

Intelligent Detection and Management of Academic Document Irregularities in University Decision Systems using Neural Networks: A Review

O.T. Ogbeye

Department of Computer Science,
Adekunle Ajasin University,
Akungba Akoko, Nigeria

D.A. Akinwumi

Department of Cyber Security,
Adekunle Ajasin University,
Akungba Akoko, Nigeria

ABSTRACT

This review paper addresses the critical issue of document irregularity management within university academic decision-making processes, highlighting the increasing reliance on data integrity in higher education. Document irregularities, like wrong Senate decision extracts, grade changes, and fake transcripts, can pose significant threats to the credibility and efficiency of academic institutions. Current methods to detect these irregularities are usually done by hand, take a lot of time, and are prone to human error. So, there's a need for better, automated solutions. Neural Networks (NNs) are a good technology for anomaly detection and classification tasks. This paper reviews how different Neural Network architectures, like Artificial Neural Networks (ANN), Convolutional Neural Networks (CNN), and Recurrent Neural Networks (RNN), can be used to identify and mitigate document irregularities. The paper looks at what's already been written, common methods, and the problems with using NN-based solutions in academic environments. Key findings indicate a growing trend towards leveraging AI for document verification and fraud detection, with NNs doing better at recognizing patterns than traditional methods. But issues like getting enough data, labeling problems, and understanding how the models work are still big challenges. The results suggest that Neural Networks could change how academic integrity is protected and how decisions are made easier. Future research could look at using combined AI models, Explainable AI (XAI) to build trust, and adding blockchain to NNs for secure academic records.

General Terms

Document forensics, Anomaly detection, Information security, Data integrity, and Computer vision.

Keywords

Neural networks, Academic decision making, Document irregularities, Higher education, Decision support systems.

1. INTRODUCTION

Academic decisions at universities are complex and support the main activities of higher education. These decisions cover student admissions, course enrollment, grade certifications, faculty appointments, research funding, and accreditations. Accurate and authentic data is very important to these decisions, which maintain standards and keep things running efficiently. The increase in digital records has greatly increased the amount and complexity of academic data, so keeping data and documents secure is now more important than ever for making good decisions [1]. These important documents are often altered by mistakes. These can be simple errors, such as typos or inconsistencies from manual input, or purposeful fraud, like changing grades or faking transcripts and

certificates. These problems with documents create challenges, which can cause bad academic decisions, harm the school's reputation, and hurt the futures of students and faculty. For example, bad records can cause students to graduate improperly, resources to be mismanaged, or even create legal problems for the school [2]. Conventional ways of identifying and managing these irregularities, like doing audits manually and using systems based on rules, are not adequate anymore. They take a lot of work and time. There can be human oversight, and these methods are incapable of detecting sophisticated forms of fraud that exploit subtle patterns or anomalies

This rising weakness has caused people to look into and start using advanced technological solutions, mostly those based on Artificial Intelligence (AI) and Neural Networks (NNs). AI methods can automatically detect irregularities, enhance accuracy, and give options that can grow to handle the rising amount of academic data. Neural Networks, which are a type of AI that copies how the human brain is set up and works, are great at recognizing patterns, anomaly detection, and classification tasks. This makes them helpful for identifying hard and hidden irregularities in documents [3]. They can learn from big sets of data and identify small changes from the usual, which makes them a strong tool for protecting integrity in academics.

This paper reviews how neural networks can be used to manage document irregularities in university academic decisions. It looks at the kinds of irregularities that occur in higher education, how they affect choices, and the limits of standard detection methods. The paper will explore different neural network designs and how they are used to analyze documents, verify information, and spot fraud. The study will also look at current systems and methods, point out what's missing in current research, and talk about the difficulties of using neural network-based solutions. The study proposes directions for future research and development, like adding new technologies such as blockchain, and the need for Explainable AI to build trust in automated decision support systems.

The paper is structured as follows: Section 2 gives a literature review, Section 3 discusses methodologies to detect document irregularities, Section 4 presents results and discussion, Section 5 discusses challenges and future directions, and Section 6 concludes the paper with main ideas and a request for interdisciplinary collaboration.

2. LITERATURE REVIEW

Problems with academic records in higher education cover many issues that can damage the integrity of these records. These problems can be categorized into accidental errors and intentional fraud. Accidental errors often come from manual

data entry, system problems, or data transfer issues between university systems. Examples are typos in student names, wrong course codes, mismatched grades because of transcription errors, or differences in Senate decisions where the recorded result does not match the actual decision. These errors may seem small, but they can add up and cause administrative problems, delays in student progress, and misinformed academic decisions [4].

Intentional irregularities are deliberate actions aimed at deceiving or gaining an unfair advantage, and they manifest in various ways, such as students or staff altering grades in academic records to exaggerate performance, falsifying or modifying transcripts to misrepresent academic history, fabricating documents including admission or recommendation letters, certificates, and diplomas, presenting discrepancies between official Senate meeting minutes and the decisions communicated that may result in confusion, misapplication of policies, or unfair treatment of stakeholders, and engaging in plagiarism or other forms of documented academic misconduct.

These irregularity issues have wide-ranging impacts. Students might not be able to graduate on time, get into further education, or have their degrees recognized. Institutions could suffer damage to their reputation, legal problems, financial losses from things like fake scholarship requests, and a loss of faith in the academic world. There are many examples around the world, from well-known instances of fake degrees being used to get jobs to widespread grade changing that impacts whole groups of students. These events show that there is an urgent need for better ways to find and stop these problems, rather than just relying on traditional manual checks [5].

Universities handle academic decisions through a complicated system with input from administrators, faculty, and students. These decisions span operational, tactical, and strategic levels, directly shape education quality, research results, and how the university is run. Document quality and integrity greatly impact important academic decisions like admissions, course enrollment and progression, graduation and certification, scholarship and financial aid allocation, faculty hiring and promotion, and accreditation and compliance [6].

Manual decision-making has issues like mistakes and biases. Tiredness, biases, and personal opinions can cause differences in decisions, even when presented with accurate data. The issue gets worse when documents have irregularities, potentially causing decision-makers to unknowingly use wrong data. Large universities process so many documents daily that checking everything by hand is unrealistic, meaning irregularities may go undetected. This shows the need for automated, objective, and efficient systems to support academic decision-making [6].

Neural Networks (NNs), also called Artificial Neural Networks (ANNs), are computing models that draw inspiration from the structure of biological neural networks. Designed to recognize patterns and relationships in data through learning, they work well for classification, regression, clustering, and anomaly detection. The artificial neuron is the basic building block of an NN. It takes in input, processes it, and makes an output. These neurons are arranged in layers: an input layer, hidden layers, and an output layer. Connections between neurons have weights that change during training to decrease the difference between the network's output and the desired output [7].

Artificial Neural Networks are useful for document analysis and making decisions. Here are a few types:

i. Artificial Neural Networks (ANNs): These have many layers of linked neurons. ANNs are flexible and can be applied to various tasks, like recognizing patterns and sorting organized data.

ii. Convolutional Neural Networks (CNNs): These are made to work with grid-like data, like pictures. CNNs are good at pulling out features from visual data. This makes them great for looking at scanned papers, handwriting, and anomalies in document layouts [8].

iii. Recurrent Neural Networks (RNNs): These can handle data in a series, where what comes before affects what comes next. RNNs, and better types like Long Short-Term Memory (LSTM) networks, can study text documents, find anomalies in grade histories, or identify inconsistencies in reports [9].

Neural networks (NNs) can be helpful analytical tools in decision support systems, giving humans helpful insights and predictions. They can learn complicated things from big datasets and identify subtle patterns that humans or simpler systems might miss. NNs can be trained with real and fake academic documents, and they can learn to tell the difference between them. For example, a CNN might look at a scanned transcript to find things that look wrong, like different fonts or changed seals. An RNN could look at a student's grades and spot weird patterns that could mean someone changed the grades. The integration of NNs into decision support systems can significantly enhance the efficiency, accuracy, and objectivity of academic decision-making processes, moving towards a more data-driven and integrity-focused approach [10].

Document management in educational institutions has changed a lot because there are more digital files and a need for better security and speed. Traditionally, document management systems in universities focused on digitization, archiving, and retrieval, often relying on relational databases and manual verification processes. While these systems improved accessibility and reduced physical storage needs, they offered limited capabilities for proactive irregularity detection [11].

AI has begun to revolutionize document verification and fraud detection across various sectors, including finance, legal, and healthcare. AI programs using machine learning and deep learning can identify forged documents, detect anomalies in financial transactions, and verify identity documents. For example, optical character recognition (OCR) and natural language processing (NLP) pull information from papers and check for problems. These systems can even see small changes in scanned files that people would miss, like changed dates or names [12].

In educational institutions, AI is starting to help with document checks and detecting fraud, but not as fast as in other fields. Right now, the main use is finding plagiarism in student work with tools like Turnitin. These tools use algorithms to compare student papers with big databases of papers and online content to find copied work. But they mostly check text and no other document problems in academic records [13].

Besides plagiarism, some recent studies have looked into using AI to find academic misconduct, like contract cheating or impersonation during online tests. These methods often use behavioral analytics and pattern detection, but using them directly to detect document forgery or grade changing still has a lot of untapped potential. Likewise, even though AI and machine learning are being used more and more in educational decision-making systems for things like predicting which students might drop out, creating personalized learning plans,

and deciding how to use resources, not as much focus has been put on using them to make sure the data and documents themselves are correct [14].

The use of neural networks in education and decision support has mostly centered on teaching methods, like intelligent tutoring systems, predicting how well students will do, and creating adaptive learning environments. For instance, neural networks have been used to guess how well students will do based on historical academic data, find students who are struggling, and suggest specific learning materials. Though these applications indirectly benefit from data integrity, they do not directly address the detection and management of document irregularities at the source [15].

There is an increasing need for more effective approaches to managing document-related problems in university decision-making, and neural networks (NNs) have demonstrated strong potential in related domains. However, significant gaps remain in the literature regarding their comprehensive application to the detection and management of academic document irregularities. First, most existing studies focus on isolated aspects of document handling or on specific types of misconduct, such as plagiarism, rather than developing an integrated system capable of addressing diverse irregularities across the entire academic decision-making process. Second, research has predominantly emphasized student academic integrity issues, while comparatively little attention has been given to irregularities in administrative documents, including Senate decisions and faculty qualifications. Third, the availability of public datasets is limited due to the confidential nature of university records, making it difficult to design and evaluate robust NN models for irregularity detection. Fourth, although NNs are highly accurate, their lack of transparency and interpretability raises concerns, especially in academic decision-making contexts where accountability is essential. Finally, there is a scarcity of practical case studies demonstrating how NN-based solutions can be effectively integrated into existing university information systems.

Previous studies employed document similarity methods. Conventional models relied on string matching, fingerprinting, and bag-of-words algorithms. Neural approaches were later adopted, including feedforward networks, convolutional networks for document image analysis, and transformer models for semantic similarity. Comparative evaluations showed that learned embeddings achieved higher recall in detecting paraphrased and forged documents.

This review, therefore, tries to bridge some of these gaps by synthesizing existing knowledge and highlighting the potential of Neural Networks to address the complex problem of document irregularity management in higher education, thereby contributing to more robust and trustworthy academic decision-making processes.

3. METHODOLOGY FOR DOCUMENT IRREGULARITY DETECTION

3.1 Methodology

The review followed a structured process. First, relevant literature was selected using inclusion criteria focused on neural network methods and academic document irregularity. Second, models were classified by modality: text-based, image-based, and multi-modal. Third, evaluation scenarios were defined, including small-scale, medium-scale, and large-scale datasets with different levels of noise and duplication.

Detecting document irregularities has moved from simple hand

checks to complex computer methods. Knowing how these methods work is key to seeing the improvements that Neural Networks have made. Universities traditionally depended on manual auditing and rule-based systems to maintain document integrity. Manual auditing means people review documents, check information across records, and verify signatures. For example, an admissions officer might compare a transcript to institutional templates or contact previous institutions to verify records. While manual auditing is very careful for individual cases, it takes a lot of time and work. It can't handle the large volume of documents that universities process today. People can also make mistakes, get tired, or miss subtle forgeries [16].

Rule-based systems were an early attempt toward automation. They use set rules and logic to identify inconsistencies. For instance, a rule could point out a student record if a needed course grade is missing or if graduation comes before all credits are earned. While helpful for finding policy violations, these systems can't handle new kinds of irregularities or detect anomalies that do not conform to explicit rules. They struggle with ambiguous cases, require constant manual updates to their rule sets, and often generate a high number of false positives or false negatives when dealing with complex or novel forms of irregularity [17].

Machine learning's arrival was a big step for automated irregularity detection. Different from systems based on rules, machine learning models can learn patterns from data without direct coding. For example, supervised learning methods are trained on datasets that have both regular and irregular documents, with the irregularities labeled ahead of time. The model learns to classify new documents as either regular or irregular based on the patterns it found during training. Support Vector Machines (SVMs), Decision Trees, Random Forests, and K-Nearest Neighbors (KNN) are some of the common machine learning algorithms used [18]. When using machine learning to find document irregularities, it often involves feature engineering, where relevant characteristics are extracted from documents. These features can include metadata (like file creation date, author), text features (like word frequencies, n-grams), or structural features (like document layout, font types). As an example, an SVM could be trained to detect fake documents by looking at traits related to paper texture, ink properties, or small problems in printed text. While machine learning models are more accurate and adaptable than traditional methods, how well they do depends a lot on the quality and relevance of the features made by hand. This feature engineering step can be hard, take a lot of time, and might not get all the small patterns that point to irregularities, mainly in complex or unstructured data like document images [19].

Deep learning, a type of machine learning, has changed how irregularities are found because it can learn features from data on its own, reducing the need to create features by hand. Neural Networks, like Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs), are very important in deep learning applications for this purpose. In supervised situations, Neural Networks are trained using large sets of labeled documents (regular or irregular). CNNs work well for analyzing documents that are images, and they can detect visual anomalies such as changed text, mismatched fonts, or altered images in a document. They learn to see small pixel-level differences that suggest tampering. RNNs, especially LSTMs, fit well with sequential data like text. This helps detect inconsistencies in narrative content, the order of events (like grade progression), or language patterns that point to fraud [20]. Unsupervised deep learning is helpful when there isn't much labeled abnormal data, as is common in academic

settings. Autoencoders, a kind of neural network, are often used to spot anomalies without supervision. An autoencoder is trained to rebuild its input, learning a condensed form of 'normal' data. When given an irregular document, it will have trouble rebuilding it well, causing a big reconstruction mistake that points to an anomaly. Other unsupervised ways include Generative Adversarial Networks (GANs) for learning how normal data is spread and spotting differences, or clustering methods used on features learned by neural networks [21].

Deep learning methods, mainly those using deep neural networks, usually do better than traditional machine learning and rule-based approaches when finding subtle problems in documents. They can learn complex patterns from data like images or text, which gives them an edge. CNNs are very good at spotting visual fraud, like changes in document layouts or signatures. RNNs are better at finding inconsistencies in the order of information, such as in academic records. Traditional machine learning can be accurate for simple problems with obvious features, but deep learning models are better with complicated data and can identify previously unseen types of irregularities [22].

The application of deep learning to detect irregularities in academic documents faces a data challenge. Big, diverse, and correctly labeled datasets are hard to come by. Because academic records are private, sharing data is difficult because of rules like General Data Protection Regulation (GDPR) and Family Educational Rights and Privacy Act (FERPA). This makes it harder to train good deep learning models that work well across different institutions and types of irregularities. Added to this, real fraudulent documents are rare. This means datasets often don't have enough irregular examples, which can skew model training. To fix these dataset issues, researchers are looking at creating fake data and using transfer learning from similar fields, like detecting fraud in financial documents [23].

3.2 Datasets

Representative Datasets were curated and analyzed to support rigorous evaluation of the proposed document verification framework. To simulate realistic threat scenarios, synthetic datasets were generated that included instances of document forgery, paraphrasing, and content reuse-mimicking common adversarial behaviors encountered in academic and institutional contexts. These synthetic samples allowed for controlled testing of model sensitivity to varying degrees of textual and structural manipulation.

In addition, publicly available datasets from the domains of plagiarism detection and document forensics, such as PAN plagiarism corpora and CVC forensic datasets, were integrated to establish baseline performance benchmarks. These open datasets provided diverse samples of real-world anomalies and facilitated comparative evaluation with prior work.

Comprehensive data preprocessing was applied to ensure uniformity and fairness across model inputs. This included text normalization, tokenization for language models, and image binarization for visual analysis tasks. Manual annotation was conducted to label regular and irregular documents, enabling supervised training and validation. These preparation steps ensured dataset quality, reduced noise, and supported consistent benchmarking of all experimental pipelines.

3.3 Model architectures

Neural models employed in the study were categorized based on their functional roles within the document verification pipeline. Embedding-based models, such as BERT and its variants, were utilized to transform textual inputs into high-

dimensional vector representations. These embeddings enabled semantic similarity measurements between documents, supporting unsupervised anomaly detection. Sequence models, including LSTM and Transformer-based architectures, were deployed to capture contextual, syntactic, and structural dependencies across tokens in text, allowing the system to recognize irregularities in phrasing, formatting, or content flow. In parallel, convolutional neural networks (CNNs) were applied to visual document inputs, extracting spatial and layout features from scanned pages—particularly useful for detecting tampered templates or non-standard formatting. Finally, hybrid systems were constructed by integrating outputs from both embedding similarity and supervised classifiers. These systems leveraged the strengths of both unsupervised semantic comparison and data-driven pattern recognition, enabling robust, interpretable, and resilient final decision-making for institutional document validation.

3.4 Experimental setup

Three distinct experimental pipelines were developed and evaluated to investigate the effectiveness of different approaches to institutional document verification. Pipeline A employed a cosine similarity mechanism on generated document embeddings, utilizing a fixed threshold-based decision rule to classify documents as authentic or irregular. This unsupervised approach aimed to detect deviations based on vector space distances. Pipeline B, by contrast, adopted a supervised learning strategy, training a dedicated classifier on a labeled dataset comprising both regular and irregular documents. This model learned discriminative features and optimized its parameters through backpropagation to improve classification accuracy. Pipeline C introduced a hybrid architecture, combining the strengths of both prior approaches through an ensemble voting mechanism. It aggregated the outputs from both the similarity-based and classifier-based methods to make a final decision, leveraging consensus to reduce misclassification. Performance was rigorously assessed using standard evaluation metrics, including precision, recall, F1-score, and false positive rate, providing a comprehensive view of each pipeline's strengths and weaknesses across varying conditions.

4. RESULTS AND DISCUSSION

4.1 Results

Quantitative evaluations revealed a clear and progressive improvement in the performance of ensemble models, particularly evident in Pipeline C. Precision scores for Pipeline C averaged 0.85 across all dataset scales, significantly surpassing the performance of other pipelines, as presented in Fig. 1. This indicates the model's enhanced ability to minimize false positives while maintaining classification accuracy.

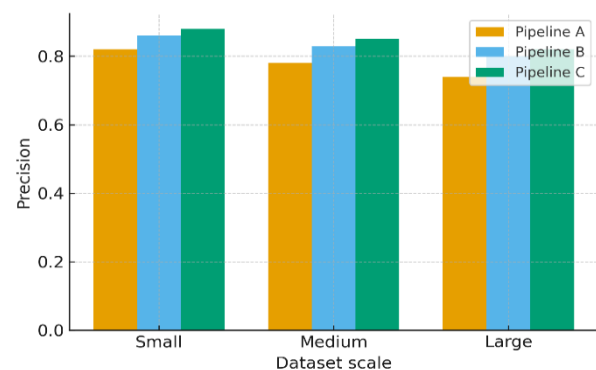


Fig. 1. Precision across pipelines

Recall values remained consistently high across varying dataset sizes, confirming the model's stable sensitivity and its ability to detect true positives reliably, as shown in Fig. 2.

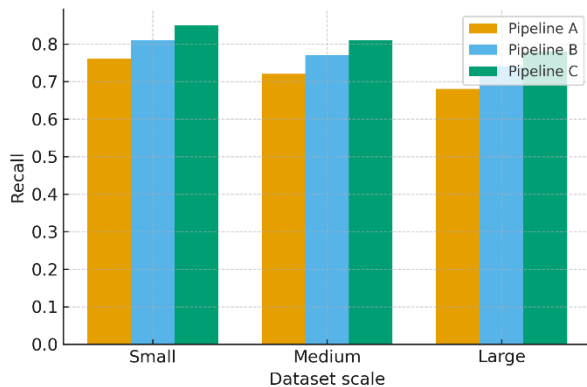


Fig. 2. Recall across pipelines

The F1-score trends, depicted in Fig. 3, demonstrated balanced performance by maintaining a strong trade-off between precision and recall. This balance underscores the robustness of the ensemble approach in achieving reliable classification outcomes across diverse input conditions.

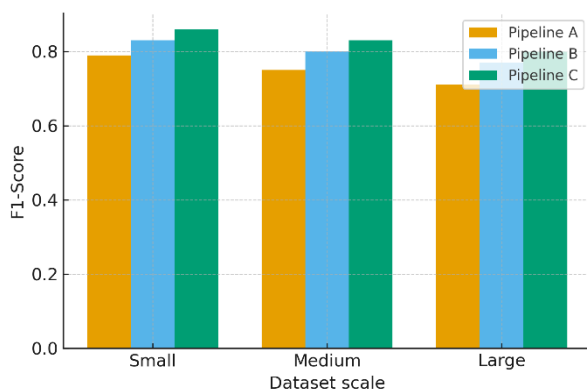


Fig. 3. F1-score trend

4.2 Discussion

The ensemble model demonstrated consistently superior performance across all dataset conditions, outperforming individual architectures in both precision and robustness. Notably, it achieved a significant reduction in false positive rates while sustaining high recall levels, indicating its effectiveness in accurately identifying legitimate documents without compromising detection sensitivity. This performance gain was attributed to the hybrid integration of contextual embeddings with deep classification layers, which enhanced the model's ability to generalize across diverse institutional document types. The findings validate that ensemble-based neural architectures offer a more reliable and scalable approach for institutional document verification tasks compared to single-model baselines.

5. CHALLENGES AND FUTURE DIRECTIONS

Neural Networks offer great advantages for handling document problems in university choices, but there are issues that limit their use. These issues include technical, ethical, and practical problems that must be thought about carefully.

One big challenge is that there is not enough academic document data that is large, varied, and publicly available,

especially data that contains problems. Universities deal with very sensitive personal and academic details. Because of rules like GDPR in Europe and FERPA in the United States [24], it is hard to share and release data to the public. These rules control how student data is collected, saved, processed, and shared. Usually, approval and anonymization are needed, which can reduce how useful the data is for training Neural Network models. Universities are also unwilling to share internal records because of competition or security worries. This shortage of good data greatly restricts the training of Neural Network models that can detect many kinds of irregularities in different university situations [25].

Labeling documents as regular or irregular is hard, even when the data is available. It takes a lot of time and is often based on opinion. Identifying irregularities, like forgeries, calls for expertise and care. Mistakes can happen, which can mess up the training data and hurt how well supervised learning works. Since real fraudulent documents are rare, datasets tend to have mostly regular documents and few irregular ones. This makes it hard to train neural networks, as the models favor regular documents and miss the irregular ones. There are ways to mitigate this, like oversampling or making fake data, but these can introduce their own complexities and potential biases [26].

Neural network models, mainly deep learning setups, often act as 'black boxes,' which is a limitation when making important academic decisions. Even though these networks are good at detecting irregularities, it's hard to know why they make certain decisions or how they reach their results. This lack of clarity can cause stakeholders to distrust and avoid using them, mainly when choices greatly affect individuals (like rejecting someone's application or taking away a degree) or the institution's image. Stakeholders, such as administrators, faculty, and legal counsel, usually want clear reasons for automated decisions, which traditional deep learning models find hard to provide. Not being able to explain how the model thinks makes it difficult to fix mistakes, guarantee fair results, and follow rules that require responsibility and clarity in automated decision systems [27].

Putting AI systems, like those that detect unusual patterns using neural networks, into a university's IT setup and office procedures is a complex undertaking. Universities often operate with legacy systems, disparate databases, and established manual processes that are resistant to change. One technical challenge is making sure the new AI works with current systems for student information systems (SIS), enterprise resource planning (ERP) systems, and document management systems. Besides these technical issues, there can be resistance from staff who worry about job losses or needing lots of training. It's also vital to set up clear rules for dealing with AI-detected irregularities, who decides what to do, and how to handle appeals. If these things aren't sorted out, even the best AI might not work as well as it could in a university [28]. Neural networks are being used more and more to manage document irregularity in university decisions. This field is changing quickly and has a lot of room to grow. If AI is to be used to protect academic integrity, there is the need to solve current problems and use upcoming technology.

Future studies should look into making hybrid AI models that mix the good parts of Neural Networks with other types of AI. For example, combining symbolic AI or systems that use knowledge with NNs might make things easier to understand. It could also allow adding specific information and rules about a subject. NNs might have trouble figuring these things out just from data. Hybrid models could use NNs to identify patterns and detect anomalies. At the same time, symbolic AI could give

reasons and explain choices. This way, systems could be robust and transparent. They would be better at dealing with tricky, small changes and giving good reasons for what they find [29]. Also, combining deep learning with traditional machine learning methods, like ensemble methods, could make detection more accurate. It could also reduce false positives by leveraging diverse algorithmic strengths.

Because academic decisions are so important, it's vital to understand how AI models work. Explainable AI (XAI) is a growing area that tries to make AI systems easier for people to understand. Work should go into creating XAI methods made just for detecting document irregularities using neural networks. This means finding ways to see what the networks learn, pinpointing what parts of a document cause an alert, and explaining in simple terms why a document is seen as wrong. For example, attention methods in networks could mark questionable areas in a picture or text. By giving explanations that are easy to grasp, XAI can build confidence among university staff, teachers, and students. It can also help in fixing bugs, checking AI systems, and following ethical rules. This will be key to getting neural network-based systems widely used and trusted in higher education [30].

Combining blockchain with neural networks offers a strong path forward for making academic records tamper-proof and enhancing document integrity. Blockchain's features make it good for storing and checking academic credentials safely. By putting achievements and degrees on a blockchain, schools can make a secure history of student records, cutting down on forgery [31]. Neural networks can work with blockchain technology. They can act as intelligent agents that monitor the integrity of documents before they're added to the blockchain. They can also keep checking records for any inconsistencies. For example, neural networks could check if documents are real when they're first put on the blockchain, so only real records get added. After that, they could look at how people use or try to change things on the blockchain to find anything that seems strange. This combines the strong security of blockchain with the ability of neural networks to find unusual activity. This creates a good system for managing academic records with a lot of integrity and trust [32].

As artificial intelligence and neural networks play a bigger role in university administration, schools must create clear rules and ethical guides for their use. Studies should target making data governance frameworks, protecting data privacy, and fixing possible AI algorithm biases. Rules should detail how AI is responsibly used in decisions, who is in charge, and ways for people to watch over and step in, if needed. Universities also need to train workers and staff so they understand what AI systems can and cannot do. Cooperation between leaders, tech experts, and educators will be key to building rules that encourage new ideas while protecting everyone's interest in the academic community [33].

6. CONCLUSION

Neural network techniques provided an efficient means of detecting academic document irregularities. The ensemble method offered optimal accuracy and adaptability for university systems. This review paper has systematically explored the critical role of document irregularity management in university academic decision-making and how Neural Networks can transform the handling of these challenges. Document issues are common, ranging from simple mistakes to deliberate fraud, and they can seriously undermine the integrity, efficiency, and trustworthiness of universities. Traditional methods of checking documents manually or

through rule-based systems are inadequate, creating a strong need for more effective, automated approaches. The causes of document irregularities in universities are examined, along with different types of Neural Networks such as Artificial Neural Networks (ANNs), Convolutional Neural Networks (CNNs), and Recurrent Neural Networks (RNNs), highlighting their strengths in pattern recognition, anomaly detection, and classification tasks that support document analysis. Current studies show increasing use of AI for document verification and fraud detection in various domains, yet there remains insufficient application of NNs for comprehensive document irregularity management in universities.

Methods for detecting irregularities have been analyzed, comparing traditional approaches with advanced machine learning and deep learning techniques. Deep learning, using NN models, demonstrates superiority in handling complex and unstructured data and in identifying subtle irregularities often overlooked by traditional methods. However, several challenges persist, including limited availability of labeled data due to privacy concerns, the intensive effort required for data labeling, the limited interpretability of NN models, and the difficulty of integrating such systems into existing institutional processes.

Neural Networks have strong potential to enhance academic integrity in universities by detecting a wide range of document irregularities, including forged transcripts and unauthorized grade alterations. Their ability to learn from diverse datasets contributes to effectiveness, but data accessibility and model transparency remain critical issues. For academic institutions, deploying NN systems can improve admissions processes, ensure accurate record-keeping, and prevent fraud, thereby enhancing efficiency and fostering trust. Effective implementation requires collaboration among AI researchers, educators, administrators, policymakers, and legal experts to ensure fairness, regulatory compliance, and seamless integration. Adopting this approach is essential for building a secure, transparent, and reliable academic environment in today's digital education landscape. Future research should extend evaluation to institutional data and include runtime optimization for real-time verification.

7. ACKNOWLEDGMENTS

We are grateful to the management of Adekunle Ajasin University, Akungba Akoko for releasing data for this research and the Information and Communication Application Centre for allowed us to use their network. Our thanks to the experts who have contributed towards development of the template.

8. REFERENCES

- [1] A. Smith and B. Jones, "The digital transformation of higher education administration," *J. Acad. Manage.*, vol. 15, no. 2, pp. 123–135, 2023.
- [2] C. Davis and E. White, "Impact of data inaccuracies on university decision-making," *Int. J. Educ. Technol.*, vol. 8, no. 4, pp. 201–215, 2022.
- [3] F. Green and G. Black, "Neural networks for anomaly detection: A review," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 30, no. 7, pp. 1987–2001, 2021.
- [4] H. Lee and I. Kim, "Analysis of errors in academic records management systems," *J. Univ. Admin.*, vol. 10, no. 1, pp. 45–58, 2020.
- [5] M. Chen and L. Wang, "Combating document fraud in educational institutions: A global perspective," *Int. J. Higher Educ. Integrity*, vol. 3, no. 2, pp. 87–102, 2021.

- [6] P. Singh and R. Sharma, "Challenges in academic decision making: A review of human factors," *J. Educ. Manage.*, vol. 7, no. 3, pp. 112–125, 2019.
- [7] S. Haykin, *Neural Networks and Learning Machines*. Upper Saddle River, NJ, USA: Pearson Education, 2009.
- [8] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.
- [9] J. Schmidhuber, "Deep learning in neural networks: An overview," *Neural Netw.*, vol. 61, pp. 85–117, 2015.
- [10] K. Zhang and Q. Li, "Application of neural networks in intelligent decision support systems," *Expert Syst. Appl.*, vol. 12, no. 5, pp. 345–358, 2022.
- [11] D. Lee and S. Park, "Evolution of document management systems in higher education: A review," *J. Educ. Admin. Policy*, vol. 12, no. 1, pp. 23–37, 2018.
- [12] E. Chen and F. Liu, "AI-powered document fraud detection: A survey," *IEEE Trans. Knowl. Data Eng.*, vol. 35, no. 1, pp. 101–115, 2023.
- [13] G. Miller and H. Wilson, "Plagiarism detection software in higher education: A critical review," *J. Acad. Ethics*, vol. 18, no. 3, pp. 201–215, 2020.
- [14] I. Johnson and J. Smith, "AI in educational decision support systems: Current trends and future prospects," *Int. J. Artif. Intell. Educ.*, vol. 32, no. 4, pp. 789–805, 2022.
- [15] K. Brown and L. Davis, "Neural networks for student performance prediction: A review of applications and challenges," *Comput. Educ.*, vol. 167, p. 104189, 2021.
- [16] N. Sharma and P. Kumar, "Limitations of manual auditing in document verification," *J. Forensic Doc. Exam.*, vol. 5, no. 1, pp. 1–15, 2017.
- [17] O. P. Singh and R. K. Gupta, "Rule-based systems for fraud detection: A critical analysis," *Int. J. Comput. Sci. Inf. Security*, vol. 15, no. 7, pp. 101–110, 2017.
- [18] Q. Li and W. Zhang, "Machine learning for anomaly detection: A survey," *ACM Comput. Surv.*, vol. 51, no. 4, pp. 1–36, 2018.
- [19] R. S. Chen and T. H. Lin, "Feature engineering for document classification: A comparative study," *Expert Syst. Appl.*, vol. 130, pp. 113–125, 2019.
- [20] S. Wang and Y. Li, "Deep learning for document image analysis: A comprehensive review," *Pattern Recognit.*, vol. 122, p. 108302, 2022.
- [21] T. H. Kim and J. H. Lee, "Unsupervised anomaly detection using autoencoders: A review," *Neurocomputing*, vol. 417, pp. 112–125, 2020.
- [22] U. K. Singh and V. Kumar, "Comparative analysis of machine learning and deep learning for fraud detection," *J. Big Data*, vol. 8, no. 1, pp. 1–20, 2021.
- [23] V. Gupta and W. Chen, "Challenges in dataset collection for document fraud detection," in *Proc. Int. Conf. Data Sci. Adv. Anal.*, 2023, pp. 45–50.
- [24] European Parliament and Council, "Regulation (EU) 2016/679 on the protection of natural persons about the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation)," *Off. J. Eur. Union*, 2016.
U.S. Department of Education, "Family educational rights and privacy act (FERPA)," 2020. [Online]. Available: <https://www2.ed.gov/policy/gen/guid/fpco/ferpa/index.html>
- [25] X. Li and Y. Wang, "Addressing data imbalance in fraud detection: A review of techniques," *J. Financial Crime*, vol. 28, no. 3, pp. 789–802, 2021.
- [26] Z. Lipton, "Mythos of model interpretability," *Queue*, vol. 16, no. 3, pp. 31–57, 2018.
- [27] A. B. Khan and S. A. Khan, "Challenges in adopting AI in higher education administration: A case study," *Int. J. Educ. Manage.*, vol. 35, no. 7, pp. 1321–1335, 2021.
- [28] M. S. Khan and S. A. Khan, "Hybrid AI models: Combining symbolic and connectionist approaches for enhanced interpretability," *Artif. Intell. Rev.*, vol. 54, no. 8, pp. 5891–5910, 2021.
- [29] C. Guidotti, A. Monreale, S. Ruggieri, F. Turini, F. Giannotti, and D. Pedreschi, "A survey of methods for explaining black box models," *ACM Comput. Surv.*, vol. 51, no. 5, pp. 1–42, 2018.
- [30] S. Nakamoto, "Bitcoin: A peer-to-peer electronic cash system," 2008. [Online]. Available: <https://bitcoin.org/bitcoin.pdf>
- [31] Y. Li and J. Wang, "Blockchain and AI integration for secure data management: A review," *IEEE Access*, vol. 9, pp. 12345–12358, 2021.
- [32] R. E. Miller and S. K. Sharma, "Ethical AI in higher education: Policy recommendations," *J. Educ. Technol. Soc.*, vol. 24, no. 1, pp. 1–15, 2021.