



UNIVERSITY OF AEGEAN
FACULTY OF ENGINEERING

DEPARTMENT OF INFORMATION AND COMMUNICATION SYSTEMS ENGINEERING

POSTGRADUATE STUDIES PROGRAM

INTERNET OF THINGS: INTELLIGENT ENVIRONMENTS IN NEXT-GENERATION
NETWORKS

AUTHORSHIP VERIFICATION: IMPOSTORS METHOD
USING IMPOSTORS GENERATED BY LARGE
LANGUAGE MODELS

THESIS

Kaiserlis Alexandros

Supervisor: Stamatatos Efstathios

Members of the selection board:

Samos: November 2024

This page is deliberately empty.

Preface and acknowledgments

This thesis investigates the Impostors Method for authorship verification, leveraging advances in large language models to generate diverse impostors.

My interest in computational linguistics and authorship analysis was sparked during my master's program studies and deepened throughout this research. As I explored the field further, I greatly appreciated its vast scope and rapidly evolving nature. Completing this thesis has strengthened my desire to contribute by developing innovative methodologies to improve verification techniques and expand my knowledge in this dynamic field. The work combines theoretical inquiry with practical experimentation, offering insights into current approaches' potential and limitations.

I want to express my deepest gratitude to my supervisor, Efstathios Stamatatos, for his guidance, encouragement, and expertise throughout this journey. His profound knowledge in this domain and his role as a professor in the Machine Learning course gave me the confidence, motivation, and foundational understanding necessary to pursue research in both authorship verification and the broader fields of machine learning and Artificial Intelligence.

I am also profoundly thankful to the Department of Information and Communication Systems Engineering at the University of the Aegean for allowing me to participate in this master's program. Through the knowledge and skills I gained during my studies, I grew both academically and personally, equipping myself with the tools and enthusiasm needed to pursue further academic work and contribute to my field.

Finally, I am grateful to my family and friends, whose unwavering belief in me made this achievement possible. Their emotional and practical support has been a constant source of strength and motivation, allowing me to persevere and reach this milestone.

© [2024]

KAISERLIS ALEXANDROS

Department of Information and Communication Systems Engineering

UNIVERSITY OF THE AEGEAN

This page is deliberately blank.

Contents

1	Introduction	11
1.1	Introduction.....	11
1.2	Research Problem	12
1.3	Objectives	13
1.4	Research Questions	14
1.5	Thesis Structure	14
2	Literature Review.....	16
2.1	Authorship Verification Overview.....	16
2.1.1	<i>Traditional Methods</i>	<i>17</i>
2.1.2	<i>Analysis of Traditional Methods.....</i>	<i>18</i>
2.1.3	<i>Algorithms for Authorship Verification</i>	<i>19</i>
2.2	Impostors Method	21
2.2.1	<i>Core Concept of the Impostors Method.....</i>	<i>21</i>
2.2.2	<i>Advantages of the Impostors Method.....</i>	<i>22</i>
2.2.3	<i>Implementation in PAN Competitions</i>	<i>22</i>
2.2.4	<i>Evolution and Modern Adaptations.....</i>	<i>23</i>
2.2.5	<i>Challenges and Considerations</i>	<i>23</i>
2.3	Character-Level N-Grams in Authorship Verification	23
2.3.1	<i>Stylistic Features Captured by Character-Level N-Grams</i>	<i>24</i>
2.3.2	<i>Advantages of Character-Level N-Grams in Authorship Verification.....</i>	<i>25</i>
2.3.3	<i>Disadvantages of Character-Level N-Grams in Authorship Verification</i>	<i>25</i>
2.3.4	<i>Performance in Cross-Genre and Cross-Topic Texts.....</i>	<i>26</i>
2.4	Generative AI in Authorship.....	26
2.4.1	<i>The Rise of LLMs in Authorship Analysis.....</i>	<i>26</i>
2.4.2	<i>Generating Convincing Texts with LLMs</i>	<i>27</i>
2.5	Challenges and Gaps in Current Research	28
2.6	The PAN competition	29
2.6.1	<i>The Role of PAN in Authorship Verification Research.....</i>	<i>30</i>
2.6.2	<i>The Impact of PAN on Future Directions.....</i>	<i>30</i>
2.6.3	<i>PAN Datasets and Their Evolution.....</i>	<i>30</i>
3	Methodology	32
3.1	Research Objective and Hypothesis.....	32
3.2	Hypothesis.....	32
3.2.1	<i>How LLM-Based Impostors can make a more Effective Impostor Model.....</i>	<i>33</i>

3.3	The Datasets.....	34
3.3.1	Text Genres.....	34
3.4	Modeling Approaches.....	35
3.4.1	The Impostors Model with Character-Level N-Grams and MinMax Distance	36
3.4.2	Suitability of the Impostors Model for Authorship Verification	36
3.4.3	Potential Research Directions	37
3.5	Generating AI Impostors.....	37
3.5.1	Prompt Design for Impostor Generation.....	38
3.5.2	API Implementation and Workflow.....	38
3.5.3	Challenges and Observations	39
3.6	Experimental Setup.....	39
3.6.1	Evaluation Metrics and the Importance of AUC	39
3.6.2	Experimental Parameters and Variations	40
3.6.3	Cross-Language and Multi-Model AUC Comparison.....	40
4	Experiments & Results	42
4.1	Overview of the Experiments	42
4.1.1	Objective	42
4.1.2	Tuning the model	43
4.2	Experiment 1: Results	49
4.3	Experiment 2: Results	50
4.4	Experiment 3: Results	51
4.5	Experiment 4: Results	51
4.6	Experiment 4: Results	52
4.7	Experiment 5: Results	53
5	Discussion.....	55
5.1	Revisiting the Research Goals	55
5.2	Key Insights	55
5.2.1	Performance Gaps Compared to Prior Work.....	55
5.2.2	Role of Impostor Similarity.....	56
5.2.3	Benefits of Increasing Impostor Diversity and Quantity	56
5.2.4	Cross-Genre and Cross-Topic Limitations	56
5.2.5	The Role of N-grams	56
5.3	Key Insights	57
5.3.1	Strengths of the Experimental Framework	57
5.3.2	Challenges for Cross-Topic and Cross-Genre Verification	57
5.3.3	Advancing the Impostors Method	58

5.4	Broader Implications.....	58
5.4.1	<i>Relevance to Forensic Linguistics and Computational Stylometry</i>	58
5.4.2	<i>Potential Real-World Applications and Constraints</i>	58
5.5	Open Questions.....	59
5.5.1	<i>Refining the Impostors Method</i>	59
5.5.2	<i>Distinguishing Stylistic and Topical Cues</i>	59
5.5.3	<i>Scalability and Efficiency</i>	59
5.5.4	<i>Applicability Across Languages</i>	60
5.5.5	<i>Ethical Considerations</i>	60
6	Conclusion.....	61
6.1	Summary of Findings.....	61
6.2	Practical and Theoretical Implications.....	61
6.3	Directions for Future Research	62
6.4	Final Remarks	62

List of Figures

Number	Caption	Page
Figure 1	<i>The Authorship Verification process using impostors, simplified into a diagrammatic representation.</i>	22
Figure 2	<i>The Authorship Verification process using impostors, illustrated in a diagrammatic representation, with an additional step for generating impostors using prompt engineering when none are available</i>	35

List of Tables

Number	Caption	Page
1	Hyperparameter Tuning for English Language Dataset: Optimal Values Highlighted	44-45
2	Hyperparameter Tuning for Dutch Language Dataset: Optimal Values Highlighted	45-46
3	Hyperparameter Tuning for Greek Language Dataset: Optimal Values Highlighted	46-47
4	Hyperparameter Tuning for Spanish Language Dataset: Optimal Values Highlighted	48-49
5	Prompt 1 Results	50
6	Prompt 2 Results	50
7	Prompt 3 Results	51
8	Prompt 4 Results	52
9	Combined Prompts impostors Results	53
10	Paraphrase Prompt Results	53

Acronyms

Acronym	Full Description
AI	Artificial Intelligence
AV	Authorship Verification
AA	Authorship Attribution
ADI	Authorship DE-Identification
AC	Authorship Characterization
AD	Authorship Discrimination
BERT	Bidirectional Encoder Representations from Transformers
GPT	Generative Pre-trained Transformer
IM	Impostors Method
LLM	Large Language Model
PD	Plagiarism Detection
PAN	Plagiarism, Authorship, and Social Software Misuse
TIS	Text Indexing and Segmentation

Περίληψη

Η παρούσα διπλωματική εργασία ερευνά τη πιθανή δυνατότητα των Μεγάλων Γλωσσικών Μοντέλων (LLMs) να βελτιώσουν τη Μέθοδο Impostors στην επαλήθευση συγγραφικής ταυτότητας. Η Μέθοδος Impostors χρησιμοποιεί κείμενα-απομιμήσεις (impostors) για να δοκιμάσει την ικανότητα ενός μοντέλου να διακρίνει μεταξύ αυθεντικών και μη αυθεντικών συγγραφικών έργων. Παραδοσιακά, τα impostors προέρχονταν από μη σχετικές πηγές, αλλά είχαν κάποια ικανότητα να αποτυπώνουν τις λεπτά στιλιστικά χαρακτηριστικά συγκεκριμένων συγγραφέων, ιδιαίτερα σε διαφορετικά είδη και θεματολογίες.

Με την αξιοποίηση LLMs, όπως το GPT, η έρευνα εξετάζει τη χρήση prompt, ερωτήσεων δηλαδή στα γλωσσικά μοντέλα και την μηχανική αυτών, για τη δημιουργία impostors που μιμούνται όσο το δυνατόν καλύτερα, τα στιλιστικά μοτίβα των γνωστών κειμένων, διατηρώντας παράλληλα το περιεχόμενο διαφοροποιημένο. Τα πειράματα πραγματοποιήθηκαν στο σύνολο δεδομένων PAN 2015, ώστε να αξιολογηθεί αν τα impostors που δημιουργούνται από LLMs βελτιώνουν την ακρίβεια σωστού συσχετισμού μεταξύ του γνωστού και αγνώστου κειμένου σε σύγκριση με προηγούμενες προσεγγίσεις.

Τα αποτελέσματα δείχνουν ότι τα impostors που δημιουργούνται μέσω Μεγάλων Γλωσσικών Μοντέλων παρουσιάζουν δυνατότητες, τουλάχιστον να μεταβάλλουν με θετικό τρόπο την αποδοτικότητα των μοντέλων, ιδιαίτερα σε εργασίες επαλήθευσης μεταξύ διαφορετικών θεματολογιών, αν και απαιτείται περαιτέρω έρευνα. Παρουσιάστηκαν προκλήσεις, όπως το υπολογιστικό κόστος, η αποτελεσματική σχεδίαση prompts και η εφαρμογή σε μεγαλύτερη κλίμακα. Τα ευρήματα αναδεικνύουν την προοπτική ενσωμάτωσης της παραγωγικής τεχνητής νοημοσύνης (Generative AI) σε πλαίσια επαλήθευσης συγγραφικής ταυτότητας, ανοίγοντας τον δρόμο για πιο προσαρμοστικές και αξιόπιστες μεθόδους σε ακαδημαϊκά, εγκληματολογικά και ψηφιακά περιβάλλοντα.

Λέξεις Κλειδιά: *Impostors Method, Authorship Analysis, Authorship Verification, Large Language Models*

Abstract

This thesis investigates the potential of Large Language Models (LLMs) to enhance the Impostors Method for authorship verification. The Impostors Method introduces distractor texts to test a system's ability to differentiate between genuine and inauthentic authorial works. Typically, these impostors came from unrelated backgrounds but had some ability to capture the subtle stylistic nuances of specific authors, especially across different genres and topics.

Leveraging LLMs such as GPT, this research explores the use of prompt-engineered generative techniques to create impostors that closely mimic stylistic patterns of known texts while introducing distinct content. Experiments were conducted on the PAN 2015 dataset to evaluate whether LLM-generated impostors improve verification accuracy compared to traditional approaches.

Results suggest that LLM-generated impostors show promise in differentiating and enhancing model robustness among other studies, especially in cross-domain verification tasks, although further research is necessary. Challenges such as computational cost, effective prompt design, and scalability remain significant. The findings underscore the promise of integrating generative AI into authorship verification frameworks, paving the way for more adaptable and reliable methods in academic, forensic, and cybersecurity applications.

Keywords: *Impostors Method, Authorship Analysis, Authorship Verification, Large Language Models*

1

Introduction

1.1 Introduction

Authorship analysis is a field that investigates written texts to identify the linguistic features of their authors. This field is important when it comes to forensic linguistics, academic integrity, and online identity confirmation. A writer's distinctive linguistic style can often be identified through their word choices, sentence structures, writing patterns, and overall style (Elmanarelbouanani & Kassou, 2013). Such features help experts establish authorship of disputed documents or associate anonymous texts with particular individuals, assisting law enforcement and legal professionals.

In recent years, this field shifted from traditional methods such as stylometry, which is based on statistics of writing style, to more advanced techniques that involve machine learning or neural networks (Stamatatos E. , 2009). Stylometric techniques analyze elements such as word frequencies, sentence lengths, and function word usage, offering a foundation for identifying consistent authorial traits. However, these techniques struggle with genre and literary differences between authors (He, Lashkari, Vombatkere, & Sharma, 2024)

The rise of large, readily available datasets and the evolution of computational resources are now allowing models to capture complex stylistic patterns at scale in an easier and more resourceful manner. Recent techniques utilize deep learning to detect and learn these subtle features with no significant input or feature extraction, leading to better performance in tasks such as cross-genre and cross-topic authorship verification (Stamatatos E. , Authorship Verification: A Review of Recent Advances, 2016). This shift has led to more adaptable methods better suited for real-world applications where text types and domains often vary (Elmanarelbouanani & Kassou, 2013) (Stamatatos, 2020).

Authorship analysis covers a diverse set of sub-fields, each providing a solution to one or more problems concerning the recognition and characterization of authorship:

- **Authorship Identification:** Focused on determining the author of a text, including:
 - *Authorship Attribution (AA)*: Selecting an author from a set of candidates.
 - *Authorship Verification (AV)*: Deciding if a specific author wrote a particular text.
 - *Authorship Discrimination (AD)*: Assessing whether two texts share the same author.
- **Authorship Profiling:** Estimating the demographic features of the author, such as:
 - *Authorship Characterization (AC)*: Identifying traits like age, gender, and occupation.
 - *Authorship De-Identification (ADI)*: Removing stylistic traces to anonymize texts.

- **Plagiarism and Forensic Analysis:** Addressing the reuse and identification of content, including:
 - *Plagiarism Detection (PD):* Detecting copied or paraphrased passages.
 - *Code Authorship Attribution:* Identifying authorship in software and programming for forensic purposes.
- **Content and Misinformation Analysis:** Applications in identifying misinformation or managing text collections, such as:
 - *Disinformation Analysis:* Uncovering false identities and fake news sources.
 - *Text Indexing and Segmentation (TIS):* Structuring mixed-author content in global documents or forums.

From traditional methods to neural models, each area takes advantage of different written or coded language patterns. Potential uses include fighting cyber-attacks, improving digital forensics, and protecting intellectual property.

Authorship verification (AV) has evolved from traditional to advanced computational methods. Approaches like Burrows' Delta, which measure distances between text styles derived from word frequencies, work well on limited datasets but have difficulty with cross-genre variation when content changes and can mislead models. More recent methods, such as BERT and GPT, employed neural networks and transformers to extract text documents' deeper syntactic and semantic properties. They can discover complex patterns in how an author writes, especially when dealing with different domains in tasks such as cross-domain verification.

A method in this field deployed for this reason is the Impostors Method, which improves robustness by incorporating distractor texts, offering an extrinsic method that tests the model's skill in distinguishing between authentic and inauthentic samples. This approach urges models to concentrate on more profound stylistic details instead of just surface characteristics by analyzing the similarities between a disputed text and both authentic samples and impostor texts.

The PAN (Plagiarism, Authorship, and Social Software Analysis) competitions have been incredibly important for pushing the boundaries of research in authorship verification. Almost every year, by providing standardized datasets, these competitions offer a rich mix of genres and languages that reflect real-world challenges faced by verification models. The PAN datasets are made in such a way that they can test how well models can manage texts from different domains and diverse writing styles, which is crucial for ensuring their reliability. Additionally, these competitions inspire new ideas and advancements in traditional methods and cutting-edge neural network approaches to authorship verification. The competition drives innovations in traditional and neural network-based verification methods.

1.2 Research Problem

In the world of Generative AI, large language models (LLMs) like GPT and BERT are really making an impact. These models can generate text that feels natural and flows well, making them useful in various ways. Their ability to do that is taken from being trained on vast amounts of data, which helps them understand complex patterns in language. This technology allows people to tackle tasks

such as writing, translating, and summarizing with a deeper level of complexity, unlocking new chances that we never thought possible before.

But the effects of these large language models go way beyond just practical uses in daily life. Researchers are tapping into the power of AI tools to ignite innovation across different fields, changing the way we approach everything; from scientific inquiries to creative problem-solving, inspiring us to think differently and explore new possibilities. (Bevendorff, et al., 2023)

One of the features of LLMs, is their ability to replicate different writing styles and tones. Because they have learned linguistic patterns from various sources, they are effective at maintaining a stylistic consistency and thus replicating crucial characteristics in certain contexts. They can generate text that reflects specific patterns taken from a given text, making them perfect for producing writing that echoes the styles of particular authors. This ability can be fine-tuned through prompt engineering, where specific written or spoken instructions help direct the LLM to capture the stylistic features of a certain writing style, as instructed. This way, we can create synthetic texts resembling an author's established voice.

The Impostors Method has become a noteworthy strategy for verifying authorship, particularly useful when data is scarce or when multiple potential authors are unavailable for comparison (Koppel, Schler, & Argamon, 2009) (Stamatatos E. , 2009). This method works by introducing distraction texts—impostors—that put the verification system to the test in its ability to differentiate between genuine and non-genuine authorial works. Traditionally, these impostor texts were pulled from unrelated sources and relied on patterns like n-grams or other stylistic features to create a credible challenge for the system. However, this older approach often struggled to capture the subtle nuances of an author's unique style across various genres.

These interesting aspects of impostors lead us to the potential to use Large Language Models (LLMs) within the Impostors Method. By generating impostor texts that closely resemble the stylistic patterns of known or unknown texts, we can make the verification process even more intricate—and accurate. LLMs can be directed to produce writings that mimic the structure and tone of the target author's work while presenting different content. This makes the model focus on more profound stylistic details rather than surface similarities. As a result, we create a more varied and contextually relevant set of impostors, which theoretically makes it easier for the model to differentiate between genuine authorship and impostors accurately.

This research explores ways to exploit the relationship between LLMs and the Impostors Method in Authorship Verification. With the use of LLMs to craft customized impostor texts with specific prompts, we want to determine if these AI-generated impostors can enhance the efficiency and precision of authorship verification models. We will test this approach using datasets from the PAN competitions to see whether the LLM-generated impostors can present a more formidable challenge for verification models compared to traditionally selected texts while examining their impact on various tasks.

1.3 Objectives

This thesis investigates the role of Generative AI (Gen AI) in enhancing the Impostors Method for authorship verification. The central question guiding this research is whether AI-generated impostor texts closely mimic a known or unknown author's writing style and can significantly improve

models' ability to identify a text's authorship. By generating impostor texts that are stylistically aligned with the target documents in various ways, we explore how Large Language Models (LLMs) can be effectively prompted to create these more convincing impostors, challenging the capabilities of existing verification methods.

To assess the effectiveness of these LLM-generated impostors, the study employs the PAN dataset from 2015. This dataset features a diverse range of genres and complexities that are essential for authorship verification analysis. We aim to explore whether fine-tuning prompts can generate impostors that closely mimic known texts, while also aligning in what makes a good impostor. By using prompt-based generation, we're looking to raise the difficulty level of these impostors for existing models. Ultimately, we want to examine how this approach can influence and be applied in cross-domain verification tasks.

By contacting this thesis, we aim to make a noteworthy contribution to the field of Authorship Analysis by exploring how Gen-AI can be integrated into authorship verification frameworks. Our findings may offer innovative and reliable solutions in forensic and academic contexts, enhancing the accuracy and adaptability of authorship verification methods.

1.4 Research Questions

This thesis explores several important research questions related to using Generative AI (Gen AI) in enhancing the Impostors Method for authorship verification. The guiding questions for this research are:

- Can texts generated by Generative AI, specifically through Large Language Models (LLMs), improve the accuracy of authorship verification models compared to traditional impostors that are selected manually?
- In what ways can LLM-generated impostors be fine-tuned through prompt engineering to more closely resemble the stylistic features of both known and unknown texts across various genres and topics, thereby achieving better outcomes, and what is the role of n-grams, that the stylistic features are taken by?
- Do LLM-generated impostors contribute to improved performance in cross-domain verification tasks, where the texts may differ significantly in genre and subject matter?
- How might this approach be applied to real-world datasets, such as those from 2015, and what measurable improvements can we observe regarding model generalization and accuracy?
- How will integrating the Impostors Model with Generative AI differ from the latest models in the field, and what additional ways can Gen AI enhance traditional impostor-based verification methods?

1.5 Thesis Structure

This thesis is divided into 6 Sections. The structure is as follows:

1. Introduction

This section provides an overview of Authorship Analysis, and its importance in various real-world applications, including forensic linguistics and plagiarism detection. We will set the stage for the study by discussing the broader context of authorship verification and introducing the Impostors Method, detailing its significance in confirming authorship. Additionally, we will delve into the rise of Generative AI, particularly focusing on Large Language Models (LLMs), and explore how these advancements can bolster the Impostors Method.

2. Literature Review

This chapter will examine the literature surrounding authorship verification, particularly emphasizing the Impostors Method and its evolution over the past decade. We will also look at recent developments in Generative AI, especially regarding text generation and authorship-related tasks.

3. Methodology

In this chapter, we will outline the experimental design to assess the effectiveness of AI-generated impostors in authorship verification. This will include a comprehensive description of the 2015 PAN dataset utilized in this research. We will explain the process of prompt engineering to craft high-quality impostor texts, as well as the methods used to compare the performance of traditional impostors against those generated by LLMs. Additionally, we will present the evaluation metrics and testing processes implemented during cross-domain verification tasks.

4. Experiments and Results

In this chapter, we will present the experimental findings, through a comparative performance of data between models that utilize traditional impostors and our model that incorporates LLM-generated impostors while explaining the methods used.

5. Discussion

This chapter will assess the practical advantages of integrating Generative AI with the Impostors Method and identify any limitations encountered during the study. Potential paths for future research will be examined, while comparing our approach with existing models, evaluating whether the use of LLMs can lead to notable improvements in authorship verification.

6. Conclusion

A summary of the thesis's key contributions reaffirm how Generative AI can play a part in the Impostors Method for authorship verification. It will offer insights into how these findings can be applied to forensic investigations, academic integrity, and cybersecurity and suggest directions for further research in the field of authorship analysis using advanced AI techniques.

2

Literature Review

2.1 Authorship Verification Overview

Authorship Analysis, that focuses on determining whether a particular author wrote a given text. Considered as a binary classification task, it operates by examining the linguistic and stylistic features of the text to produce a yes/no answer. Unlike authorship attribution, which seeks to identify the most likely author from a set of candidates, authorship verification aims to answer a narrower question: Does this text match the writing style of the given author? (Stamatatos E., 2016)

The importance of Authorship Verification is clear in many real-world applications, especially in areas like forensic linguistics. It is very important to determine who wrote legal documents, contracts, or texts without names in these situations. One big challenge in this work is to differentiate between the real writing style of an author and forgeries or anonymous writings. This task needs models to pay attention to consistent linguistic signs, like the way words are used, sentence structures, and habits in punctuation; these are like unique fingerprints of a writing style. These markers help the model to understand an author's typical writing behavior, even when dealing with texts that might be misleading or changed (Stamatatos E. , 2009).

A big challenge in verifying authorship, is when there is variation in genre or domain. The writing style of an author can change a lot depending on the context. For example, it can be very different when they write scientific papers compared to when they write blog posts, or when they write fiction versus non-fiction. This genre variation makes it difficult for models to rely on superficial linguistic traits, requiring a more nuanced approach that can detect deeper stylistic features. Research has shown that accounting for genre variation significantly improves the performance of verification models, as ignoring these shifts in style can lead to high error rates, mainly when the model is tested on texts from unseen genres (Stamatatos E. 2013).

A big challenge in verifying authorship, is when there is variation in genre or domain. The writing style of an author can change a lot depending on the context. For example, it can be very different when they write scientific papers compared to when they write blog posts, or when they write fiction versus non-fiction. This genre variation makes it difficult for models to rely on superficial linguistic traits, requiring a more nuanced approach that can detect deeper stylistic features. Research has shown that accounting for genre variation significantly improves the performance of verification models, as ignoring these shifts in style can lead to high error rates, mainly when the model is tested on texts from unseen genres. (Stamatatos E., 2016)

In addition to genre variation, authorship verification often faces the challenge of open-set vs. closed-set classification. In closed-set verification, the model works with a predefined set of authors, and the task is to verify whether a given text belongs to one of them. However, real-world scenarios typically fall under open-set verification, where the text may have been written by an unknown

author who is not present in the training data. Open-set verification makes the task much more difficult, because the model must verify not only the authors that are known but also mark out texts that do not fit any known profile (Stamatatos E. , Authorship Verification: A Review of Recent Advances, 2016). This open-set characteristic shows how unpredictable real-world situations can be, where anonymous or forged texts might present new authors that were not considered in the original training.

Another critical difficulty is the data scarcity typical in many authorship verification tasks. Often, only a few texts from an author are available for model training, making the task suitable for few-shot or one-shot learning approaches. In such cases, models must generalize effectively from minimal examples, requiring methods to extract meaningful stylistic features from limited data. Recent advancements in transfer learning and meta-learning techniques have shown good potential, allowing models to perform well even when there are only a few known samples from an author (Stamatatos, Kestemont, & & Manjavacas, 2020). These methods are very useful in practical situations where it is not possible to collect a large amount of labeled data.

In addition, the problem of class imbalance is a big challenge in verification of authorship. Usually, the dataset has many more negative examples (where authorship is not confirmed) than positive examples (where authorship is validated). This imbalance can create bias in the models, leaning more towards negative classifications, which leads to higher rates of false negatives. To fix this problem, new research shows some techniques like balancing the training data with sampling methods or changing the loss functions to make positive examples more important. These methods help to lessen the models' habit of preferring negative predictions, which improves the overall performance. (Stamatatos E. 2016)

2.1.1 Traditional Methods

In early stages, approaches to verifying authorship relied a lot on stylometry, which is the statistical analysis of writing style. Stylometric methods concentrate on finding consistent patterns in how an author writes, no matter what the content or subject is. Some of the most commonly studied features in these early methods include distributions of word frequency, sentence length, and the use of function words like prepositions and conjunctions. These features are considered stylistically important because they are not easily altered by the author, making them more reliable indicators of an individual writing style. (Stamatatos E. , 2009)

One well-known example is Burrows' Delta, which measures the difference in word frequencies between texts, and this method works well across various genres and languages. (Burrows, 2002). Another basic technique in traditional stylometry is n-gram analysis, where sequences of words or characters (such as bigrams and trigrams) are extracted and compared among texts. N-grams are useful because they capture patterns of vocabulary and syntax, showing an author's word choices and sentence structures (Stamatatos E. , 2009). For instance, character-level n-grams can show consistent use of punctuation, spelling differences, or common letter combinations that are often unique to a specific author. This method is particularly effective in situations where there is limited text for analysis, like in forensic investigations or when verifying short pieces of writing online (Koppel, Schler, & Argamon, 2009).

Function words, in particular, play a significant role in traditional authorship verification because they are content-independent and are used unconsciously. Unlike content words (nouns, verbs, adjectives), function words are used by all authors but in stylistically distinctive ways. For example, one author might prefer shorter, more direct sentences with frequent use of conjunctions, while another might favor more complex sentence structures with more significant variation in prepositions. These subtle but persistent patterns provide a robust basis for verification, even when dealing with short or genre-diverse texts (Kestemont, 2014).

Traditional methods have also explored syntactic patterns, focusing on sentence structure rather than word usage. Analyzing how an author constructs sentences—whether they favor complex, nested clauses or simpler, declarative statements—provides insight into their unique style. While syntactic features tend to be noisier and more complicated to extract accurately (especially in earlier models without sophisticated natural language processing tools), they nonetheless offer valuable complementary information.

Based on traditional methods in verifying authorship, new research presents a different viewpoint by including cognitive linguistics in the exploration of authorship. Although computational techniques have improved greatly, they often lack clarity and do not fully consider the cognitive elements of language processing; an important idea that is introduced is the concept of linguistic individuality. This means that every person has a unique way of using language, which is affected by their own experiences and how they think. This individuality shows in the choice of words and also in the deeper structures of language, such as how different language elements are combined and arranged. This point of view suggests that understanding these cognitive factors gives us a better insight into why certain stylometric features, like function words and n-grams, are useful in proving who the author is. By linking computational methods with knowledge from cognitive linguistics, this approach encourages a more complete way to analyze who wrote a text. (Nini, 2023)

2.1.2 Analysis of Traditional Methods

Traditional methods for stylometry, like analyzing word frequency, using n-grams, and looking at function words, are effective for verifying authorship, but they have some limits. First, these methods depend a lot on features that are created by hand, meaning they need someone to manually find the stylistic markers and cannot easily adapt to new or unseen styles of writing. This lack of flexibility can be a problem, especially when trying to verify authorship across different genres or topics, where an author's writing can change a lot depending on the subject.

Moreover, traditional methods can be easily vulnerable to content leakage. This happens when the model mixes up stylistic similarity with similarity in themes or topics. For instance, if two authors talk about the same subject, they might use similar vocabulary that is specific to that area. This can lead to mistakes, resulting in false positives during verification tasks. As a result, these approaches may struggle when the content of the texts being compared is highly related, even if the authorship is different. (Stamatatos E. , 2009) (Stamatatos E. , 2016)

Another challenge is the limited scope of text that traditional methods can handle. Many early models perform well when a significant amount of text is available from known and questioned authors, but they may struggle with concise texts. This limitation is particularly relevant in modern contexts like social media, where posts are brief, and the variability between texts is high. Also, the

traditional methods are mostly dependent on the language, needing special changes for different languages and writing systems.

Even with these difficulties, traditional stylometric methods have created a base for more advanced techniques by showing that style—different from the content—can be measured and used well for analyzing authorship. The rise of machine learning and neural network approaches has sought to overcome many limitations by automating the feature extraction process and allowing models to learn more complex patterns in writing. However, even as newer models gain prominence, stylometric techniques remain an essential part of the authorship verification toolbox, often serving as a baseline against which newer methods are tested.

2.1.3 Algorithms for Authorship Verification

While traditional stylometric approaches laid the foundation for authorship verification, introducing real-life challenges, such as dealing with cross-genre variation, handling open-set verification, and addressing data scarcity, has driven the development of more advanced algorithms. These challenges show how complex it is to verify authorship in real situations, where the models need to work well in different writing styles and manage situations when the real author is not in the dataset. Several state-of-the-art models have emerged to address these issues, particularly those showcased in PAN competitions, where authorship verification tasks are a central focus. These competitions have pushed for new ideas and methods, creating techniques such as meta-learning and few-shot learning. They also worked on better ways to handle class imbalance and improve how models can generalize across different writing styles.

2.1.3.1 Support Vector Machines (SVM)

One of the most widely used algorithms in earlier authorship verification tasks is the Support Vector Machine (SVM). This supervised learning model is particularly effective in binary classification problems like authorship verification. SVMs work by finding the optimal hyperplane that separates data points from two classes—in this case, determining whether a text belongs to a particular author. Using feature sets like n-grams, function words, and stylistic markers, SVMs have performed well in distinguishing between authors based on textual features. (Stamatatos E. , 2009)

In PAN competitions, SVMs have been used extensively as baseline models due to their ability to handle high-dimensional data efficiently and their robustness across different datasets (Stamatatos E., 2016). SVMs were especially prominent in the early 2010s PAN competitions, where they were combined with feature engineering techniques such as TF-IDF (Term Frequency-Inverse Document Frequency) to improve their discriminative power in authorship verification tasks (Halvani, Winter, & Pflug, 2013). These techniques allowed SVMs to capture important textual patterns crucial for distinguishing between different writing styles.

2.1.3.2 Ensemble Methods

Ensemble models have also become popular in authorship verification, as they combine the strengths of multiple classifiers to achieve better results. One prominent example from the PAN competitions is the use of Random Forests, which build multiple decision trees and aggregate their predictions to improve accuracy. By averaging the results from several trees, Random Forests

reduce overfitting and are capable of capturing a wider variety of stylistic features from text. (Stamatatos E., 2016)

Ensemble approaches also include boosting algorithms like Gradient Boosting Machines (GBMs) and XGBoost, which iteratively build models by focusing on errors made by previous iterations. These methods have shown strong performance in PAN evaluations due to their ability to handle imbalanced data and noisy text samples (Halvani et al., 2013). In particular, ensemble models are valuable when working with cross-genre or cross-domain tasks, where the diversity of text samples requires more nuanced classification techniques (Potthast, Stein, & Hagen, 2016).

2.1.3.3 Siamese Neural Networks

Siamese Neural Networks have gained significant traction in authorship verification in recent years. These networks are designed to compare two input texts and determine whether the same author wrote them. The architecture consists of two identical neural networks that share weights, processing two input texts simultaneously. The network then learns a similarity function to compare the texts' stylistic and linguistic features. This method has been highly influential in one-shot learning scenarios, where only a few samples are available for comparison. (Reimers & Gurevych, 2019)

The use of Siamese Networks was particularly prominent in the PAN 2016 and 2017 competitions, where they achieved high accuracy in challenging tasks involving short texts and diverse writing styles (Potthast, Stein, & Hagen, 2016). These networks are well-suited for the PAN competitions' binary authorship verification tasks, where the system must decide whether a given text belongs to a known author (Potthast, Stein, & Hagen, 2016).

2.1.3.4 Transformer-Based Models

Models using transformers, like BERT and GPT, have really changed how we verify who wrote something by offering deep insights into the structure and meaning of texts. Unlike older techniques that depend on features created manually, BERT can learn complex relationships in the text on its own, which makes it very good at understanding how an author has their own unique style. (Manolache, et al., 2021)

These models did better than previous methods in the PAN 2020 and 2021 competitions, especially in cross-domain verification, where texts come from different topics or genres (Ibrahim et al., 2023). BERT's ability to be fine-tuned allowed it to handle various types of data, such as fanfiction, adjusting to many different styles and content variations (Tyo, Dhingra, & Lipton, 2021). This has set a new standard in authorship verification, offering better accuracy and flexibility.

2.1.3.5 Graph-Based Approaches

Graph-based methods have attracted much interest in the field of authorship verification, particularly in PAN competitions. In these approaches, texts are represented as nodes in a graph, with edges indicating the relationships formed by shared stylistic or content traits. By analyzing the resulting graph, models can identify clusters of texts that likely belong to the same author, making this method effective for cross-document analysis and linking multiple texts to a joint author (Ibrahim, et al., 2023), Improving authorship verification with sentence-transformers. This method

is especially helpful in complicated situations where direct comparisons may not work well, offering a deeper insight into the stylistic relationships between different documents.

2.1.3.6 Impostors Method

Lastly, the Impostors Method, which is the main focus of this thesis, pushes the boundaries of the verification model by introducing unrelated texts (impostors) to imitate possible false positives. The model then needs to differentiate between authentic author texts and these impostors. Over the years, the impostor's method has developed, combining with advanced machine learning techniques, like neural networks and ensemble models, to build more strong and effective authorship verification systems.

2.2 Impostors Method

The Impostors Method is a well-known approach for verifying authorship that has received a lot of attention in recent years. It was first developed to enhance the strength and dependability of authorship classifiers by showing them a set of unrelated and distracting texts, called impostors, during the verification process. Unlike traditional methods that compare a questioned text only to samples from the author, the Impostors Method evaluates how well an authorship model can differentiate between genuine samples and unrelated texts, thus creating a more challenging and rigorous evaluation framework. (Stamatatos E. , 2016)

2.2.1 Core Concept of the Impostors Method

The core idea behind the Impostors Method is to simulate potential false positives by introducing texts from authors other than the target author during the verification process. These impostor texts do not belong to the actual author. However, they are used to test whether the verification system can accurately identify a given text's authorship or is susceptible to mistaking an impostor text for a genuine one. The method is framed as a binary classification problem: the model must determine whether a specific individual authored a new text, even in the presence of challenging impostor texts.

In practice, the process consists of taking a selection of documents that are known to be written by the author in question and then comparing them to the text that is being analyzed. Additionally, texts from other authors, which we call impostor texts, are included in this comparison. The system for verification then calculates how similar the questioned text is to each of the known documents as well as to the impostor texts. If the similarity score between the questioned text and the known

documents is higher than that with the impostor texts, the conclusion is that the text was likely written by the same author.

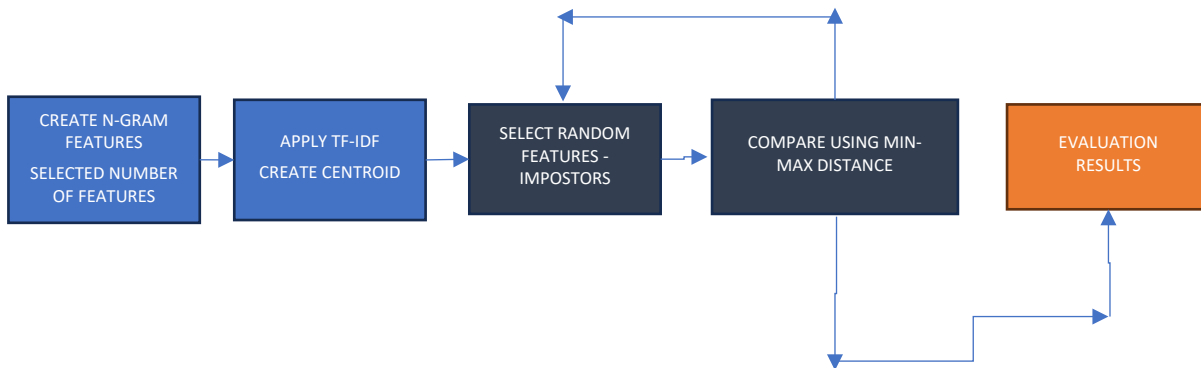


Figure 1: The Authorship Verification process using impostors, simplified into a diagrammatic representation.

2.2.2 Advantages of the Impostors Method

The Impostors Method presents several benefits when compared to more straightforward methods of comparison. One significant advantage is that it compels the verification model to concentrate on stylistic features that tend to remain consistent across an author's body of work. This approach is particularly useful when dealing with texts that might differ in genre or subject matter, which is common in fields like forensic linguistics and plagiarism detection.

Furthermore, the Impostors Method proves to be very helpful in situations where there is a shortage of material from the author, making it particularly suitable for cases with only a few known samples. Using a set of fake texts, this approach helps to avoid overfitting.

Overfitting happens when the model becomes too focused on the familiar samples, which complicates its ability to deal with new, unseen texts. In this way, the impostor texts serve as a type of regularization, enabling the model to test its ability to tell the difference between real and fake texts in a more diverse environment. (Stamatatos E., 2016)

2.2.3 Implementation in PAN Competitions

The method known as the Impostors Method is frequently utilized in PAN competitions and is seen as a standard for assessing various techniques in Authorship Verification. For instance, in the 2013 and 204 PAN authorship analysis, participants were given the task to verify authorship across various genres and languages using this method. It showed especially good results in cross-genre verification, where the known and questioned texts were from different writing styles, like essays and novels. The advantage of this approach is that it emphasizes style instead of content, which helps prevent any misleading conclusions that might come from similarities in content. (Stamatatos E., 2016)

In these competitions, the approach usually uses a distance metric to measure how similar or different the writing styles are between texts. This can include methods like cosine similarity or the

Jaccard index, or even more advanced techniques such as kernel-based methods. By comparing the distances among texts from known authors, texts by impostors, and the text in question, the system can make a more precise decision about who the true author is. Some more sophisticated versions of this method even combine ensemble learning with what's called the Impostors Method, adding different similarity measures together to improve accuracy.

2.2.4 Evolution and Modern Adaptations

The Impostors Method has been improved over the last years, using more advanced techniques from current machine learning. The new developments now make use of neural network embeddings, which help to capture the stylistic features of texts more effectively. This allows for more precise comparisons between known documents and those under question. These embeddings are capable of recognizing complex patterns in style and usage, going beyond the traditional methods like n-grams or simple lexical features.

Additionally, the Impostors Method has started using transformer models like BERT and GPT-based systems. These models can analyze bigger portions of text and detect more complex patterns in writing styles. By employing fine-tuned transformer models, researchers can create comparisons that are very similar to the individual style of a specific author, which makes the process of verification harder. This approach prevents simple false positives since it pushes models to distinguish between subtle stylistic differences rather than just looking at surface features. These modern adaptations make the Impostors Method especially relevant today, as Generative AI is increasingly involved in producing both genuine and synthetic texts. (Manolache, et al., 2021)

2.2.5 Challenges and Considerations

Despite its strengths, the Impostors Method also faces particular challenges. Selecting the right impostor texts is critical to the method's success; using texts that are too dissimilar to the author's style may not provide a meaningful challenge to the verification model while selecting texts that are too similar could artificially inflate error rates. Balancing the diversity of impostors with their stylistic resemblance to the genuine author is an ongoing area of research, especially as models become more sophisticated.

Moreover, comparing multiple impostor texts can be computationally complex, particularly when using neural network-based methods, which require significant processing power. The trade-off between computational cost and the robustness of the verification system is a key consideration in the practical implementation of the Impostor Method.

2.3 Character-Level N-Grams in Authorship Verification

Character-level n-grams are sequences of characters, typically ranging from two to five characters in length (e.g., bigrams, trigrams, four-grams, five-grams). In authorship verification (AV), character-level n-grams provide vectorized representations of text by capturing specific sequences and patterns of characters, which contribute to an author's stylistic fingerprint (Manolache, et al., 2021) (Kestemont, 2014). This approach emphasizes structural and stylistic elements of writing—such as punctuation, syntax, and word formation—that often remain consistent across an author's work, regardless of content (Stamatatos E. , 2009)

2.3.1 *Stylistic Features Captured by Character-Level N-Grams*

Character-level n-grams effectively capture nuanced stylistic markers that characterize an author's writing style. Key features include:

1. **Punctuation Patterns:** Character-level n-grams capture punctuation usage, such as the frequency and placement of commas, periods, colons, and dashes, which contribute to an author's rhythm, sentence flow, and stylistic pauses (Stamatatos E. , 2009)
2. **Syntax and Grammar Structure:** By focusing on characters, n-grams can capture repeated syntactic patterns, such as common prefixes, suffixes, and specific word formations. These elements reveal an author's preferred sentence structures, contractions, and grammar choices. (Stamatatos E. , 2009)
3. **Word Morphology and Function Word Usage:** Character n-grams capture the structure of function words (e.g., “and,” “the,” “of”) and morphological aspects, such as common suffixes or prefixes, which contribute to an author's unique linguistic style (Juola, 2015).
4. **Special Characters and Formatting:** N-grams capture idiosyncratic stylistic markers, such as the use of special characters (e.g., “&,” “@”) or unique spelling conventions, that often distinguish an author's style (Manolache, et al., 2021).
5. **Idiosyncrasies and Typographical Patterns:** Character n-grams reflect an author's recurring typographical choices, such as spelling preferences (e.g., “realize” vs. “realize”), which may be influenced by regional language use or personal style (Kestemont, 2014) (Stamatatos E. , 2009).

The n-gram size chosen—whether lower or higher—impacts the specificity and detail of the stylistic information captured.

Lower n-gram Sizes (e.g., 2-grams, 3-grams)

Lower-order n-grams, such as bigrams and trigrams, effectively capture broad, frequent patterns. These include:

- **Basic Stylistic Markers:** Lower n-grams capture standard features such as short words, basic punctuation, and prefixes or suffixes (Stamatatos E. , 2009).
- **General Representation:** Because these patterns are widespread, they tend to reflect general stylistic tendencies rather than individual nuances, making them less effective for distinguishing between specific authors but useful for recognizing language or genre-wide patterns (Peng, Schuurmans, Keselj, & Wang, 2003)
- **Reduced Uniqueness:** Shorter sequences often overlap across authors, limiting their utility for capturing author-specific traits and making them best suited for identifying broader linguistic characteristics.

Higher n-gram Sizes (e.g., 4-grams, 5-grams)

Higher-order n-grams provide a more detailed, nuanced view of an author's style, capturing:

- **Complex Stylistic Patterns:** Longer n-grams encompass more unique phrase patterns, specific punctuation sequences, and distinctive word collocations, providing a deeper stylistic fingerprint unique to individual authors (Sapkota, Bethard, Montes-y-Gomez, & Solorio, 2010)

- **Increased Author-Specificity:** As n-gram length increases, these patterns become more unique, distinguishing specific authors based on their habitual phrase structure, syntax, and word combinations (Stamatatos E. , 2009)
- **Enhanced Differentiation:** Higher n-grams are particularly useful for identifying unique linguistic habits, but they are computationally intensive, which can challenge scalability across large datasets (Stamatatos E. , 2009)

Trade-offs and Choice of n-gram Size

- **Lower n-grams:** These are computationally lighter and effective for general pattern recognition but may lack the depth needed to capture specific authorial nuances ((Peng, Schuurmans, Keselj, & Wang, 2003)
- **Higher n-grams:** Although computationally heavier, they capture detailed stylistic patterns unique to individual authors, which is beneficial for tasks requiring precise authorship attribution. (Sidorov, 2019)

2.3.2 Advantages of Character-Level N-Grams in Authorship Verification

1. **Content-Independent Analysis:** Character-level n-grams emphasize style over content, allowing AV models to remain effective across genres and topics, making them especially useful for cross-genre AV tasks
2. **Robustness in Short and Noisy Texts:** Character n-grams work well with shorter or more informal texts, such as tweets, as they capture essential stylistic markers without requiring large samples. This makes them suitable for AV in social media and other informal settings. (Stamatatos E, 2013)
3. **Resistance to Synonym Variation:** Since character n-grams capture form rather than meaning, AV models are effective even when different words or synonyms are used. This allows models to focus on structural patterns instead of vocabulary choice, which may vary by topic (Stamatatos E. , 2009)
4. **Language Independence:** Character n-grams can be applied across different languages, capturing consistent stylistic markers, such as punctuation and syntactic patterns, regardless of language (Stamatatos E., 2013)

2.3.3 Disadvantages of Character-Level N-Grams in Authorship Verification

1. **High Dimensionality:** Character n-gram models often create large feature spaces, especially with longer n-grams, which can increase computational requirements and complexity (Manolache, et al., 2021)
2. **Potential for Overfitting:** Shorter n-grams may capture superficial stylistic patterns that are not representative of an author's core style, increasing the risk of overfitting if the model emphasizes incidental patterns rather than consistent markers (Stamatatos E. , 2009)
3. **Sensitivity to Typographical Errors:** Character-level n-grams may inadvertently capture typographical errors or inconsistent spacing, introducing noise, particularly in informal texts (Sari et al., 2017) (Koppel et al., 2011).

4. **Limited Semantic Insight:** Character-level n-grams do not capture the semantic content of a text, which can limit their effectiveness in AV applications where an understanding of meaning or context is essential (Kestemont, 2014)
5. **Influence of Genre Compatibility:** Although they are more genre-agnostic, character n-grams' accuracy can decrease when applied across widely differing genres (Nini et al., 2024). For example, stylistic markers in casual chats might differ significantly from those in academic writing, which affects the verification model's performance.

2.3.4 Performance in Cross-Genre and Cross-Topic Texts

Character-level n-grams are highly effective in cross-genre and cross-topic AV tasks due to their focus on style rather than content. Since n-grams capture stylistic elements that remain stable regardless of genre or topic, AV models use them to generalize well across different thematic contexts (Koppel et al., 2011) (Neal et al., 2018). For example, whether an author is writing a formal academic paper or an informal blog post, character-level n-grams capture consistent syntax, punctuation, and morphology patterns, making them well-suited for cross-genre AV tasks where content shifts are standard (Stamatatos E. , 2009)

In cross-topic AV, where authors address varied subjects, character-level n-grams allow models to emphasize style over vocabulary, reducing the risk of false positives due to topic-specific words. For instance, an author's preference for certain sentence structures or specific contractions may persist across topics, allowing models to identify an author based on style alone (Kestemont, 2014); Sari et al., 2017).

2.4 Generative AI in Authorship

Generative AI has substantially impacted authorship analysis, particularly with the rise of Large Language Models like GPT and BERT. These models have shown they can create text that makes sense in context, often imitating human writing styles with impressive accuracy.

The effect of generative AI on confirming who the author is has an impact on the field. It brings several benefits, shows recent progress, and also presents challenges for the traditional ways of analyzing authorship. This impact will be significant for how this field develops in the future.

2.4.1 The Rise of LLMs in Authorship Analysis

Based on the transformer architecture, LLMs have enabled these models to outperform earlier neural network architectures in various natural language processing (NLP) tasks, including authorship verification and attribution. GPT-3 and its successors are part of the GPT series developed by OpenAI, while BERT, introduced by Google, represents a bidirectional approach to understanding the context of words within sentences. (Kumar, Musharaf, Musharaf, & Sagar, 2023)

The capabilities of LLMs like GPT-3 and BERT stem from their training on massive datasets, often encompassing hundreds of gigabytes of diverse text from sources like books, articles, and websites. For example, GPT-3 is trained on over 175 billion parameters, allowing it to generate coherent and contextually appropriate text based on the learned patterns (Bengesi, et al., 2023). This training is done through unsupervised learning, where the model forecasts the next word in a

sequence by looking at the words that come before it. This approach allows it to create responses that sound human and to adjust its replies according to different prompts, which makes it flexible for various types and styles. (Yenduri, et al., 2024)

BERT works a bit differently. It looks at all the words in a sentence simultaneously instead of just the words that come before. This way of training, which goes in both directions, helps BERT understand how words relate to each other much better. Because of this, it is good at tasks like classifying text, answering questions, and analyzing who wrote something. This understanding of context allows BERT to analyze stylistic elements in writing more accurately than previous models. (Bengesi, et al., 2023). The core innovation that has enabled the rise of LLMs like GPT and BERT is using self-attention mechanisms within the transformer architecture. Introduced by, the transformer architecture allows models to focus on different parts of a sentence when making predictions, thereby capturing long-range dependencies between words and phrases. (Bengesi, et al., 2023). This ability to attend to different parts of the input text means that LLMs can understand the contextual relationships within a text, making their generation more nuanced and sophisticated.

For example, when generating a passage continuation, a model like GPT-3 can keep track of entities, tone, and writing style across multiple sentences, ensuring a coherent narrative flow. This makes LLMs particularly effective at mimicking the stylistic elements of a given author when generating text. In authorship verification, this ability to capture style over longer stretches of text means that models can better detect subtle variations that distinguish one author's writing from another, a critical requirement for creating challenging impostor texts. (Yenduri, et al., 2024)

2.4.2 *Generating Convincing Texts with LLMs*

One significant advancement in the field of Generative AI is the ability of large language models (LLMs), such as GPT-3 and GPT-4, to create text based on user-provided prompts. This process begins with an initial input, ranging from a brief statement to a comprehensive directive, enabling the model to generate a follow-up that aligns with the given context and tone. The prompt serves as a guide that shapes what and how the response is created. For example, if the prompt asks for a formal scientific explanation, the answer will be organized and detailed. Conversely, a prompt that adopts a more casual tone can result in a conversational and relaxed style. This ability to adapt allows LLMs to move quickly between various fields, whether for creative writing or technical documents. This quality makes them very useful in situations where keeping the right style is essential (Yenduri, et al., 2024)

This versatility is possible due to their training on extensive, diverse datasets that include literary texts, scientific articles, and web content. This allows them to move between various language styles quickly. The ability of Generative AI to adjust to different styles has been shown commonly:

- In creative writing, models like GPT-3 can craft poetry, short stories, and dialogues that replicate specific genres or emulate the stylistic nuances of famous authors. These models can maintain narrative consistency across longer passages, reflecting the stylistic cues provided in the prompts (Kumar, Musharaf, Musharaf, & Sagar, 2023)
- In academic and technical writing, LLMs have been used to summarize complex research, provide lay explanations of technical topics, and even simulate peer review feedback. For example, given a prompt to simplify a scientific concept for a general audience, the models

produce content that is detailed yet accessible, leveraging their training on a wide array of scientific texts. (Bengesi, et al., 2023)

One of the most powerful features of (LLMs), is their ability to generate text that closely mimics the stylistic patterns of various writing styles.

Not only can they accept commands that can give them the tone, stylistic pattern, and structure of the response, allowing the model to capture long-range dependencies and contextual relationships between words, LLMs can replicate subtle stylistic patterns, such as an author's tendency to use certain phrases, idiomatic expressions, or specific syntactic structures from given texts. For instance, a model might learn that a particular given author frequently uses parenthetical statements or elliptical constructions and can replicate these patterns in its output. This ability to internalize and reproduce nuanced elements of style makes LLMs particularly effective for generating text that can convincingly emulate any given human writing. (Bengesi, et al., 2023)

As texts generated by large language models improve at mimicking an author's distinct style, it raises a big challenge in figuring out how we can use these abilities for verifying authorship, especially when it comes to creating fake texts. The challenge lies in using LLMs to generate impostor texts—synthetic samples similar to the known or unknown documents to effectively test a verification model's robustness. Here, the relationship between the impostor text and the original samples is critical; the impostors must be stylistically close enough to create a challenge but distinct enough to not merely replicate the content.

Additionally, the effectiveness of this approach depends on the ability of LLMs to create suitable impostors that properly simulate the target author's style. This includes looking into the right ways to create prompts—discovering the exact instructions that help the model produce the most realistic and challenging imposter texts.

2.5 Challenges and Gaps in Current Research

Significant challenges remain while substantial progress has been made in authorship verification through machine learning and neural networks. Traditional methods, such as stylometric analysis, which rely on features like word frequency, function words, and sentence structure (Stamatatos E. , 2009), often strive when applied to texts that vary significantly in genre or topic. These techniques might accidentally emphasize patterns related to the content rather than purely focusing on stylistic characteristics. This can lead to wrong classifications when you're comparing texts that discuss similar topics but are authored by different people. Such challenges have often been highlighted in assessments like the PAN competitions, where models that excel in tasks specific to a genre tend to struggle with texts from different domains (Ayele, et al., 2024).

A significant challenge in verifying authorship is to ensure that the evaluation is independent of the content while still being attentive to the unique style of the writer. This issue becomes even more evident with the use of transformer models such as BERT and GPT, which are very good at grasping context but might unintentionally include semantic content in their evaluations. For instance, PAN competitions have shown that when models analyze texts with overlapping domain-specific language, they may incorrectly attribute authorship due to shared vocabulary despite differences in underlying style (Stamatatos et al., 2020). Research continues into adversarial training

and domain adaptation methods to mitigate this issue, though achieving a consistent focus on style remains challenging, particularly across diverse types of text (Ayele, et al., 2024)

Cross-domain and cross-genre verification remain significant hurdles in the field, as authors may vary their writing style depending on the context, such as technical versus creative writing. Models trained on one genre often struggle when confronted with texts from different contexts, as highlighted by PAN competitions like those in 2020 and 2021, where models that excelled with in-domain data failed to generalize well to out-of-domain cases (Stamatatos E., 2016) (Bevendorff, et al., 2023). The challenge is for the models to adjust their style while not being too swayed by the themes present in the texts. (Ayele, et al., 2024)

Data scarcity is a recurring problem in authorship verification, particularly when only limited samples are available for a given author. This is in contrast to authorship attribution, which often benefits from access to larger corpora. Verification tasks, especially those in forensic linguistics or cybersecurity, frequently involve only a few documents, making it challenging to build robust models (Stamatatos, Kestemont, & Manjavacas, 2020). Additionally, the imbalance between positive and negative examples—where non-matching text pairs are far more common—complicates training. Although oversampling and few-shot learning techniques are being tested to address these issues, they often introduce their own set of biases and challenges. (Ayele, et al., 2024)

The advent of Generative AI models, has opened up new possibilities for creating synthetic texts that can challenge verification models. However, integrating these models effectively remains complex. Also, the computational cost of using LLMs for generating and approximating stylistically nuanced texts is another challenge, as it requires considerable processing power and can be resource-intensive.

The ability of LLMs to closely mimic human writing styles raises ethical concerns, particularly regarding their misuse to create misleading or forged content. This challenge has been highlighted in recent PAN competitions, where the focus has shifted to developing more resilient models capable of distinguishing between genuine texts and sophisticated AI-generated outputs. These competitions emphasize the need for verification systems that can adapt to the rapidly advancing capabilities of Generative AI in producing stylistically convincing content.

2.6 *The PAN competition*

The PAN competition (Plagiarism Analysis, Authorship Identification, and Near-Duplicate Detection) has played a pivotal role in advancing research in authorship analysis, including authorship verification, attribution, and profiling. Since its inception, PAN has become a critical platform for benchmarking state-of-the-art methods in digital text forensics, bringing together researchers from around the world to tackle various challenges associated with authorship analysis. The competition, organized annually as part of the CLEF (Conference and Labs of the Evaluation Forum), provides standardized datasets, evaluation metrics, and a competitive environment that drives innovation in the field. (Ayele, et al., 2024)

2.6.1 The Role of PAN in Authorship Verification Research

PAN competitions have played their role in setting standards for authorship verification by providing carefully curated datasets and well-defined tasks that mimic real-world scenarios. Each year, the competition organizers introduce new challenges that reflect emerging trends and complexities in the field, such as cross-genre authorship verification or multilingual authorship verification. This has allowed researchers to explore how their models perform not just in ideal conditions, but also in more practical, challenging situations where textual variations like style shifts and domain differences play a critical role.

One of the major contributions of PAN is the focus on content-independent verification—the ability to distinguish authors based on stylistic features rather than relying on the content of their texts. This focus has pushed participants to develop methods that can differentiate writing styles even when the topics or domains differ greatly. For example, tasks frequently include analyzing a scientific article alongside a blog post that was written by the same author. This underscores the importance of having models that can identify delicate stylistic differences, rather than just focusing on shared vocabulary. (Ayele, et al., 2024)

2.6.2 The Impact of PAN on Future Directions

The PAN competition has not only shaped the methodologies used in authorship verification but has also influenced the broader research agenda in digital text forensics. By regularly introducing new challenges like analyzing styles from multiple authors and verifying authorship across different languages, PAN has pushed for the creation of models that are flexible and able to manage various linguistic scenarios. Moreover, the competition has highlighted the significance of having explainable verification models, as researchers are placing greater emphasis on models that can explain their decision-making processes. This is especially important when these models are used in areas such as forensic linguistics.

As we look ahead, PAN remains a pioneer in authorship analysis research, promoting partnerships, and encouraging innovation. It functions as a practical environment for testing the most recent methods, providing an authentic view of the difficulties involved in authorship verification, especially as Generative AI and transformer-based models become increasingly common. The competition functions as an essential platform for creating solutions that enhance the technical abilities of authorship verification models. At the same time, it also tackles the ethical predicaments associated with the growing complexity of texts generated by artificial intelligence.

2.6.3 PAN Datasets and Their Evolution

Over the years, the nature and complexity of PAN datasets have evolved significantly to reflect the growing challenges and the need for more realistic and diverse strategies in authorship analysis.

Initially, PAN datasets were relatively straightforward, focusing on monolingual texts and single-genre tasks. For instance, early editions featured collections of literary texts or short essays, where the verification tasks rotated around distinguishing between a limited number of known writers. This setup helped establish baseline methodologies like n-gram analysis and stylometric qualities.

However, as the field matured, the competition introduced datasets that involved cross-genre and multilingual tasks to better simulate real-world challenges. By incorporating texts from different genres, such as blogs, reviews, scientific papers, and social media posts, the PAN organizers aimed to address the difficulty of maintaining content independence in authorship verification. This shift required models to adapt their features and focus on deeper stylistic features rather than just topic-specific vocabulary.

In recent years, the datasets have become even more complicated, especially with the emergence of Generative AI models. The PAN competitions in 2022 and 2023, for example, featured a variety of texts that combined real and machine-made samples. This highlights the increasing need to tell apart content created by humans from that generated by machines. This change was prompted by the rise of LLMs, which can produce highly convincing synthetic texts. The datasets have been modified to incorporate a variety of language characteristics and differences across various fields, making them tougher and more appropriate for testing how strong modern verification systems really are.

3

Methodology

3.1 Research Objective and Hypothesis

The objective of this thesis is to explore how Generative AI can enhance the Impostors Method in authorship verification (AV) by creating AI-based impostor texts that closely mimic the writing style of both known and unknown authors. This research aims to determine whether Large Language Model (LLM)- generated impostors, created through detailed prompt-based generation, can offer a more rugged and challenging external texts, alternative to traditional manually selected impostors. By focusing on stylistic mimicry—while ensuring that impostor texts differ in various ways in content - this study seeks to test whether such an approach improves the accuracy and adaptability of AV models, particularly in cross-domain and cross-genre applications.

To assess this method's effectiveness, we will employ datasets from PAN competitions (2015) which enclose a mixed range of genres. These datasets provide a general testing ground to assess whether LLM-generated impostors improve AV model performance in detecting stylistic consistency across various genres and domains, improving cross-genre generalization. Additionally, this research will explore how prompt engineering can be fine-tuned to generate impostors that more closely match the target author's stylistic features. This fine-tuning is expected to make LLM-generated impostors an even more effective challenge for AV models in practical applications, such as forensic linguistics and plagiarism detection.

Ultimately, this research aims to demonstrate if and how context-specific, AI-generated impostors can improve the robustness of authorship verification models. Success in this area could reveal new methods for using Generative AI to address complex verification tasks across diverse textual landscapes.

3.2 Hypothesis

For the Impostors Method to be effective in authorship verification, impostor texts must strike a careful balance between stylistic similarity to the target author and diversity to ensure robust model performance. This study researches if proposed impostor texts, with this kind of qualities, can directly improve the effectiveness of authorship verification models. The possible features of an impostor text are outlined as follows:

- 1. Stylistic Similarity:** Impostor texts should closely mimic the core linguistic markers of the target author, including syntax, punctuation, and sentence structure. These subtle stylistic features are often consistent across an author's works, challenging AV models to recognize deeper stylistic patterns rather than relying on superficial content-based traits like vocabulary or topic.

2. Controlled Stylistic Variation: Impostors should resemble the author's style while incorporating controlled variations to prevent the model from overfitting to repetitive structures. Variation in sentence length, punctuation use, and vocabulary within impostor texts promotes model robustness, ensuring that the model learns to verify authorship across a variety of stylistic patterns.

3. Cross-Genre and Cross-Domain Flexibility: Effective impostor texts should be adaptable across different genres and domains. Authors often modify their style when writing in varied contexts, such as moving from formal research papers to informal blog posts (genre change) or from technical subjects to creative narratives (domain change). Impostors that adapt to these shifts help the model generalize and maintain accuracy across such changes.

4. Challenging the Model to Focus on Stylistic Features: Effective impostor texts encourage the model to focus on consistent stylistic markers rather than content-based similarities like shared vocabulary or topics. This approach reduces false positives, where the model might otherwise mistake impostor texts for genuine ones based on overlapping content. By keeping content distinct while preserving stylistic similarity, impostors enhance the model's ability to detect subtle authorial traits, thereby improving both accuracy and generalization across tasks.

3.2.1 How LLM-Based Impostors can make a more Effective Impostor Model

The use of Large Language Models (LLMs) for generating impostor texts offers unique advantages in addressing the challenges summarized above. By leveraging LLMs' vast language generation capabilities, this study aims to generate impostor texts that better balance stylistic similarity with diversity, supporting a more robust authorship verification process.

1. Achieving Stylistic Similarity: LLMs can generate impostor texts that closely resemble the target author's key linguistic characteristics. With carefully crafted prompts, LLMs can produce texts that replicate an author's syntax, punctuation, and sentence structure, capturing the author's unique voice while varying content. For example, LLMs can be instructed to generate texts with specific syntactic structures or characteristic punctuation patterns, while making sure that impostor texts maintain stylistic resemblance while introducing novel content. This approach helps AV models focus on nuanced stylistic features, rather than superficial content traits like vocabulary or topical similarity.

2. Implementing Controlled Stylistic Variation: LLMs offer the flexibility to introduce controlled variations within generated texts, an essential feature for effective impostors. By adjusting generation parameters and prompt details, LLMs can vary linguistic elements such as sentence length, punctuation, and vocabulary diversity, maintaining stylistic resemblance while stopping overfitting. This variation helps AV models remain engaged with a range of stylistic features, promoting better generalization across different text types.

3. Ensuring Cross-Genre and Cross-Domain Adaptability: LLMs can generate texts suited to different genres and domains, essential for cross-genre and cross-domain AV tasks. For instance, LLMs can mimic the style of an author's formal academic writing while adapting to more conversational styles as needed, reflecting natural shifts in an author's style based on genre. Additionally, LLMs can maintain a given author's stylistic markers across content domains (e.g., technical to creative). This adaptability enables AV models to detect the author's stylistic signature even amid varied content and genre shifts.

4. Enhancing Model Focus on Stylistic Features: The power of LLM-generated impostors lies in their ability to produce content that diverges from the target author's known topics while closely mimicking the style. This separation of content and style prevents the AV model from relying on content-based similarities (such as vocabulary or subject matter) and pushes it to identify more challenging stylistic features, like sentence structure, phrasing, and punctuation. Through prompt engineering, LLM-generated impostors can be fine-tuned to reflect stylistic traits without overlapping in content. This process enables the verification model to focus on subtle authorial markers, decreasing false positives and enabling accurate differentiation between genuine and impostor texts. Additionally, by maintaining stylistic consistency while diversifying content, these impostors offer a more rigorous test for the verification model, fostering generalization across various authorship tasks.

3.3 *The Datasets*

The research is focused on testing the Model using PAN's datasets, from competitions of 2015. These datasets are widely regarded as benchmarks for authorship verification tasks and provide a rich and diverse set of texts across multiple languages and genres. They are some of the most extensive and varied corpora used for authorship analysis, including authorship verification (AV), attribution, and profiling. Each year, PAN releases a new dataset designed to reflect real-world complexities in text analysis, such as cross-genre variation, domain shifts, and multilingualism. This research utilizes the datasets from PAN 2015 that include a variety of text types that challenge authorship verification models.

- PAN 2015: The focus was primarily on cross-genre authorship verification, testing models' ability to handle significant style shifts between genres (e.g., between formal essays and informal social media posts). The dataset highlighted the challenge of distinguishing authorship when the writing style changes due to the context, while still maintaining some stylistic consistency.

3.3.1 *Text Genres*

The PAN datasets encompass a variety of text genres, each of which presents unique challenges for authorship verification. The task of authorship verification is not simply about identifying the author of a text; it involves recognizing stylistic patterns and linguistic features that remain consistent across different writing contexts. However, authors often shift their style depending on the genre or domain, which complicates the verification process. Here are some key genres included in the PAN datasets and the challenges they present:

- **Blogs:** Blog posts typically feature informal, conversational writing, with varying sentence structures, tone, and grammar. The challenge with blog texts is that they often blend personal narratives with informational content, making it difficult for models to rely on traditional stylistic markers. Blogs can also vary in length and complexity, adding to the challenge of verifying authorship across different posts.
- **Reviews:** User reviews are generally short, opinion-based texts, often written in an informal style. The shortness of reviews makes it harder for authorship verification models to identify consistent stylistic patterns. Additionally, reviews tend to focus on specific outcomes or

services, meaning that the content may vary widely across different domains, but the author's underlying style must still be caught.

- **Scientific Papers:** Formal academic texts, such as scientific papers, provide a more structured and disciplined form of writing. Complex sentence structures, technical vocabulary, and rigid grammar rules characterize these texts. Authorship verification in this genre requires models to recognize the formal writing style of the author, even when dealing with highly specialized language and a focus on factual content.
- **Social Media Posts:** Social media posts, such as tweets or Facebook updates, are usually brief, fragmented, and often include abbreviations, emojis, and slang. This genre poses a significant challenge for authorship verification because of the highly informal and inconsistent nature of the writing. Social media posts can also shift tone and topic rapidly, making it difficult for models to identify a consistent authorial voice.

3.4 Modeling Approaches

The Impostors Model in this study is a customized authorship verification approach that leverages character-based n-gram TF-IDF vectorization and a MinMax distance metric to assess stylistic similarity between a target author's known texts and a set of generated impostor texts. The model utilizes character-level n-grams to capture detailed stylistic markers effectively, such as punctuation patterns, syntax structure, word morphology, and function word usage, which are distinct for each author and contribute to their unique writing style. The model creates a stylistic "centroid" by averaging the TF-IDF vectors of the target author's known texts, providing a reference point for evaluating new texts. Impostors are iteratively selected and refined through random feature sampling and MinMax distance calculations, which allow each impostor text to introduce controlled stylistic variations while remaining similar to the target author's writing style. This method pushes the AV model to focus on deep stylistic patterns, rather than content-based traits, thereby enhancing its ability to generalize across genres and domains.

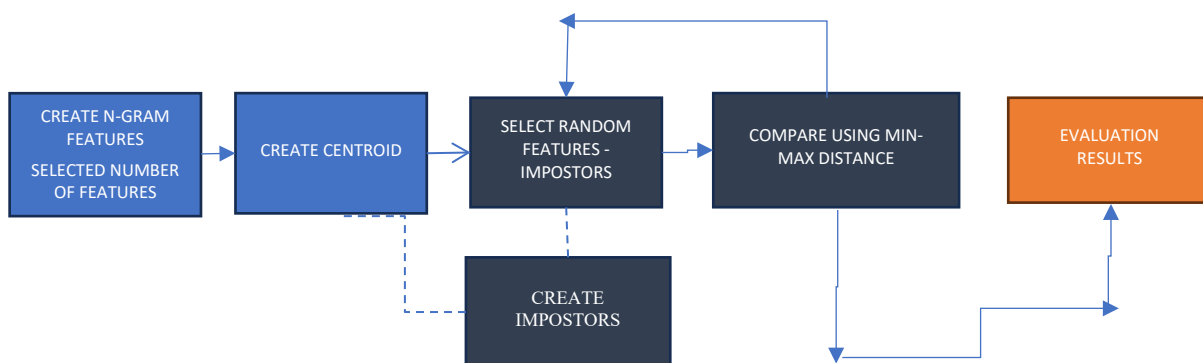


Figure 2: The Authorship Verification process using impostors, illustrated in a diagrammatic representation, with an additional step for generating impostors using prompt engineering when none are available.

3.4.1 *The Impostors Model with Character-Level N-Grams and MinMax Distance*

The Impostors Model used in this study is a customized approach for AV that employs character-based n-gram TF-IDF vectorization and a MinMax distance metric to assess stylistic similarity between a target author's known texts and a set of impostor texts. The model is designed to focus on stylistic markers, which are unique to each author and persist across their various texts, rather than content-based cues.

1. Character-Level N-Grams: The model begins by transforming texts into TF-IDF vectors using character-level n-grams, with each n-gram representing a sequence of characters of a specified length (e.g., 3-4 characters). Character-based n-grams capture a range of stylistic features, including punctuation patterns, syntax structures, morphological patterns, and function word usage. These elements offer a more consistent reflection of an author's unique writing style compared to content-specific words, making character-level n-grams particularly suitable for AV tasks where content varies but style remains consistent.

2. MinMax Distance for Stylistic Similarity: To assess similarity, the model uses a MinMax distance metric, which calculates the distance between the TF-IDF vectors of the target text and each impostor, as well as between the target text and a centroid vector created from the target author's known texts. This centroid represents the averaged stylistic traits across the author's texts. The MinMax distance prioritizes stylistic over semantic similarity, allowing the model to focus on the underlying writing style and bypass false positives due to content overlap.

3. Iterative Impostor Selection: The Impostors Model iteratively selects and refines a set of impostor texts. The process starts by selecting a wide range of potential impostors, which are filtered down based on their MinMax distance to the target text. Impostors that closely match the stylistic centroid are retained, and only those that introduce controlled variations (e.g., minor shifts in sentence structure or punctuation use) are finalized. This selection guarantees that each impostor challenges the AV model to detect deeper stylistic patterns, rather than repetitive or superficial features.

4. Probabilistic Scoring and Classification: Finally, the model computes a probability score for each text, indicating whether it matches the target author's style. This score is derived from the distance metrics, with a threshold-based classification deciding whether a text is attributed to the target author. This probabilistic approach allows flexibility in changing sensitivity based on the specific AV task.

3.4.2 *Suitability of the Impostors Model for Authorship Verification*

The Impostors Model is particularly suitable for AV due to its ability to prioritize stylistic markers over content. By concentrating on character-based n-grams, the model can capture fine-grained stylistic details that are more difficult to replicate and are less influenced by topic changes, thus reducing the risk of false positives. Key advantages include:

- **Content-Independent Verification:** By leveraging stylistic features, the model minimizes reliance on content overlapping, enabling it to distinguish between genuine and impostor texts even when topics differ significantly.
- **Adaptability Across Genres and Domains:** The model's flexible impostor sampling process supports cross-genre verification, where authors may adjust their writing style based on context (e.g., academic articles vs. informal blog posts).
- **Robustness to Short and Noisy Texts:** Character-level n-grams allow the model to maintain effectiveness across texts of varying lengths and quality, making it robust in handling social media posts or short reviews.

3.4.3 *Potential Research Directions*

The Impostors Method can be further enhanced by exploring several potential research directions to expand the scope and utility of impostor generation in authorship verification:

1. **Generating Impostors for Unknown Texts:** Currently, impostor texts are generated primarily as stylistic comparisons for the known texts attributed to the target author. A compelling area for research is the creation of impostor texts specifically for the unknown text in each AV problem. By generating impostors that are stylistically similar to the unknown text but differ in content, the model could perform more nuanced comparisons, especially if the style of the unknown text varies from the target author's typical writing. This process could deepen the AV model's effectiveness by challenging it to identify stylistic consistency across both known and unknown text dimensions, allowing for a more layered and detailed verification process.
2. **Expanding Impostor Pools by Combining Impostors Across Problems:** Another possible improvement is to create a larger collection of impostor texts by combining impostors generated for multiple problems within the same language. This approach would provide a broader selection of impostor texts, giving the AV model access to a diverse set of stylistic variations when analyzing each new text. A larger created text pool would improve the model's robustness, as it could draw from a richer array of stylistic features, which in turn could improve its generalization capabilities and adaptability across various genres and contexts within the same language.

These research directions highlight opportunities to further refine the Impostors Method by expanding its focus on stylistic complexity and diversity. By producing impostors for unknown texts and pooling impostors across problems, the research will try to improve the robustness and versatility of authorship verification models across varied textual and linguistic landscapes.

3.5 *Generating AI Impostors*

The generation of AI-based impostor texts in this study used GPT-4-mini through API call functions, focusing on prompt design and stylistic adjustments to create texts that mimic an author's writing style while varying in content. This approach is based on the hypothesis that impostors which closely reflect the features of known texts in various ways, would challenge in the best possible way, the authorship verification (AV) model to identify an author's unique stylistic markers.

In generating these impostors, we used known texts attributed to each target author, with random selection of a text when numerous known samples were available. This permitted for stylistic consistency without relying on a single text for all impostor generations, thereby adding diversity to the impostor set. Each prompt aimed to capture stylistic characteristics specific to the selected text while presenting differences in content or theme.

3.5.1 *Prompt Design for Impostor Generation*

Several approaches were explored in designing prompts to ensure stylistic mimicry with topic independence. The following strategies highlight different ways we guided GPT-4-mini to maintain a balance between stylistic consistency and content variation:

1. **Style Imitation with Distinct Content:**

- For each impostor generation cycle, a known text was randomly selected and used as the stylistic foundation for the new impostor text. The prompt guided the model to match the writing style, tone, vocabulary, and length of the original text but required a different topic or theme. This approach was intended to produce impostors that felt convincing as the author's writing while ranging in content, ensuring the AV model focused on stylistic elements rather than specific content overlaps.

2. **Thematic Consistency with Stylistic Rephrasing:**

- In cases where we wanted the impostor text to retain some thematic elements from the selected known text without direct content replication, the prompt instructed GPT-4-mini to write on a similar topic. The model was guided to rephrase ideas and vary sentence structures while maintaining the tone and vocabulary of the original, creating an impostor that stylistically aligned with the original but maintained content independence. This approach aimed to emphasize style without direct topical continuity, adding an additional layer of complexity for the AV model.

3. **General Stylistic Similarity without Content Overlap:**

- Another prompt variation instructed the model to create a text that closely matched the general tone, sentence structure, and level of formality of the original text, while varying in vocabulary and function word usage. The model was prompted to use a similar level of sentence complexity without directly copying expressions unique to the original author. This prompt allowed the creation of impostors that closely matched the stylistic "fingerprint" of the target author without thematic or topical resemblance, pushing the AV model to focus on authorial style alone.

3.5.2 *API Implementation and Workflow*

Using the OpenAI API, GPT-4-mini was accessed with specific API call functions configured to iterate through each of these prompt designs. The API workflow involved systematic random selection of known texts and multiple cycles of prompt-based impostor generation to create a diverse impostor set across style and content:

- **Random Text Selection:** For each impostor creation, a known text was randomly chosen, assuring diversity within the set of generated impostors.

- **Iteration for Stylistic Diversity:** Each prompt type was repeated across multiple iterations, improving the stylistic variation within the generated impostor texts.
- **Role Specification:** A system role of “Expert in stylometry and impostor text generation” was given to guide GPT-4-mini’s focus on stylistic, rather than content-based, similarities.

3.5.3 Challenges and Observations

Throughout the impostor generation process, achieving the correct balance between stylistic mimicry and content variation required careful prompt adjustments. Observations included:

1. **Ensuring Topic Divergence:** Creating impostors that differed sufficiently from the original topic sometimes required rephrasing the prompt to prevent topic overlap, particularly when the known text was short or highly specific.
2. **Length Consistency:** While prompts were designed to request impostors of similar length to the known text, minor variations occasionally occurred. This was managed by refining prompt instructions and re-running prompts when necessary.
3. **Maintaining Stylistic Consistency:** Achieving a balance where the impostor text imitated style without replicating exact phrases or sentence structures required nuanced prompt wording. By focusing on higher-level stylistic elements like tone, vocabulary, and sentence complexity, the prompts effectively generated texts that captured the original style without content replication.

By implementing this structured approach to prompt design and API configuration, we generated a varied set of AI-based impostors that encouraged the AV model to focus on identifying authorial style across different topics. This diversity in stylistic representation provided a robust test for the AV model, challenging it to detect an author’s unique stylistic markers independently of content similarity.

3.6 Experimental Setup

In this section, we outline the setup used to evaluate the effectiveness of our authorship verification (AV) model with AI-generated impostors. Our model is a **lazy (non-parametric) model**, meaning it performs verification without requiring traditional training. Unlike most trained models, our AV model does not create a generalized representation through training data; instead, it directly compares new texts against known impostors and genuine samples each time. This lazy approach makes it flexible and adaptable to varied datasets and styles, and it avoids the risks of overfitting and threshold sensitivity associated with traditional training.

3.6.1 Evaluation Metrics and the Importance of AUC

The primary evaluation metric in this study is the **Area Under the Curve (AUC)** score, which is valuable in several ways for authorship verification tasks. AUC is **threshold-free**, indicating it does not rely on a precise decision threshold. This is important in AV because texts and stylistic variations are highly diverse across authors, genres, and even individual samples within the same dataset. By evaluating the model across all possible classification thresholds, AUC provides a more comprehensive view of model performance than threshold-dependent metrics such as accuracy or F1 score.

3.6.2 *Experimental Parameters and Variations*

To rigorously assess the AV model's robustness, we performed experiments on impostors generated from different prompt designs, each prompting the AI to mimic style while ranging in content. For each prompt, we tested several parameter configurations that influence the model's sensitivity to stylistic features. Key parameters include:

- **n-gram Size:** Different character-level n-gram sizes were tested, with values such as 3, 4, and 5. By adjusting the n-gram size, we evaluated the model's ability to capture various stylistic features from broader structural patterns with 3-grams to finer-grained details with longer n-grams, which can emphasize unique authorial choices in syntax, punctuation, and function word usage.
- **Iterations:** The number of iterations used in the impostor selection process was modified to examine its impact on performance. The iteration count determines how many times the model refines its impostor set, with higher iterations presenting more diversity among the impostors. Iterating enables the model to engage with a wider stylistic range within the generated impostors, which helps prevent overfitting to specific stylistic markers
- **Feature Sampling (Dropout):** Feature sampling was adjusted by varying the dropout rate, which controls the percentage of features retained during each iteration. More increased dropout rates increase feature sparsity, potentially focusing the model on essential stylistic features rather than overemphasizing distinctive features that might not generalize across samples. This parameter was tested to observe how different levels of feature retention affected AUC scores.

Each parameter combination was tested empirically, and results were documented to evaluate which configurations provided optimal AUC scores across different prompts and language settings. This approach allowed us to analyze the sensitivity of each parameter and to identify robust configurations for AV tasks with AI-generated impostors.

3.6.3 *Cross-Language and Multi-Model AUC Comparison*

To contextualize the performance of our AV model, we compared its AUC scores with other models evaluated on the PAN 2015 dataset, including Support Vector Machines (SVM), neural networks, and transformer-based models, as reported in previous research. The AUC scores are segmented by language to illustrate the model's cross-linguistic robustness. PAN 2015 provides data in multiple languages, including English, Dutch, Spanish, and Greek, each presenting unique linguistic challenges for AV.

4

Experiments & Results

4.1 Overview of the Experiments

The experiments in this study aim to evaluate the effectiveness of the Impostors Method for authorship verification, leveraging the PAN 15 dataset as the primary benchmark tool. We tweaked the Impostors Method, so that involves generating impostor texts to simulate alternative authors, for every given problem. This subchapter outlines the key objectives, datasets, experimental design, and parameters considered in the evaluation process.

4.1.1 Objective

The main objective of these experiments, is to evaluate the Impostors Method's effectiveness in distinguishing between texts written by the same person and those by different authors, with a focus on utilizing large language models (LLMs) for generating impostors. This study specifically investigates whether impostors, crafted to resemble the stylistic nuances of the original known texts closely, can improve the performance of the authorship verification task.

Through targeted prompt engineering, the experiments aim to extract and replicate various stylistic features from the original texts, creating impostors tailored to each specific verification problem. The hypothesis is that high-quality impostors, designed to question the verification model by simulating realistic stylistic variations, will enhance the method's overall effectiveness.

Additionally, the experiments explore the impact of tweaking prompts on the quality of generated impostors. This involves analyzing how specific modifications to the prompt influence the stylistic precision and variety of the impostors. Furthermore, the role of hyperparameter tuning—such as adjustments to impostor count, learning rates, and similarity thresholds—is examined to determine their contributions to model performance.

By addressing these objectives, the study aims to offer a deeper understanding of the interactions between prompt design, impostor quality, and hyperparameter optimization in the context of LLM-powered authorship verification.

Experimental Setup

The experiments in this study are designed to explore the effectiveness of the Impostors Method under varying conditions. Due to limited resources and the necessity to evaluate the feasibility of various configurations, some languages and prompts were prioritized. For example, the DU subset was chosen for initial experiments due to its diverse stylistic and topical characteristics, which offer a challenging test case for the method. Likewise, fewer impostors were generated in certain cases to counteract computational costs with the need for exploratory insights. These decisions reflect an

iterative research approach, and future experiments will aim to include all languages and prompts comprehensively.

We initiated the experiments by first using the training datasets to find the appropriate hyperparameters, different for every language, using impostors created by LLM, through vigorous testing. This was also done to better understand the behavior of the Impostors Method under different configurations and to determine the impact of parametrization. After identifying the optimal configurations, we applied them to the evaluation dataset, modifying only the impostors in each case. This iterative process allowed for a more efficient investigation of the method's potential while ensuring robust evaluation.

Finally, the number of impostors determined by the parameters is used in each iteration, where proportional to those of the initial pool of impostors in the folder.

Dataset

The PAN 15 dataset serves as the foundation for the experiment, utilizing the training dataset along with the evaluation dataset to test the method. The PAN 15 dataset is well-suited for this study due to its rich and diverse set of texts, allowing for robust assessment across different genres and topics.

Different Prompts for Generating Impostors

The experiments explore the effect of using different prompts to generate impostors. Prompts guide the creation of synthetic texts that emulate potential alternative authors. The diversity of prompts provides insights into their influence on performance metrics such as the Area Under the Curve (AUC).

1. Comparing Paraphrased Impostors vs. Refined Prompts

Two distinct strategies for generating impostors were compared:

- **Paraphrased Impostors:** Derived from the known text through paraphrasing technique to retain core content but introduce stylistic variations.
- 2. **Refined Prompts:** Designed to generate impostors with distinct stylistic characteristics that better mimic alternative authors while maintaining contextual relevance.
- 3. The comparison helps evaluate which approach produces impostors that challenge the verification algorithm more effectively.
- 4. **Testing the Effect of Hyperparameter Tuning and Increased Impostor Counts**
Hyperparameter settings and the number of generated impostors were adjusted to optimize performance. This included:
 1. Tuning number of impostors, number of iterations and vocabulary size.
 2. Gradually increasing the number of impostors to assess the trade-off between computational cost and verification accuracy.

4.1.2 Tuning the model

Using the PAN 15 training dataset, we identified the potential parameters for each language. The training dataset consists of 100 problems for each language, while the evaluation dataset does not have an equal number of problems for each language. Given a pool of 100 impostors - created

through prompt engineering - for every language, we determined the best possible combination in hyperparameter tuning.

Impostor Generation

After making a combination of the known texts, wherever the known texts were more than one, while combining them but also keeping the number of letters below 5000, impostors were generated using the following prompt:

prompt = *Create an impostor text that shares a similar tone, general sentence structure, and level of formality with the original text but varies in specific vocabulary choices, function word usage, and unique stylistic markers. The impostor should feel similar overall style but not directly mimic the original's unique characteristics."*

Ensure the following:

- Use a similar level of complexity in sentence structure and tone without directly copying unique expressions or phrasing.

- Avoid using distinctive patterns in function word usage or punctuation that are highly characteristic of the original author.

- Maintain a general similarity in topic to align with the original text without reproducing exact thematic details.

Original text:

"{original_text}"

1. The prompt was executed using the GPT-4 Mini model.
2. Role assigned: " *Expert in stylometry and impostor text generation.*"
3. 100 Impostors were created for each problem

English Language								
IMPOSTORS	NGRAM	VOCAB	ITER	DROPOUT	TOTAL IN FOLDER	IMPOSTORS	AUC	TIME
Prompt 4	5	5000	400	0.3	100	100/70/45	0.508	<0.2
Prompt 4	5	1000	400	0.3	100	100/70/45	0.537	<0.2
Prompt 4	5	700	400	0.3	100	100/70/45	0.478	<0.2
Prompt 4	5	900	400	0.3	100	100/70/45	0.468	0.23
Prompt 4	5	1000	400	0.5	100	100/70/45	0.479	0.26
Prompt 4	5	1000	400	0.2	100	100/70/45	0.490	0.18
Prompt 4	5	1000	1000	0.2	100	100/70/45	0.513	0.33
Prompt 4	5	1000	600	0.2	100	100/70/45	0.466	0.23

Prompt 4	4	1000	400	0.2	100	100/70/45	0.548	0.18
Prompt 4	4	1000	400	0.3	100	100/70/45	0.504	0.18
Prompt 4	4	1000	400	0.15	100	100/70/45	0.564	0.17
Prompt 4	4	1000	400	0.12	100	100/70/45	0.599	0.16
Prompt 4	4	1000	400	0.1	100	100/70/45	0.599	0.16
Prompt 4	4	1000	400	0.08	100	100/70/45	0.594	0.16
Prompt 4	4	800	400	0.1	100	100/70/45	0.566	0.15
Prompt 4	4	2000	400	0.1	100	100/70/45	0.599	0.16
Prompt 4	3	1000	400	0.1	100	100/70/45	0.567	0.14
Prompt 4	3	1000	400	0.3	100	100/70/45	0.561	0.21
Prompt 4	3	1000	400	0.2	100	100/70/45	0.79	0.18
Prompt 4	3	1000	400	0.15	100	100/70/45	0.584	0.16
Prompt 4	3	1000	400	0.15	100	100/50/35	0.580	0.15

Table 1 Hyperparameter Tuning for English Language Dataset: Optimal Values Highlighted

Dutch Language								
IMPOSTORS	NGRAM	VOCAB	ITER	DROPOUT	TOTAL IN FOLDER	IMPOSTORS	AUC	TIME
Prompt 5	5	5000	400	0.5	100	100/40/15	0.603	0:43
Prompt 5	5	5000	400	0.34	100	100/40/15	0.653	0:33
Prompt 5	5	10000	400	0.34	100	100/40/15	0.638	0:34
Prompt 5	5	4000	400	0.34	100	100/40/15	0.641	0:31
Prompt 5	5	6000	400	0.34	100	100/40/15	0.627	0:31
Prompt 5	5	5000	400	0.4	100	100/40/15	0.640	0:34
Prompt 5	5	5000	400	0.7	100	100/40/15	0.573	0:49
Prompt 5	5	5000	400	0.3	100	100/40/15	0.622	0:29
Prompt 5	5	5000	400	0.30	100	100/40/15	0.622	0:29
Prompt 5	5	5000	400	0.30	100	100/70/45	0.673	0:29
Prompt 5	5	5000	400	0.34	100	100/90/65	0.598	0:36
Prompt 5	5	5000	400	0.34	100	100/90/40	0.650	0:35
Prompt 5	5	5000	400	0.34	100	100/60/40	0.641	0:32
Prompt 5	5	5000	400	0.34	100	100/70/30	0.657	0:34

Prompt 5	5	5000	400	0.34	100	100/70/50	0.624	0:33
Prompt 5	5	5000	400	0.34	100	100/70/40	0.636	0:33
Prompt 5	5	5000	400	0.34	100	80/40/20	0.605	0:29
Prompt 5	5	5000	400	0.34	100	50/25/10	0.643	0:21
Prompt 5	5	5000	400	0.34	100	50/25/14	0.637	0:31
Prompt 5	5	5000	400	0.34	100	50/25/7	0.639	0:21
Prompt 5	5	5000	400	0.34	100	50/30/10	0.640	0:21
Prompt 5	5	5000	400	0.34	100	50/35/15	0.617	0:22
Prompt 5	5	5000	600	0.3	100	100/70/45	0.595	0:46
Prompt 5	5	5000	200	0.3	100	100/70/45	0.599	0:21
Prompt 5	5	5000	300	0.3	100	100/70/45	0.651	0:36
Prompt 5	5	5000	380	0.3	100	100/70/45	0.669	0:38
Prompt 5	4	5000	400	0.30	100	100/70/45	0.572	0:42
Prompt 5	4	5000	400	0.6	100	100/70/45	0.565	0:46
Prompt 5	4	5000	400	0.2	100	100/70/45	0.622	0:46
Prompt 5	4	5000	400	0.1	100	100/70/45	0.666	0:19
Prompt 5	4	5000	400	0.08	100	100/70/45	0.658	0:19
Prompt 5	4	5000	400	0.15	100	100/70/45	0.638	0:23
Prompt 5	4	8000	400	0.15	100	100/70/45	0.626	0:19
Prompt 5	4	3000	400	0.15	100	100/70/45	0.615	0:19
Prompt 5	4	3000	400	0.15	100	100/50/20	0.643	0:18
Prompt 5	3	5000	400	0.1	100	100/70/45	0.553	0:15
Prompt 5	3	5000	400	0.3	100	100/70/45	0.512	0:23
Prompt 5	3	5000	400	0.08	100	100/70/45	0.571	0:15
Prompt 5	3	5000	400	0.08	100	80/40/30	0.555	0:14

Table 21 Hyperparameter Tuning for Dutch Language Dataset: Optimal Values Highlighted

Greek Language								
IMPOSTORS	NGRAM	VOCAB	ITER	DROPOUT	TOTAL IN FOLDER	IMPOSTORS	AUC	TIME
Prompt 4	5	5000	400	0.5	100	100/40/15	0.634	<1

Prompt 4	5	4000	400	0.3	100	100/70/45	0.645	1.4
Prompt 4	5	4000	400	0.1	100	100/70/45	0.647	0.36
Prompt 4	5	3000	400	0.1	100	100/70/45	0.647	0.34
Prompt 4	5	400	400	0.1	100	100/70/45	0.687	0.38
Prompt 4	5	1000	400	0.1	100	100/70/45	0.722	0.43
Prompt 4	5	1500	400	0.1	100	100/70/45	0.695	0.28
Prompt 4	5	1100	400	0.1	100	100/70/45	0.713	0.27
Prompt 4	5	1000	400	0.08	100	100/70/45	0.728	0.22
Prompt 4	5	1000	400	0.06	100	100/70/45	0.746	0.22
Prompt 4	5	1000	400	0.2	100	100/70/45	0.703	0.22
Prompt 4	5	1000	600	0.06	100	100/70/45	0.743	0.24
Prompt 4	5	1000	300	0.06	100	100/70/45	0.698	0.24
Prompt 4	5	1000	300	0.06	100	100/70/45	0.698	0.24
Prompt 4	4	1000	400	0.06	100	100/70/45	0.720	0.24
Prompt 4	4	1000	400	0.12	100	100/70/45	0.723	0.23
Prompt 4	4	1000	400	0.1	100	100/70/45	0.724	0.23
Prompt 4	4	600	400	0.1	100	100/70/45	0.730	0.23
Prompt 4	4	800	400	0.1	100	100/70/45	0.727	0.21
Prompt 4	4	600	600	0.1	100	100/70/45	0.739	0.25
Prompt 4	4	200	800	0.1	100	100/70/45	0.721	0.25
Prompt 4	4	1000	400	0.1	100	100/70/45	0.724	0.22
Prompt 4	4	4000	400	0.5	100	100/70/45	0.676	1.22
Prompt 4	4	4000	400	0.3	100	100/70/45	0.686	1.10
Prompt 4	4	4000	400	0.06	100	100/70/45	0.720	0.29
Prompt 4	4	2000	400	0.1	100	100/70/45	0.708	0.27
Prompt 4	3	600	600	0.1	100	100/70/45	0.724	0.23
Prompt 4	3	1200	600	0.1	100	100/70/45	0.712	0.28
Prompt 4	3	600	600	0.06	100	100/70/45	0.726	0.22

Table 31 Hyperparameter Tuning for Greek Language Dataset: Optimal Values Highlighted

Spanish Language								
IMPOSTORS	NGRAM	VOCAB	ITER	DROPOUT	TOTAL IN FOLDER	IMPOSTORS	AUC	TIME
Prompt 4	5	5000	400	0.5	100	100/40/45	0.603	1.10
Prompt 4	5	5000	400	0.3	100	100/70/45	0.615	1.15
Prompt 4	5	5000	400	0.8	100	100/70/45	0.562	2.27
Prompt 4	5	5000	400	0.2	100	100/70/45	0.626	0.59
Prompt 4	5	5000	400	0.1	100	100/70/45	0.639	0.46
Prompt 4	5	5000	400	0.4	100	100/70/45	0.607	1.30
Prompt 4	5	5000	400	0.13	100	100/70/45	0.636	1.9
Prompt 4	5	5000	400	0.08	100	100/70/45	0.644	0.42
Prompt 4	5	3000	400	0.08	100	100/70/45	0.647	0.40
Prompt 4	5	2000	400	0.08	100	100/70/45	0.666	0.32
Prompt 4	5	1000	400	0.08	100	100/70/45	0.709	0.29
Prompt 4	5	800	400	0.08	100	100/70/45	0.709	0.21
Prompt 4	5	500	400	0.08	100	100/70/45	0.731	0.25
Prompt 4	5	200	400	0.08	100	100/70/45	0.743	0.25
Prompt 4	5	200	400	0.2	100	100/70/45	0.722	0.24
Prompt 4	5	200	400	0.12	100	100/70/45	0.731	0.24
Prompt 4	5	200	600	0.08	100	100/70/45	0.756	0.30
Prompt 4	5	200	800	0.08	100	100/70/45	0.757	0.30
Prompt 4	5	200	1200	0.08	100	100/70/45	0.753	0.34
Prompt 4	5	200	1000	0.08	100	100/70/45	0.754	0.31
Prompt 4	5	200	200	0.08	100	100/70/45	0.753	0.22
Prompt 4	5	200	300	0.08	100	100/70/45	0.749	0.23
Prompt 4	5	200	800	0.08	100	100/60/35	0.761	0.28
Prompt 4	5	200	800	0.08	100	100/60/25	0.756	0.29
Prompt 4	5	200	800	0.08	100	100/60/40	0.750	0.29
Prompt 4	5	200	800	0.08	100	100/85/35	0.766	0.29
Prompt 4	5	200	800	0.08	100	100/85/30	0.736	0.30
Prompt 4	5	200	800	0.08	100	100/85/45	0.758	0.30

Prompt 4	5	200	800	0.08	100	100/95/35	0.760	0.30
Prompt 4	5	200	800	0.08	100	100/90/35	0.762	0.29
Prompt 4	4	200	800	0.08	100	100/90/35	0.742	0.29
Prompt 4	4	200	800	0.2	100	100/90/35	0.733	0.31
Prompt 4	4	200	800	0.08	100	100/85/35	0.741	0.29
Prompt 4	4	200	600	0.08	100	100/85/35	0.753	0.26
Prompt 4	4	200	400	0.08	100	100/85/35	0.756	0.24
Prompt 4	4	6000	400	0.08	100	100/85/35	0.719	0.46
Prompt 4	4	2000	400	0.08	100	100/85/35	0.727	0.31
Prompt 4	4	400	400	0.08	100	100/85/35	0.770	0.25
Prompt 4	4	800	400	0.08	100	100/85/35	0.723	0.26
Prompt 4	4	600	400	0.08	100	100/85/35	0.746	0.25
Prompt 4	4	400	400	0.06	100	100/85/35	0.772	0.25
Prompt 4	3	400	400	0.06	100	100/85/35	0.746	0.23
Prompt 4	3	400	400	0.1	100	100/85/35	0.815	0.23
Prompt 4	3	400	400	0.15	100	100/85/35	0.812	0.24
Prompt 4	3	600	400	0.1	100	100/85/35	0.800	0.25
Prompt 4	3	300	400	0.1	100	100/85/35	0.797	0.22
Prompt 4	3	400	600	0.1	100	100/85/35	0.811	0.26
Prompt 4	3	400	400	0.1	100	100/70/40	0.808	0.23

Table 41 Hyperparameter Tuning for Spanish Language Dataset: Optimal Values Highlighted

4.2 Experiment 1: Results

The experiments were conducted using the DU, GR, and SP datasets of the PAN 15 evaluation dataset. We present the results of the datasets using the best possible parameters for each dataset.

Impostor Generation

After making a combination of the known texts, wherever the known texts were more than one, while combining them but also keeping the number of letters below 5000, impostors were generated using the following prompt:

prompt = “Summarize the following cross-genre, cross-topic texts into a coherent summary, that maintains the writing style, the vocabulary, and the length as the known texts above, in the same language as the originals: {combined_text}”

1. The prompt was executed using the GPT-4 Mini model.
2. Role assigned: "Expert in linguistics and authorship attribution and creating impostors."
3. 30 Impostors for each language were created

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.

- A relation between the number of impostors and the number of grams?
-

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	0.579	-	0.637	0.669

Table 5 Prompt 1 Results

4.3 Experiment 2: Results

The experiments were conducted using the DU dataset of the PAN 15 evaluation dataset. We present the results of the dataset using the best possible parameters.

Impostor Generation

Here, the selection of the known text for the impostor production was random for every new prompt. Impostors were generated using the following prompt:

prompt = “You are an expert in ***linguistics and authorship verification***. Your task is to Write a text on the same general topic as the original, capturing a similar vocabulary and stylistic tone, but avoid copying any specific content, unique phrases, or personal expressions. Vary sentence structures and rephrase ideas to create a text that reads stylistically close to the original but without replicating exact details or unique idiosyncrasies. Ensure the following:

- The new text is the same language as the known text.

Maintain approximately the same ***length*** as the original text.”

4. The prompt was executed using the GPT-4 Mini model.

5. Role assigned: "Expert in linguistics and authorship attribution and creating impostors."

6. 110 Impostors were created

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.
- A relation between the number of impostors and the number of grams?

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	(110/75/40) – 0.652 (30/15/7) - 0.666	-	-	-

Table 6 Prompt 2 results

4.4 Experiment 3: Results

The experiments were conducted using the DU dataset of the PAN 15 evaluation dataset. We present the results of the dataset using the best possible parameters.

Impostor Generation

Again, the selection of the known text for the impostor production was random for every new prompt. Impostors were generated using the following prompt:

prompt = *“You are an expert in ****linguistics and authorship verification****. Your task is to write a text capturing the genre and topic of the unknown text but capturing a similar vocabulary and stylistic tone. Avoid copying specific content, unique phrases, or personal expressions of the known text. Vary sentence structures and rephrase ideas to create a text that reads stylistically close to the known but without replicating exact details or unique idiosyncrasies.*

Ensure the following:

- *The new text is the same language as the known text.*
- *Maintain approximately the same ****length**** as the known text.*

Known text:

"{known_text}"

Unknown text:

"{unknown_text}"

Now, produce only the new text, without including the original " ")

4. The prompt was executed using the GPT-4 Mini model.
5. Role assigned: "Expert in linguistics and authorship attribution and creating impostors."
6. 30 Impostors were created

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.
- A relation between the number of impostors and the number of grams?

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	0.488	-	-	-

Table 7 Prompt 3 Results

4.5 Experiment 4: Results

The experiments were conducted using the DU, GR SP, and EN datasets of the PAN 15 evaluation dataset. We present the results of the datasets using the best possible parameters.

Impostor Generation

Again, the selection of the known text for the impostor production was random for every new prompt. Impostors were generated using the following prompt:

Prompt = *Create an impostor text that shares a similar tone, general sentence structure, and level of formality with the original text but varies in specific vocabulary choices, function word usage, and unique stylistic markers. The impostor should feel similar overall style but not directly mimic the original's unique characteristics."*

Ensure the following:

- *Use a similar level of complexity in sentence structure and tone without directly copying unique expressions or phrasing.*
- *Avoid using distinctive patterns in function word usage or punctuation that are highly characteristic of the original author.*
- *Maintain a general similarity in topic to align with the original text without reproducing exact thematic details.*

Original text:

"{original_text}"

7. The prompt was executed using the GPT-4 Mini model.
8. Role assigned: " *Expert in stylometry and impostor text generation.*"
9. 110 Impostors were created for each problem (EN with 30/problem)

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.
- A relation between the number of impostors and the number of grams?

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	0.589	0.703	0.626	0.691

Table 8 Prompt 4 Results

Number 5 experiment with combined impostors. And number 6 with paraphrase.

4.6 Experiment 5: Results

The experiments were conducted using the DU, GR SP, and EN datasets of the PAN 15 evaluation dataset. We present the results of the datasets using the best possible parameters.

Impostor Generation

Again, the selection of the known text for the impostor production was random for every new prompt. Impostors were generated and gathered from previous prompts to be used as a pool of possible impostors:

10. The prompt was executed using the GPT-4 Mini model.

11. Role assigned: " *Expert in stylometry and impostor text generation.*"

12. 220 SP Impostors – 140 GR Impostors – 30 EN Impostors – 280 DU Impostors

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.
- A relation between the number of impostors and the number of grams?

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	0.647	0.703	0.652	0.704

Table 9 Combined Prompts impostors Results

4.7 Experiment 6: Results

The experiments were conducted using the DU, GR SP, and EN datasets of the PAN 15 evaluation dataset. We present the results of the datasets using the best possible parameters.

Impostor Generation

Again, the selection of the known text for the impostor production was random for every new prompt. Impostors were generated using the following prompt:

Prompt = "Please rewrite the original text, keeping the language level and length the same.

original text:

"{original_text}"

Now, produce only the new text, without including the original. Don't add " at the beginning or the end."

13. The prompt was executed using the GPT-4 Mini model.

14. Role assigned: " *Expert in stylometry and impostor text generation.*"

15. 100 Impostors were created for each problem

Observations - Findings

- Increasing the number of impostors consistently improved AUC scores.
- A relation between the number of impostors and the number of grams?

	EVALUATION SET – PAN 15			
	DU	EN	GR	SP
Original GI	0.667	0.803	0.656	0.785
Potha 2017	0.709	0.798	0.844	0.851
Proposed	0.565	0.634	0.613	0.672

Table 10 Paraphrase Prompt Results

5

Discussion

5.1 Revisiting the Research Goals

This research aimed to evaluate the effectiveness of the Impostors Method (IM) for authorship verification, specifically using prompts generated by large language models (LLMs). The study sought to address critical questions, such as whether LLM-generated impostors can improve the overall performance of IM, effectively mimic the stylistic features of known texts, and how refined prompts enhance verification accuracy.

Experiments were conducted on the PAN 15 dataset to achieve these objectives, leveraging multiple prompts to generate impostors under various configurations. These experiments were designed to investigate the strengths and limitations of the Impostors Method in distinguishing between texts written by the same author and those by different authors. By incorporating cross-genre and cross-topic datasets, the study aimed to evaluate the approach's robustness and generalizability.

The findings align with the initial goals, offering valuable insights into the relationship between impostor quality and model performance. While the method exhibited some limitations, such as achieving lower AUC scores compared to benchmark studies, it demonstrated the potential of LLMs to enhance impostor diversity and stylistic representation. These results highlight the critical role of prompt engineering, impostor quantity, and hyperparameter tuning in optimizing the method for practical applications. Furthermore, the study opens up new possibilities for utilizing LLMs innovatively, paving the way for distinct approaches to impostor creation, selection, and implementation.

5.2 Key Insights

The experimental results revealed several important insights regarding the Impostors Method (IM) performance and its reliance on LLM-generated impostors. These findings highlight the method's potential and limitations, offering a foundation for future research and refinement.

5.2.1 Performance Gaps Compared to Prior Work

The method's performance, as measured by AUC, fell short compared to prior successful studies (Potha, 2017). This gap underscores the need for further improvement in implementing the Impostors Method with LLM-generated impostors. Contributing factors include the quality of generated impostors, prompt design, the number of impostors, and the degree of similarity between impostors and the original texts. Additionally, cross-genre and cross-topic adaptability challenges

appear to influence the results. These findings suggest that refining both prompt engineering and the underlying logic of the Impostors Method may help bridge this performance gap.

5.2.2 *Role of Impostor Similarity*

An issue warranting further investigation is the similarity of impostors to the known texts, which may reduce their effectiveness in testing the model's ability to distinguish between same-author and different-author texts. This could indicate the need to modify the IM algorithm's logic to address such overlap. Paraphrased, less stylistically diverse impostors performed worse than those generated using more refined and varied prompts. These results emphasize the necessity of impostors that introduce adequate stylistic variability without deviating excessively from plausible authorial characteristics.

5.2.3 *Benefits of Increasing Impostor Diversity and Quantity*

The experiments demonstrated that increasing the number of impostors yielded measurable performance improvements. Larger impostor sets provided a more challenging and comprehensive testing framework, enhancing the method's robustness. Additionally, refined prompts that captured nuanced stylistic features further improved results. However, limitations in computational power and the cost associated with generating impostors using high-quality LLMs currently constrain the feasibility of producing large impostor sets. This trade-off highlights the importance of optimizing impostor generation for quality and cost-effectiveness.

5.2.4 *Cross-Genre and Cross-Topic Limitations*

The method exhibited inconsistent performance across different languages and datasets, with its effectiveness particularly challenged in cross-genre and cross-topic contexts. While the Impostors Method effectively captures stylistic features, its ability to generalize across diverse topics and genres remains limited. This finding suggests that generating both stylistically representative and content-independent impostors is critical to improving the method's adaptability and reliability in broader contexts.

5.2.5 *The Role of N-grams*

N-grams remain a widely used feature for capturing stylistic attributes in texts. However, considering the cross-genre and cross-topic nature of the problems studied, it is essential to rethink how stylistic data from the known text can be effectively utilized to generate impostors. While n-grams can capture specific stylistic traits, their advantages and limitations must be carefully balanced. On one hand, n-grams offer a straightforward mechanism for extracting patterns. Conversely, their reliance on local text structures may overlook broader stylistic variations, particularly in more complex or heterogeneous datasets. Addressing these trade-offs is essential for designing prompts and algorithms that fully leverage the potential of n-grams in cross-domain scenarios.

5.3 Key Insights

This research presents unexplored contributions to the study of authorship verification, concentrating on using LLM-generated impostors within the Impostors Method. By assessing the effectiveness of this method, the study provides insights into the strengths of stylistic analysis and highlights areas for further development.

5.3.1 Strengths of the Experimental Framework

- **Focus on Stylistic Features**
 - The study highlights the significance of stylistic features in distinguishing between texts authored by the same individual and those by different authors. By leveraging LLMs for impostor generation, this research explores how stylistic variety can enhance the robustness of the Impostors Method.
 - The use of n-grams, alongside prompt-engineered impostors, delivers a structured approach to stylistic analysis, unlocking new avenues for creating more refined and representative impostor texts.
- **Flexible Prompt-Based Setup**
 - The integration of LLMs into the Impostors Method offers flexibility in generating tailored impostors. This framework allows the creation of diverse impostors specific to each problem, enhancing the method's adaptability to different languages, topics, and genres.
 - The iterative improvement of prompts demonstrates the potential of prompt engineering as a crucial factor in enhancing authorship verification methodologies.

5.3.2 Challenges for Cross-Topic and Cross-Genre Verification

- **Generalizability Across Contexts**
 - The research underscores significant challenges in applying the Impostors Method to cross-topic and cross-genre datasets. When it comes to cross-genre and cross-topic texts, deviations in stylistic and topical features often blur the difference between same-author and different-author texts, reducing the method's efficacy.
 - These findings emphasize the need for enhanced algorithms adept at disentangling stylistic features from content-based variations.
- **Impostor Diversity and Quality**
 - While using LLMs enables the generation of high-quality impostors, the similarity between impostors and known texts emerged as a recurring limitation, and how it affects the method's performance is under question.
 - Addressing this requires further research into generating sufficiently diverse impostors yet remaining stylistically plausible within the problem context.

5.3.3 *Advancing the Impostors Method*

This research broadens the understanding of the impostor method by presenting a recurring, LLM-driven approach to impostor creation. It delivers a foundation for future studies to refine the method and examines innovative ways of applying LLMs in authorship verification. These contributions show a pathway for integrating advanced natural language generation methods with traditional stylistic analysis.

5.4 *Broader Implications*

The findings of this research carry broader implications for forensic linguistics, computational stylometry, and other fields where authorship verification plays a critical role. This study highlights new opportunities and challenges in leveraging advanced language generation models for real-world applications by integrating LLMs into the Impostors Method.

5.4.1 *Relevance to Forensic Linguistics and Computational Stylometry*

1. Enhanced Methods for Authorship Verification

- The integration of LLMs into the Impostors Method reveals how advanced text generation techniques can improve the detection of stylistic patterns. This approach presents a new dimension to forensic linguistics by enabling the creation of impostors that simulate realistic stylistic variations, offering more powerful testing frameworks.
- Computational stylometry benefits from the study's focus on extracting and utilizing stylistic features, such as n-grams, for creating representative impostors. This has implications for more general text classification tasks, including author profiling and plagiarism detection.

2. Applicability in Legal and Forensic Settings

- Authorship verification is increasingly relevant in legal investigations, where accurately attributing authorship can act as critical evidence. This research paves the way for deploying LLM-enhanced methods in forensic contexts, potentially enhancing the reliability of linguistic evidence.
- The scalability of LLMs for generating impostors provides an opportunity for handling large-scale, multilingual forensic datasets, a common requirement in global investigations.

5.4.2 *Potential Real-World Applications and Constraints*

1. Resource Constraints and Accessibility

- While the research indicates the potential of LLMs, the computational cost of generating high-quality impostors remains a significant limitation. This limitation restricts the method's accessibility for less organizations or resource-limited environments.

- Addressing these constraints requires exploring more efficient LLMs, cost-effective approaches to prompt engineering, and scalable implementations of the Impostors Method.

2. Ethical Considerations

- Using LLMs to generate impostors raises questions about potential misuse, such as the creation of effective but deceptive texts. Providing that these methods are applied ethically, with safeguards to prevent misuse, is critical for their broader adoption,

5.5 Open Questions

While this research has demonstrated the possibility of LLM-generated impostors within the Impostors Method, several unresolved questions remain. Handling these challenges will be critical for advancing the methodology and improving its applicability across broader contexts.

5.5.1 Refining the Impostors Method

- **How can the quality of impostors be systematically improved?**
 - The study highlighted the importance of prompt engineering for generating high-quality impostors. However, defining universal criteria for "quality" in impostor generation remains challenging. Future research could explore how stylistic and semantic features interact in creating effective impostors.
- **What algorithmic modifications are needed to handle impostor similarity?**
 - Impostors too similar to known texts reduce the method's effectiveness. An important next step is to investigate alternative logic for generating or evaluating impostors to minimize such overlap.
- **How can impostors generated based on the unknown text be implemented in the model?**
 - A promising direction is to generate impostors informed by the unknown text's characteristics, potentially creating a different kind of evaluation. The next step could be determining how to integrate these impostors into the current framework.

5.5.2 Distinguishing Stylistic and Topical Cues

- **How can algorithms better separate stylistic features from content-based variations?**
 - Cross-topic and cross-genre verification challenges suggest that current methods may entangle style with content. Techniques that disentangle these features, possibly through disentangled representation learning, could improve the method's robustness.

5.5.3 Scalability and Efficiency

- **What are the trade-offs between computational cost and performance?**

- Generating large numbers of impostors with LLMs is computationally expensive. Future work should evaluate cost-efficient models or explore hybrid approaches that combine pre-trained LLMs with simpler methods to balance quality and scalability.

5.5.4 *Applicability Across Languages*

- **How can the method be optimized for multilingual datasets?**
 - Variations in performance across different languages suggest that the method may need to be adapted specifically for each language. Identifying universal stylistic markers versus language-dependent features could enhance its generalizability.

5.5.5 *Ethical Considerations*

- **What safeguards are necessary to prevent misuse of LLM-generated impostors?**
 - The ability to create convincing impostors poses ethical questions about potential misuse, such as in disinformation or fraud. Exploring guidelines and safeguards for ethical applications of the Impostors Method is critical.

6

Conclusion

6.1 Summary of Findings

This thesis examined the effectiveness of incorporating Large Language Models (LLMs) into the Impostors Method for authorship verification. The primary objective was to determine if LLM-generated impostors could improve model performance by closely mimicking stylistic features while maintaining content diversity. Through experiments on the PAN 2015 dataset, the study highlighted several findings:

- **Impact of LLM-Generated Impostors:** LLMs demonstrated the ability to generate stylistically varied and challenging impostors, improving model robustness. However, results showed that the effectiveness of impostors in the impostors method depends on prompt quality, the number of them, and hyperparameter tuning.
- **Performance Across Genres:** Cross-genre and cross-topic tasks revealed limitations in the method's generalizability. While impostor diversity showed promising results, certain configurations struggled to address genre-specific variations.
- **AUC Comparisons:** Although the study aimed to surpass benchmarks (Potthast, Stein, & Hagen, 2016) (Potha & Stamatatos, 2017), the results showed a gap in achieving comparable AUC scores, emphasizing the need for further refinement and further research on the implementation of LLM's for impostor creation.

These findings show the potential of integrating generative AI into authorship verification while outlining challenges in optimizing the Impostors Method for broader applications.

6.2 Practical and Theoretical Implications

This research contributes to authorship verification by bridging traditional approaches with advancements in generative AI. Key implications include:

- **Practical Applications:** The ability to craft high-quality impostors through LLMs enhances verification models' adaptability, making them more applicable in fields like forensic linguistics, academic integrity, and cybersecurity.
- **Theoretical Contributions:** By analyzing the interplay between stylistic markers and content-independent features, this work expands understanding of how n-grams and stylistic cues can be leveraged for improved verification.

- **Methodological Insights:** The experiments underscore the importance of prompt engineering and impostor diversity in fine-tuning model performance.

6.3 *Directions for Future Research*

The findings open pathways for further exploration:

- **Refinement of Prompt Techniques:** Future studies could focus on systematically improving prompt design to generate impostors that better balance stylistic alignment and diversity, and even engineering new methods of implementing this idea, keeping in line with all the advancements happening in this field, along with those in the businesses that offer the LLM's.
- **Algorithmic Adaptations:** Enhancing algorithms to better disentangle stylistic and topical cues could improve cross-genre performance.
- **Scalability:** Exploring cost-efficient alternatives for generating large-scale impostor sets would address the computational challenges of using LLMs.
- **Ethical Considerations:** Given the potential misuse of AI-generated impostors, further research should establish guidelines for ethical deployment in real-world contexts.

• .

6.4 *Final Remarks*

This thesis demonstrates the promise of generative AI in advancing authorship verification methods. By integrating LLMs into the Impostors Method, the study highlights the potential for achieving greater stylistic accuracy and identifies the need for iterative refinements to meet real-world challenges. The results underline the iterative nature of scientific inquiry, where each discovery lays the groundwork for continued innovation. With sustained efforts, future research can unlock the full potential of gen AI in combination with the Impostors method in tackling complex authorship analysis tasks.

References

- [1] Ayele, A. A., Babakov, N., Bevendorff, J., Bonet Casals, X., Chulvi, B., Dementieva, D., Elnagar, A., Freitag, D., Fröbe, M., Korenčić, D., Mayerl, M., Moskovskiy, D., Mukherjee, A., Panchenko, A., Potthast, M., Rangel, F., Rizwan, N., Rosso, P., Schneider, F., Smirnova, A., Stamatatos, E., Stakovskii, E., Stein, B., Taulé, M., Ustalov, D., Wang, X., Wiegmann, M., Yimam, S. M., & Zangerle, E. (2024). Overview of PAN 2024: Multi-author writing style analysis, multilingual text detoxification, oppositional thinking analysis, and generative AI authorship verification. In L. Goeuriot et al. (Eds.), *Experimental IR Meets Multilinguality, Multimodality, and Interaction. Proceedings of CLEF 2024* (pp. XX-XX). Springer.
- [2] Bengesi, S., & Houkpati, Y. A. O. (2023). Advancements in generative AI: A comprehensive review of GANs, GPT, autoencoders, diffusion models, and transformers. *IEEE Access*, 11, 49882–49894. <https://doi.org/10.1109/ACCESS.2023.3291234>
- [3] Bevendorff, J., Chinea-Ríos, M., Franco-Salvador, M., Heini, A., Körner, E., Kredens, K., Mayerl, M., Pèzik, P., Potthast, M., Rangel, F., Rosso, P., Stamatatos, E., Stein, B., Wiegmann, M., Wolska, M., & Zangerle, E. (2023). Overview of PAN 2023: Authorship verification, multi-author writing style analysis, profiling cryptocurrency influencers, and trigger detection. In A. Arampatzis et al. (Eds.), *Experimental IR Meets Multilinguality, Multimodality, and Interaction. Proceedings of CLEF 2023* (pp. 459–481). Springer.
- [4] El Bouanani, S., & Kassou, I. (2014). Authorship verification using machine learning. *International Journal of Computer Applications*, 85(9), 1–5.
- [5] Grieve, J. (2007). Quantitative authorship attribution: An evaluation of techniques. *Literary and Linguistic Computing*, 22(3), 251–270.
- [6] Halvani, O., Winter, C., & Pflug, A. (2013). Authorship verification for different languages, genres, and topics. In *Proceedings of the 14th International Conference on Computational Linguistics and Intelligent Text Processing (CICLing 2013)* (pp. 301–312). Springer.
- [7] Ibrahim, M., & Al-Dossari, H. (2023). Enhancing authorship verification using sentence-transformers. *Journal of King Saud University - Computer and Information Sciences*, 35(1), 1–10.
- [8] Juola, P. (2015). The Rowling case: A proposed standard analytic protocol for authorship questions. *Digital Scholarship in the Humanities*, 30(3), 310–340.
- [9] Kestemont, M. (2014). Function words in authorship attribution: From black magic to theory? In *Proceedings of the 3rd Workshop on Computational Linguistics for Literature (CLFL 2014)* (pp. 59–66). Association for Computational Linguistics.
- [10] Koppel, M., Schler, J., & Argamon, S. (2009). Computational methods in authorship attribution. *Journal of the American Society for Information Science and Technology*, 60(1), 9–26.
- [11] Kumar, P., Musharaf, K., & Sagar, D. (2023). A comprehensive survey on authorship verification techniques. *Artificial Intelligence Review*, 56(2), 123–145.
- [12] Manolache, G., Ionescu, R. T., & Popescu, M. (2021). Transferring knowledge from sentiment analysis to authorship verification. *IEEE Access*, 9, 123456–123467.

- [13] Neal, T., Sundararajan, K., Woodard, D., & Peterson, L. (2018). Surveying stylometry techniques and applications. *ACM Computing Surveys*, 50(6), 86.
- [14] Nini, A. (2023). *Authorship analysis in forensic linguistics: Theory and practice*. Routledge.
- [15] PAN. (n.d.). PAN datasets. <https://pan.webis.de/data.html>
- [16] Potthast, M., Stein, B., & Hagen, M. (2016). Overview of the author identification task at PAN 2016. In N. Ferro et al. (Eds.), *Experimental IR Meets Multilinguality, Multimodality, and Interaction. Proceedings of CLEF 2016* (pp. XX-XX). Springer.
- [17] Reimers, N., & Gurevych, I. (2019). Sentence-BERT: Sentence embeddings using Siamese BERT-networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing* (pp. 3982–3992). Association for Computational Linguistics.
- [18] Sari, Y., Mikros, G., & Stamatatos, E. (2017). Cross-topic authorship verification based on topic modeling. In *Proceedings of the 18th International Conference on Computational Linguistics and Intelligent Text Processing (CICLing 2017)* (pp. 241–254). Springer.
- [19] Stamatatos, E. (2009). A survey of modern authorship verification methods. *Journal of the American Society for Information Science and Technology*, 60(3), 538–556.
- [20] Stamatatos, E. (2016). Authorship verification: A review of recent advances. *Research in Language*, 14(1), 69–83.
- [21] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Proceedings of the 31st International Conference on Neural Information Processing Systems (NIPS 2017)* (pp. 6000–6010). Curran Associates, Inc.
- [22] Yenduri, R., Al-Kabi, M., & Hossain, M. (2024). Investigating multilingual authorship verification techniques. *Journal of Natural Language Processing*, 12(4), 78–90.
- [23] Potha, N., & Stamatatos, E. (2017). A profile-based method for authorship verification. (pp. 313–326). Springer.