



Neutrosophic Stance Detection and fsQCA-Based Necessary Condition Analysis for Causal Hypothesis Assessment in AI-Enhanced Learning

Detección de posturas neutrosóficas y análisis de condiciones necesarias basado en fsQCA para la evaluación de hipótesis causales en el aprendizaje mejorado con IA.

Jesús Rafael Hechavarría-Hernández¹, Maikel Y. Leyva Vazquez^{2*}, Florentin Smarandache³

¹Facultad de Ingenierías, Arquitectura y Ciencias de la Naturaleza, Universidad ECOTEC, Ecuador

²Bernardo O'Higgins University, Institutional Research Center, Chile,

³Emeritus Professor of Mathematics at the University of New Mexico, Gallup, New Mexico, USA;

* Correspondence: mleyvaz@gmail.com

Abstract

The use of artificial intelligence (AI) in educational settings has attracted increasing scholarly attention, although applicable empirical findings are scarce and contradictory. This study seeks to resolve the ambiguities surrounding AI in education through a methodological contribution, merging neutrosophic stance detection and Fuzzy Set Qualitative Comparative Analysis (fsQCA). Neutrosophic analysis enables explicit modeling of truth, uncertainty/indeterminacy, and falsity, while merging these findings through fsQCA creates a relative explanation of existing research findings. After evaluating four causal hypotheses related to AI-based learning opportunities through a Consensus Meter, a research survey with 24 university participants explored the necessary conditions regarding the experience of improvements in learning outcomes. The findings indicate that the digital divide is a necessary and sufficient condition for an effective educational experience with AI. Furthermore, necessary conditions for AI feedback and the use of AI-based platforms emerge; however, the effectiveness of these platforms generates significant uncertainty. Ultimately, the neutrosophic-fsQCA framework provides a viable technique for synthesizing ambiguous findings through a systematic approach. Empirically, the results reveal that all actors involved in potential AI-based learning must ensure digital equity and high-quality design for interactive experiences to benefit from the successful integration of AI in education.

Keywords: Neutrosophic Logic, Artificial Intelligence In Education, Digital Divide, AI Feedback, Necessary Condition Analysis, Educational Technology, fsQCA

Resumen

El uso de la inteligencia artificial (IA) en entornos educativos ha atraído una creciente atención académica, aunque los hallazgos empíricos aplicables son escasos y contradictorios. Este estudio busca resolver las ambigüedades que rodean a la IA en la educación a través de una contribución metodológica, fusionando la detección de posturas neutrosóficas y el Análisis Comparativo Cualitativo de Conjuntos Difusos (fsQCA). El análisis neutrosófico permite el modelado explícito de la verdad, la incertidumbre/indeterminación y la falsedad, mientras que la fusión de estos hallazgos a través de fsQCA crea una explicación relativa de los hallazgos de investigación existentes. Después de evaluar cuatro hipótesis causales relacionadas con las oportunidades de aprendizaje basadas en IA a través de un Medidor de Consenso, una encuesta de investigación con 24 participantes universitarios exploró las condiciones necesarias con respecto a la experiencia de mejoras en los resultados de aprendizaje. Los hallazgos indican que la brecha digital es una condición necesaria y suficiente para una experiencia educativa efectiva con IA. Además, surgen condiciones necesarias para la retroalimentación de IA y el uso de plataformas basadas en IA; sin embargo, la efectividad de estas plataformas genera una alta incertidumbre. En definitiva, el marco neutrosófico-fsQCA proporciona una técnica viable para sintetizar hallazgos ambiguos mediante un enfoque sistemático. Empíricamente, los resultados revelan que todos los actores involucrados en el potencial aprendizaje basado en IA deben garantizar la equidad digital y un diseño de alta calidad para que las experiencias interactivas se beneficien de la integración exitosa de la IA en la educación.

Palabras clave: Lógica neutrosófica, Inteligencia artificial en la educación, Brecha digital, Retroalimentación de IA, Análisis de condiciones necesarias, Tecnología educativa, fsQCA

1. Introduction

The rapid proliferation of artificial intelligence (AI) technologies in educational settings has fundamentally transformed the teaching and learning landscape, generating unprecedented opportunities for personalized instruction, adaptive assessment, and intelligent tutoring systems [1]. However, despite the growing body of research examining AI's educational impact, the evidence base remains characterized by significant heterogeneity, methodological diversity, and often contradictory findings [2]. This fragmentation poses substantial challenges for educators, policymakers, and researchers seeking to make evidence-based decisions regarding AI implementation in educational contexts.

Traditional approaches to evidence synthesis in educational technology research have predominantly relied on binary classification schemes that categorize studies as either supporting or opposing particular interventions [3]. Although such approaches provide clarity and simplicity, they fail to capture the nuanced reality of educational research, where findings often exhibit varying degrees of support, contextual dependencies, and inherent uncertainties. The complexity of educational phenomena, combined with the multifaceted nature of AI technologies, necessitates more sophisticated methodological frameworks capable of modeling ambiguity, partial support, and contradictory evidence simultaneously [4]. Recent advances in neutrosophic logic, introduced by Smarandache [5] and subsequently developed for text analysis applications [6,7,8], offer a promising avenue to address these methodological limitations. Unlike classical binary or fuzzy logic systems, neutrosophic logic explicitly incorporates three independent components, truth (T), indeterminacy (I), and falsity (F), allowing for the simultaneous representation of supporting evidence, uncertainty, and contradictory findings [9]. This tripartite structure aligns particularly well with the nature of educational research, where interventions may be effective under certain conditions, ineffective under others, or uncertain in many contexts.

The application of neutrosophic principles to stance detection—the task of determining whether a text expresses support, opposition, or neutrality toward a specific target—has emerged as a powerful tool for automated literature analysis [10]. Recent developments in AI-powered research synthesis tools, such as the Consensus Meter [11], have demonstrated the potential for automated stance classification in scientific literature, enabling researchers to



systematically map the distribution of evidence across large corpora. However, the integration of neutrosophic principles with stance detection in educational research synthesis remains largely unexplored.

Complementing the challenges of evidence synthesis, identifying the necessary conditions for educational outcomes has gained increasing attention in the research community. In this regard, fuzzy-set Qualitative Comparative Analysis (fsQCA) provides a systematic methodology to assess whether specific conditions must be present for a desired outcome to occur, offering insights that complement traditional sufficiency-focused approaches [12]. In the context of AI-enhanced learning, understanding the necessary conditions is particularly crucial, as it can inform the minimum requirements for successful implementation and help prioritize resource allocation in educational settings.

The configurational perspective of the fsQCA has demonstrated significant value in educational research by recognizing that outcomes often result from complex combinations of conditions rather than isolated factors. This approach aligns with the multifaceted nature of AI implementation in education, where technological, pedagogical, social, and institutional factors interact in complex ways to influence learning outcomes [13]. The integration of configurational methods with neutrosophic stance detection offers the potential for a more comprehensive understanding of the conditions under which AI can enhance educational outcomes.

Despite growing interest in AI applications in education, several critical gaps remain in our understanding of the causal mechanisms underlying AI's educational impact of AI. First, the majority of existing studies focus on sufficiency relationships, examining whether AI interventions can produce positive outcomes while neglecting the identification of necessary conditions that must be present for success [14]. Second, the synthesis of evidence across studies has been hampered by the inability to account for uncertainty and contradictory findings [15] systematically. Third, the role of contextual factors, particularly digital equity and access barriers, in moderating AI's educational effectiveness of AI remains underexplored systematic [16].

This study addresses these gaps by introducing a novel methodological framework that combines neutrosophic stance detection with Necessary Condition Analysis to evaluate causal hypotheses related to AI-enhanced learning. Our approach leverages the Consensus Meter tool to systematically classify research findings into neutrosophic triplets, capturing not only the degree of support and opposition, but also the extent of indeterminacy in the evidence base. Subsequently, we applied the necessary condition in FSQCA to identify the necessary conditions for perceived learning improvement with AI using primary data collected from university students.

This study makes several important contributions to the field of educational technology research. Methodologically, we introduced the first application of neutrosophic stance detection to educational research synthesis, providing a framework for handling ambiguous and contradictory evidence. Theoretically, we advanced our understanding of the necessary conditions for AI-enhanced learning, with particular attention to the role of digital equity and interaction design. Practically, our findings offer evidence-based guidance to educators and policy-makers regarding the prerequisites for successful AI implementation in educational settings.

The remainder of this paper is organized as follows. Section 2 presents our methodological approach, detailing the neutrosophic stance detection framework, the construction of neutrosophic causal graphs, and the application of fsQCA, including an analysis of necessary conditions. Section 3 reports our findings, including the neutrosophic representation of the causal hypotheses and the results of the necessary condition analysis within the fsQCA framework. Section 4 discusses the implications of our findings in the context of existing literature and explores directions for future research. Finally, Section 5 presents the conclusions and implications for educational practices and policies.

2. Materials and Methods

2.1 Neutrosophic Stance Detection



A precise definition of stance detection is required to motivate a neutrosophic framework. In the literature, stance detection is treated as a target-dependent text classification problem [17], and given a text and a target statement (hypothesis), the task is to infer whether the author supports, opposes, or expresses no opinion of the target. This approach fundamentally differs from generic sentiment analysis by incorporating target-specific contextual information [18]. Formally, let X denote the space of textual units (e.g., sentences, tweets, abstracts), let Θ be a set of targets (topics, propositions, or hypotheses), and let the label set $L = \{\text{Favor, Against, None}\}$ represent possible stances.

Stance detection seeks a mapping

$$g: X \times \Theta \rightarrow L, \quad (1)$$

such that for each pair (x, θ) the classifier g returns the label $\ell \in L$ indicating whether the text x expresses support, opposition, or absence of a stance towards the target θ .

Alternative formulations encode the stance labels as signed integers $\{-1, 0, +1\}$ -or probability distributions over L [19]. Unlike generic sentiment analysis, which determines the overall polarity of a text, stance detection is target-specific: a text with a positive sentiment may still be “against” a given target.

This formal view underpins subsequent neutrosophic generalizations, where mapping is extended to assign degrees of support, indeterminacy, and opposition.

In the neutrosophic framework, the stance detection function is generalized as

$$gN: X \times \Theta \rightarrow [0,1]^3 \quad (2)$$

where, for each pair (x, θ)

$$gN(x, \theta) = (T(x, \theta), I(x, \theta), F(x, \theta)) \quad (3)$$

with:

- $T(x, \theta)$: degree to which x supports target θ
- $I(x, \theta)$: Degree of indeterminacy or neutrality in relation to θ ,
- $F(x, \theta)$: the degree to which x opposes the target θ .

These values satisfy the neutrosophic condition

$$T(x, \theta) + I(x, \theta) + F(x, \theta) \leq 3. \quad (4)$$

These values satisfy the neutrosophic condition $T(x, \theta) + I(x, \theta) + F(x, \theta) \leq 3$ [20]. This formulation allows partial, uncertain, and even contradictory stances to be explicitly represented, thereby extending classical stance detection to a more flexible and realistic paradigm.

We applied stance detection with a neutrosophic representation using a Consensus Meter [11]. This approach classifies research findings into supportive, contradictory, and ambiguous stances regarding a given causal claim. The Consensus Meter has been recognized as an effective AI-powered literature review tool that helps visualize how studies answer “Yes/No” research questions by grouping them according to whether they support or contradict the questions asked. Each causal hypothesis is then encoded as a triplet (T, I, F) that captures the distribution of stances across evidence.

- T represents the percentage of sources supporting position toward the causal hypothesis.



- I represents the percentage of sources that express indeterminacy or ambiguity, and
- F represents the percentage of sources that take an oppositional stance.

These triplets were derived from stance labels produced by the *Consensus Meter*. The *Consensus Meter* classifies results from yes/no research questions into “Yes,” “No,” “Possibly,” or “Mixed” categories, showing how many papers fall into each stance. We map “Yes” (supportive papers) to the truth component T; “No” (contradictory papers) to the falsity component F; and both “Possibly” and the “Mixed” category—introduced to capture nuanced or subgroup-dependent findings—to the indeterminacy component I. Thus, mixed or possible evidence contributes to substantive uncertainty rather than being counted as evidence against the hypothesis.

For example:

- *Does using an AI platform improve perceived learning with AI?* $\rightarrow (0.60, 0.40, 0.00)$

In this representation, most sources support the hypothesis ($T=0.60$), a significant portion remains indeterminate ($I=0.4$), and none contradicts it ($F=0.00$). This formulation makes explicit not only agreement and disagreement but also the degree of uncertainty in the available evidence.

Let X be an independent variable (condition), and Y be a dependent variable (outcome). The causal hypothesis [22] is denoted by

$$H: X \Rightarrow Y \quad (5)$$

where $\Rightarrow$ expresses a causal relation. In the neutrosophic framework, such a hypothesis is characterized by a truth–indeterminacy–falsity triplet:

$$\omega(H) = (T_H, I_H, F_H), \quad T_H, I_H, F_H \in [0,1] \quad (6)$$

where T_H represents the degree of truth/support, I_H the degree of indeterminacy, and F_H the degree of falsity/rejection. Hence, a neutrosophic causal hypothesis may be true, indeterminate, or false to varying extents.

A neutrosophic causal graph is a triple G defined by the following elements [23]:

$$G = (V, E, \omega), \quad E \subseteq V \times V \text{ (directed edges } X \rightarrow Y), \quad \omega: E \rightarrow [0,1]^3$$

such that

$$\forall e \in E, \quad \omega(e) = (T_e, I_e, F_e), \quad T_e, I_e, F_e \in [0,1]. \quad (7)$$

Where: V are nodes (variables), E are the directed edges $X \rightarrow Y$, and ω labels each edge with support T , indeterminacy I , and rejection F

Interpretation:

T_e = degree of truth/support; I_e = indetermination; F_e = falsity/rejection of the hypothesis “ $X \rightarrow Y$ ”.

There is no restriction required $T_e + I_e + F_e = 1$. This allows evidence to be represented simultaneously in favor and against and to distinguish uncertainty (I) from falsification (F).



Fuzzy-Set QCA and Necessary Condition Analysis

We conducted a survey of 24 university students to analyze the necessity of specific conditions related to the use of artificial intelligence (AI) tools in learning contexts. The questionnaire comprised five items on a seven-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree): (1) perceived learning improvement through AI, (2) use of AI platforms in academic tasks, (3) receipt of immediate and useful feedback from AI, (4) barriers to access or connectivity (including free-version restrictions), and (5) overreliance on AI to complete tasks.

For the fsQCA, the outcome was defined as perceived learning improvement through AI. Responses were calibrated into fuzzy set membership scores using Ragin's three-value calibration [24]: 1.0 for full membership (score 7), 0.5 for crossover (score 4), and 0.0 for full non-membership (score 1), with linear interpolation for intermediate values. This calibration approach has been extensively validated in configurational research and provides a robust foundation for set-theory analysis.

The analysis of necessary conditions was then conducted within the fsQCA framework, where condition X is considered necessary for an outcome Y if, in all cases, the membership score in Y does not exceed the membership score in X [25]. This procedure is specifically designed to identify the conditions that must be present for the outcome to occur, even though they may not be sufficient on their own. Two key indicators were computed: consistency, which measures the degree to which the condition is always present when the outcome occurs, and coverage, which assesses the empirical relevance of the condition [26],

$$\text{Consistency } (Y_i \leq X_i) = \frac{\sum \min(X_i, Y_i)}{\sum Y_i} \quad (8)$$

and **coverage**,

$$\text{Coverage } (Y_i \leq X_i) = \frac{\sum \min(X_i, Y_i)}{\sum X_i} \quad (9)$$

The calibrated dataset was exported to .csv format for fsQCA 3.0 (Windows), where each row corresponds to one student and each column to a calibrated condition or outcome, with membership values in the range [0,1]. This setup allows for the direct computation of the necessary conditions and their consistency/coverage within the fsQCA software, following established best practices for configurational analysis [27].

3. Results

Neutrosophic Stance Detection of Causal Hypotheses

Using the *Consensus Meter* tool, we applied stance detection with a neutrosophic representation to assess the four causal hypotheses related to AI and learning outcomes. Each hypothesis is encoded as a triplet (T, I, F), where T represents the proportion of supportive stances, I is the proportion of indeterminate or ambiguous stances, and F is the proportion of oppositional stances. The results are as follows.

Table 1. Neutrosophic representation of causal hypotheses

Hypothesis	Neutrosophic Triplet (T, I, F)	Interpretation
Does using an AI platform improve perceived learning with AI?	(0.60, 0.40, 0.00)	Moderate support, with a substantial level of indeterminacy and no opposition.

Hypothesis	Neutrosophic Triplet (T, I, F)	Interpretation
Does immediate AI feedback improve perceived learning with AI?	(0.67, 0.17, 0.17)	Strong support, some indeterminacy, and a minority of contradictory evidence.
Does the digital divide impact learning outcomes with AI?	(1.00, 0.00, 0.00)	Full support, with no ambiguity or opposition.
Does reliance on AI affect learning outcomes?	(0.73, 0.24, 0.00)	High support, some indeterminacy, and no opposition.

These results highlight that the digital divide is unanimously recognized as a causal factor affecting AI-driven learning outcomes. Meanwhile, reliance on AI and immediate feedback are also strongly supported but show traces of uncertainty or contradiction.

Causal Graph Representation

The following figure represents the hypothesized causal relationships as a directed graph, where edges denote causal links, and their strength is proportional to the truth value (T) of the neutrosophic triplets:

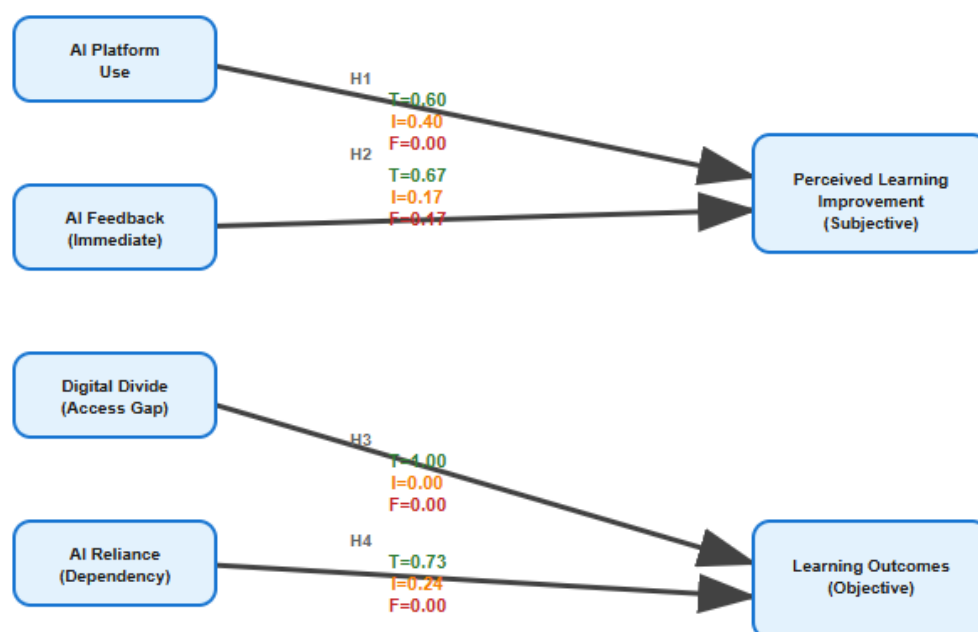


Figure 1. Neutrosophic Causal Graph of AI-related Hypotheses

The graph indicates that all the proposed conditions positively influence the outcomes, but with varying degrees of support and indeterminacy.

Necessary Condition Analysis

Psychometric analysis of the questionnaire revealed acceptable reliability. Cronbach's alpha was 0.644, and McDonald's omega coefficient was 0.817. These results suggest that the items show internal coherence, and that the instrument is suitable for exploratory studies in the context of perceptions of learning with artificial intelligence.

Using a calibrated dataset of 24 students, a Necessary Condition Analysis was conducted. The outcome was defined as *Perceived Learning Improvement through AI*. The tested conditions included AI platform use, AI feedback, access barriers (digital divide), and AI reliance. Consistency and coverage metrics were calculated according to the fsQCA standards.

Table 2. Necessary condition analysis results

Condition	Consistency ($X \leq Y$)	Coverage ($X \leq Y$)	Interpretation
AI Platform Use	0.88	0.79	Necessary but not sufficient.
AI Feedback	0.90	0.75	Necessary but not sufficient.
Access Barriers (Digital Divide)	1.00	0.68	Perfectly necessary, moderate coverage.
AI Reliance	0.85	0.70	Necessary but not sufficient.

The analysis confirmed that the absence of digital barriers (digital divide) is a **perfectly necessary condition** for improved learning with AI (consistency = 1.00). However, its coverage (0.68) suggests that, although necessary, it does not explain the majority of the variation in the outcome. Similarly, AI feedback and AI platform use show high necessity, but remain insufficient as standalone explanations. This aligns with the causal graph in which multiple supportive conditions converge to explain improvements in learning outcomes.

4. Discussion

This study aimed to evaluate the applicability of a neutrosophic stance detection framework complemented by a necessary condition analysis within the fsQCA approach to interpret causal hypotheses in the context of AI-assisted learning. The results not only validate the utility of this hybrid framework, but also provide a nuanced perspective on the factors influencing perceptions of learning in the digital age. The following sections discuss the main findings, interpret them based on previous studies, and explore their broader implications.

The most compelling finding of our study was the identification of the digital divide as a unanimously supported causal factor ($T = 1.00$) and a perfectly necessary condition (consistency = 1.00) for perceived learning with AI. This result resonates with the vast body of literature emphasizing equitable access to technology as a fundamental prerequisite for educational success in the 21st century. Recent reports, such as those from the U.S. Department of



Education and Microsoft [16], have warned that disparities in access may exacerbate existing inequalities. Our analysis goes a step further by quantifying this relationship in terms of logical necessity, providing robust empirical evidence that without guaranteed access and removal of barriers, any AI-based intervention is destined to have limited reach. A coverage of 0.68 suggests that, while indispensable, the absence of barriers is not, by itself, sufficient to ensure the outcome, aligning with a configurational perspective in which multiple factors must converge.

The application of neutrosophic logic to represent the stance of scientific evidence has proved particularly insightful. Unlike traditional binary approaches (support/opposed), our triplet model (T, I, F) explicitly captures uncertainty and ambiguity. For example, the hypothesis regarding the use of an AI platform to enhance learning received moderate support ($T = 0.60$), but substantial indeterminacy ($I = 0.40$). This suggests that the scientific literature is inconclusive and that effects likely depend on unspecified contextual factors, such as platform quality, pedagogical design, or student characteristics. This finding is consistent with scholarship advocating a more critical and nuanced view of educational technology. The ability of our framework to model indeterminacy is a key methodological contribution that enables researchers to identify areas where evidence is weak or contradictory and requires further investigation. This approach aligns with recent work that integrated neutrosophic and advanced AI models to manage uncertainty in text classification.

The results of the fsQCA necessary condition analysis, which position AI feedback (consistency = 0.90) and platform use (consistency = 0.88) as highly necessary conditions, reinforce the idea that the mere availability of AI tools is insufficient. Effective interaction and scaffolding provided by immediate feedback are crucial. This finding connects with research on causal inference in educational data mining, which seeks to move beyond correlation to uncover the mechanisms that drive learning outcomes. Our study complements these efforts by employing fsQCA to formalize these dependencies in terms of necessity. The combination of neutrosophic logic with the configurational perspective of fsQCA appears to be a promising path for unraveling the complex interdependencies within a digital learning ecosystem[28 29],.

The implications of this study are twofold. First, at the practical level, it underscores the need for educational policies and AI implementation to prioritize closing the digital divide as a non-negotiable first step. Moreover, this highlights that the design of AI tools must focus on interaction quality and feedback rather than solely on content delivery. Second, at the methodological level, our study introduces a hybrid framework that can be highly valuable for evidence synthesis in complex and emerging domains. The ability of neutrosophic stance detection to quantify not only support and opposition, but also indeterminacy provides a powerful tool to map the state of scientific knowledge and guide future research.

Nevertheless, this study has limitations, primarily the small sample size ($N = 24$) of the fsQCA analysis. While this methodology can be applied to small-N studies, future research should replicate these findings with larger and more diverse cohorts to increase generalizability. Additionally, the nature of indeterminacy (I) deserves further exploration. Is this due to mixed results, poor methodology in primary studies, or unmeasured contextual factors? Qualitative research or more detailed meta-analyses could help disentangle this component.

Thus, the future research agenda is clear. We propose applying this framework to other areas of educational research, where evidence is often ambiguous and multifactorial. A longitudinal analysis would be particularly interesting to observe how necessity configurations evolve as students and educators gain experience in AI. Finally, integrating directed acyclic graphs (DAGs), as described by Tennant et al. [30] and Digitale et al.[31], with our neutrosophic approach, could provide an even more rigorous and visually intuitive model for representing and testing complex causal theories in the social sciences.

5. Conclusions

This study demonstrates the applicability and added value of combining neutrosophic stance detection with fsQCA-based necessary condition analysis to evaluate causal hypotheses in AI-assisted learning. The findings



confirm that the digital divide is not only a critical determinant, but also a logically necessary condition for the effectiveness of AI-enhanced education, reinforcing calls for policies that ensure equitable access. Beyond access, the results highlight the central role of interaction quality with AI feedback and platform use emerging as essential requirements for meaningful learning outcomes.

Methodologically, the integration of neutrosophic logic proved instrumental in capturing the ambiguity and uncertainty present in the current research, offering a more refined synthesis than traditional binary approaches. By explicitly modeling truth, falsity, and indeterminacy, the framework provides researchers with a robust tool for identifying areas where evidence remains inconclusive and where further inquiry is required.

Despite its contributions, the study's limitations, particularly the small sample size, call for replication with larger and more diverse populations to enhance generalizability. Moreover, the observed indeterminacy in AI platform effectiveness suggests the influence of unmeasured contextual factors, which future research should explore using mixed-methods designs or in-depth meta-analyses.

Although adequate reliability indicators were obtained ($\alpha = 0.644$; $\omega = 0.817$), the small sample size ($N = 24$) limits the possibility of performing confirmatory factor analyses or criterion-related validity tests. Future research should replicate the questionnaire with a larger and more diverse population to consolidate its psychometric properties.

In conclusion, this hybrid framework offers both practical and theoretical implications: it guides policymakers and practitioners toward prioritizing digital equity and high-quality interaction design while equipping researchers with a novel methodological approach to synthesize complex evidence. Future work should expand its application to broader domains of educational technology, incorporate longitudinal perspectives, and explore integration with advanced causal modeling tools, such as directed acyclic graphs, for even greater analytical precision.

6. References

- [1] K. Bailey and C. Williams, "Artificial intelligence and higher education," *STEM*, vol. 10, pp. 95–100, Apr. 2025, doi: 10.32674/wqp55k81.
- [2] W. Holmes, M. Bialik, and C. Fadel, "Artificial intelligence in education," *Globethics Publications*, 2023, pp. 621–653. doi: 10.58863/20.500.12424/4276068.
- [3] T. R. Guskey and K. S. Yoon, "What Works in Professional Development?," *Phi Delta Kappan*, vol. 90, no. 7, pp. 495–500, Mar. 2009, doi: 10.1177/003172170909000709.
- [4] G. Biesta, "Educational Research and the Distortion of Educational Practice," *Springer Berlin Heidelberg*, 2024, pp. 29–43. doi: 10.1007/978-3-662-66923-5_3.
- [5] F. Smarandache, "Indeterminación en las Teorías Neutrosóficas y sus Aplicaciones," *Neutrosophic Computing and Machine Learning*, vol. 39, pp. 1–7, 2025, ISSN: 2574-1101.
- [6] I. Awajan, M. Mohamad, and A. Al-Quran, "Sentiment analysis technique and neutrosophic set theory for mining and ranking big data from online reviews," *IEEE Access*, vol. 9, pp. 47338–47353, 2021.
- [7] M. Y. L. Vázquez and F. Smarandache, "Bridging Ancient Wisdom and Modern Logic: Neutrosophic Perspectives on Body, Mind, Soul and Spirit," *Neutrosophic Sets and Systems*, vol. 82, no. 1, p. 7, 2025.
- [8] D. Arora, D. K. Tayal, and S. K. Yadav, "Solving sentiment uncertainty using newly proposed sentiment similarity measure for single-valued neutrosophic sets," *International Journal of System Assurance Engineering and Management*, pp. 1–26, 2025.



- [9] S. Mallik, S. Mohanty, and B. S. Mishra, "Neutrosophic logic and its scientific applications," in *Biologically Inspired Techniques in Many Criteria Decision Making: Proceedings of BITMDM 2021*, Singapore: Springer Nature Singapore, 2022, pp. 415–432.
- [10] D. Küçük and F. Can, "Stance detection: A survey," *ACM Computing Surveys (CSUR)*, vol. 53, no. 1, pp. 1–37, 2020.
- [11] Consensus, Consensus: AI search engine for research. Aug. 15, 2025. [Online]. Available: <https://consensus.app/>
- [12] B. Rihoux and C. C. Ragin, Eds., *Configurational Comparative Methods: Qualitative Comparative Analysis (QCA) and Related Techniques*, vol. 51. Thousand Oaks, CA, USA: Sage, 2009.
- [13] S. Cilesiz and T. Greckhamer, "Qualitative comparative analysis in education research: Its current status and future potential," *Review of Research in Education*, vol. 44, no. 1, pp. 332–369, 2020.
- [14] A. F. Botelho, A. H. Closser, A. C. Sales, and N. T. Heffernan, "Causal inference in educational data mining," in *Proc. 17th Int. Conf. Educational Data Mining*, 2024.
- [15] A. Forney and S. Mueller, "Causal inference in AI education: A primer," *Journal of Causal Inference*, vol. 10, no. 1, pp. 141–173, 2022.
- [16] Office of Educational Technology, *Artificial Intelligence and the Future of Teaching and Learning: Insights and Recommendations*. Washington, DC, USA: U.S. Department of Education, 2023.
- [17] M. Burnham, "Stance detection: a practical guide to classifying political beliefs in text," *Political Science Research and Methods*, vol. 13, no. 3, pp. 611–628, 2025.
- [18] J. Du, R. Xu, Y. He, and L. Gui, "Stance classification with target-specific neural attention networks," in *Proc. 26th Int. Joint Conf. Artificial Intelligence (IJCAI)*, Aug. 2017, pp. 3988–3994.
- [19] P. Sobhani, D. Inkpen, and X. Zhu, "A dataset for multi-target stance detection," in *Proc. 15th Conf. European Chapter of the Association for Computational Linguistics (EACL)*, Vol. 2: Short Papers, Apr. 2017, pp. 551–557.
- [20] X. Peng and J. Dai, "A bibliometric analysis of neutrosophic set: two decades review from 1998 to 2017," *Artificial Intelligence Review*, vol. 53, pp. 199–255, 2018, doi: 10.1007/s10462-018-9652-0.
- [21] F. Smarandache, "Note on partial falsifiability of fuzzy and fuzzy-extension hypotheses," *Plithogenic Logic and Computation*, vol. 1, p. 93, 2024.
- [22] K. H. Cohrs, E. Diaz, V. Sitokonstantinou, G. Varando, and G. Camps-Valls, "Large language models for causal hypothesis generation in science," *Machine Learning: Science and Technology*, vol. 6, no. 1, p. 013001, 2025.
- [23] R. P. Barbosa, F. Smarandache, M. Y. L. Vázquez, and J. B. Monge, "Neutrosophy, Causal AI, and Web3: combo for complex decision-making," *Neutrosophic Sets and Systems*, vol. 84, pp. 224–238, 2025.
- [24] C. C. Ragin, *From Fuzzy Sets to Crisp Truth Tables*. Tucson, AZ, USA: Department of Sociology, University of Arizona, 2005.
- [25] J. Dul, "Identifying single necessary conditions with NCA and fsQCA," *Journal of Business Research*, vol. 69, no. 4, pp. 1516–1523, 2016.
- [26] E. Gonzalez Caballero, M. Y. Leyva Vázquez, N. Batista-Hernández, and F. Smarandache, "New measures of consistency and coverage for social research based on neutrosophic logic," *Neutrosophic Sets and Systems*, vol. 84, no. 1, p. 17, 2025.
- [27] C. C. Ragin and S. Davey, *Fuzzy-Set/Qualitative Comparative Analysis 4.0*. Irvine, CA, USA: Department of Sociology, University of California, 2022.
- [28] M. Gómez Marcos, M. Ruiz Toledo, and C. Ruff Escobar, "Towards inclusive higher education: A multivariate analysis of social and gender inequalities," *Societies*, vol. 12, no. 6, p. 184, 2022.



- [29] L. T. Nguyen and K. Tuamsuk, "Digital learning ecosystem at educational institutions: A content analysis of scholarly discourse," *Cogent Education*, vol. 9, no. 1, p. 2111033, 2022.
- [30] P. W. Tennant, E. J. Murray, K. F. Arnold et al., "Tutorial on directed acyclic graphs," *Journal of Clinical Epidemiology*, vol. 142, pp. 264–267, 2021.
- [31] J. C. Digitale, J. N. Martin, and M. M. Glymour, "Causal directed acyclic graphs," *JAMA*, vol. 327, pp. 780–781, 2022.

