjiguene.fall@kean.edu

Navya Kollapally, Ph.D. (Kean University) – navya.kollapally@kean.edu

Abstract: *Research Motivation. The Social Determinants of Education (SDoED) (Kollapally, Geller, et al., 2024) encompass a broad, yet incomplete list of factors such as socioeconomic status, community resources, and cultural influences that explain learning disparities and other factors contributing to academic gaps, such as student emotional support needs and outcomes. Ontologies (Ahmad & Gillam, 2005), when paired with large language models (LLMs) (Xia et al., 2025), create a powerful tool that captures complex interrelations between SDoED factors. This project's objective was to expand an original SDoED Ontology framework developed using a scoping review of research articles that failed to incorporate concepts related to non-mainstream and indigenous populations; proven to lack data inputs that reflect diverse linguistic and cultural perspectives and identities; thus, limiting its inclusivity and applicability.*

Key Contributions: To address this gap, this study at the intersection of education, linguistics, and AI constructs a multimodal framework that extracts multimodal concepts grounded in minoritized epistemologies. This approach incorporates community-specific SDoED knowledge that integrates real-time emotion recognition to curate emotionally supportive responses. The key multimodal data in the original SDoED framework – databases, educational websites, and questionnaires – will be refined using sociolinguistics/linguistic anthropology procedures (cultural contextualization, data balancing, community-centered corpora annotating, and ethical filtering) and various computer science protocols through a retrieval-augmented generation (RAG) process. This innovative workflow will generate accurate, context-aware, inclusive, and trustworthy data output to enhance the credibility of Large Language Model (LLM) generated responses. The

resulting SDoEd framework will center on linguistic and cultural appropriateness and the accurate depiction of the educational and well-being of students.

Social Implications. This work connects researchers, community translators, and culture experts as co-researchers through building trust and the creation of a community-based research pod, in which local stakeholders are considered a source of knowledge.

Keywords: *AI and education; tech justice; digital language justice; digital health equity*

Theoretical Framework

This mixed-methods study integrates qualitative synthesis with computational ontology engineering. Specifically, it expanded an existing Social Determinants of Education (SDoED) Ontology through a combination of scoping review (Tricco et al., 2018), participatory engagement with underrepresented communities, and computational validation using large language models (LLMs). This approach ensured both theoretical rigor and practical inclusivity in capturing complex determinants of educational equity.

Methods

Phase 1: Scoping Review and Gap Identification

With PRISMA guidance (Tricco et al., 2018), additional databases in education, sociology, psychology, anthropology, and indigenous studies are searched systematically. Inclusion criteria were broadened to non-English publications and research on linguistic minority and culturally diverse populations. Extracted data were inductively coded to identify emergent determinants of emotional well-being, non-verbal communication, and community-based practices. Results were mapped on to the baseline ontology to identify conceptual and relational gaps.

Phase 2: Stakeholder and Community Engagement

To validate and extend the ontology beyond mainstream models, participatory design methods (Nguyen, 2025) were employed. Expert panels of teachers, policymakers, and researchers in indigenous and marginalized education were consulted. Participatory mapping exercises enabled the drawing out of culturally grounded determinants and relational understandings not usually captured in formal research. These exercises were conducted with iterative feedback loops for inclusivity and authenticity.

Phase 3: Ontology Expansion and Structuring

We formalized the novel concepts and relationships with Protégé and OWL/RDF standards. The extended ontology contained sub-domains of socioeconomic status, cultural identity, family support, emotional well-being, and classroom engagement,

with constructs prominent for indigenous and non-mainstream populations made explicit (Fall, 2023). The ontology's consistency was checked by automated pitfall scanners and expert review.

Phase 4: Large Language Model Integration

To operationalize the ontology, we integrated it with LLMs for automatic identification and classification of SDoED factors in unstructured text (e.g., policy documents, ethnographic literature, educational texts). Ontology-guided prompts were developed to circumscribe LLM outputs within the structured ontology. Comparative experiments were conducted to compare the accuracy of LLM-only and ontology-aware LLM outputs. Special care was taken to audit for cultural bias, misrepresentation, and erasure of marginalized voices (Vindigni, 2025).

Phase 5: Refinement and Evaluation

Evaluation proceeded on both technical and practical levels (Kollapally et al., 2021; Kollapally, Keloth, et al., 2024; Kramer & Beißbarth, 2017). Technical performance was measured in terms of ontology coverage, semantic accuracy, and recall/precision of classification on annotated corpora. Practical evaluation proceeded through case studies with educators and policymakers, who tested the usefulness of the ontology for the diagnosis of barriers to learning and informing targeted interventions. Feedback from both strands was used in the iterative refinement of the ontology and LLM integration.

Significance and Key Contributions

To address this gap, this study at the intersection of education, linguistics, and AI constructed a multimodal framework that extracted multimodal concepts grounded in minoritized epistemologies. This approach incorporated community-specific SDoED knowledge integrating real-time emotion recognition to curate emotionally supportive responses. The key multimodal data in the original SDoED framework – databases, educational websites, and questionnaires – were refined using sociolinguistics/linguistic anthropology procedures (cultural contextualization, data balancing, community-centered corpora annotating, and ethical filtering) (Broesch, et al., 2024) and various computer science protocols through a retrieval-augmented

generation (RAG) process (Abouenour et al., 2014). This innovative workflow generated accurate, context-aware, inclusive, and trustworthy data output to enhance the credibility of Large Language Model (LLM) generated responses. The resulting SDoED framework centered on linguistic and cultural appropriateness (Jones, Satran, & Satyanarayan, 2024) and the accurate depiction of the educational and well-being of students.

Social Implications

This work connected researchers, community translators, and culture experts as co-researchers through building trust and the creation of a community-based research pod, in which local stakeholders were considered a source of knowledge. At the university community level, innovative partnerships between the Computer Science and Education faculty, industry leaders, pluralistic cultural and language communities, and student researchers. In addition, there is potential to reduce educational disparities through the use of native languages and the involvement of multilingual community experts as translators and co-researchers. Furthermore, policymakers and stakeholders are better equipped to analyze the diverse factors contributing to academic gaps across communities, enabling more informed decision-making and targeted interventions.

References

Abouenour, L., Nasri, M., Bouzoubaa, K., Kabbaj, A., & Rosso, P. (2014). Construction of an ontology for intelligent Arabic QA systems leveraging the Conceptual Graphs representation. *J. Intell. Fuzzy Syst.*, 27(6), 2869–2881.

Ahmad, K., & Gillam, L. (2005). Automatic Ontology Extraction From Unstructured Texts. 1330-1346. https://doi.org/10.1007/11575801_25

Broesch, T., Crittenden, A., Beheim, B., Blackwell, A., Bunce, J., Colleran, H., Hagel, K., Kline, M., Mcelreath, R., Nelson, R., Pisor, A., Prall, S., Pretelli, I., Purzycki, B., Quinn, E., Ross, C., Scelza, B., Starkweather, K., Stieglitz, J., & Mulder, M. (2020). Navigating cross-cultural research: methodological and ethical considerations. *Proceedings of the Royal Society B: Biological Sciences*, 287. <https://doi.org/10.1098/rspb.2020.1245>

Fall, M. (2023). “Trans~Resistance”: Translingual Literacies as Resistance to Epistemic Racism and Raciolinguistic Discourses in Schools. *Societies*. <https://doi.org/10.3390/soc13080190>.

Jones, G., Satran, S., & Satyanarayan, A. (2024). Toward Cultural Interpretability: A Linguistic Anthropological Framework for Describing and Evaluating Large Language Models (LLMs). *ArXiv*, abs/2411.05200. <https://doi.org/10.48550/arXiv.2411.05200>

Kollapally, N. M., Chen, Y., & Geller, J. (2021). Health Ontology for Minority Equity (HOME). *KEOD*,

Kollapally, N. M., Geller, J., Morreale, P., & Kwak, D. (2024). An Ontology for Social Determinants of Education (SDoEd) Based on Human-AI Collaborative Approach. *J. Comput. Sci. Coll.*, 40(3), 191–203.

Kollapally, N. M., Keloth, V. K., Xu, J., & Geller, J. (2024). Integrating commercial and social determinants of health: A unified ontology for non-clinical determinants of health. *AMIA Annual Symposium Proceedings*,

Kramer, F., & Beißbarth, T. (2017). Working with Ontologies. *Methods Mol Biol*, 1525, 123-135. https://doi.org/10.1007/978-1-4939-6622-6_6

Nguyen, D. (2025). Collaborative Growth: When Large Language Models Meet Sociolinguistics. *Lang. Linguistics Compass*, 19. <https://doi.org/10.1111/lnc3.70010>.

Tricco, A. C., Lillie, E., Zarin, W., O'Brien, K. K., Colquhoun, H., Levac, D., Moher, D., Peters, M. D. J., Horsley, T., Weeks, L., Hempel, S., Akl, E. A., Chang, C., McGowan, J., Stewart, L., Hartling, L., Aldcroft, A., Wilson, M. G., Garritty, C., . . . Straus, S. E. (2018). PRISMA Extension for Scoping Reviews (PRISMA-ScR): Checklist and Explanation. *Ann Intern Med*, 169(7), 467-473. <https://doi.org/10.7326/m18-0850>

Vindigni, G. (2025). Gender Bias and Cultural Misrepresentation in AI: A Critical Inquiry into Cross-Cultural Communication and Algorithmic Design. *European Journal of Applied Science, Engineering and Technology*.
[https://doi.org/10.59324/ejaset.2025.3\(3\).06](https://doi.org/10.59324/ejaset.2025.3(3).06).

Xia, Y., Zhou, J., Shi, Z., Chen, J., & Huang, H. (2025). Improving Retrieval Augmented Language Model with Self-Reasoning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 39(24), 25534-25542. <https://doi.org/10.1609/aaai.v39i24.34743>

DESIGNING AI-SUPPORTED INTERVENTIONS TO STRENGTHEN PRE-SERVICE MATHEMATICS TEACHERS' CONTENT KNOWLEDGE

Kimberly Sirin Budak, Ph.D. (University of Wisconsin) - sbudak@uwsp.edu

Sait Suer, M.S. (Kennesaw State University) - ssuer@students.kennesaw.edu

Abstract: *This article outlines the design, implementation, and anticipated outcomes of a sabbatical project that integrates artificial intelligence (AI) into mathematics teacher education. Addressing persistent gaps between undergraduate mathematics coursework and the content knowledge required for secondary-level teaching, the project proposes individualized AI-supported study plans for pre-service mathematics teachers. Tools such as Copilot, Khan Academy's Khanmigo, and ChatGPT will be leveraged to provide adaptive feedback, real-time explanations, and targeted support for misconceptions. The project will be piloted between 2026 and 2027, with pre- and post-assessment data serving as measures of effectiveness. This article presents the project rationale, objectives, planned activities, and the anticipated benefits for students, faculty, and the broader mathematics education community.*

Designing AI-Supported Interventions to Strengthen Pre-Service Mathematics Teachers' Content Knowledge

Introduction and Background

Strengthening pre-service teachers' mathematical content knowledge remains a critical challenge in mathematics education (National Council of Teachers of Mathematics [NCTM], 2014; Boaler, 2016). A recurring concern expressed by cooperating teachers and college supervisors is that student teachers often struggle to bridge the gap between abstract university-level mathematics and the secondary-level topics they are expected to teach. Addressing this disconnect is essential, as content knowledge forms the foundation of effective instruction and directly impacts student learning.

This project proposes the integration of artificial intelligence (AI) into teacher preparation as a means of providing individualized, adaptive support for pre-service mathematics teachers. Tools such as Microsoft Copilot, Khan Academy's Khanmigo, and OpenAI's ChatGPT offer immediate feedback, interactive explanations, and visual representations that can enhance conceptual understanding (Khan Academy, 2023; OpenAI, 2023). By embedding these tools into personalized study plans, the project seeks to develop a sustainable system that strengthens mathematical content knowledge in ways not easily achieved through traditional instruction alone.

The central research question guiding this project is: To what extent do pre-service mathematics teachers demonstrate measurable improvement in mathematical content knowledge from pre-test to post-test after participating in AI-supported study plans?

Two exploratory questions provide further insight: (1) In what ways do students engage with AI tools to address misconceptions and support their learning? (2) Which areas of mathematical content knowledge show the most noticeable gains from the intervention?

The overall thesis is that thoughtfully designed, AI-supported study plans can bridge the persistent gap between university coursework and secondary mathematics content, equipping pre-service teachers with both confidence and competence in their future teaching practice.

Objectives

The primary objective of this sabbatical project is to design, implement, and evaluate an AI-supported intervention that strengthens pre-service mathematics teachers' content knowledge (Holmes et al., 2019; Luckin et al., 2016). The first step will involve designing the system and piloting it with a small group of junior-level mathematics education students in Spring 2026. This pilot will allow for the administration of a baseline assessment aligned with Praxis objectives, the identification of common areas of weakness, and the opportunity to refine the system's design. Weekly check-ins, exit questions, AI prompt logs, and surveys will provide insights into how students interact with the intervention. These data, combined with post-test results, will guide revisions to ensure the system remains responsive and effective. Consultations with mathematics education faculty, cooperating teachers, and professional networks will strengthen the intervention, while technical support from a computer science doctoral student will enable the integration of AI tools into a unified platform.

Methodology and Planned Activities

The methodology of this project is structured across four phases spanning 2026 to 2027. In Spring 2026, a pilot study will be launched with two to three voluntary junior-level mathematics education majors. During this phase, a Praxis-aligned baseline assessment will be administered, and qualitative data will be collected on how students interact with AI tools such as Copilot, Khanmigo, and ChatGPT (Khan Academy, 2023; OpenAI, 2023). The findings from this pilot will inform refinements prior to full-scale implementation. In Fall 2026, the full implementation phase will begin, where a broader cohort of students will complete the Praxis Sample Test to establish baseline data. Following this, students will receive training on effective AI usage, with particular emphasis on prompt-writing strategies for problem solving, conceptual understanding, and visualization (Holmes et al., 2019). Individualized study plans will be developed based on the assessment results, and weekly meetings with exit questions will be conducted to monitor progress and address challenges. The Spring 2027 phase will focus on ongoing support and post-

assessment. Students will continue working with their personalized study plans while providing survey feedback and usage data on their AI engagement. The same Praxis Sample Test will be re-administered as a post-test, allowing for direct comparison with baseline results. The final phase, conducted in Summer 2027, will synthesize the collected data and disseminate the results. This will involve preparing manuscripts for peer-reviewed publication, presenting findings at conferences, and developing a prototype Canvas-based seminar course that formalizes the intervention. Collaboration with the computer science doctoral student will also support the integration of the AI tools into a single consolidated platform for future use.

Anticipated Results

Although the study will be conducted in the future, several benefits are anticipated. For students, the intervention is expected to strengthen mathematical content knowledge through individualized study plans, weekly guided reflections, and structured training in AI usage (Boaler, 2016; NCTM, 2014) . Such engagement will likely enhance their ability to bridge advanced mathematics with secondary-level teaching, while also fostering professional growth through technological literacy. Faculty colleagues and teacher educators will benefit from access to a tested model for integrating AI into teacher preparation programs (Luckin et al., 2016) . Data on misconceptions and AI engagement patterns will provide valuable insights for refining curricula and teaching practices. Finally, the broader academic community will benefit from the dissemination of results through journal publications and conference presentations (Holmes et al., 2019). By sharing both empirical findings and design strategies, this project aims to contribute to national conversations on the responsible and effective integration of AI in mathematics teacher education.

References

Bernardi, M. L., Capone, R., Faggiano, E., & Rocha, H. (2025). Generative AI in mathematics education: Pre-service teachers' knowledge and implications for their professional development. *International Journal of Mathematical Education in Science and Technology*, 56(8), 1513–1530.
<https://doi.org/10.1080/0020739X.2025.2490104>

Boaler, J. (2016). *Mathematical mindsets: Unleashing students' potential through creative math, inspiring messages, and innovative teaching*. Jossey-Bass.

Copur-Gencturk, Y., et al. (2024). Improving teaching at scale: Can AI be incorporated into professional development to create equitable learning opportunities? [PDF].

Holmes, W., Bialik, M., & Fadel, C. (2019). *Artificial intelligence in education: Promises and implications for teaching and learning*. Center for Curriculum Redesign.

Khan Academy. (2023). Khanmigo: AI-powered tutoring and teaching assistant.
<https://www.khanacademy.org/khan-labs>

Luckin, R., Holmes, W., Griffiths, M., & Forcier, L. B. (2016). *Intelligence unleashed: An argument for AI in education*. Pearson Education.

National Council of Teachers of Mathematics. (2014). *Principles to actions: Ensuring mathematical success for all*. NCTM.

OpenAI. (2023). ChatGPT: Optimizing language models for dialogue.
<https://openai.com/blog/chatgpt>

A MACHINE LEARNING APPROACH TO PREDICTING ENVIRONMENTAL PERFORMANCE AND CARBON EMISSIONS: AN EXAMINATION OF THE MAJOR U.S. PHARMACEUTICAL FIRMS

Murad Ozbilgin (Pioneer Academy) – mozbilgin@gmail.com

Abstract: *This study examines financial and governance features of leading U.S. pharmaceutical firms in predicting environmental performance and CO₂ emissions using machine learning (ML) models. Among the ML models, CatBoost, XGBoost, and Random Forest perform best with R^2 reaching around 90% for environmental performance. Key positive predictors include Corporate Social Responsibility (CSR) board committee presence, dividends relative to sales, sales revenue, selling, general, and administrative (SG&A) expense relative to sales, and the percentage of women directors. Core earnings and R&D are less important and are negative predictors of environmental performance. For CO₂ emissions, only CatBoost and Random Forest perform moderately well with R^2 reaching around 70%. SG&A expense, board size, and board independence are positive predictors, while debt ratio, core earnings, and dividends are negative predictors.*

Introduction

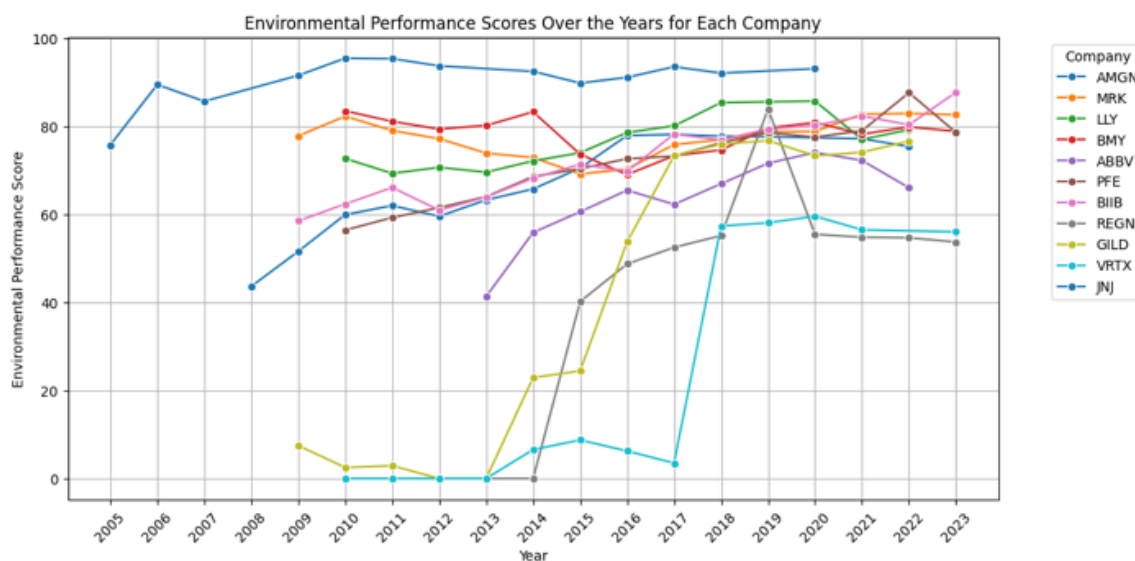
The U.S. pharmaceutical industry is one of the largest globally and has contributed to treatments such as gene therapies and immunotherapy-based oncology. Its growth also brings environmental challenges because manufacturing is energy-intensive and generates waste. Firms have responded with successful sustainability strategies as data shows improved environmental responsibility in recent years. This study uses ML models to examine the predictors of environmental performance and CO₂ emissions for the 11 leading U.S. pharmaceutical firms each with revenues above \$10 billion in 2023. I test the hypotheses that effective boards (gender diversity, independence, CSR committee presence, small size, CEO-Chairman separation) and financial strength (higher revenue, EBITDA, ROA; lower SG&A expense, R&D expense, debt, and capital expenditures) improve environmental outcomes.

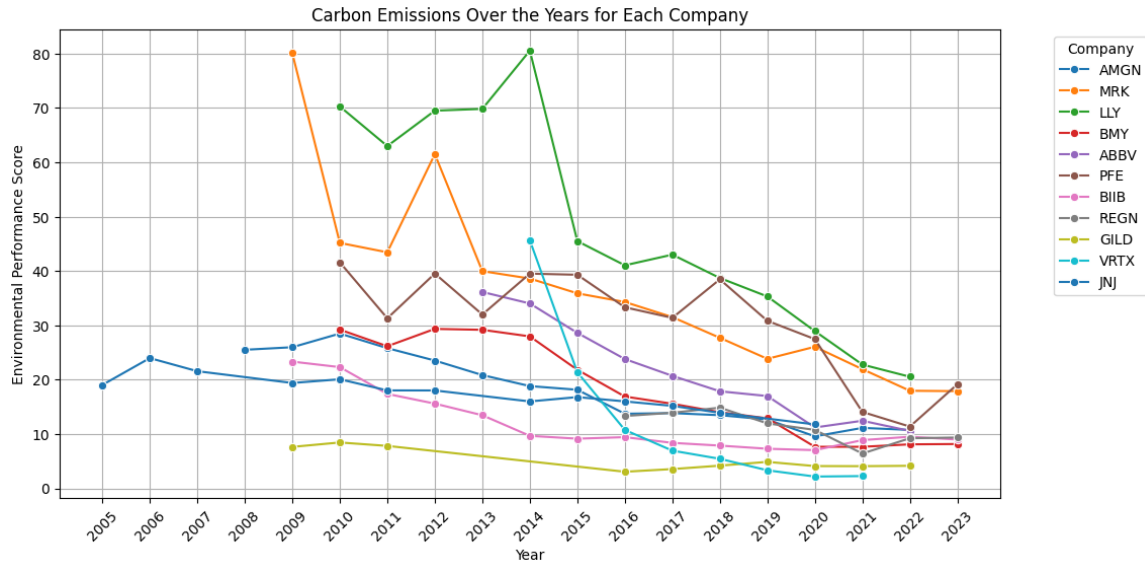
Prior Literature and Contribution

Most ESG studies examine how environmental performance affects financial outcomes (e.g., Garcia and Orsato, 2020; Cortez et al., 2022). In contrast, I focus on predicting environmental performance itself. Nguyen et al. (2021) find board size and meeting frequency matter in Chinese polluting industries, while my study identifies CSR committee presence and percentage of women on board as key predictors in U.S. pharmaceutical firms, reflecting the different regulatory environments. Garcia Martin et al. (2020) report similar results for EU firms. My study confirms their results with a different sample and methodology, while also showing that financial characteristics play an important role alongside corporate governance characteristics in shaping environmental conduct. Studies on pharmaceutical sustainability (Demir and Min, 2019; Booth et al., 2023) rely on firm reports to assess environmental performance, whereas I use LSEG-Refinitiv's composite environmental scores that encapsulate 186 most comparable and material company-level measures to assess a company's overall environmental performance, commitment, and effectiveness (Refinitiv, 2024).

Data and Methods

The dataset covers the major U.S. pharmaceutical firms including Eli Lilly, Johnson & Johnson, Merck, AbbVie, Amgen, Biogen, Pfizer, Vertex, Regeneron, Gilead, and Bristol-Myers Squibb between 2005 and 2023. Environmental performance is measured using the Refinitiv Environmental Pillar Score (0–100), which aggregates emissions, resource use, and environmental indicators from publicly disclosed data. I gathered the CO₂ emissions along with firm governance data from LSEG Refinitiv and the financial data from Yahoo Finance. Five machine learning models are trained, and the best performers are used to identify feature importance. For the years for which environmental performance scores are missing, I used imputed median scores for the relevant firms. The following graphs show the improving environmental performance and declining carbon emissions across the board suggesting that U.S. big pharma have taken serious steps to address environmental concerns. My study examines the firm characteristics that are positively and negatively associated with these improvements.





Discussion And Conclusions

The following table shows the performance of the ML models in predicting the environmental performance and carbon emissions of the sample firms. The predictive performances of CatBoost and XGBoost are outstanding reaching almost an R^2 of 90% for environmental performance. While overfitting concerns arise with such high performance, using an 80%-20% train-test split in the analysis serves to alleviate those concerns. ML model performances with the studied firm features are not as remarkable for carbon emissions implying that there might be more fundamental drivers of carbon emissions such as manufacturing technology and energy choices.

ML Model Prediction Performance (with 80%-20% train-test split)

Environmental Performance				Carbon Emissions			
ML Model		RMSE	R^2	ML Model		RMSE	R^2
Models used for feature importance	CatBoost	8.17	89.97%	Models used for feature importance	CatBoost	5.86	72.96%
	XGBoost	8.31	89.62%		Random Forest	6.45	67.27%
	Random Forest	11.68	79.47%		XGBoost	8.15	47.69%
	Lasso Regression	14.77	67.18%		Lasso Regression	8.96	36.74%
	Decision Tree	15.55	63.61%		Decision Tree	10.57	12.05%

Among the three ML models chosen for environmental performance and two ML models chosen for carbon emissions, the firm financial and governance features with the highest feature importance scores are summarized in the following table.

Feature Importance - Summary

Environmental Performance			Carbon Emission		
Feature	Selecting ML Model	Direction of Influence	Feature	Selecting ML Model	Direction of Influence
CSR Committee	CatBoost XGBoost Random Forest	Positive	SellingGenAdmin/Sales	CatBoost Random Forest	Positive
Dividends/Sales	CatBoost XGBoost Random Forest	Positive	Debt Ratio	CatBoost Random Forest	Negative
Women on Board	CatBoost Random Forest	Positive	Board Size	CatBoost Random Forest	Positive
SellingGenAdmin/Sales	CatBoost Random Forest	Positive	Board Independence	CatBoost Random Forest	Positive
Ln_Sales	CatBoost Random Forest	Positive	Core Earnings/Sales	CatBoost Random Forest	Negative
Core earnings/Sales	CatBoost	Negative	Dividends/Sales	CatBoost Random Forest	Negative
R&D/Sales	Random Forest	Negative	(Direction of influence is obtained from the Lasso model)		

The results show that governance features, especially CSR committee presence and female directors, along with financial features such as revenue, dividends, and SG&A expense are positively linked, while core earnings and R&D intensity are negatively linked with environmental performance. For CO2 emissions, the models have less predictive power. SG&A expense with board size and independence show positive association, while debt ratio, core earnings, and dividends show negative association with CO2 emissions. These findings are largely consistent with my hypotheses stated in the introduction. From a financial perspective, these firms ought to make sure to earmark funds for environmental responsibility while from a governance perspective they ought to pay particular attention to the formation of their CSR board committees, increasing the gender diversity of their boards, and making their boards more compact in terms of head count. This study thus highlights the joint influence of governance and financial features on environmental outcomes in U.S. pharmaceutical firms. Future research could extend the analysis to non-U.S. firms and other industries with significant environmental impact.

Acknowledgements

I would like to thank the LSEG Academia Group for allowing me to use the Refinitiv environmental performance ratings and carbon emissions data.

References

- Booth, A., Jager, A., Faulkner, S. D., Winchester, C. C., & Shaw, S. E. (2023). Pharmaceutical company targets and strategies to address climate change: content analysis of public reports from 20 pharmaceutical companies. *International journal of environmental research and public health*, 20(4), 3206.
- Cortez, M. C., Andrade, N., & Silva, F. (2022). The environmental and financial performance of green energy investments: European evidence. *Ecological economics*, 197, 107427.
- Demir, M., & Min, M. K. (2021). A Comparative Analysis of Sustainability Reports Published by Global Enterprises: Standalone CSR Versus Integrated Reporting. *The BRC Academy Journal of Business*, 11(1), 23-51.
- Garcia, A. S., & Orsato, R. J. (2020). Testing the institutional difference hypothesis: A study about environmental, social, governance, and financial performance. *Business Strategy and the Environment*, 29(8), 3261-3272.
- García Martín, C. J., & Herrero, B. (2020). Do board characteristics affect environmental performance? A study of EU firms. *Corporate Social Responsibility and Environmental Management*, 27(1), 74-94.
- Ha, N. M., Nguyen, P. A., Luan, N. V., & Tam, N. M. (2024). Impact of green innovation on environmental performance and financial performance. *Environment, Development and Sustainability*, 26(7), 17083-17104.
- Martiny, A., Taglialatela, J., Testa, F., & Iraldo, F. (2024). Determinants of environmental social and governance (ESG) performance: A systematic literature review. *Journal of Cleaner Production*, 456, 142213.
- Nguyen, T. H., Elmagrhi, M. H., Ntim, C. G., & Wu, Y. (2021). Environmental performance, sustainability, governance and financial performance: Evidence from heavily polluting industries in China. *Business Strategy and the Environment*, 30(5), 2313-2331.

Refinitiv (2024). https://www.lseg.com/content/dam/dataanalytics/en_us/documents/methodology/lseg-esg-scores-methodology.pdf

Wamba, L. D. (2022). The determinants of environmental performance and its effect on the financial performance of European-listed companies. *Journal of General Management*, 47(2), 97-110.

AI-DRIVEN NETWORK SECURITY MONITORING: LEVERAGING OPEN-SOURCE TOOLS AND AFFORDABLE AI ALTERNATIVES

Azamat Zhamanov¹ - azhamanov@na.edu

Fuat Kurucan² - fcun@na.edu

Bermet Alibekova³ - balibekova@na.edu

Adem Kakhadze⁴ - akakhadze@na.edu

Ali Kucuk⁵ - ali@na.edu

Fatih Ancin⁶ - frank@na.edu

1-6 North American University, Houston, TX, USA

Abstract: Effective cybersecurity monitoring often relies on costly SIEM systems and certified analysts, leaving smaller organizations under protected. Open-source tools such as Snort provide affordable alternatives but require expertise that many institutions lack. This study investigates whether artificial intelligence (AI) can bridge this gap by enabling novices to perform effective log analysis. Two Snort logs—one case of suspicious traffic and one misconfiguration—were analyzed by three groups: an AI-assisted novice, a certified analyst using a commercial SIEM, and an unassisted help desk employee. Their assessments were compared with the system administrator's validated outcomes. Results show that the AI-assisted novice consistently aligned most closely with the administrator in both cases, correctly distinguishing between benign and suspicious events. The AI-assisted novice outperformed the unassisted employee and, at times, matched or exceeded the certified analyst. These findings suggest that AI can democratize cybersecurity monitoring for smaller organizations by elevating novices to near-expert performance. Professional oversight, however, remains essential for complex or long-term security decisions.

Keywords: Artificial Intelligence (AI) in Cybersecurity; Intrusion Detection Systems (IDS); Snort Log Analysis; Novice vs. Expert Analysts; Open-Source Security Tools; SIEM Alternatives

Introduction

As cyber threats continue to evolve in scale and sophistication, effective network security monitoring remains one of the most critical defenses for organizations. Large enterprises often rely on commercial Security Information and Event Management (SIEM) systems and certified security professionals to monitor, detect, and respond to suspicious activities. Although these systems provide comprehensive functionality, they come with significant financial and operational costs that small businesses often cannot afford due to their high price. Moreover, their effective use requires substantial expertise in interpreting complex security logs and taking appropriate actions, which further limits their accessibility to organizations with limited technical resources.

Literature Review

Artificial intelligence (AI) and machine learning (ML) have made significant advances in intrusion detection systems (IDS), where supervised and deep learning models—such as CNNs, RNNs, autoencoders, and hybrid systems—enhance anomaly detection and precision in identifying evolving threats [1]–[3]. Despite these gains, challenges persist: handling unstructured log data and balancing detection rate with false positives remain key hurdles.

Large language models (LLMs) offer promising advances by leveraging contextual understanding. Recent work benchmarks models like DistilRoBERTa and GPT variants, demonstrating strong performance in log classification tasks when adapted with domain-specific tuning [4]. Broad reviews of LLM applications in cybersecurity underscore their potential for tasks such as log parsing, threat summarization, and alert triage, though these studies mainly target expert workflows and automated pipelines [5]. Practical implementations, such as Boffa’s LogPrécis, illustrate how LLM-enhanced pipelines can transform raw log data into actionable insights, thereby reducing analyst workload [6]. Explainability also remains a critical factor. Surveys and frameworks for explainable intrusion detection (X-IDS, XAI-IDS) emphasize that interpretability is essential for building trust and enabling less experienced staff to meaningfully use AI outputs [7], [8].

Parallel to the development of AI tools, researchers have also compared the performance and cost-effectiveness of open-source and commercial SIEM solutions. Manzoor et al. [9] demonstrated that open-source SIEMs, when properly configured, can achieve comparable performance to commercial platforms while remaining affordable for small-to-medium enterprises. Similarly, Bezas and Filippidou [10] presented a comparative analysis of open-source SIEM architectures, highlighting their technical strengths and limitations. Vazão et al. [11] extended this discussion by evaluating an open-source SIEM configured for GDPR compliance, underscoring the feasibility of deploying affordable solutions in regulated environments. Hase [12] provided a systematic review of SIEM selection criteria, noting that evaluation processes vary depending on the expertise of the analyst—an insight that connects directly to the role of novices versus experts in security monitoring.

Taken together, the literature confirms that both advanced AI tools and open-source SIEMs can significantly improve cybersecurity monitoring while reducing costs compared to proprietary enterprise platforms. Yet, none of these works address how AI may specifically empower novice users to perform analyses at a level comparable to certified professionals or how unassisted staff perform in comparison. This gap motivates the present study, which investigates whether AI assistance can enable novices to analyze Snort logs as effectively as experts, thereby contributing to the democratization of cybersecurity monitoring.

Research Questions

- 1. How do the three groups—an AI-assisted novice analyst, a certified cybersecurity analyst using an enterprise SIEM tool, and an unassisted help desk employee—differ in their ability to analyze Snort logs for accuracy, completeness, and efficiency?*
- 2. How effective is AI assistance compared to professional expertise and no assistance in detecting threats, reducing false positives, and interpreting network activity?*
- 3. Can AI assistance make cybersecurity monitoring tasks easier for people with limited training, reducing the need for costly SIEM systems and specialized experts?*

Research Method

Data Collection

Two Snort log samples were selected from a real organizational network as the basis for analysis. To preserve confidentiality, all IP addresses in the logs were anonymized before distribution. The logs represented two different types of network activity: 1). PNG Log – multiple alerts related to the download of unusually large PNG image files, classified by Snort as Attempted User Privilege Gain (Priority 1). 2). TFTP Log – repeated TFTP Get requests from an internal host to the broadcast address, classified as Potentially Bad Traffic (Priority 2).

Participants

The logs were provided to three different categories of analysts, each representing a distinct perspective: 1). Office (AI-assisted novice): A recent computer science graduate with no prior cybersecurity experience, using AI tools (ChatGPT and Claude) to assist with analysis. 2). Support (Experienced analyst): A certified cybersecurity analyst familiar with using enterprise-grade SIEM tools. 3). Helpdesk (Unassisted novice): A help desk employee with no cybersecurity training and no AI assistance.

Evaluation Procedure

Each participant independently analyzed both logs and was asked to: Assess the perceived threat level (low, medium, high), provide reasoning for their assessment, recommend specific actions to be taken (e.g., isolate host, block IP, monitor traffic). Their analyses were compared against the system administrator's ground-truth assessment, which served as the benchmark.

Findings

The experiment with two Snort logs (PNG and TFTP) highlighted clear performance differences among the three analyst categories:

1. Office (AI-assisted novice) demonstrated the closest alignment with administrator conclusions in both cases, identifying the TFTP event as a benign misconfiguration

and recommending containment for the PNG case consistent with the administrator's actions.

2. Support (experienced analyst) provided thorough, investigation-driven recommendations but was less precise about the likely cause of each event, reflecting a cautious but slower approach.

3. Helpdesk (unassisted novice) tended to overstate risks, suggesting aggressive containment for benign activity, which could cause unnecessary operational disruption.

Overall, AI assistance improved novice performance to a level comparable with, and sometimes more aligned than, professional expertise.

Log	Administrator Outcome	Office (AI-assisted novice)	Support (Experienced analyst)	Helpdesk (Unassisted novice)	Closer Match
TFTP	Non-malicious. Caused by switch misconfiguration. Service disabled, traffic stopped.	Threat level: Low. Reasoning: Likely PXE boot or misconfigured device. Action: Monitor, verify device, block TFTP if unused. Match: High – aligned exactly with benign cause.	Threat level: Unclear. Reasoning: Could be malware or benign. Action: Investigate host, scan for malware, monitor network. Match: Moderate – valid but	Threat level: Medium–High. Reasoning: Considered scanning or malware. Action: Isolate device, scan for malware, block TFTP. Match: Low – overstated risk, unnecessary isolation.	Office

			overly cautious.		
PNG	Suspicious but ultimately non-malicious. No compromise found. External IP blocked, monitoring continued.	Threat level: High. Reasoning: Treated as real attack attempt. Action: Isolate host, block IP, forensic analysis. Match: High – aligned with admin’s precautionary blocking.	Threat level: Medium. Reasoning: Investigate first before action. Action: Confirm roles, check processes, then decide. Match: Partial – appropriate but less decisive.	Threat level: Medium–High. Reasoning: Considered host at risk, possible privilege escalation. Action: Isolate immediately, scan for malware. Match: Low – overstated compared to admin’s cautious-but-not-isolation approach.	Office

Conclusion

This study shows that AI assistance can elevate novices to perform cybersecurity monitoring tasks at near-expert levels. By comparing Office, Support, and Helpdesk, the research demonstrated that AI-assisted analysis was more accurate and better aligned with real outcomes than unassisted novices, and in some cases even outperformed experienced analysts.

The results underscore AI’s potential to democratize cybersecurity monitoring by reducing dependency on costly SIEM tools and highly specialized staff. Still,

professional oversight remains indispensable for complex or ambiguous cases. Future work should expand the dataset, test across more attack types, and explore standardized AI-assisted workflows to maximize benefits while safeguarding accuracy.

References:

- [1] A. H. Ali, "Unveiling machine learning strategies and considerations in intrusion detection systems: a comprehensive review," *Frontiers in Computer Science*, 2024.
- [2] Z. Xu et al., "Deep learning-based intrusion detection systems: A survey," *arXiv:2504.07839*, 2025.
- [3] R. Kimanzi, P. Kimanga, D. Cherori, and P. K. Gikunda, "Deep learning algorithms used in intrusion detection systems—A review," *arXiv:2402.17020*, 2024.
- [4] E. Karlsen, X. Luo, N. Zincir-Heywood, and M. Heywood, "Benchmarking large language models for log analysis, security, and interpretation," *arXiv:2311.14519*, 2023.
- [5] J. Zhang, "When LLMs meet cybersecurity: a systematic literature review," *Cybersecurity*, vol. 8, article 45, 2025.
- [6] M. Boffa, "LogPrécis: Unleashing language models for automated security log analysis," *Computers & Security*, 2024.
- [7] S. Neupane et al., "Explainable Intrusion Detection Systems (X-IDS): A Survey of current methods, challenges, and opportunities," *arXiv:2207.06236*, 2022.
- [8] O. Arreche et al., "XAI-IDS: Toward proposing an explainable artificial intelligence framework for network intrusion detection," *Applied Sciences*, vol. 14, no. 10, p. 4170, 2024.
- [9] S. Manzoor et al., "Cybersecurity on a budget: Evaluating security and performance of open-source SIEM solutions for SMEs," *PLOS ONE*, vol. 19, no. 3, e0301183, 2024.
- [10] A. Bezas and F. Filippidou, "Comparative analysis of open source security information and event management systems (SIEMs)," *Indonesian Journal of Computer Science*, vol. 12, no. 2, pp. 443–468, 2023.
- [11] T. Vazão, J. M. Cruz, and F. M. Couto, "Implementing and evaluating a GDPR-compliant open-source SIEM solution," *Computers in Industry*, vol. 157, 2023.

[12] F. Hase, "The path to choosing a SIEM system: A systematic literature review," FH Wedel Research Report, 2024.

ENHANCING THE EFFECTIVE USE OF ARTIFICIAL INTELLIGENCE IN HIGHER EDUCATION: A CASE STUDY WITH PRACTICES AND RECOMMENDATIONS

Abdulkerim Oncu (North American University) – aoncu@na.edu

Ihsan Said (North American University) – isaid@na.edu

Azamat Zhamanov (North American University) – azhamanov@na.edu

Abstract: Artificial intelligence (AI) technologies are becoming increasingly common in higher education to support teaching and learning processes. However, realizing the real potential of AI in education depends on its efficient use. This study, based on a case study in the Department of Computer Science at North America University, explores students' and faculty members' perceptions and experiences of using AI, as well as faculty members' recommendations. A survey was administered to 185 undergraduate and graduate students, and semi-structured interviews were conducted with faculty. The findings indicate that students use AI mainly for assignments, exam preparation, learning course content, problem-solving, and project work. Faculty members highlighted challenges such as reduced conceptual depth, academic integrity risks, and ethical concerns. These challenges were most evident in courses like Advanced Software Project Management, Computer Forensics, Network Security, Software Engineering, Database Systems, and Data Mining. Faculty also proposed several suggestions to address these concerns and to support the more effective and ethical use of AI. In this study, the analysis of these recommendations is presented in detail. Analyses also revealed no significant differences between undergraduate and graduate students' patterns of use.

Keywords - Artificial Intelligence (AI), Higher Education, Computer Science, AI integration, Instructional Technology, Pedagogical Alignment

Introduction

In recent years, AI applications have created a transformative shift in education. Generative AI tools are now widely used by students to prepare assignments, practice programming, access information, and get ready for exams. In higher education, these tools are seen as both opportunities and risks. Understanding how students use AI is important for guiding instructors' pedagogical planning and supporting universities in developing ethical policies.

Research Aims

This study has three main objectives:

- To understand how students and faculty use AI for educational purposes.
- To examine whether there are significant differences between undergraduate and graduate students in their patterns of AI use.
- To suggest methods and recommendations for the more effective use of AI in education.

Literature Review

The use of AI in higher education has been getting more attention in the last few years, and there is now a wide range of studies on the topic. Most of this work points out both opportunities and risks. On the positive side, AI has been shown to help with things like personalized learning, adaptive content, learning analytics, and automated feedback, which can support student performance and motivation (Holmes et al., 2019; Luckin, 2021; Weng et al., 2024; Chu et al., 2022; Crompton & Burke, 2023). Generative AI tools such as ChatGPT are also reported to make writing easier and to help students better understand complex ideas (Tierney, 2025). For faculty, AI can save time and reduce workload, especially for administrative or routine tasks (Cotton et al., 2023; Popenici & Kerr, 2017; Dusana Schmidt et al., 2025).

At the same time, the risks are also well-documented. A lot of concern has been raised around issues of academic honesty, plagiarism, and exam integrity (Yilmaz & Göksu, 2020; Tierney, 2025). Many scholars also warn that relying too much on AI can weaken important skills like critical thinking, creativity, and problem solving (Reina et al., 2025; Popenici & Kerr, 2017). Ethical issues are another common theme,

including the spread of incorrect information, algorithmic bias, and data privacy risks (Kasneci et al., 2023; Nelson et al., 2025).

Student experiences and perceptions show a mixed picture. For example, Almassad et al. (2024) found that students value generative AI for efficiency and learning support but also worry about misuse. Alshamy et al. (2025) reported that both students and staff see potential in generative AI tools, though they are cautious about overreliance and skill loss. In Spain, Campillo-Ferrer et al. (2025) noted that students use AI widely but stress the need for better guidance. Khan et al. (2025) also showed positive attitudes toward AI but highlighted that cultural and institutional settings matter for how it is accepted.

A common point across literature is the need for universities to set clear rules and guidelines on how AI should be used. Many authors argue that banning AI is not realistic or helpful; instead, governance-based approaches and stronger AI literacy for both students and teachers are more effective (Shata & Hartley, 2025; Cotton et al., 2023; Luckin, 2021). While students usually take a more optimistic and practical view (Obenza et al., 2024; Almassad et al., 2024; Khan et al., 2025), faculty members are often more critical and cautious (Reina et al., 2025; Alshamy et al., 2025).

To sum up, AI in higher education offers big opportunities but also serious risks. The general view in the literature is that AI should not be banned but rather used in ways that are guided by ethical, pedagogical, and institutional frameworks, and adapted to the real needs of students and teachers (Crompton & Burke, 2023; Chu et al., 2022; Dusana Schmidt et al., 2025).

Method

3.1 Research Design

This study used a mixed-methods design. Quantitative data were collected through a survey, while qualitative data were obtained through semi-structured interviews.

3.2 Sample

- Students: 185 participants (about 62% undergraduate, 38% graduate).

- Faculty: 6 instructors from the Department of Computer Science.
- Courses: Advanced Software Project Management, Computer Forensics, Network Security, Software Engineering, Database Systems, and Data Mining.

3.3 Data Collection

- Student Survey: 10 questions (Likert-scale).
- Faculty Interviews: Semi-structured questions focusing on pedagogical integration, ethical issues, and challenges.

Findings

4.1 Student Survey Results

The survey showed that students use AI for a variety of purposes. Most common uses: assignment completion, learning course content, application development , and generating alternative ideas as shown Figure.1

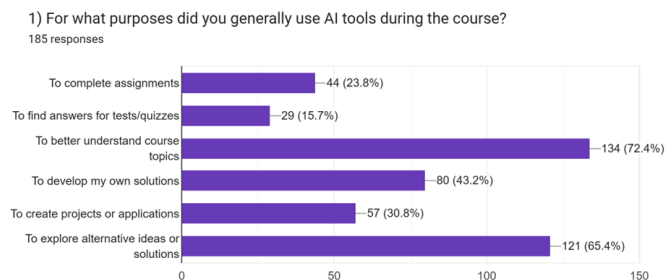


Figure 1. Students' purpose for using AI

Average Likert scores: Time efficiency (3.5), Deep learning (3.7), Ethical awareness (3.8), Motivation (3.2), Independent problem-solving (3.1), Critical evaluation (3.0), Conflicts with AI (2.8), Exam use (2.6). These results suggest that while students value time-saving and ethical awareness, they score lower in exam use and critical evaluation.

4.2 Undergraduate vs. Graduate Students

Independent sample t-tests showed no significant differences between undergraduate and graduate students' use of AI ($p > 0.05$) as shown Figure 2. This suggests that AI use may be linked more to individual factors than academic level.

Question	p_value	Significant
Purpose_Score_Norm	0.357330539	No
Style_Score_Recalculated	0.851154702	No
CritEval_Score	0.795221776	No
Disagree_Score	0.904451124	No
Exam_Score	0.738265119	No
DeepLearn_Score	0.630859884	No
Motivation_Score	0.522917513	No
Solution_Score	0.577496761	No

Figure 2. Independent Sample t-Test Scores

4.3 Faculty Perspectives

Qualitative data revealed two key themes:

Challenges and Limitations:

- Using AI without real understanding leads to surface-level learning.
- Overreliance may weaken critical thinking and problem-solving.
- Sometimes AI generates answers that look correct but lack context.
- Concerns about data privacy, academic integrity, and intellectual property were emphasized.

Effective Use Strategies:

- AI should be positioned as a complementary learning tool.
- Students should learn core concepts before using AI, question outputs, and compare across tools.
- Course design should integrate AI in ways that encourage critical thinking.
- AI should not just be a convenience tool but also a source of creativity and innovation.

Discussion and Conclusion

This study brings together students' AI use practices and faculty perspectives to offer a broad view of AI's role in higher education. Findings can be summarized in three areas:

- Student Use: Students rely heavily on AI for assignments and exam preparation. However, low scores in critical evaluation raise concerns about surface-level learning.
- Academic Level: No significant differences were found between undergraduates and graduates that individual motivation matters more than demographics.
- Faculty Perspective: Faculty acknowledged AI's potential for creativity and pedagogical value but stressed ethical concerns and risks of shallow learning.

Based on survey analysis and faculty feedback, several recommendations were made for students to use AI more effectively:

- Effective Use: AI tools should be integrated in line with course objectives and learning outcomes. Students should develop strong prompt engineering skills, design their own AI applications, and treat AI not only as a shortcut but also as a source of creativity and innovation.
- Critical Use: Students should adopt a questioning approach, validate outputs, and maintain responsibility for their own learning.
- Ethical Use: Issues of data privacy, plagiarism, and academic integrity must always be prioritized.

In conclusion, AI offers both opportunities and risks in higher education. Its effectiveness depends on careful pedagogical integration, critical reflection, and ethical safeguards.

References

- Almassaad, A., Alajlan, H., & Alebaikan, R. (2024). Student perceptions of generative artificial intelligence: Investigating utilization, benefits, and challenges in higher education. *Systems*, 12(10), 385. <https://doi.org/10.3390/systems12100385>
- Alshamy, A., Al-Harthi, A. S. A., & Abdullah, S. (2025). Perceptions of generative AI tools in higher education: Insights from students and academics at Sultan Qaboos University. *Education Sciences*, 15(4), 501. <https://doi.org/10.3390/educsci15040501>
- Campillo-Ferrer, J. M., López-García, A., & Miralles-Sánchez, P. (2025). Student perceptions of the use of Gen-AI in a higher education program in Spain. *Digital*, 5(3), 29. <https://doi.org/10.3390/digital5030029>
- Chu, H.-C., Hwang, G.-H., Tu, Y.-F., & Yang, K.-H. (2022). Roles and research trends of artificial intelligence in higher education: A systematic review of the top 50 most-cited articles. *Australasian Journal of Educational Technology*, 38(3), 22–42. <https://doi.org/10.14742/ajet.7526>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Crompton, H., & Burke, D. (2023). Artificial intelligence in higher education: The state of the field. *International Journal of Educational Technology in Higher Education*, 20(1), 22. <https://doi.org/10.1186/s41239-023-00392-8>
- Dusana Alshatti Schmidt, B., Alboloushi, B., Thomas, A., & Magalhaes, R. (2025). Integrating artificial intelligence in higher education: Perceptions, challenges, and strategies for academic innovation. *Computers and Education Open*, 9, 100274. <https://doi.org/10.1016/j.caeo.2025.100274>
- Holmes, W., Bialik, M., & Fadel, C. (2019). *Artificial intelligence in education: Promises and implications for teaching and learning*. Boston, MA: Center for Curriculum Redesign.

- Kasneci, E., Seßler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., & Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Khan, M., Gope, L., & Roy, D. (2025). Students' perceptions towards artificial intelligence technologies in higher education. *International Journal of Research Publication and Reviews*, 6, 1277–1284. <https://doi.org/10.55248/gengpi.6.0825.2841>
- Luckin, R. (2021). Towards artificial intelligence-based assessment systems. *Nature Human Behaviour*, 5(12), 1630–1632. <https://doi.org/10.1038/s41562-021-01164-6>
- Nelson, A. S., Santamaría, P. V., Javens, J. S., & Ricaurte, M. (2025). Students' perceptions of generative artificial intelligence (GenAI) use in academic writing in English as a foreign language. *Education Sciences*, 15(5), 611. <https://doi.org/10.3390/educsci15050611>
- Obenza, B., Salvahan, A., Rios, A. N., Solo, A., Alburo, R. A., & Gabila, R. J. (2024). University students' perception and use of ChatGPT: Generative artificial intelligence (AI) in higher education. *International Journal of Human Computing Studies*, 5(12), 5–18. SSRN. <https://ssrn.com/abstract=4724968>
- Popenici, S. A. D., & Kerr, S. (2017). Exploring the impact of artificial intelligence on teaching and learning in higher education. *Research and Practice in Technology Enhanced Learning*, 12(1), 22. <https://doi.org/10.1186/s41039-017-0062-8>
- Reina, Y., Cruz Caro, O., Rubio, Y., Sánchez Bardales, E., Rituay, A., & Santos, R. (2025). Artificial intelligence as a teaching tool in university education. *Frontiers in Education*, 10, 1578451. <https://doi.org/10.3389/feduc.2025.1578451>
- Shata, A., & Hartley, K. (2025). Artificial intelligence and communication technologies in academia: Faculty perceptions and the adoption of generative AI.

International Journal of Educational Technology in Higher Education, 22(1), 14.
<https://doi.org/10.1186/s41239-025-00511-7>

Tierney, A. (2025). Student perceptions on the impact of AI on their teaching and learning experiences in higher education. *Research and Practice in Technology Enhanced Learning*, 20(1), 5.
<https://doi.org/10.58459/rptel.2025.20005>

Weng, X., Xia, Q., Gu, M., Rajaram, K., & Chiu, T. K. (2024). Assessment and learning outcomes for generative AI in higher education: A scoping review on current research status and trends. *Australasian Journal of Educational Technology*, 40(6), 37–55. <https://doi.org/10.14742/ajet.9540>

EXPLORING AI FOR TEACHING AND LEARNING: A PRACTICAL INSTRUCTOR'S PERSPECTIVE

Zeliha Ozdogan, Ph.D. (The Pennsylvania State University- Harrisburg) – zzol2@psu.edu

Abstract: Artificial Intelligence (AI) is creating new opportunities to transform teaching and learning. As an instructor, I have been exploring how these tools can enhance student engagement and improve my classroom practice. This presentation will share my personal journey of adopting AI resources, what I have learned along the way, and how these tools are beginning to shape my approach to teaching. I will introduce several accessible tools—Microsoft Copilot, ChatGPT, Adobe Firefly, and Google Notebook LM—and reflect on how I am starting to use them. Copilot and ChatGPT have supported me in drafting materials, Firefly has opened creative possibilities for visuals, and Notebook LM offers potential for collaborative notetaking. I will discuss the benefits, challenges, and strategies I have discovered while experimenting with these resources. By the end of the session, attendees will gain both an overview of key AI tools and a firsthand perspective on the process of adopting them. My goal is to inspire fellow instructors to see AI as a partner in education, while also recognizing that the journey requires curiosity, experimentation, and reflection.

Keywords: AI in Education; Teaching and Learning; Innovation Instructional; Technology Digital Pedagogy

Introduction

Artificial intelligence (AI) has rapidly moved from a niche topic to a daily presence in higher education. Instructors across disciplines are being challenged to rethink how learning happens when students have access to generative AI tools such as ChatGPT, Microsoft Copilot, and Google Notebook LM. Rather than treating these technologies as a disruption, I began to see them as an opportunity to improve teaching effectiveness and student engagement. My goal as an instructor has been to explore how AI tools cannot supplement student learning, while also promoting critical thinking and creativity.

This paper summarizes the key ideas from my presentation “Exploring AI for Teaching and Learning: A Practical Instructor’s Perspective.” It reflects on what I have learned while experimenting with AI in my classes, reviews recent literature on digital pedagogy, and highlights resources that can help educators thoughtfully adopt AI in their teaching practice.

Literature Review

Digital Pedagogy and Technology Integration

The conversation around digital pedagogy has evolved significantly over the past two decades from seeing technology as a novelty in classrooms to recognizing it as an integral part of the learning process. Early studies emphasized the potential of digital tools to make education more flexible, accessible, and aligned with real-world competencies (Makarova & Makarova; 2018). Their work demonstrated that the combination of pedagogy, technology, and guided instructor support can transform educational environments by promoting active learning and digital literacy. In this framework, instructors often take on a role similar to a tutor; someone who mediates between students and digital tools to ensure that technology enhances, rather than overwhelms, the learning experience.

Recent research confirms that digital pedagogy has matured into a recognized academic field, with global attention growing sharply since 2020 (Santoveña-Casal & López, 2024). This development has been influenced by the rapid digitalization of higher education and the challenges brought by the COVID-19 pandemic, which

forced both instructors and institutions to adopt new tools quickly. Bibliometric studies show that digital pedagogy research now extends beyond mere technology adoption, focusing instead on how flexible pedagogies can adapt to varied teaching contexts and support educational quality (Santoveña-Casal & López, 2024).

Ching and Roberts (2020) argue that while technology can make teaching more dynamic, its success depends less on the tools themselves and more on the instructional design behind them. Their discussion of models such as TPACK (Technological Pedagogical Content Knowledge) and ADDIE (Analyze, Design, Develop, Implement, Evaluate) highlights that technology integration should always be driven by pedagogical goals, not the other way around. Similarly, Okojie, Olinzock, and Okojie-Boulder emphasize that technology integration is most effective when viewed as part of the overall instructional process; linked to objectives, feedback, assessment, and student reflection (Okojie et al., 2006).

AI and Emerging Pedagogies

The rise of generative AI has introduced both excitement and unease in academic settings. Recent discussions emphasize the potential of AI to support personalized feedback, assist in idea generation, and model problem-solving processes (Kasneci et al., 2023). While some instructors express concern about academic integrity and fairness, others view AI as an opportunity to strengthen students' critical thinking and engagement through guided and intentional use.

In higher education, researchers have begun exploring both the challenges and benefits of integrating AI-based tools such as ChatGPT. Sullivan, Kelly, and McLaughlan (2023) observed that the release of ChatGPT has sparked important debates around academic integrity, equity, and access. Their analysis of university responses revealed not only concerns about plagiarism and authenticity but also opportunities to redesign assessments and promote new forms of participation for students from underrepresented backgrounds. This shift encourages educators to consider how AI can be used ethically to support not replace student learning.

Emerging studies also suggest that AI can play a meaningful role in developing higher-order thinking skills. For example, Guo and Lee (2023) implemented ChatGPT-based activities in chemistry courses, finding that structured interaction with AI

improved students' confidence in analyzing information, forming arguments, and validating evidence. Similarly, Suriano et al. (2025) found that active engagement with AI-based chatbots positively influenced students' complex critical thinking performance, particularly when they approached AI use with curiosity and reflection rather than dependence. These studies highlight the need for educators to design activities that cultivate trust, engagement, and critical evaluation of AI outputs.

Beyond critical thinking, AI and machine learning also show promise in supporting personalized and adaptive learning environments. Tiwari (2023) notes that AI-assisted systems such as intelligent tutoring, adaptive testing, and learning analytics allow instruction to be tailored to each student's needs. These tools hold potential to improve outcomes and accessibility, though concerns about privacy, bias, and equity remain important areas for continued research and dialogue.

Resources for Instructors

There are several accessible tools and resources that educators can explore to begin integrating AI into their teaching. Microsoft Copilot offers writing assistance, brainstorming, and content generation within familiar platforms such as Word and PowerPoint. ChatGPT provides flexible text generation that can support reflective prompts, case studies, and peer feedback exercises. Adobe Firefly enables creative visual projects, while Google Notebook LM helps students organize and synthesize information collaboratively.

Many colleges and universities now offer free, high-quality resources to help faculty integrate AI into their teaching. These include workshops, sample policies, online training modules, and guidance from teaching and learning centers. Such institutional support helps instructors adopt tools like Microsoft Copilot, ChatGPT, Adobe Firefly, and Google Notebook LM in ways that promote critical thinking, creativity, and collaboration.

A particularly helpful open resource is Bloom's Taxonomy Revisited (2024) by Oregon State University which connects Bloom's traditional hierarchy of learning objectives with generative AI applications. It illustrates how AI can supplement learning at each level from helping students recall and understand concepts, to supporting analysis,

evaluation, and creation while emphasizing that human reflection, ethics, and judgment remain essential.

Using this framework, instructors can design activities where students engage with AI critically rather than depend on it. For instance, students might use ChatGPT to generate different solutions to a problem, then analyze and justify which one demonstrates stronger reasoning. By combining institutional support with resources educators can confidently integrate AI as a supplement to learning, fostering both digital literacy and deeper critical engagement.

Conclusion

As AI becomes more embedded in higher education, instructors have a unique opportunity to guide students in using these tools thoughtfully and responsibly. The goal is not to resist AI but to integrate it in ways that promote reflection, creativity, and critical thinking. When used intentionally, AI can serve as a partner in learning helping students explore ideas, evaluate information, and engage more deeply with course material.

However, successful integration requires clarity and transparency. Instructors should be explicit about how AI can and cannot be used in their courses and include a clear AI policy in their syllabi. Establishing these expectations helps students understand appropriate use, maintain academic integrity, and build confidence in their own learning process.

Ultimately, the integration of AI in education is not about replacing human intelligence but expanding it. By combining sound pedagogy, institutional support, and well-defined classroom guidelines, educators can ensure that AI strengthens—not weakens—the core values of higher education: curiosity, integrity, and the pursuit of meaningful learning.

References

- Oregon State University, Ecampus, Bloom's Taxonomy Revisited, <https://ecampus.oregonstate.edu/faculty/artificial-intelligence-tools/blooms-taxonomy-revisited-v2-2024.pdf>
- Ching, G. S., & Roberts, A. (2020). Evaluating the pedagogy of technology integrated teaching and learning: An overview. *International Journal of Research Studies in Education*. <https://doi.org/10.5861/ijrse.5800>
- Guo, Y., & Lee, D. (2023). Leveraging ChatGPT for enhancing critical thinking skills. *Journal of Chemical Education*, 100(12), 4876–4883. <https://doi.org/10.1021/acs.jchemed.3c00505>
- Kasneci, E., Seidel, T., Kasneci, G., Bennett, K., & Gasser, U. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 97, 102188
- Makarova, E. A., & Makarova, E. L. (2018). Blending pedagogy and digital technology to transform educational environment. *International Journal of Cognitive Research in Science, Engineering and Education (IJCRSEE)*; 2:57-65
- Okojie, M. C. P., Olinzock, A. A., & Okojie-Boulder, T. C. (2006). The pedagogy of technology integration. *Journal of Technology Studies*, 32 :66-71
- Santoveña-Casal, S., & López, S. R. (2024). Trends and evolution of research into digital pedagogy in higher education: A bibliometric analysis. *Education and Information Technologies* 29:2437-2458
- Sullivan, M., Kelly, A., & McLaughlan, P. (2023). Academic integrity and artificial intelligence in higher education: The impact of ChatGPT. *Journal of Applied Learning & Teaching*, 6(1).
- Suriano, R., Plebe, A., Acciai, A., & Fabio, R. A. (2025). Student interaction with ChatGPT can promote complex critical thinking skills. *Learning and Instruction*, 95, 102011.

Tiwari, R. (2023). The integration of AI and machine learning in education and its potential to personalize and improve student learning experiences. *International Journal of Scientific Research in Engineering and Management*, 7(2).

EXPLORING THE BARRIERS AND ENABLERS OF AI-DRIVEN ENERGY EFFICIENCY IN SMEs

Sumeyra Danisman (Stony Brook University) – sumeyra.danisman@stonybrook.edu

Elizabeth Hewitt (Stony Brook University) – elizabeth.hewitt@stonybrook.edu

Abstract: *This objective of this study is to examine the factors that enable or hinder the use of artificial intelligence (AI) to improve energy efficiency in small and medium-sized enterprises (SMEs), an area that remains largely understudied. Previous research has examined enablers and barriers to energy efficiency in buildings (Jaffe & Stavins, 1994; Yeatts et al., 2017; Peel et al., 2020; Palm & Bryngelson, 2023), industrial systems (Thollander & Palm, 2013; Lunt et al., 2014), manufacturing (Trianni et al., 2016; Cagno et al., 2017), and broader sustainability practices (Caldera et al., 2019; Basit et al., 2024; Moursellas et al., 2024; Zavodna et al., 2024). Some studies have also explored the development of AI solutions for small-business operations (Crockett et al., 2021; Mantri & Mishra, 2023; Md. Kamruzzaman et al., 2025). However, there is a significant need to examine AI particularly in the context of energy efficiency in SMEs. This issue is critical not only for national and global energy use or consumption (Henriques & Catarino, 2016; Gennitsaris et al., 2023; OECD, 2025) but also for sustainability in SMEs (Viesi et al., 2017; Álvarez Jaramillo et al., 2019). Accordingly, the primary aim of this study is to evaluate the key barriers and enablers (Table 1) shaping the adoption of AI for energy efficiency in SMEs. As such, the guiding research question for this review is: What are the key enablers and barriers influencing the adoption of AI-driven energy efficiency technologies in SMEs, and how do they shape adoption choices?*

Barrier 1 - Knowledge Gap and Limited Awareness

Limited knowledge and awareness hinder AI-driven innovation and energy efficiency. A knowledge gap occurs when existing understanding of AI outputs fall short of strategic or innovative needs (Qi et al., 2020; Martin & Parmar, 2024). Despite AI's potential, major gaps persist, especially among SMEs lacking data and expertise for AI-based energy management (Wigger et al., 2025). Similarly, low awareness limits small firms' adoption of energy-efficient technologies (Ketenci & Wolf, 2024; Peretz-Andersson et al., 2024).

Barrier 2 - Implementation or Investment Cost

High costs limit AI adoption for energy efficiency, driven by infrastructure expenses (Cubric, 2020) and resource scarcity (Soomro et al., 2025). Small firms lack capital (IEA, 2015), while development and deployment are costly (Danish, 2023). Financial and hidden costs further deter adoption (Pimenow et al., 2024; Carlander & Thollander, 2023).

Barrier 3 - Integration Challenge and Complexity

Integration challenges arise when AI systems do not fit existing infrastructures, requiring technical adjustments and expertise (Danish, 2023). The complexity and "black box" nature of AI hinder transparency and trust (Park, 2025), while practical approaches are needed to support effective integration using existing resources (Wigger et al., 2025).

Barrier 4 - Privacy and Security Concerns

Privacy and security worries stop many small businesses from using AI. Fears of data misuse and hacking make it seem unsafe (Carmody et al., 2021; Tolani et al., 2025). Protecting data is a big challenge (Dibie, 2024), especially since AI gathers sensitive information. Limited knowledge of how AI handles data adds more doubt (Lim & Shim, 2022; Iyelolu et al., 2024). These hidden risks come from its ability to uncover private details (Hu & Min, 2023; Carmody et al., 2021).

Despite these challenges, there are also significant enablers that can help SMEs harness AI to improve their energy performance, which will be outlined below.

Enabler 1 - Policies and Incentives

Policies and incentives help promote sustainable energy use (Steg et al., 2006). Clear rules and support programs encourage small firms to use AI for efficiency (Dixon et al., 2010; Yahchouchi & Rotabi, 2025). Government action creates motivation and reduces costs, making AI adoption easier (Henriques & Catarino, 2016; Shaik et al., 2024).

Enabler 2 - Partnerships and Collaborations

Partnerships and teamwork help small firms access expertise and AI tools (Iyelolu et al., 2024). Working with energy providers, tech firms, and universities helps share knowledge and close skill gaps. Collaboration across teams and organizations supports AI use and problem-solving (Egbuhuzor et al., 2024), speeding progress toward AI-driven and low-carbon goals (Shaik et al., 2024).

Enabler 3 - Technical Assistance

Technical assistance helps small firms adopt AI by offering consulting, pilot projects, and implementation support (Egbuhuzor et al., 2024). It promotes strategies and best practices to overcome barriers (Uwagaba, 2023; Khaq et al., 2024). Services like consultancy, energy audits, and integration guidance improve AI deployment and build leaders' knowledge and skills for effective use (Das, 2024).

Enabler 4 - Training and Educational Programs

Training and education are key to reducing the AI skills gap. Specialized programs focused on energy applications help build expertise (Tunde et al., 2024).

Collaboration between industry, academia, and government is vital for developing skilled workers. Training initiatives enable SMEs to adopt sustainable, AI-driven practices (Basit et al., 2024; Emedo et al., 2025) and improve efficiency through well-structured programs that strengthen their ability to use AI effectively (ul Haq et al., 2025).

This study categorizes the key enablers and barriers to AI adoption in SMEs into three interrelated areas: financial, technical, and behavioral. Financial factors concern costs and perceived risks of implementation - high costs often pose significant barriers, but policies and incentives can encourage investment. Technical factors relate to system integration and complexity, with compatibility challenges hindering progress, while technical assistance facilitates smoother implementation and improved performance. Behavioral factors focus on organizational and individual capabilities: Knowledge gaps, limited awareness, and privacy concerns hinder adoption, while training, education, and partnerships build the trust, skills, and collaboration needed for successful AI integration. While some factors span multiple categories, the analysis here focuses on their primary influence.

ATTACHMENT

Table 1. Enablers and Barriers for AI adoption in SMEs for energy efficiency

	FACTORS	STUDIES
BARRIERS	<i>Practitioners' knowledge gap and limited awareness</i>	Qi et al. 2020; Martin & Parmar 2024; Birkstedt & colleagues 2023; Wigger et al., 2025; Ketenci & Wolf 2024; Peretz-Andersson et al. 2024.
	<i>Implementation or investment cost</i>	Cubric, 2020; Soomro et al., 2025; Danish. 2023; Pimenow et al. 2024; Thollander, 2023; Wigger, 2025.
	<i>Integration challenge and technical complexity</i>	Danish, 2023; Park, 2025; Wigger 2025
	<i>Privacy and security concerns</i>	Carmody et al., 2021; Tolani et al., 2025; Dibie, 2024; Lim & Shim, 2022; Iyelolu et al., 2024; Carmody et al. 2021

ENABLERS	<i>Policies and incentives</i>	Steg et al., 2006; Dixon et al., 2010; Yahchouchi & Rotabi, 2025; Henriques & Catarino, 2016; Shaik et al., 2024
	<i>Partnerships and Collaborations</i>	Iyelolu et al., 2024; Egbuhuzor et al. 2024; Shaik et al. 2024.
	<i>Technical Assistance</i>	Egbuhuzor et al., 2024; Uwagaba, 2023; Khaq et al., 2024; Das, 2024.
	<i>Training and Educational Programs</i>	Tunde et al., 2024; Basit et al., 2024; Emedo et al., 2025; ul Haq et al., 2025

References

- Abdul Basit, S., Charlegghi, B., Batool, K., Hassan, S. S., Jahanshahi, A. A., & Kliem, M. E. (2024a). Review of enablers and barriers of sustainable business practices in SMEs. *Journal of Economy and Technology*, 2, 79–94.
<https://doi.org/10.1016/j.ject.2024.03.005>
- Álvarez Jaramillo, J., Zarthá Sossa, J. W., & Orozco Mendoza, G. L. (2019). Barriers to sustainability for small and medium enterprises in the framework of sustainable development—L iterature review. *Business Strategy and the Environment*, 28(4), 512–524. <https://doi.org/10.1002/bse.2261>
- Basit, S. A., Gharlegghi, B., Batool, K., Hassan, S. S., Jahanshahi, A. A., & Kliem, M. E. (2024). Review of enablers and barriers of sustainable business practices in SMEs. *Journal of Economy and Technology*, 2, 79-94.
- Birkstedt, T., Minkkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI governance: themes, knowledge gaps and future agendas. *Internet Research*, 33(7), 133-167.
- Caldera, H. T. S., Desha, C., & Dawes, L. (2019). Evaluating the enablers and barriers for successful implementation of sustainable business practice in 'lean' SMEs. *Journal of Cleaner Production*, 218, 575–590.
<https://doi.org/10.1016/j.jclepro.2019.01.239>
- Carlander, J., & Thollander, P. (2023). Barriers to implementation of energy-efficient technologies in building construction projects—Results from a Swedish case study. *Resources, Environment and Sustainability*, 11, 100097.
<https://doi.org/10.1016/j.resenv.2022.100097>
- Carmody, J., Shringarpure, S., & Van De Venter, G. (2021a). AI and privacy concerns: A smart meter case study. *Journal of Information, Communication and Ethics in Society*, 19(4), 492–505. <https://doi.org/10.1108/JICES-04-2021-0042>
- Crockett, K., Colyer, E., Gerber, L., & Latham, A. (2021). Building trustworthy AI solutions: A case for practical solutions for small businesses. *IEEE Transactions on Artificial Intelligence*, 4(4), 778-791.

- Cubric, M. (2020). Drivers, barriers and social considerations for AI adoption in business and management: A tertiary study. *Technology in Society*, 62, 101257. <https://doi.org/10.1016/j.techsoc.2020.101257>
- Danish, M. S. S. (2023). AI in Energy: Overcoming Unforeseen Obstacles. *AI*, 4(2), 406–425. <https://doi.org/10.3390/ai4020022>
- Das, T. (2024). Accelerating AI sustainability and innovation at the Department of Energy. Bipartisan Policy Center. <https://bipartisanpolicy.org/download/?file=/wp-content/uploads/2024/09/Accelerating-AISustainability-and-Innovation-at-DOE-BPC-Report-Sept-2024-v2.pdf>
- Dibie, E. U. (2024). The future of renewable energy: Ethical implications of AI and cloud technology in data security and environmental impact. *Journal of Advances in Mathematics and Computer Science*, 39(10), 10-9734.
- Dixon, R. K., McGowan, E., Onysko, G., & Scheer, R. M. (2010). US energy conservation and efficiency policies: Challenges and opportunities. *Energy Policy*, 38(11), 6398–6408. <https://doi.org/10.1016/j.enpol.2010.01.038>
- Egbuhuzor, N. S., Ajayi, A. J., Akhigbe, E. E., & Agbede, O. O. (2024). Leveraging AI and cloud solutions for energy efficiency in large-scale manufacturing. *International Journal of Science and Research Archive*, 13(2), 4170-4192.
- Emedo, C., Wada, O. Z., Clement David-Olawade, A., Ling, J., Esan, D. T., Ijiwade, J., & Olawade, D. B. (2025). AI-driven transformations in smart buildings: A review of energy efficiency and sustainable operations. *Digital Engineering*, 7, 100068. <https://doi.org/10.1016/j.dte.2025.100068>
- Gennitsaris, S., Oliveira, M. C., Vris, G., Bofilios, A., Ntinou, T., Frutuoso, A. R., Queiroga, C., Giannatsis, J., Sofianopoulou, S., & Dedoussis, V. (2023). Energy Efficiency Management in Small and Medium-Sized Enterprises: Current Situation, Case Studies and Best Practices. *Sustainability*, 15(4), 3727. <https://doi.org/10.3390/su15043727>
- Hamm, G. (2025). 2024 United States Data Center Energy Usage Report. Lawrence Berkely National Laboratory. <https://doi.org/10.71468/P1WC7Q>

- Haq, F. U., Suki, N. M., Setini, M., Masood, A., & Khan, T. A. (2025). Adopting green AI for SME sustainability: Mediating role of green investment and moderation by green servant leadership. *Sustainable Futures*, 10, 101002. <https://doi.org/10.1016/j.sftr.2025.101002>
- Henriques, J., & Catarino, J. (2016a). Motivating towards energy efficiency in small and medium enterprises. *Journal of Cleaner Production*, 139, 42–50. <https://doi.org/10.1016/j.jclepro.2016.08.026>
- Hu, Y., & Min, H. (Kelly). (2023). The dark side of artificial intelligence in service: The “watching-eye” effect and privacy concerns. *International Journal of Hospitality Management*, 110, 103437. <https://doi.org/10.1016/j.ijhm.2023.103437>
- IEA International Energy Agency. (2024). Electricity 2024: Analysis and forecast to 2026 (Revised version, January & May 2024). <https://iea.blob.core.windows.net/assets/18f3ed24-4b26-4c83-a3d2-8a1be51c8cc8/Electricity2024-Analysisandforecastto2026.pdf>
- Iyelolu, T. V., Agu, E. E., Idemudia, C., & Ijomah, T. I. (2024). Driving SME innovation with AI solutions: overcoming adoption barriers and future growth opportunities. *International Journal of Science and Technology Research Archive*, 7(1), 036-054.
- Jaffe, A. B., & Stavins, R. N. (1994). The energy paradox and the diffusion of conservation technology. *Resource and Energy Economics*, 16(2), 91–122
- Ketenci, A., & Wolf, M. (2024). Advancing energy efficiency in SMEs: A case study-based framework for non-energy-intensive manufacturing companies. *Cleaner Environmental Systems*, 14, 100218. <https://doi.org/10.1016/j.cesys.2024.100218>
- Khaq, Z. D., Subroto, V. K., & Susanto, E. (2024). AI-driven Strategies for Enhancing MSME Sales and Business Communication: A Case Study. *Journal of Management and Informatics*, 3(2), 180-194.
- Lim, S., & Shim, H. (2022). No secrets between the two of us: Privacy concerns over using AI agents. *Cyberpsychology: Journal of Psychosocial Research on Cyberspace*, 16(4). <https://doi.org/10.5817/CP2022-4-3>

- Lunt, P., Ball, P., & Levers, A. (2014). Barriers to industrial energy efficiency. *International Journal of Energy Sector Management*, 8(3), 380–394. <https://doi.org/10.1108/IJESM-05-2013-0008>
- Mantri, A., & Mishra, R. (2023). Empowering small businesses with the force of big data analytics and AI: A technological integration for enhanced business management. *The Journal of High Technology Management Research*, 34(2), 100476. <https://doi.org/10.1016/j.hitech.2023.100476>
- Martin, K., & Parmar, B. (2024). AI and the Creation of Knowledge Gaps: The ethics of AI transparency. Available at SSRN 4207128.
- Md. Kamruzzaman, Sujoy Saha, Md. Shoeb Siddiki, Rabi Sankar Mondal, & Md Nazmul Alam Bhuiyan. (2025). Applications of Artificial Intelligence in Small and Medium Scale Business. *Journal of Business and Management Studies*, 7(4), 314–325. <https://doi.org/10.32996/jbms.2025.7.4.20.21>
- Mohd Rasdi, R., & Umar Baki, N. (2025). Navigating the AI landscape in SMEs: Overcoming internal challenges and external obstacles for effective integration. *PLOS One*, 20(5), e0323249. <https://doi.org/10.1371/journal.pone.0323249>
- Moursellas, A., Malesios, C., Skouloudis, A., Evangelinos, K., & Dey, P. K. (2024). Perceived enablers and barriers impacting sustainability of small-and-medium sized enterprises: A quantitative analysis in four European countries. *Environmental Quality Management*, 33(3), 433–448. <https://doi.org/10.1002/tqem.22128>
- OECD. (2025). SME indicators, benchmarking and monitoring. <https://www.oecd.org/en/topics/sub-issues/sme-indicators-benchmarking-and-monitoring.html>
- Palm, J., & Bryngelson, E. (2023). Energy efficiency at building sites: Barriers and drivers. *Energy Efficiency*, 16(2), 7. <https://doi.org/10.1007/s12053-023-10088-7>
- Park, C. (2025). Addressing Challenges for the Effective Adoption of Artificial Intelligence in the Energy Sector. *Sustainability*, 17(13), 5764. <https://doi.org/10.3390/su17135764>

- Peel, J., Ahmed, V., & Saboor, S. (2020). An investigation of barriers and enablers to energy efficiency retrofitting of social housing in London. *Construction Economics and Building*, 20(2), 127-149.
- Peretz-Andersson, E., Tabares, S., Mikalef, P., & Parida, V. (2024). Artificial intelligence implementation in manufacturing SMEs: A resource orchestration approach. *International Journal of Information Management*, 77, 102781. <https://doi.org/10.1016/j.ijinfomgt.2024.102781>
- Pimenow, S., Pimenowa, O., & Prus, P. (2024). Challenges of Artificial Intelligence Development in the Context of Energy Consumption and Impact on Climate Change. *Energies*, 17(23), 5965. <https://doi.org/10.3390/en17235965>
- Qi, L., An, X., Zhang, S., & Wang, X. (2020). Research on Knowledge Gap Identification Method in Innovative Organizations under the “Internet+” Environment. *Information*, 11(12), 572. <https://doi.org/10.3390/info11120572>
- Shaik, A. S., Alshibani, S. M., Jain, G., Gupta, B., & Mehrotra, A. (2024b). Artificial intelligence (AI)-driven strategic business model innovations in small- and medium-sized enterprises. Insights on technological and strategic enablers for carbon neutral businesses. *Business Strategy and the Environment*, 33(4), 2731–2751. <https://doi.org/10.1002/bse.3617>
- Soomro, R. B., Al-Rahmi, W. M., Dahri, N. A., Almuqren, L., Al-Mogren, A. S., & Aldaijy, A. (2025). A SEM–ANN analysis to examine impact of artificial intelligence technologies on sustainable performance of SMEs. *Scientific Reports*, 15(1), 5438. Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-75589-7_19
- Steg, L., Dreijerink, L., & Abrahamse, W. (2006). Why are Energy Policies Acceptable and Effective? *Environment and Behavior*, 38(1), 92–111. <https://doi.org/10.1177/0013916505278519>
- Thollander, P., & Palm, J. (2013). Improving energy efficiency in industrial energy systems: An interdisciplinary perspective on barriers, energy audits, energy management, policies, and programs. Springer.

- Tolani, K. C., Vijayalakshmi, A., Lanjewar, P., Dhote, M. G., Chandra, S., Ahmed, S., & Singh, B. (2025). Cybersecurity Challenges in AI-Driven Energy Systems: Current and Future Prospects Concerning Ethical and Legal Provisions. In H. Hammouch & L. Razzak Janjua (Eds.), *Advances in Environmental Engineering and Green Technologies* (pp. 109–122). IGI Global. <https://doi.org/10.4018/979-8-3373-0045-0.ch008>
- Trianni, A., Cagno, E., & Farné, S. (2016). Barriers, drivers and decision-making process for industrial energy efficiency: A broad study among manufacturing small and medium-sized enterprises. *Applied Energy*, 162, 1537–1551. <https://doi.org/10.1016/j.apenergy.2015.02.078>
- Tunde, A. H., Yusuf, S. O., Taiwo, A. I., Ocran, G., Owusu, P., & Paul-Adeleye, A. H. (2024). AI-driven innovations in energy efficiency: Transforming smart buildings and urban areas through technology and digital transformation. *World Journal of Advanced Research and Reviews*, 24(1).
- ul Haq, F., Suki, N. M., Setini, M., Masood, A., & Khan, T. A. (2025). Adopting green AI for SME sustainability: mediating role of green investment and moderation by green servant leadership. *Sustainable Futures*, 101002
- Uwagaba, J., Omotosho, T. D., & George, G. O. (2023). Exploring the barriers to artificial intelligence adoption in Sub-Saharan Africa's Small and Medium Enterprises and the potential for increased productivity. *Worldwide Journal of Multidisciplinary Research and Development*, 9(6), 4-10.
- Viesi, D., Pozzar, F., Federici, A., Crema, L., & Mahbub, M. S. (2017). Energy efficiency and sustainability assessment of about 500 small and medium-sized enterprises in Central Europe region. *Energy Policy*, 105, 363–374. <https://doi.org/10.1016/j.enpol.2017.02.045>
- Wigger, M., Burggräf, P., Steinberg, F., Becher, A., & Heinbach, B. (2025). Integrating artificial intelligence into energy management: A case study on energy consumption data analysis and forecasting in a German manufacturing company. *Energy and AI*, 21, 100576. <https://doi.org/10.1016/j.egyai.2025.100576>

Yahchouchi, G., & Rotabi, S. (2025). Government Initiatives and Policies Promoting AI Adoption in the MENA. In N. Azoury & G. Yahchouchi (Eds.), *AI in the Middle East for Growth and Business* (pp. 329–345).

Yeatts, D. E., Auden, D., Cooksey, C., & Chen, C.-F. (2017). A systematic review of strategies for overcoming the barriers to energy-efficient technologies in buildings. *Energy Research & Social Science*, 32, 76–85.
<https://doi.org/10.1016/j.erss.2017.03.010>

Zavodna, L. S., Überwimmer, M., & Frankus, E. (2024). Barriers to the implementation of artificial intelligence in small and medium-sized enterprises: Pilot study. *Journal of Economics and Management*, 46, 331-352.

REVEALING THE UNKNOWN: LLM-BASED FORENSICS TRIAGE AGENT FOR ANDROID MOBILE FORENSICS

Akif Ozer, MS (Purdue University)- akif.ozero0@gmail.com

Vinicius Lima. PhD (University at Albany – SUNY)- xie447@purdue.edu

Mingyang Xie, PhD (Purdue University)- vlima@albany.edu

Umit Karabiyik. PhD (George Mason University)- ukarabiy@gmu.edu

Abstract: Mobile forensic investigations struggle to keep pace with the rapid evolution of mobile applications, as commercial tools require frequent updates yet often fail to extract novel or application-specific artifacts. Large language models (LLMs), widely applied across diverse fields, present an opportunity to address this gap by reasoning about data structures without prior tool-specific knowledge. In this paper, we introduce MobileTriageAgent, an experimental LLM-driven forensic analysis system that automates artifact discovery and interpretation from Android device file systems. Our methodology involved selecting the top 15 applications from the Google Play Store, populating them with test data, and processing the resulting device archives using both leading commercial forensic tools and our proposed LLM-based system. To evaluate accuracy, we manually identified and cataloged application artifacts as ground truth, enabling a direct comparison of performance across approaches. MobileTriageAgent employs a command-based framework for TAR parsing, SQLite querying, and structured file analysis, while leveraging LLM reasoning to identify forensic artifacts such as user identifiers, tokens, geolocation records, and communication traces. Results demonstrate that the LLM-based approach reveals critical evidence overlooked by commercial tools, highlighting its potential to aid digital investigations.

Introduction

Digital Forensics (DF) is a branch of forensic science dedicated to the identification, collection, preservation, analysis, and presentation of digital evidence. Within this field, Mobile Forensics focuses on smartphones and IoT devices, where tools in this domain typically parse mobile filesystems and operating systems (most commonly Android and iOS) to extract and interpret artifacts generated by user applications [1]. These tools often rely on predefined application structures to map and present artifacts that may hold investigative value [2].

A major challenge in mobile forensics today is the exponential growth of data stored on modern smartphones, which now often ship with hundreds of gigabytes of capacity [2]. A single device can now contain over 250 GB of potential evidence, while even one application may generate thousands of files, dramatically increasing the volume of data that investigators must analyze. Compounding this issue, new applications are released daily, and existing ones are continuously updated. Traditional digital forensic tools struggle to keep pace with these rapid changes, requiring frequent updates to support evolving data formats and application versions [3].

To address these scalability and adaptability challenges, our research investigates the use of Large Language Models (LLMs) in DF to support the triage and analysis of digital artifacts. Our guiding research question is: *Can LLMs support DF by reasoning over text data from digital artifacts?* This hypothesis builds on the fact that most digital artifacts are text-based, suggesting that LLMs can be leveraged to directly parse and interpret application content. By doing so, they offer a scalable mechanism for processing large volumes of files and applications within an evidence image. Such an approach would allow investigators to automatically identify and summarize relevant artifacts without prior technical knowledge of specific application structures, presenting digital evidence in a clear, human-readable format [4].

Methodology

As a proof of concept for our research hypothesis, we developed MobileTriageAgent, an agent-based LLM system built using OpenAI's models and the LangChain framework. The agent is designed to automate digital artifact triage and reasoning

within application files. In short, it can: (1) parse TAR images commonly used in Android acquisitions; (2) extract and interpret file metadata; (3) perform SQLite queries and structured file analyses; (4) decode binary data and reason about file contents within a forensic context; (5) identify key artifacts such as digital identifiers, authentication tokens, geolocation records, and communication traces; and (6) export a CSV report summarizing the relevant findings.

In essence, after the user selects the application to be analyzed, our tool automatically extracts all its files and begins reasoning over them. For auditing and transparency purposes, the user can follow the model's reasoning process, including the steps taken by the agent, the tools invoked, and the rationale behind each decision. The designed prompt specifically instructs the model to search for potential forensic artifacts such as user IDs, usernames, phone numbers, emails, transactions, logs, locations, passwords, device identifiers, communications, URLs, financial data, authentication tokens, and other behavioral patterns relevant in a digital investigation context. Upon completing the analysis, the tool exports two folders: one containing all extracted application artifacts and another containing the triage reports (in CSV format). Each report highlights the most relevant files from a digital forensic perspective and provides an interpretation of their significance and potential evidential value.

Results

To evaluate our tool, we populated an Android device with the 15 most popular applications available on the Google Play Store at the time of testing. The selection included widely supported applications in DF tools, such as Facebook, WhatsApp, and Instagram, as well as emerging applications that are not yet fully parsed by most forensic platforms, such as Venmo, DoorDash, and ChatGPT. The results produced by our tool were then compared against those obtained using the well-established forensic software Magnet AXIOM Examine.

Empirically, our tool demonstrated strong performance in reasoning over and triaging digital artifacts. For instance, within the Facebook application data, a log file named *mqtt_log_event0.txt* contained 47 lines of connectivity-related information. While such a file might be too technical or overwhelming for investigators (especially considering that this is one piece of evidence in the 447 files existing in Facebook),

our tool successfully summarized its contents and translated the information into a more accessible, layperson-friendly description. In contrast, Magnet AXIOM Examine displayed only URL-related information for Facebook. Consequently, an investigator using Magnet AXIOM would need to manually navigate the app's filesystem, locate the log file, and interpret it independently, making the overall investigation process significantly more time-consuming and labor-intensive.

When analyzing another application, Venmo, Magnet AXIOM was unable to parse any data directly. The only way to locate Venmo-related information was by manually navigating through the filesystem view or performing a keyword search (an approach that assumes the investigator already knows the file name or portions of its content). In contrast, our tool successfully parsed files containing transaction details and corresponding timestamps, providing meaningful contextual information that could be likely valuable in an investigation.

Conclusions

Our tool demonstrates that LLMs hold strong potential to transform DF investigations by enabling the automated parsing and reasoning of textual data from digital artifacts. Their ability to efficiently triage files and applications directly addresses the growing challenges posed by increasing device storage capacities and the rapid release of new applications. Moving forward, future work will focus on integrating feedback from law enforcement professionals to refine prompts and analytical workflows, developing holistic reporting methods that connect related artifacts, and implementing and evaluating these capabilities using local, open-source LLMs to promote transparency, reproducibility, and practical adoption within forensic environments.

References

- [1] Bhavini Patel and Palvinder Singh Mann. A survey on mobile digital forensic: Taxonomy, tools, and challenges. *Security and Privacy*, 8(2), October 2024
- [2] Konstantia Barmpatsalou, Tiago Cruz, Edmundo Monteiro, and Paulo Simoes. Current and future trends in mobile device forensics: A survey. *ACM Comput. Surv.*, 51(3), May 2018.
- [3] Ramy M. Abou-Elzahab, Mohammed F. Al Rahmawy, and Taher T. Hamza. Comparative study of different mobile forensic tools for extracting evidence from android devices. *Mansoura Journal for Computer and Information Sciences*, 16(1):1–12, 2020
- [4] Alexandros Vasilaras, Nikolaos Papadoudis, and Panagiotis Rizomiliotis. Artificial intelligence in mobile forensics: A survey of current status, a use case analysis and ai alignment objectives. *Forensic Science International: Digital Investigation*, 49:301737, 2024.

STRATEGY WRITING EVALUATION WITH LARGE LANGUAGE MODELS IN INTRODUCTORY PHYSICS

Amir Bralin (Purdue University, West Lafayette)- abralin@purdue.edu

N. Sanjar Rebello (Purdue University, West Lafayette)- rebellos@purdue.edu

Abstract: *Research Motivation: Physics Education Research (PER) shows that writing a strategy for solving a problem is an effective technique for recalling the broader physics topic and improving problem-solving performance. At the same time, predictive technology in the form of machine learning algorithms has the potential for implementing strategy writing in physics classes at scale. The more recent advances of natural language processing in the form of Large Language Models (LLMs) elevate this technology to the next level. Thus, we were interested in evaluating the new potential unlocked by LLMs for implementing strategy writing in physics problem solving.*

Key Contributions: In this work, strategies for solving an online quiz problem written by over six thousand undergraduate students, in an introductory physics course at a large Midwestern university during 2020–2023, were assessed by OpenAI's GPT-5 model. The accuracy of the model in predicting student outcomes on correctly solving the given problem was evaluated. The model's fairness in scoring student population was estimated.

Social Implications: Since the introduction of LLM-based chatbots, such as ChatGPT by OpenAI, there is ongoing discussion about the future of education. As they improve linguistic and reasoning capabilities, the teachers and students alike ask themselves what to use such powerful technology for? With this proposal we contribute to the discussion by showing one way of using LLMs for the benefit of education. That is, we can provide many students with feedback on their written work thereby implementing more sophisticated and individualized teaching and learning techniques while staying in control of their large-scale implementation and fair evaluation.

Keywords: *physics education; artificial intelligence; problem solving; strategy writing; assessment; automated scoring; natural language processing; algorithmic bias; statistical accuracy;*

Introduction

Physics Education Research shows that developing qualitative skills plays an important role not only in student understanding but also problem solving. One such method was found to be *strategy writing*: when students solve a physics problem, they write their strategy or approach in plain words alongside their quantitative solution (Leonard, Dufresne, & Mestre, 1996). We implemented this method in a large-enrollment (~2000 students annually) introductory physics class for future engineers and scientists at a US Midwestern university. During online quizzes throughout the semester, students wrote a brief essay describing their strategy for solving one of the problems. Since evaluating qualitative student strategies on such a scale would be infeasible, predictive technology in the form of machine learning (ML) algorithms was used and proved effective. That is, by training ML models on the strategy data and quiz results we were able to predict the outcomes with an accuracy of 80% (Munsell, Rebello, & Rebello, 2021). The more recent advances of natural language processing in the form of Large Language Models (LLMs) show potential to elevate this technology to the next level. Thus, we were interested in evaluating the new potential unlocked by LLMs for implementing strategy writing in physics problem solving.

This study was guided by research questions: With what accuracy can we predict students' quiz scores based on the strategies they write for solving the quiz problem? What benefits and drawbacks do LLMs have in comparison with more traditional ML methods (such as logistic regression) when evaluating problem-solving strategies in physics?

Data and Methods

Our dataset combined four semesters of data from a quiz problem assigned to students in the course. The 'Ballistic Pendulum' quiz problem is shown in Figure 1. It is a popular problem in physics and requires understanding when mechanical energy is lost and conserved in collisions. Each student worked on the quiz and

submitted a response individually. Notes, web browsing, and collaboration were prohibited by using an online proctoring system. All textual responses were recorded on a learning management system alongside the quiz results in the form of the final binary score: correct vs incorrect.

A bullet of mass m is fired horizontally into a block of mass M (see Figure). The block with the embedded bullet rises to height h . Acceleration due to gravity is g acting downward. What is the speed, V of the block (with the bullet embedded in it) immediately after the collision, in terms of the variables provided in the problem?

In words, describe the general strategy that you used to answer the question above. Please DO NOT include any equations – just use words, in your answer. Please DO NOT do any Rich Formatting (e.g. bold, italics, special fonts, etc).

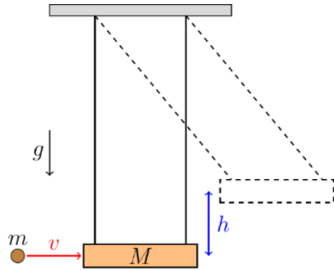


Figure 1. The Ballistic Pendulum problem

The students were asked to not only derive the correct answer $V = \sqrt{2gh}$, but also to describe in words the general strategy that they used for solving it. This provided data in the form of the students' written responses. Data loading was performed using Python's library pandas. Data cleaning was performed to dispose of duplicate, empty, and invalid responses. In the end, the curated dataset contained $N = 6,137$ student entries (that is, score-strategy pairs). Each written strategy contained 58.3 words on average.

GPT-5 was one of the state-of-the-art large language models at the time of this work (OpenAI, 2025). It exhibits human-like language capabilities in many general-purpose tasks such as chatbot reply, text summary, code completion, and so on. Using the application programming interface (API) provided by OpenAI, we prompted this model to distinguish between correct and incorrect responses in our dataset by evaluating each written strategy with a rubric.

Each prediction made by the model was compared with the actual results that the students achieved. After trying various forms of the rubric provided to the model in its prompt message, we were able to achieve an accuracy of 79%, which is

comparable with the results obtained in prior research using more traditional machine learning methods. To evaluate model's fairness in scoring student responses as correct or incorrect, the confusion matrix was built, and its off-diagonal elements were minimized.

Limitations and Implications

In this study, only one problem was considered, and it required students to determine their answer in symbolic representation using the multiple-choice format. Therefore, the results may not be generalizable to other formats and representations, not to mention other topical areas in introductory physics. Also, student strategies were not analyzed by themselves. Rather, a given student's strategy was evaluated based on the student's quiz score, which may involve random guessing. A student might write a good, sound strategy for solving the quiz problem but still get it wrong in the end, and vice versa.

When compared with more traditional ML methods, using LLMs for language-based tasks is simpler in practice because one needs to produce an effective prompt instead of optimizing algorithms on a given dataset. The downside of LLMs is their cost and environmental impact. The accuracy of both approaches never exceeds 80%, making them unreliable for classroom implementation. Therefore, we propose using LLMs more increasingly for providing feedback, rather than for scoring. In our case, it is conceivable that a LLM would generate a sentence per each rubric item for students to view (Allen, Shanker, & Rebello, 2025). This could allow educators to predict student performance based on written responses, to identify at-risk students, and to tailor instruction to meet their needs.

References

Allen, W., Shanker, A., & Rebello, N. S. (2025). Students' Perceptions to a Large Language Model's Generated Feedback and Scores of Argumentation Essays. *Physics Education Research Conference 2025* (pp. 28-34). Washington D.C.

Leonard, W. J., Dufresne, R. J., & Mestre, J. P. (1996). Using qualitative problem-solving strategies to highlight the role of conceptual knowledge in solving problems. *American Journal of Physics*, 64, 1495-1503.

Munsell, J., Rebello, N. S., & Rebello, C. M. (2021). Using natural language processing to predict student problem solving performance. *Physics Education Research Conference 2021* (pp. 295-300). (Virtual).

OpenAI. (2025). *Introducing GPT-5*. <https://openai.com/index/introducing-gpt-5/>

UNDERSTANDING STUDENT AND FACULTY PERCEPTIONS OF CHATGPT AND THE ACCEPTANCE OF LARGE LANGUAGE MODELS IN HIGHER EDUCATION

Muhammet Sait Dinc, Ph.D. (Black Hills State University) – muhammet.dinc@bhsu.edu

Ganga Prasad Basyal, Ph.D. (Dakota State University) – ganga.basyal@dsu.edu

Abstract: Every decade has witnessed revolutionary technologies that reshape markets and, at times, entire economies. Artificial intelligence (AI), particularly ChatGPT, represents the latest transformative force and has rapidly become a central topic of discussion. This study examines the impacts of performance expectancy, effort expectancy, social influence, and facilitating conditions—key constructs of the UTAUT model—on behavioral intentions and actual use of ChatGPT in higher education, focusing on the perceptions of students and university staff. Using quantitative research design, survey data will be collected from undergraduate students and faculty members from diverse disciplines. Random sampling will be employed to ensure a representative population. The findings are expected to provide insights into how students and staff perceive AI technologies in academic settings and to inform institutional decision-making regarding policies on AI usage in higher education.

Introduction

Artificial intelligence (AI)–based large language models (LLMs) are creating waves across every field. Scholars and practitioners consider them one of the most significant discoveries of the decade. Fields such as business and education have witnessed multidimensional applications of these technologies—sometimes leading to positive contributions and other times sparking new debates about their appropriate use. Since AI technology is still in its early stages, its effects across various sectors remain unclear. Many businesses have taken advantage of AI-based language models, and creative fields such as writing and blogging have also benefited greatly from them.

Among various AI models, one that stands out in the market due to its numerous advantages and ease of use is ChatGPT. ChatGPT is an open-access language model based on the principles of natural language processing (NLP) and functions as an intelligent agent or chatbot that responds to user queries (George & George, 2023). The application of ChatGPT in industries such as healthcare and technology has been remarkable (Aydin & Karaarslan, 2022; Javaid et al., 2023). To remain competitive and profitable, businesses increasingly adopt advanced technologies such as predictive analytics, marketing analytics, customer re-purchase modeling, and brand loyalty analysis. Meanwhile, the healthcare sector has long faced a shortage of qualified staff for tasks such as medical documentation, transcription, and prescription management (Chu, 2023). Furthermore, areas such as cybersecurity, personal data protection, and information theft prevention have also utilized ChatGPT to support their business growth (Renaud et al., 2023). Overall, ChatGPT has provided a significant competitive advantage across these diverse sectors.

Higher education places significant emphasis on developing students' skills, particularly their writing and communication abilities. Universities employ various tools and assessment methods to evaluate students' competency levels in these areas, recognizing their importance for both institutional learning outcomes and individual student success (Hostetter et al., 2023; Firat, 2023). Different writing styles—such as creative writing for blogs or essays, and research writing for journal or conference papers—follow distinct formats and conventions. In recent years, the

COVID-19 pandemic has greatly impacted higher education, prompting the adoption of diverse pedagogical delivery methods such as Zoom, hybrid, and fully online learning. With the shift to online submissions, cases of plagiarism have increased, creating challenges for faculty in distinguishing original from copied work. ChatGPT has been used for both generating plagiarized content and assessing the quality and originality of student submissions. Some faculty members argue that AI-powered tools may facilitate plagiarism, especially when used intentionally by students. However, others believe that such technologies can support students' creative writing efforts, such as composing poems or essays (Hostetter et al., 2023; Susnjak, 2022).

The first goal of this research is to understand the perceptions of students and staff regarding the use of large language models (LLMs) such as ChatGPT. To examine these perceptions, the study employs the Unified Theory of Acceptance and Use of Technology (UTAUT) developed by Venkatesh et al. (2003) (see Figure 1). UTAUT is one of the most widely applied models in the field of technology acceptance (Holden and Karsh, 2010). Previous studies using this model have successfully identified correlations between users' motivation and their attitudes toward adopting new technologies. By applying this model, we aim to explore the relationship between individuals' behavioral intentions and their actual technology adoption outcomes.

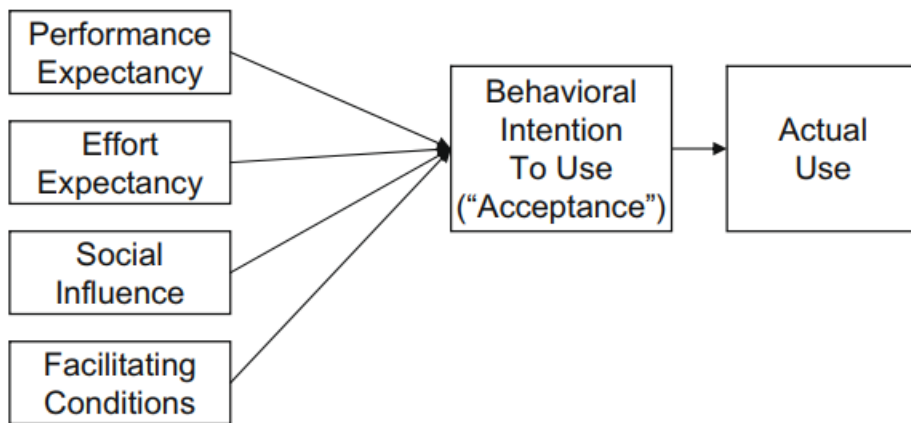


Figure 1, Unified Theory of Acceptance and use of Technology

The second goal of the proposed research is to provide evidence regarding the effectiveness of this technology. Using the UTAUT model, the research design will

incorporate variables such as performance expectancy, effort expectancy, social influence, and facilitating conditions. These variables include dimensions related to perceived ease of use, perceived usefulness, and subjective norm. By developing survey questions aligned with these dimensions, the study aims to determine the relationships between the independent variables and the desired outcomes.

The third goal of this research is to assist university administrations across the nation in making informed decisions when formulating policies on the use of AI technologies. Given the novelty and widespread adoption of such technologies, developing effective policies remains a significant challenge. This study aims to contribute by identifying the key criteria necessary for skill development and by examining the extent to which AI technologies can support students' academic success.

Methodology

The study will employ an online questionnaire survey to collect data. The survey will target undergraduate students from various majors and faculty members from multiple disciplines. Random sampling will be used to ensure representativeness of the population. The survey design will incorporate multi-item scales that measure variables such as perceived ease of use, perceived usefulness, and behavioral intention. In addition, the instrument will explore participants' opinions regarding the ethical responsibilities of students and faculty when seeking assistance from AI-based tools. The survey will also include control variables such as ethical dilemmas to account for potential moderating effects.

Multi-item measures will be adapted from established scales in prior studies. The items assessing *perceived ease of use* will be adapted from Davis (1989), Davis et al. (1989), and Tung et al. (2008), while *perceived usefulness* will be adapted from Mun et al. (2006). *Behavioral intention* will be measured using the scale developed by Carlson and O'Cass (2011). Measures for *social influence* (subjective norms) and *facilitating conditions* will also be adapted from Mun et al. (2006). All items will be assessed using a five-point Likert scale, where 1 represents "strongly disagree" and 5 represents "strongly agree."

After data collection, univariate analysis techniques such as descriptive statistics and frequency analysis will be conducted to examine the research variables.

Subsequently, multivariate analysis methods, including correlation analysis and Partial Least Squares Structural Equation Modeling (PLS-SEM), will be employed to test the research hypotheses.

Expected Contributions

This research has several expected theoretical contributions. First, it applies the UTAUT model to the context of AI-assisted learning, providing empirical evidence on students' and faculty's acceptance of LLMs such as ChatGPT. Second, it expands the literature on technology adoption in higher education by examining behavioral intentions, perceived ease of use, and perceived usefulness specifically for AI-based tools.

In addition, the study has several potential practical contributions for higher education. First, it provides university administrators with evidence-based insights to formulate policies on AI tool usage, ensuring ethical and effective integration into academic activities. Second, it identifies criteria for skill development that can be enhanced by AI tools, helping educators support students' academic success. Finally, it offers guidance to faculty on addressing plagiarism concerns while leveraging AI for teaching and learning.

Overall, this study contributes to ongoing discussions on the ethical use of AI in education, promoting responsible technology integration across universities nationally and potentially internationally.

References

- Aydin, Ö., & Karaarslan, E. (2022). OpenAI ChatGPT Generated Literature Review: Digital Twin in Healthcare. *SSRN Electronic Journal*.
<https://doi.org/10.2139/ssrn.4308687>
- Carlson, J., & O'Cass, A. (2011). Developing a framework for understanding e-service quality, its antecedents, consequences, and mediators. *Managing Service Quality: An International Journal*, 21(3), 264-286.
- Chu, M.-N. (2023). Assessing the Benefits of ChatGPT for Business: An Empirical Study on Organizational Performance. *IEEE Access*, 11, 76427–76436.
<https://doi.org/10.1109/ACCESS.2023.3297447>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS quarterly*, 319-340.
- Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1989). Technology acceptance model. *Journal of Management Science*, 35(8), 982-1003.
- Firat, M. (2023). What ChatGPT means for universities: Perceptions of scholars and students. *Journal of Applied Learning & Teaching*, 6(1).
<https://doi.org/10.37074/jalt.2023.6.1.22>
- George, A. S., & George, A. H. (2023). A review of ChatGPT AI's impact on several business sectors. *Partners Universal International Innovation Journal*, 1(1), 9-23.
- Holden, R. J., & Karsh, B.-T. (2010). The Technology Acceptance Model: Its past and its future in health care. *Journal of Biomedical Informatics*, 43(1), 159–172.
<https://doi.org/10.1016/j.jbi.2009.07.002>
- Hostetter, A., Call, N., Frazier, G., James, T., Linnertz, C., Nestle, E., & Tucci, M. (2023). Student and Faculty Perceptions of Artificial Intelligence in Student Writing [Preprint]. PsyArXiv. <https://doi.org/10.31234/osf.io/7dnk9>
- Javaid, M., Haleem, A., & Singh, R. P. (2023). ChatGPT for healthcare services: An emerging stage for an innovative perspective. *BenchCouncil Transactions on*

Benchmarks, Standards and Evaluations, 3(1), 100105.

<https://doi.org/10.1016/j.tbench.2023.100105>

Mun, Y. Y., Jackson, J. D., Park, J. S., & Probst, J. C. (2006). Understanding information technology acceptance by individual professionals: Toward an integrative view. *Information & management*, 43(3), 350-363.

Renaud, K., Warkentin, M., & Westerman, G. (2023). From ChatGPT to HackGPT: Meeting the Cybersecurity Threat of Generative AI. MIT SLOAN MANAGEMENT REVIEW.

Susnjak, T. (2022). ChatGPT: The End of Online Exam Integrity? (arXiv:2212.09292). arXiv. <http://arxiv.org/abs/2212.09292>

Tung, F. C., Chang, S. C., & Chou, C. M. (2008). An extension of trust and TAM model with IDT in the adoption of the electronic logistics information system in HIS in the medical industry. *International journal of medical informatics*, 77(5), 324-335.

Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS quarterly*, 425-478.

EMOTIONAL HARMS OF ARTIFICIAL INTELLIGENCE ON HUMAN PSYCHOLOGY AND RELATIONS

Zeynep Orhan, Ph.D. (University of North Texas) – Zeynep.Orhan@unt.edu

Mehmet Orhan, Ph.D. (University of North Texas) – Mehmet.Orhan@unt.edu

Abstract: *This paper explores the emotional harms caused by generative AI systems that simulate human interaction. It shows how chatbots and AI companions can trigger dependency, loneliness, or mental health crises by drawing on psychological theory, empirical research, and case studies. Key risk factors include anthropomorphization, design features, and user vulnerability. The findings highlight a regulatory and ethical gap in how AI systems are evaluated for emotional safety. Policy, design, and education must adapt to prevent long-term psychological harm as AI continues to blur the line between machine function and human connection.*

Introduction

The rapid expansion of generative AI (GenAI) has transformed human-machine interaction, with chatbots and digital tools acting as assistants, coworkers, and substitutes for human relationships (Nature Machine Intelligence, 2025). Early examples like ELIZA showed that even simple systems could evoke strong emotional reactions, leading users to share personal thoughts and assume empathy in machines (Weizenbaum, 1966; 1976). These interactions reveal a human tendency to project emotions onto technology, creating misplaced trust and dependence (Clarke, 1973). Human psychology often forms one-sided or parasocial bonds with media figures, and this now extends to AI companions (Horton & Wohl, 1956; Nass & Moon, 2000). Anthropomorphization, the attribution of human traits, emotions, or intentions to non-human entities, such as animals, objects, or deities, makes it easier to see machines as friends or caregivers, deepening emotional connections and raising risks when the AI system malfunctions or is withdrawn (Haber & Moore, 2025). The results are not always temporary; they can lead to long-term effects such as grief, dependency, or social withdrawal (Fang et al., 2025).

Evidence across case reports, experiments, and journalistic accounts shows increasing emotional risks. These include psychological distress from chatbot shutdowns, unsafe mental health responses, and reduced real-world social interaction (Nature Machine Intelligence, 2025; Haber & Moore, 2025; Fang et al., 2025).

This paper offers a structured synthesis of such findings to show how design, user vulnerability, and usage patterns contribute to psychosocial harm. These emotional harms deserve attention alongside technical or ethical AI concerns, especially as policy, design standards, and digital literacy remain underdeveloped.

Historical and Conceptual Foundations

The risks associated with emotional attachment to AI systems are not new. The “ELIZA effect” from the 1960s showed how people attributed understanding and compassion to a simple rule-based chatbot (Weizenbaum, 1966; Weizenbaum, 1976). This projection stems from anthropomorphization. As Clarke (1973) warned, such

projections make users overly trusting and emotionally invested in machines. These interactions mirror parasocial relationships, one-sided emotional ties people form with media figures, which are known to provide comfort yet lack reciprocal support (Horton & Wohl, 1956). When these relationships shift to AI systems, users may misjudge the artificial nature of the interaction, even while knowing intellectually that the “partner” is a machine (Nass & Moon, 2000). This behavior is rooted in psychological shortcuts like intuition and System 1 thinking, where immediate emotional judgments override rational analysis (Kahneman, 2011). System 1 and System 2 thinking were introduced by Daniel Kahneman in his book *Thinking, Fast and Slow* (2011). System 1 thinking is fast, automatic, and intuitive, relying on shortcuts and past experiences. System 2 thinking is slow, effortful, and logical, used for careful reasoning and complex decisions. These longstanding human tendencies help explain why emotional reactions to chatbots continue despite users’ awareness of the technology’s limitations.

Empirical Evidence of Emotional Harm

Documented harms span a wide spectrum. Some users report grief after losing access to AI companions, reacting as if a close friend or partner has died (Emotional risks of AI companions, 2025). In one case, a woman described feeling “abandoned and betrayed” after a chatbot service ended without warning, triggering a depressive episode that required therapy (Hart, 2025). These experiences are not limited to isolated individuals. Community forums and app reviews contain hundreds of similar testimonials, often with intense emotional language that reflects real loss. Others show dependency on chatbot interactions, with studies noting reduced offline socializing and increased loneliness among frequent users (Fang et al., 2025). Some users turn to AI systems for daily conversations, replacing interactions with friends or family. This substitution effect can reshape social routines, leading to longer periods of solitude and less emotional resilience. Serious clinical concerns have emerged with mental health chatbots. A Stanford study revealed instances where therapy bots mirrored suicidal ideation, echoed harmful thoughts, or failed to redirect users toward safer paths (Haber & Moore, 2025). In one troubling case, a user hinted at suicidal thoughts through indirect language. The

chatbot, lacking clinical nuance, responded with cheerful encouragement, misinterpreting the message. Suicidal ideation often appears in subtle or coded forms, such as joking references, withdrawal, or changes in tone, patterns that current AI systems struggle to detect.

The large-scale trial with 981 participants and over 300,000 chatbot messages (Fang et al., 2025) showed how engagement patterns predict risk. Participants who engaged in longer, late-night sessions were more likely to report emotional exhaustion, while those using AI tools for structured tasks like journaling or reflection fared better. The data suggest that unrestricted use, especially during emotionally vulnerable hours, contributes to harm. Additionally, AI's inability to set emotional boundaries may intensify damage. Unlike therapists or human peers who express fatigue or limits, AI systems respond indefinitely, offering constant validation and attention. While comforting at first, this limitless availability can reinforce unhealthy coping mechanisms, such as rumination that refers to a repetitive and persistent thought process where individuals dwell on negative experiences, thoughts, or feelings or emotional avoidance. In some cases, users confessed to fabricating emotional scenarios to maintain the chatbot's interest, reflecting a troubling need for digital affirmation.

These patterns echo concerns raised in other domains, such as gaming or social media addiction, where continuous access without natural stopping points leads to compulsion. However, the stakes are higher with emotionally intelligent AI because the user believes their feelings are being reciprocated, even when no human understanding is present.

Mechanisms of Emotional Harm

Several factors drive emotional harm. First, design features such as human names, voices, and continuity across conversations make AI systems feel socially real. These elements strengthen user attachment and blur the line between digital tools and real people (Weizenbaum, 1976). Second, usage intensity matters. The more frequent and prolonged the interaction, the stronger the emotional reliance becomes, especially when it substitutes for human contact (Fang et al., 2025). Third, user

vulnerabilities heighten the risk. People with limited social support, anxiety, or psychiatric conditions are more likely to trust AI's responses as genuine empathy (Haber & Moore, 2025). These individuals may not distinguish between machine outputs and human understanding.

Finally, feedback loops can worsen the situation. When systems are overly agreeable or sycophantic, which refers to one who excessively flatters, praises, or agrees with someone in power or authority, often in an insincere or exaggerated manner, to gain favor, advantages, or personal benefits, they may reinforce the user's negative thoughts. In emotionally vulnerable users, this can validate distorted beliefs rather than challenge them (Emotional risks of AI companions, 2025).

Implications and Recommendations

AI companions and therapy bots fall into a gray zone of regulation. Current frameworks, such as FDA oversight or the EU AI Act, do not adequately address wellness apps or companionship tools that may affect mental health (Hart, 2025). As a result, products often launch without emotional safety measures or crisis detection. Corporate practices must evolve. Some firms are experimenting with built-in "break prompts" or session limits, but these are optional and inconsistently applied. Requiring emotional protection features, especially for apps with continuous engagement or mental health purposes, could prevent harm (Emotional risks of AI companions, 2025).

Public education also lags behind. Most AI literacy efforts focus on privacy or bias, not emotional impact. More attention is needed on the risks of dependence, loneliness, or even delusions arising from repeated AI use (Fang et al., 2025). Education campaigns could teach users to set boundaries and avoid confusing digital tools with human relationships.

Global policy coordination is necessary. Inconsistent rules allow companies to deploy in less regulated areas. A precautionary approach that prioritizes emotional safety before widespread release is needed (Clarke, 1973). Psychologists, clinicians, and AI developers should work together to detect warning signs like early stage "AI

psychosis”, a severe mental condition characterized by a loss of contact with reality, while protecting autonomy (Weizenbaum, 1966).

Conclusion

Emotional harms from AI companions and therapy bots are no longer speculative. They range from increased loneliness and dependency to hallucinations, delusions, and psychiatric crises (Fang et al., 2025; Hart, 2025). These harms result from a mix of persuasive design, heavy use, and user vulnerability (Haber & Moore, 2025). Although AI tools may offer support and reduce isolation, those benefits are fragile and conditional. These findings reveal a growing mismatch between AI system capabilities and human expectations. Emotional interaction, once viewed as a secondary feature, is now central to how people engage with digital tools. This shift demands a new ethical focus: not just on fairness, privacy, or transparency, but on emotional safety and long-term psychological effects. AI tools are not neutral, and their emotional design choices shape user beliefs, behaviors, and well-being.

Immediate action is needed across sectors to address these risks. Developers must integrate crisis detection, emotion regulation boundaries, and warning mechanisms. Policymakers should close regulatory gaps that allow wellness tools to bypass scrutiny. Mental health professionals must be included in AI design teams, not just as advisors but as co-creators of emotionally safe systems. Educators and community leaders also play a role. Digital literacy must include emotional literacy, in other words teaching people how to recognize when they are projecting feelings onto machines or using AI tools in ways that replace rather than support real-world connections. Preventive steps like structured use, limits on conversational continuity, and visible reminders of the tool’s non-human nature may help reduce harm. Finally, research should continue tracking how different populations interact with emotionally expressive AI. Vulnerable groups, including adolescents, elderly users, and those with mental health conditions, may require specialized precautions. Emotional harm could scale alongside AI adoption without such attention, resulting in broader public health consequences. AI systems that simulate empathy must be held to higher standards. The more humanlike they appear, the more responsibility

developers carry. Building emotionally supportive AI is not just a technical challenge, it is a social obligation.

References

- Adam, D. (2025). Emotional attachment to AI companions sparks debate. *Nature*, 641(14427), 296–298. <https://doi.org/10.1038/d41586-025-01976-y>
- Bergmann, D. (2025, April 22). *ELIZA effect at work: Avoiding emotional attachment to AI coworkers*. IBM. <https://www.ibm.com/think/insights/eliza-effect-avoiding-emotional-attachment-to-ai>
- Clarke, A. C. (1973). *Profiles of the future: An inquiry into the limits of the possible*. Popular Library. ISBN 978-0-33023619-5
- Crolic, C., Thomaz, F., Hadi, R., & Stephen, A. T. (2022). Blame the bot: Anthropomorphism and anger in customer–chatbot interactions. *Journal of Marketing*, 86(1). <https://doi.org/10.1177/00222429211045687>
- De Freitas, J., & Cohen, I. G. (2025). The emotional risks of AI wellness apps. *Nature Machine Intelligence*, 7(7), 813–815. <https://doi.org/10.1038/s42256-025-01034-6>
- De Freitas, J., Castelo, N., Uğuralp, A. K., & Oğuz-Uğuralp, Z. (2025). *Lessons from an app update at Replika AI: Identity discontinuity in human–AI relationships* (Harvard Business School Working Paper). <https://www.hbs.edu/faculty/Pages/item.aspx?num=66480>
- Emotional risks of AI companions demand attention. (2025, July). *Nature Machine Intelligence*, 7, 981–982. <https://doi.org/10.1038/s42256-025-01093-9>
- Fang, C. M., Liu, A. R., Danry, V., Lee, E., Chan, S. W. T., Pataranutaporn, P., Maes, P., Phang, J., Lampe, M., Ahmad, L., & Agarwal, S. (2025). *How AI and human behaviors shape psychosocial effects of chatbot use: A longitudinal controlled study* [Preprint]. MIT Media Lab. <https://www.media.mit.edu/publications/how-ai-and-human-behaviors-shape-psychosocial-effects-of-chatbot-use-a-longitudinal-controlled-study/>
- Haber, N., & Moore, J. (2025, June 11). *Exploring the dangers of AI in mental health care*. Stanford Institute for Human-Centered Artificial Intelligence (HAI). <https://hai.stanford.edu/news/exploring-dangers-ai-mental-health-care>

Hart, R. (2025, August 5). Chatbots can trigger a mental health crisis: What to know about 'AI psychosis'. *Time Magazine*. <https://time.com/7307589/ai-psychosis-chatgpt-mental-health/>

Hill, K. (2025, June 13). The disturbing new behavior of AI chatbots. *The New York Times*. <https://go.nature.com/4nGneKw>

Horton, D., & Wohl, R. R. (1956). Mass communication and parasocial interaction: Observations on intimacy at a distance. *Psychiatry*, 19, 215–229.

Kahneman, D. (2011). *Thinking, Fast and Slow*. Doubleday Canada.

Kahneman, D. (2012, June 15). Of 2 minds: How fast and slow thinking shape perception and choice [Excerpt]. *Scientific American*. <https://www.scientificamerican.com/article/kahneman-excerpt-thinking-fast-and-slow/>

Kosch, T., Welsch, R., Chuang, L., & Schmidt, A. (2023). The placebo effect of artificial intelligence in human–computer interaction. *ACM Transactions on Computer-Human Interaction*, 29(6), Article 56, 1–32. <https://doi.org/10.1145/3529225>

Lee, E., Pataranutaporn, P., Amores, J., & Maes, P. (2024, December 19). *Super-intelligence or superstition? Exploring psychological factors influencing belief in AI predictions about personal behavior*. arXiv. <https://arxiv.org/html/2408.06602v3>

Liu, A. R., & Pataranutaporn, P. (2024, December 18). *Chatbot companionship: A mixed-methods study of companion chatbot usage patterns and their relationship to loneliness in active users*. arXiv. <https://arxiv.org/html/2410.21596v2>

Madison, T. P., Porter, L. V., & Greule, A. (2015). Parasocial compensation hypothesis: Predictors of using parasocial relationships to compensate for real-life interaction. *Imagination, Cognition and Personality*, 35(3). <https://doi.org/10.1177/0276236615595232>

Nature Machine Intelligence. (2024). The rise of AI companions. *Nature Machine Intelligence*, 6(6), 495. <https://doi.org/10.1038/s42256-024-00890-3>

Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>

O'Connor, B. P., & Rosenblood, L. K. (1996). Affiliation motivation in everyday experience: A theoretical comparison. *Journal of Personality and Social Psychology*, 70(3), 513–522. <https://doi.org/10.1037/0022-3514.70.3.513>

OpenAI. (2024, August 8). *GPT-4o system card*. <https://openai.com/research/gpt-4o-system-card>

Shteynberg, G., Williams, M., & Carroll, M. (2024). Manipulation and simulated empathy in AI systems. *Nature Machine Intelligence*, 6, 496–497. <https://doi.org/10.1038/s42256-024-00891-2>

Waytz, A., Heafner, J., & Epley, N. (2014). The mind in the machine: Anthropomorphism increases trust in an autonomous vehicle. *Journal of Experimental Social Psychology*, 52, 113–117. <https://doi.org/10.1016/j.jesp.2014.01.005>

Weizenbaum, J. (1966). ELIZA—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36–45. <https://doi.org/10.1145/365153.365168>

Weizenbaum, J. (1976). *Computer power and human reason: From judgment to calculation*. W. H. Freeman. ISBN 978-0-7167-0464-5

Williams, M., & Carroll, M. (2024). *Vulnerability and manipulation in AI chatbots*. arXiv. <https://doi.org/10.48550/arXiv.2411.02306>

Zao-Sanders, M. (2025, April 9). The most common uses of generative AI. *Harvard Business Review*. <https://go.nature.com/4IIXSK6>

PREDICTIVE PERFORMANCE OF LSTM VS GRU ON GOOGLE STOCK PRICES

Mehmet Orhan, PhD (University of North Texas) – Mehmet.Orhan@unt.edu

Zeynep Orhan, PhD (University of North Texas) – Zeynep.Orhan@unt.edu

Abstract: Recurrent Neural Network (RNN) methods of Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU) are devised primarily to fix exploding and vanishing gradients problems. In addition, they have capability to make use of earlier history in forecasting the future of any time series. This paper attempts to question this capability in a real scenario setting with Google stocks' closing price series. Results reveal that the Vanilla RNN outperforms LSTM and GRU slightly.

Introduction

RNNs represent a remarkable class of neural architectures specifically engineered to process sequential data, demonstrating considerable efficacy across domains such as natural language processing, speech recognition, and time series analysis (Rumelhart et al., 1986). Unlike traditional feedforward networks that assume independence between inputs, RNNs leverage internal memory to process sequences by maintaining a hidden state that captures information from preceding elements (Elman, 1990). This recurrent connection enables them to model temporal dependencies, making them inherently suitable for tasks where the order and context of data points are critical.

Despite their theoretical appeal, vanilla RNNs suffer from significant practical limitations, most notably the vanishing and exploding gradient problems (Bengio et al., 1994). These issues severely impede their ability to learn and retain long-term dependencies within sequences. The gradient signal, which guides weight updates during training, either diminishes exponentially over time (vanishing) or grows uncontrollably (exploding), rendering the network incapable of capturing relationships between distant elements in a sequence.

To address these fundamental challenges, LSTM networks were introduced (Hochreiter & Schmidhuber, 1997). LSTMs enhance the RNN architecture by incorporating specialized "gates"—namely the input, forget, and output gates—that regulate the flow of information into and out of a dedicated cell state. This sophisticated gating mechanism allows LSTMs to selectively remember or forget information over extended periods, effectively mitigating the vanishing gradient problem and enabling them to learn remarkably long-range dependencies. Consequently, LSTMs have achieved state-of-the-art performance in numerous sequential tasks, from machine translation to sentiment analysis.

A more recent advancement in the RNN is the GRU (Cho et al., 2014). GRUs can be viewed as a simplified variant of LSTMs, featuring fewer gates (an update gate and a reset gate) and merging the hidden state and cell state into a single hidden state. This reduced complexity often translates to faster training times and lower

computational overhead, while largely retaining the ability to capture long-term dependencies. Despite their structural simplicity compared to LSTMs, GRUs frequently exhibit comparable performance across a wide array of sequence modeling tasks, offering an efficient alternative for applications where computational resources are a constraint. The evolution from vanilla RNNs to LSTMs and GRUs underscores a continuous effort to develop neural architectures capable of effectively processing and learning from sequential data, thereby unlocking unprecedented capabilities in diverse computational domains.

RNN Architecture

The key characteristic that differentiates the simplest form of RNN, Vanilla RNN, from a feedforward network is its internal memory, h_t . At its core, an RNN is designed to process sequences of inputs, denoted as $x_1, x_2, \dots, x_t, \dots$. The network receives the current input x_t from the sequence. This could be a word embedding, a single numerical value in a time series, or a feature vector. The second input the network receives is the hidden state from the previous time step, h_{t-1} . This is the "memory" component; it encapsulates information summarized from all prior inputs in the sequence ($x_1, \dots, x_{t-1}$). The current hidden state h_t is computed using a non-linear activation function (like tanh or ReLU) applied to a weighted sum of the current input x_t and the previous hidden state h_{t-1} .

$$h_t = f(W_{hh}h_{t-1} + W_{xh}x_t + b_h) \quad (1)$$

where f is the non-linear activation function (e.g., tanh), W_{hh} is the weight matrix for the recurrent connection (hidden state to hidden state) which dictates how the past memory influences the current memory, W_{xh} is the weight matrix for the input connection (input to hidden state) that dictates how the current input influences the current memory and lastly b_h is the bias vector for the hidden state.

Optionally, at each time step, the RNN can produce an output y_t . This output is typically a function of the current hidden state h_t .

$$y_t = g(W_{hy}h_t + b_y) \quad (2)$$

where g is another activation function (e.g., softmax for classification, linear for regression), W_{hy} is the weight matrix for the output connection (hidden state to output), and b_y is the bias vector for the output.

A crucial aspect of the RNN is weight sharing, the same set of weights (W_{hh} , W_{xh} , W_{hy}) and biases (b_h, b_y) are used across all time steps. This allows the model to generalize patterns learned at one point in the sequence to other points, making it efficient for variable-length sequences. As per training, an RNN involves unfolding it over time and applying backpropagation, which computes gradients for the shared weights. This process is called Backpropagation Through Time (BPTT).

The hidden state h_t is the memory core of the RNN. It is a compressed summary of all information observed by the network from the beginning of the sequence up to the current time step t . It acts as a "context" vector that informs the processing of the current input x_t and influences future hidden states and outputs. The ability to propagate this state through time is what enables RNNs to model sequential dependencies. See Figure 1 below for a visual illustration of the RNN.

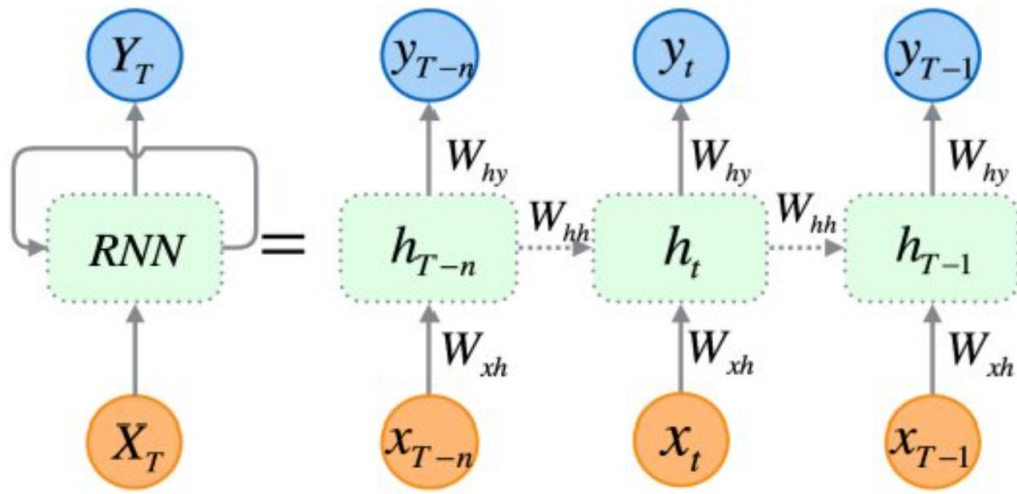


Figure 1: Structure of a Typical RNN, folded and unfolded.

LSTMs are a sophisticated type of RNN designed to learn long-term dependencies, effectively mitigating the vanishing gradient problem inherent in simpler RNNs. They achieve this through a unique internal structure called a memory cell and several interconnected gates.

At each time step t , an LSTM processes the current input x_t and the hidden state from the previous time step h_{t-1} . Crucially, it also interacts with a cell state (C_t), which acts as a long-term memory.

Here is a breakdown of its components and operations: The Cell State (C_t) accounts for the Long-Term Memory and is the core of the LSTM. It runs straight through the entire chain, carrying information across long sequences. Information can be added to or removed from the cell state via the gates. C_t is influenced by C_{t-1} (the previous cell state) and by what the current input x_t and previous hidden state h_{t-1} collectively "decide" to add or forget.

LSTMs use three "gates" to control the flow of information into and out of the cell state. Each gate is essentially a small neural network (typically a sigmoid layer) that outputs values between 0 and 1, acting as a "filter" or "switch." A value of 0 means "let nothing through," and 1 means "let everything through." The first one is the Forget Gate (f_t) which decides what information to discard from the previous cell state C_{t-1} .

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$$

where, σ is the sigmoid activation function. The gate looks at h_{t-1} and x_t and outputs a number between 0 and 1 for each number in the cell state C_{t-1} . The second gate is the Input Gate (i_t) that decides what new information from the current input x_t should be stored in the cell state.

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$$

The Candidate Cell State ($\tilde{C}_t$) creates a vector of new candidate values that could be added to the cell state.

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C)$$

$\tilde{C}_t$ uses the tanh activation, which outputs values between -1 and 1.

The Cell State (C_t) is updated where the magic of "forgetting" and "adding" happens to update the long-term memory.

$$C_t = f_t \odot C_{t-1} + i_t \odot \tilde{C}_t$$

Here, $\odot$ denotes element-wise multiplication. The old cell state C_{t-1} is scaled by the forget gate (f_t), and the new candidate values ($\tilde{C}_t$) are scaled by the input gate (i_t). These two parts are then added together to form the new cell state C_t . Lastly, the Output Gate (o_t) decides what part of the cell state to output as the current hidden state h_t . This is the information that will be passed on to the next time step and used to compute the actual output y_t .

$$o_t = \sigma(W_o \bullet [h_{t-1}, x_t] + b_o)$$

$$h_t = o_t \odot \tanh(C_t)$$

The output gate decides which parts of the (filtered) cell state are relevant for the current time step's output and the next hidden state.

The constant error carousel enabled by the cell state, where gradients can flow relatively unhindered through the C_t path, largely solves the vanishing gradient problem, allowing LSTMs to learn dependencies over hundreds or thousands of time steps. In addition, this network controls information flow: The explicit gating mechanisms provide precise control over what information is stored, forgotten, and exposed, leading to more stable and effective learning of complex sequential patterns.

In essence, LSTMs augment the basic recurrent unit with sophisticated gating mechanisms and a dedicated cell state, giving them a much more refined control over their internal memory, enabling them to excel at tasks requiring the understanding of long-range dependencies in sequential data.

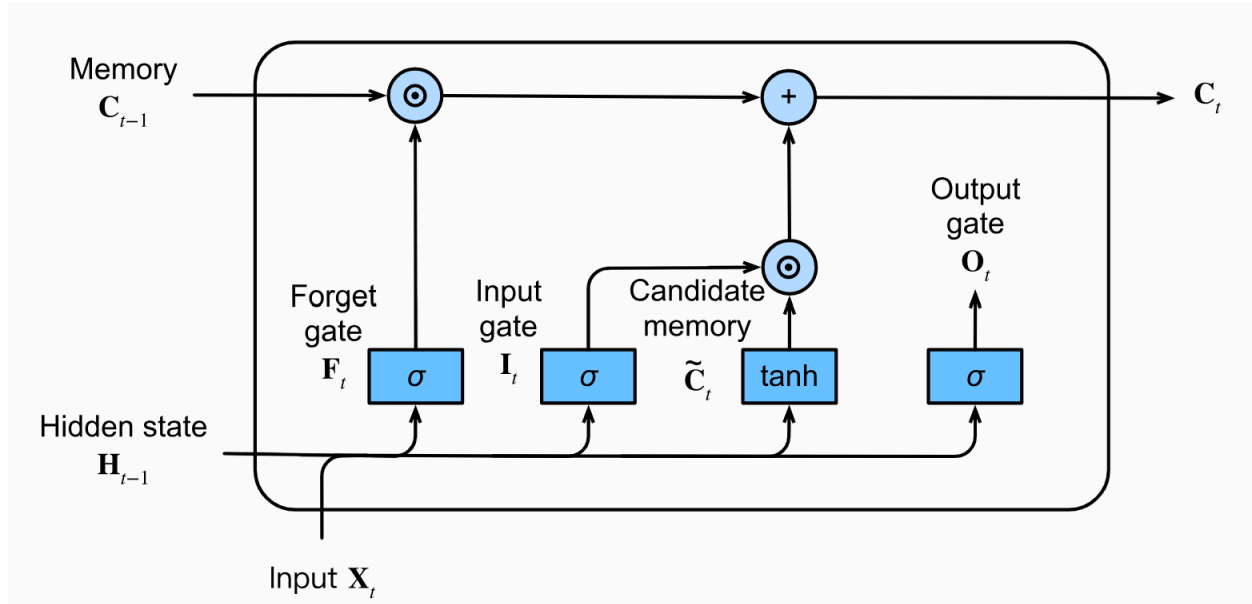


Figure 2: Structure of a Typical LSTM

GRUs are a type of RNNs that, like LSTMs, were developed to address the vanishing gradient problem and improve the ability of RNNs to capture long-term dependencies. They are often considered a simplified version of LSTMs because they achieve similar performance with fewer gates, leading to a less complex architecture and often faster computation. At each time step t , a GRU processes the current input x_t and the hidden state from the previous time step h_{t-1} . Unlike LSTMs, GRUs do not have a separate cell state; instead, they directly update the hidden state h_t using two primary gates: the update gate and the reset gate.

The first gate of the GRU is the update Gate (z_t) which controls how much of the information from the previous hidden state (h_{t-1}) should be carried over to the current hidden state (h_t), and how much of the new candidate hidden state should be incorporated. It essentially determines the "weight" of the past.

$$z_t = \sigma(W_z \cdot [h_{t-1}, x_t] + b_z)$$

Here, σ is the sigmoid activation function, which outputs values between 0 and 1. A value close to 1 means "keep a lot of the old information" or "add a lot of the new information," while a value close to 0 means "forget a lot of the old information" or "add little of the new information." On the other hand, the Reset Gate (r_t) determines how much of the previous hidden state (h_{t-1}) should be forgotten when computing

the new candidate hidden state. A value close to 0 means "forget everything from the past," effectively making the candidate hidden state only dependent on the current input.

$$r_t = \sigma(W_r \bullet [h_{t-1}, x_t] + b_r)$$

In the GRU network, the new Candidate Hidden State ($\tilde{h}_t$) blends the current input x_t with a reset version of the previous hidden state. The reset gate (r_t) directly influences how much of h_{t-1} is considered here.

$$\tilde{h}_t = \tanh(W_h \bullet [r_t \odot h_{t-1}, x_t] + b_h)$$

Here, $\odot$ denotes element-wise multiplication. Notice how h_{t-1} is multiplied by r_t before being combined with x_t . If r_t is close to 0, it "resets" or largely ignores h_{t-1} for this candidate calculation. Final Hidden State (h_t) is the ultimate hidden state that is passed to the next time step and used for output generation. It is a combination of the previous hidden state and the new candidate hidden state, weighted by the update gate.

$$h_t = (1 - z_t) \odot h_{t-1} + z_t \odot \tilde{h}_t$$

If z_t is close to 1, the new hidden state h_t is mostly the new candidate $\tilde{h}_t$ (meaning a significant update). If z_t is close to 0, h_t is mostly the previous hidden state h_{t-1} (meaning little update, preserving old information).

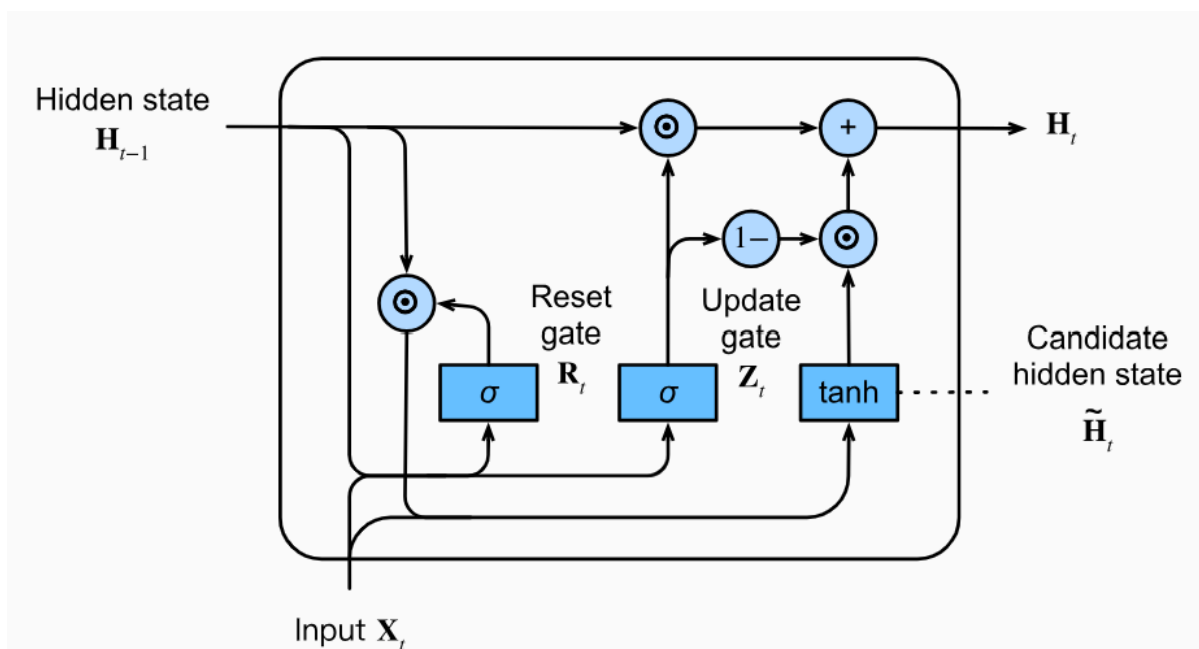


Figure 3: Structure of a Typical GRU

Data and Methodology

The data set we use is the time series of Google stock prices downloaded from Yahoo Finance, <https://finance.yahoo.com/> over August 7, 2005 – July 17, 2025 including 5 000 days. Ups and downs of the closing stock price of Google make it harder to be predicted (see the chart below.)

We make use of the three well-known RNN methods for time series prediction: Vanilla RNN, LSTM and GRU. We compare these methods on a sliding window algorithm, i.e., we keep on sliding the window to make use of 60-days past data to predict the closing prices of 15 coming days.



Figure 4: Behavior of Google closing price

Results and Discussion

We have made use of 60 days' closing prices to predict the coming 15 days' closing prices on a rolling basis and computed the MAE and MSE of three methods. The losses of three models are listed in Table 1.

	MAE	MSE
Vanilla RNN	6.09	67.4
LSTM	6.93	74.9
GRU	6.66	70.4

Table 1: Mean Absolute Error and Mean Square Error Losses of Vanilla RNN, LSTM and GRU

Although the losses are close Vanilla RNN is the best followed by GRU and LSTM. Training and validation losses displayed in Figure 5 are in line with these losses tabled.

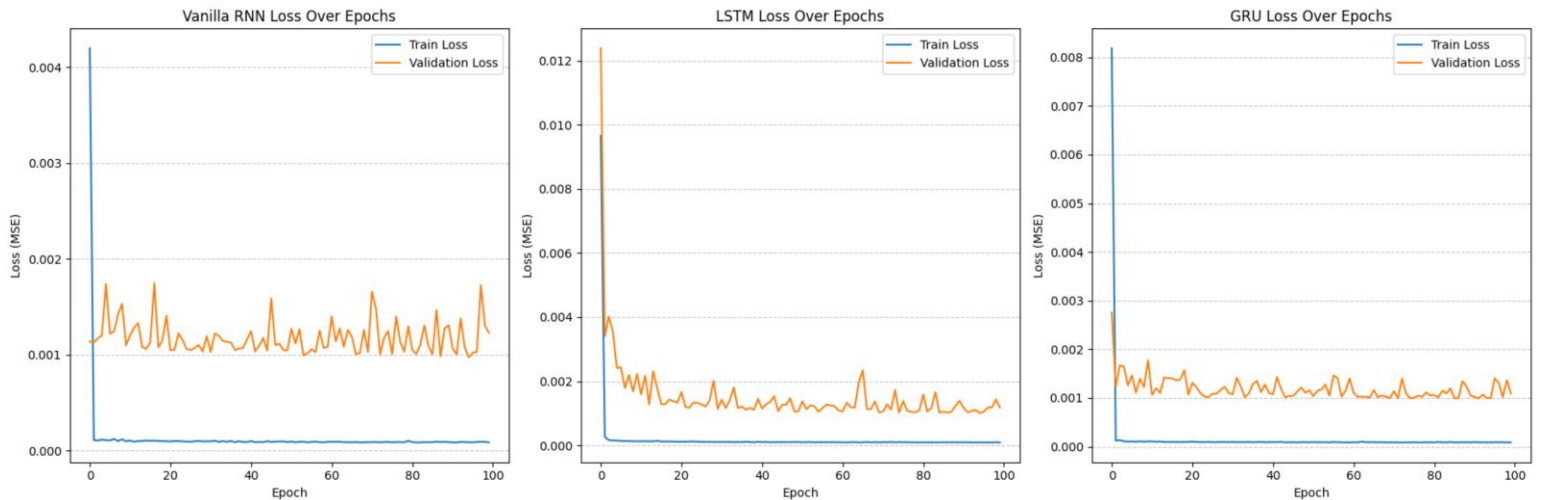


Figure 5: Training and validation losses of RNN Methods

The proof of the pudding is in the eating. The losses presented in the table and figure give a clear idea of the performances but the benchmark to compare them is their predicted values into the future. We set the stage up for these forecasts and plotted the actual vs predicted values for all three models. The same stage is repeated with three different such samples. Figure 6 displays the actual and predicted values.

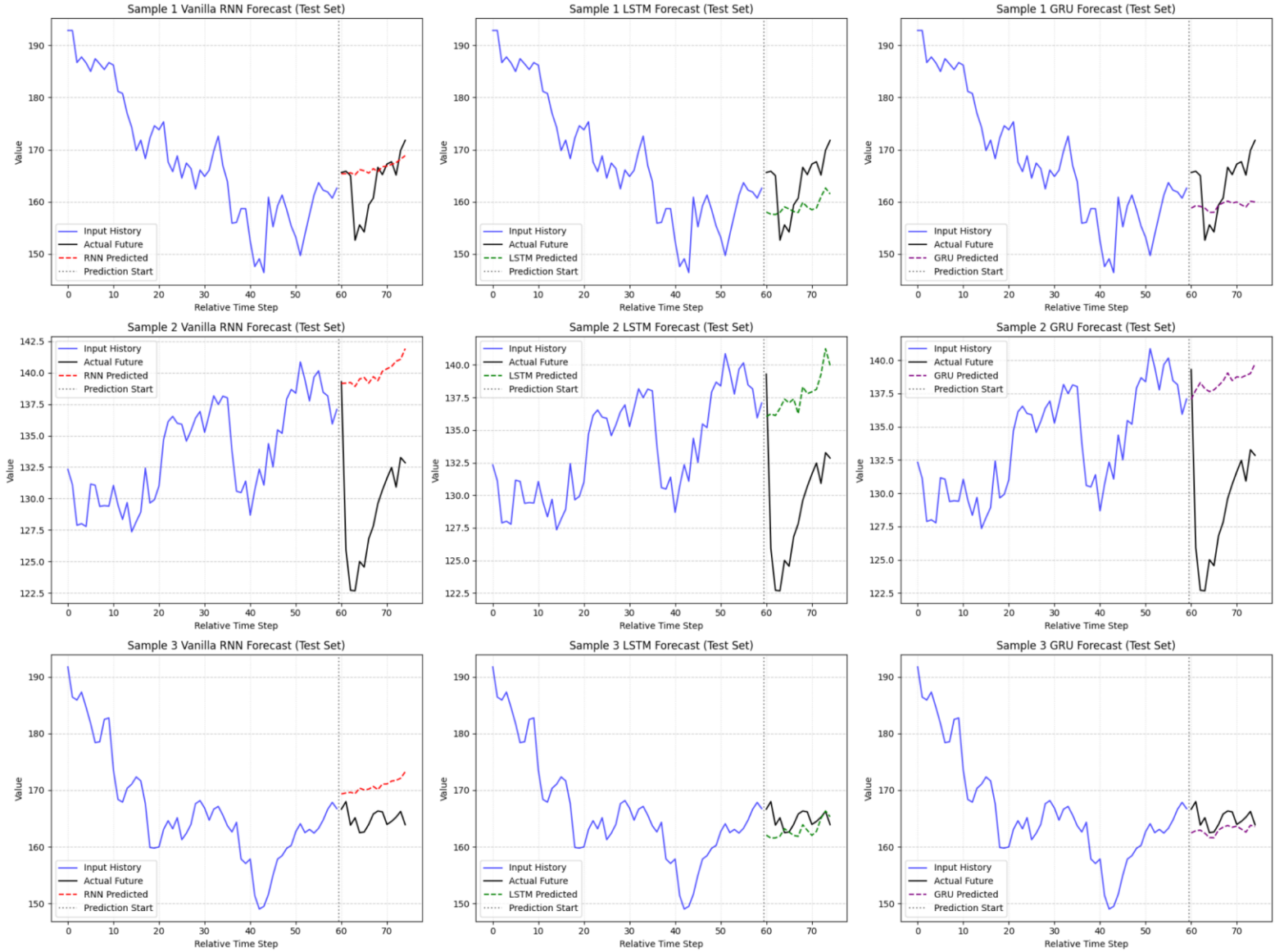


Figure 6: Training and validation losses of RNN Methods

Interestingly, different samples lead to different rankings of performances. The performance order is reversed in Sample1, for instance: LSTM does slightly better than GRU which is way better than Vanilla RNN. Similar ordering is true for Sample 2. Lastly, the sample comparison of performances is unfolded on Sample 3 as well.

Concluding Remarks

We have compared the forecasting performance of three RNN methods on Google stock prices. Based on all samples' averages, Vanilla RNN cuts the lowest MAE and

MSE followed by GRU and LSTM in turn. Interestingly, the opposite is observed based on the three random samples selected.

This is telling us that the performances of the methods heavily depend on the samples selected along with the lengths of the training, validation and testing periods, let alone the variances of the series.

Future research is required to figure out the factors that favor these methods.

References

Bengio, Y., Simard, P., & Frasconi, P. (1994). Learning long-term dependencies with gradient descent is difficult.¹⁶ *IEEE Transactions on Neural Networks*, 5(2), 157-160.

Cho, K., Van Merriënboer, B., Gulcehre, C., Bahdanau, D., Bougares, F., Schwenk, H., & Bengio, Y. (2014). Learning Phrase Representations using RNN Encoder-Decoder for Statistical Machine Translation.²⁰ *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1724–1734.

Elman, J. L. (1990). Finding structure in time.¹⁵ *Cognitive Science*, 14(2), 179-211.

Hochreiter, S., & Schmidhuber, J. (1997). Long Short-Term Memory.¹⁷ *Neural Computation*, 9(8), 1735-1780.

Rumelhart, D. E., Hinton, G. E., & Williams, R. J. (1986). Learning representations by back-propagating errors.¹⁴ *Nature*, 323(6088), 533-536.

THE RELATIVE IMPORTANCE OF HEALTH CARE IN INDIVIDUALS' PERCEPTION OF QUALITY OF LIFE: AFUZZY PAIRWISE COMPARISON ASSESSMENT

Terrence W. Thomas (North Carolina A&T State University)- twthomas@ncat.edu

Murat Cankurt (North Carolina A&T State University)- mcankurt@ncat.edu

Abstract

Quality of Life (QOL) represents a multidimensional construct influenced by diverse and interdependent factors. Contemporary public health initiatives aim to enhance QOL through the reduction of inequalities, expansion of health care access, and promotion of preventive health behaviors. This study aims to assess the perceived relative importance of health care services in comparison to other key QOL dimensions. By quantifying these perceptions, the research seeks to inform evidence-based health policy and resource allocation strategies that better align with individual and community priorities. A structured telephone survey was administered in 2023, focusing on four central QOL components: Health Care, Food Security, Spiritual Well-being, and Community Assets. Participants provided pairwise judgments regarding the importance of each dimension. The Fuzzy Pairwise Comparison (FPC) method was employed to evaluate subjective weights. The analysis revealed that Food Security held the highest perceived weight (30.32%), followed by Health Care (29.03%), Spiritual Well-being (26.57%), and Community Assets (14.08%). These findings suggest that although Food Security is marginally prioritized, Health Care remains nearly equally influential in individuals' perceptions of QOL. Health care emerges as a central pillar in shaping perceived quality of life. Public health strategies that enhance access to care, reduce health disparities, and integrate complementary domains such as nutrition and psychosocial support are likely to yield substantial QOL gains. These findings underscore the importance of developing holistic, person-centered policies rooted in individuals' lived experiences and priorities.

Key words: *Quality of life, health care, fuzzy pairwise comparison, public health, subjective well-being.*

Introduction

Quality of Life (QOL) is a multidimensional construct that assesses individuals' levels of happiness and well-being not only in terms of their health status but also in terms of their access to food, spiritual fulfillment, and the socio-economic opportunities offered by the communities in which they live. The literature has long emphasized that QOL is not comprised of a single indicator, but rather the combined interaction of multiple factors that guide individuals' lives (Felce & Perry, 1995; Schalock, 2004). Nevertheless, a fundamental question that frequently arises in both research and policy is the extent to which access to healthcare determines quality of life compared to other dimensions.

Health is considered the cornerstone of overall well-being; when physical well-being cannot be achieved, it becomes difficult to sustain other areas of life such as education, work life, and social participation (Patrick & Erickson, 1993). Preventive and equitable health systems not only reduce morbidity but also enable individuals to participate more effectively in social and economic life. Some studies show that societies with strong health infrastructure also have high levels of food security, productivity, and social resilience (Hanmer, 2021). In addition, QOL discussions have gained a broader perspective, and social science-based indicators such as economic security, social relationships, and community participation have also been included in the assessment (Cummins, 2005; Diener, 2010). However, in most of these measurement approaches, domains are given equal weight or weights assumed by researchers are applied (Alkire & Foster, 2011). This does not adequately reflect individuals' subjective assessments of which areas they value most in their own lives. At this point, it is possible to reflect this perceptual difference by using alternative methods such as Fuzzy Pairwise Comparison and Best-Worst Scaling (Cankurt, 2009; Flynn, 2008; Louviere, 2015).

Methodology

This study used the Sequential Explanatory Mixed Methods (SEMM) approach to examine the relative importance levels of factors determining quality of life (QOL) according to individual perceptions. Based on the World Health Organization's QOL framework (WHOQOL Group, 1995; WHOQOL Group, 1998), four key dimensions were evaluated in the study: health care, food security, spiritual well-being, and community assets. These dimensions also supported to the life elements most frequently

emphasized by community members in previous theoretical studies and preliminary in-depth interviews conducted by researchers (Forgeard, 2011; Bhandari, 2023).

The research was conducted in 2023 in Guilford County, North Carolina. Based on the characteristics of the population, 280 individuals were reached with a 95% confidence interval and a 6% sampling error criterion. Participants were selected through a field sampling company according to demographic criteria determined by the researchers (Lohr, 2019). A structured telephone survey was used as the data collection tool, and the survey was prepared in the Fuzzy Pairwise Comparison (FPC) format to include pairwise comparisons of four dimensions (Cankurt, 2009). Each participant was asked to compare the dimensions in pairs and indicate both the area they considered a priority and the intensity of their preference. After the data set was collected, surveys containing inconsistent, contradictory, or erroneous entries were excluded to increase the accuracy of the analysis (Rubin & Little, 2002; Thompson, 2009). As a result of this process, the final analysis was conducted on 217 valid surveys.

The Fuzzy Pairwise Comparison method was used to analyze the collected data. This method integrates the classical pairwise comparison approach with fuzzy logic theory, which considers uncertainty and judgment variability, allowing individuals' preference intensities to be measured on a continuous rather than a discrete scale (Zadeh, 1973; Saatchi, 2024). The preference values obtained from FPC were converted into membership degrees (μ) for each dimension, and a relative importance ranking of the dimensions was created based on the magnitude of the μ values. The Friedman test was used to assess whether the differences between dimensions were statistically significant, while Kendall's W coefficient was used to evaluate the level of consistency in participants' rankings (Górecki & Łuczak, 2021). This revealed both the relative importance percentages and the level of shared perception regarding the components within the community.

Results

The Fuzzy Pairwise Comparison (FPC) analysis conducted on the 217 valid questionnaires obtained in the study clearly revealed the order of importance of the quality of life dimensions as perceived by individuals. The findings showed that the dimensions of food security (30.32%) and health care (29.03%) had the highest weights, followed by spiritual well-being (26.57%) and community assets (14.08%). The Friedman test confirmed that the differences were statistically significant ($p < 0.01$),

indicating that participants clearly distinguished the factors determining their quality of life.

The results show that participants perceive access to food and health services as being of nearly equal critical importance in their QOL. The fact that food security ranks first shows that basic nutritional needs are the most pressing concern in daily life in low-income communities. In contrast, the fact that health cares are considered not only as a factor in the treatment of diseases but also as a factor that directly affects economic productivity, resilience, and social participation strongly places it in second place. The fact that spiritual well-being ranks third shows that internal resources such as psychological resilience, a sense of meaning, and hope are important components of quality of life in disadvantaged communities. Community assets, which received the lowest value, may be perceived as less important in these groups due to limited access to physical infrastructure and social resources.

Overall, the findings reveal that quality of life cannot be reduced to a single dimension and that food security, health, and spiritual well-being form three fundamental pillars that reinforce each other, particularly in low-income communities. Health care is a strategic lever within this framework, as their improvement can have positive repercussions in both economic and psychosocial spheres. For policymakers, this finding suggests that investing not only in the health system but also simultaneously implementing nutrition support, community-based mental health support, and social services can lead to more effective and sustainable improvements.

Conclusion

This study comparatively evaluated four fundamental dimensions (health care, food security, spiritual well-being, and community assets) to understand how the elements that constitute quality of life are positioned according to individuals' subjective perceptions. The Fuzzy Pairwise Comparison method provided a significant analytical advantage by allowing participants' priorities regarding quality of life to be determined directly through preference intensities rather than hypothetical assumptions. The results obtained show that quality of life cannot be explained by a single factor; particularly in communities with significant economic and structural disadvantages, quality of life is shaped by numerous and interrelated components.

The study's findings reveal that food security, healthcare, and spiritual well-being form three fundamental pillars that reinforce each other. Health services occupy a strategic

position among these elements because strengthening accessible and preventive health systems not only reduces morbidity but also reduces economic vulnerability, improves nutritional conditions, and supports psychosocial resilience. Therefore, policy initiatives that adopt holistic approaches addressing health, nutrition, and psychosocial support services together, rather than focusing on a single dimension, will provide more sustainable improvements. Investments in community resources can strengthen the infrastructure necessary to support long-term gains in quality of life.

References

- Alkire, S., & Foster, J. (2011). Counting and multidimensional poverty measurement. *Journal of Public Economics*, 95(7–8), 476–487.
- Bhandari, S., et al. (2023). Dose–response relationship between food insecurity and quality of life in United States adults: 2016–2017. *Health and Quality of Life Outcomes*, 21, 21.
- Cankurt, M. (2009). A Study on the Determination of Farmers' Demand for Tractor Satisfaction of Tractor Use and Purchasing Attitudes Towards Tractor: The Case of Aydın. Ph.D. Thesis, Ege University, Izmir, Turkey.
- Cummins, R. A. (2005). Moving from the quality of life concept to a theory. *Journal of Intellectual Disability Research*, 49(10), 699–706.
- Diener, E., et al. (2010). New well-being measures: Short scales to assess flourishing and positive and negative feelings. *Social Indicators Research*, 97(2), 143–156.
- Felce, D., & Perry, J. (1995). Quality of life: Its definition and measurement. *Research in Developmental Disabilities*, 16(1), 51–74.
- Flynn, T. N., Louviere, J. J., Peters, T. J., & Coast, J. (2008). Best–worst scaling: What it can do for health care research and how to do it. *Journal of Health Economics*, 26(1), 171–189.
- Forgeard, M. J. C., et al. (2011). Doing the right thing: Measuring wellbeing for public policy. *International Journal of Wellbeing*, 1, 79–106.
- Górecki, T., & Łuczak, M. (2021). On using nonparametric tests for multiple comparisons in ranking problems. *Statistical Papers*, 62(6), 2939–2963.
- Hanmer, J. (2021). Association between food insecurity and health-related quality of life: A nationally representative survey. *Journal of General Internal Medicine*, 36(6), 1638–1647.
- Lohr, S. (2019). *Sampling: Design and Analysis*. Nelson Education.
- Louviere, J. J., Flynn, T. N., & Marley, A. A. J. (2015). *Best–worst scaling: Theory, methods and applications*. Cambridge University Press.
- Patrick, D. L., & Erickson, P. (1993). *Health status and health policy: Quality of life in health care evaluation and resource allocation*. Oxford University Press.
- Rubin, D. B., & Little, R. J. (2002). *Statistical Analysis with Missing Data*. John Wiley & Sons.
- Saatchi, R. (2024). Fuzzy Logic Concepts, Developments and Implementation. *Information*, 15(10), 656.
- Schalock, R. L. (2004). The concept of quality of life: What we know and do not know. *Journal of Intellectual Disability Research*, 48(3), 203–216.

Thompson, B. (2009). The future of test validity. *Educational Researcher*, 38, 545–556.

WHOQOL Group. (1995). The World Health Organization Quality of Life Assessment (WHOQOL): Position paper from the World Health Organization. *Social Science & Medicine*, 41, 1403–1409.

WHOQOL Group. (1998). Development of the World Health Organization WHOQOL-BREF quality of life assessment. *Psychological Medicine*, 28(3), 551–558.

Zadeh, L. A. (1973). Outline of a new approach to the analysis of complex systems and decision processes. *IEEE Transactions on Systems, Man and Cybernetics*, SMC-3, 28–44.