

Mystical Literature Under the Scrutiny of Artificial Intelligence: The Case of Maria Valtorta

RIZZIERO CAROFANO

Marelli Suspension Systems, Torino
rizziero.carofano@fastwebnet.it
ORCID: 0009-0005-5386-6216

LIBERATO DE CARO

Istituto di Cristallografia, Consiglio Nazionale delle Ricerche, Bari
liberato.decaro@cnr.it
ORCID: 0000-0002-7927-6178

Abstract. Maria Valtorta, a mystic lived in Italy in the XX century, claimed to have received mystical visions about the life of Jesus, the Virgin Mary, and some saints. She reported in her writings texts and dictations attributed to Jesus, her guardian angel, Azaria, and the Paraclete. Additionally, she wrote an autobiography. These texts underwent binary and multiclass supervised classification using Machine Learning algorithms. Classification metrics and confusion matrices have shown that the four classes of texts (Maria Valtorta, Jesus, Azaria, Paraclete) are so well separated in the deep structure of the language as if the texts were attributable to four different authors. Possible causes of this unexpected result are discussed.

Keywords: Machine Learning, Authorship Identification, Text Attribution, Artificial Intelligence, Maria Valtorta, Mystical visions.

Contribution. This study demonstrates that Machine Learning can provide an objective and quantitative approach to the analysis of mystical literature. It contributes to computational linguistics, stylometry, and digital humanities. We introduce a novel methodology for investigating complex literary and religious corpora and encourage interdisciplinary dialogue between Artificial Intelligence, linguistics and the study of mystical writings.

Use of AI. Artificial intelligence was used as an integral part of the research methodology. Specifically, supervised Machine Learning algorithms were employed to perform binary and multiclass classification of the textual corpus. The AI-based analyses generated the classification results, performance metrics, and confusion matrices discussed in the article. The research design, data preparation, methodological choices, interpretation of the results, discussion, and conclusions were developed and critically evaluated by the authors. No artificial intelligence tool was used to generate the scientific content or conclusions of the manuscript.

Introduction

The history of Christianity is rich in numerous mystical writings, produced by individuals who claimed to have had visions about the life of Jesus or Mary, and to have received messages, dictations, prophecies, which theology groups into the category of private revelations, to distinguish them from biblical writings, defined as Public Revelation, the only one to bind the faith of the Catholic Christians. Nevertheless, these mystical writings sometimes play an important role in the life of the Catholic Church, even influencing the liturgical calendar, as happened with the Feast of Divine Mercy, which falls on the first Sunday after Easter, following the ecclesiastical approval of the private revelations of Saint Faustina Kowalska. While the verification of the supernatural origin of mystics' visions pertains wholly to the Church, nevertheless, within the potentialities offered by Artificial Intelligence (AI) nowadays in analyzing vast databases of texts, it is legitimate to pose the following question: can AI be of any assistance in the study of mystics' writings?

As a case study, the possibility of analyzing the writings of Maria Valtorta (1897–1961) with a scientific approach has been the focus of

some recent studies, in the last decade (Matricciani and De Caro 2017a; Matricciani and De Caro 2017b; Matricciani and De Caro 2018; De Caro et al. 2020; Matricciani and De Caro 2020; Battisti 2020). In particular, the astronomical analysis of the numerous narrative details she reported in her main work, “The Gospel as Revealed to Me” (Valtorta 2015), which describes the life of Jesus and his mother, the Virgin Mary, has led to a coherent framework of dates that has allowed associating a precise chronology with all events described in the canonical Gospels (Matricciani and De Caro 2017a; Matricciani and De Caro 2017b) a research that subsequently inspired a series of studies on biblical chronology (La Greca and De Caro 2017; La Greca and De Caro 2019; La Greca and De Caro 2020; De Caro et al. 2021; Faro 2022).

Maria Valtorta claimed to have received visions about the life of Jesus, the Virgin Mary, and some saints. She also claimed to have received many dictations from Jesus, spiritual suggestions from her guardian angel Azaria, and comments on Sacred Scripture from the Holy Spirit (Paraclete). In the case of Maria Valtorta’s writings, there has been no official ecclesiastic pronouncement regarding the authenticity of these mystical visions. Indeed, between 1943 and 1951, Maria Valtorta produced a huge number of writings, the reading of which was forbidden to Christians because the main work was published, under the title “The Poem of the Man-God”, without prior ecclesiastical authorization, a formal constraint that has now expired after the reform of the ecclesiastical *Imprimatur*.

Maria Valtorta spent much of her life isolated from the outside world, as she became paralyzed from the waist down starting in 1934, about thirteen years before the beginning of the mystical experiences she claimed to have received. She received a medium-high education, corresponding to what was taught in the early years of a classical high school in Italy in the first half of the 20th century, but had no access to specialized libraries, nor to previously written/revealed stories of Jesus’ lives before putting her mystical visions into writings. Furthermore, Maria Valtorta reported everything in writing directly with her handwriting, sitting on her bed, without intermediaries. And she wrote everything simultaneously with her alleged mystical experiences, as emerges from her own writings, without letting time pass.

The large quantity of literary works written by Maria Valtorta has already been studied with a scientific approach, to evaluate the similarities and differences in style present in her writings (Matricciani and De Caro 2018). Maria Valtorta wrote an autobiography in two months, at the beginning of 1943 (Matricciani and De Caro 2018; Valtorta 2009), just before her mystical experiences began, on 23 April, 1943. From this, it is possible to deduce her writing style, which reflects her culture, her level of education, her linguistic skills. Maria Valtorta's writing style was characterized through appropriate statistical mathematical tools. Language statistics are generally calculated by counting characters, words, sentences, and word classifications. Some of these parameters are also the main "ingredients" of readability formulas, mathematical tools that measure quantifiable textual characteristics, mainly the length of words and sentences (Matricciani and De Caro 2018). According to these formulas, different texts can be automatically compared to evaluating their differences. Readability indices allow matching texts to expected readers, avoiding excessive difficulties and texts that are inaccessible or excessively simplified, and are based on quantities that any writer (or reader) can calculate directly. However, each readability formula gives a partial measurement of reading difficulty because its result is mainly linked to the length of words and sentences. It provides no clue about the correct use of words, the variety and richness of literary expression, its beauty, or effectiveness. Understanding a text is the result of many other factors, the most important being the reader's culture and sensitivity (Matricciani and De Caro 2018).

The analysis of Valtorta's texts was carried out using mathematical data that were not known or perceived, a priori, by the writer, in the sense that she is unaware of them. If we assume what Maria Valtorta says is true, that Jesus, the Virgin Mary, the Guardian Angel, and the Holy Spirit have all spoken the different texts found in her writings, then it has been interesting to examine them using impartial mathematical and statistical tools (Matricciani and De Caro 2018). In fact, for Maria Valtorta considered a direct author, we can analyze her autobiography (Valtorta 2009) —written before any mystical vision and dictation—and the many detailed descriptions of landscapes, roads, cities, etc., written in her main

work, which narrates the life of Jesus. These latter texts define the peculiar writing style of Maria Valtorta, and they can be taken as a reference to which any other text could be compared. The hypothesis that has been verified is whether there are texts not directly attributable to her personal linguistic and stylistic abilities, both measured mathematically (Matricciani and De Caro 2018). Starting from 23 April 1943, she wrote without making any corrections, incessantly, almost every day until 1947, and later intermittently till 1951 all her mystical writings. In the end, she filled 122 notebooks (over 13,000 pages), excluding the 7 notebooks of her Autobiography previously cited, written at the beginning of 1943 (Matricciani and De Caro 2018). A great amount of the texts related to the character Jesus (dictations of Jesus), in her mystical writings, were written from 1943 to 1944 (Valtorta 2006a; 2006b). Further few texts were written in the period 1945–1950 (Valtorta 2006c). The dictates of Maria Valtorta's guardian angel, Azariah, about theological and spiritual comments on the readings of 58 holiday masses, were written in year 1946 (Valtorta 2007). In the *Lessons on the Epistle of St. Paul to the Romans*, the Holy Author (Paraclete) dictates her, 48 lessons on the *Epistle to the Romans* in the years 1948–1950 (Valtorta 2008). Therefore, almost all of Maria Valtorta's writings were carried out in 7 years, with a peak of production from 1943 to 1945, a period short enough to reliably assume that her writing style could not vary significantly (Matricciani and De Caro 2018; Can and Patton 2004). Details on the timeline of Maria Valtorta's writings are given in (Matricciani 2022). In fact, without any new systematic studies, e.g. university courses, which justify the acquisition of levels of culture and greater mastery of the language (Crossley 2020), the deep language structure of an author's writing style cannot significantly change in few years. Conversely, a very long period of at least 10–15 years, characterizing the literacy production, could cause measurable style changes (Can and Patton 2004). Other causes that could justify the change of the writing style of an author are: the learning of a second language during the literacy production, in the third age (Pfenniger and Polz 2018); degenerative mental diseases, such as Alzheimer's disease (Hirst and Feng 2012), or dementia (Le et al. 2011). All these causes can be excluded for the considered period of Maria Valtorta's literacy production.

Furthermore, the emotional state of an author or mood disorders (e.g. depressive disorder) can influence the beauty and effectiveness of literary texts in transmitting emotions, but they cannot influence the deep structure of an author's writing style of very long texts. In these cases, Deep-Learning models could be used to aid medical diagnoses, distinguishing "psychologically healthy" voices from "critical" ones (Du and Sun 2022). In synthesis, research has demonstrated that the style of a single author remains stable, with few exceptions, previously mentioned, while writing style systematically differs between different authors, with an accuracy $\approx 98\%$ (Ramnial et al. 2016).

Since Maria Valtorta's literacy production is not influenced by the above-mentioned exceptions, a single writing style should characterize all her literacy texts. But from the study already conducted, a surprising fact emerged: the range of Maria Valtorta's writings, in terms of syntactic and semantic indices, extends almost as much as the entire Italian literature of several Italian writers, from the 1300s to today (Matriccioni and De Caro 2018). The texts of Maria Valtorta, in this statistical analysis, were grouped into homogeneous subsets: the parables of Jesus; the sermons and speeches of Jesus; passages beginning with "Jesus says"; passages beginning with "Mary says"; the commentary of the Holy Spirit on the Epistle of Saint Paul to the Romans (Valtorta 2008); the Book of Azaria (Valtorta 2007); the descriptions of landscapes, roads, cities, etc., which narrates the life of Jesus, written by Maria Valtorta in her main work (Valtorta 2015); her autobiography (Valtorta 2009). Probability tests showed that for most of her writings, especially for those she claims to have received by dictation, they are markedly distinct from all other texts, as if they were written by different authors. Moreover, the literary works explicitly attributable to Maria Valtorta (autobiography and descriptions) have probability distributions that differ significantly from those of literary works attributable to the alleged characters of Jesus and Mary (Matriccioni and De Caro 2018). The same mathematical approach, to the analysis of the deep language structure of texts, has recently allowed evidence that at least three authors wrote the texts attributed to Galen of Pergamum (129-216 Anno Domini), the well-known medical doctor and philosopher, a giant in the history of medicine (La Greca et al. 2024).

Given these interesting results, already scientifically demonstrated (Matricciani and De Caro 2018) and just summarized, the objective of the present study is to further investigate, through a different approach, through AI, the authorship of Maria Valtorta's writings, to verify if also AI-based approaches confirm the unexpected result already proved in (Matricciani and De Caro 2018). In fact, in the field of AI, the so-called text mining aims to highlight just the identification of the author, defining its literary profile, verifying the authorship of the text attributed to him. The author's profile is studied by examining texts to determine gender, age group, etc. The verification of text paternity aims to determine if two or more texts were written by the same author by analyzing linguistic patterns. Author identification, instead, aims to predict the possible author of an anonymous text from a predefined set of candidate authors. In particular, text paternity identification is used in many fields, including literary studies, history, and forensic linguistics. It is an interdisciplinary research field that uses computational methods to determine the author's identity, usually for an anonymous or contested text. This discipline lies at the intersection of computational linguistics, AI, and philology, using Machine Learning (ML) techniques to analyze the stylistic and linguistic characteristics of texts. The goal is to identify unique and recurring patterns in the author's writing style that can distinguish their works from those of other authors (Stamatatos 2000). The original idea presented in this study is to apply these author identification techniques to Maria Valtorta's literacy texts, particularly to the texts she claims to be "dictated by Jesus" (Valtorta 2006a; 2006b; 2006c; 2006d), comments from the Holy Spirit on biblical texts (Valtorta 2008), spiritual instructions from her guardian angel, Azaria (Valtorta 2007), as well as to the texts of her autobiography (Valtorta 2009), written just before the alleged mystical experiences began.

The traditional approach to author attribution through ML focuses on identifying distinctive patterns in the text that are unconsciously used by the authors (Koppel et al. 2002). The author attribution process begins with the collection of a corpus of texts whose author is known, which serves as a training set for the ML model. We introduced four distinct categories of texts, each attributable to a specific character in her writings: Maria

Valtorta, Jesus, the Paraclete, and Azaria, as if they were texts produced by different authors. Using ML techniques, we adopted two classification approaches: one that analyzes author pairs by directly comparing them; and another that considers all four authors simultaneously, to discern and quantify the stylistic differences between the texts.

The goal of this analysis is to verify if ML consistently identifies the same author for all four classes of texts or if, on the contrary, different writing styles are evidenced. In the former case, it could be concluded that it is plausible that Maria Valtorta simulated mystical experiences in which she claims to have received private revelations. In the latter case, it should be concluded that her mystical experiences lead still to modifying writing style. For theology, the eventual presence of different writing styles in the literacy production of an alleged mystic can be considered a useful indication to further investigate his/her alleged private revelations to verify their eventual authenticity.

Given the above general framework, the following study has been structured as follows. In the second section, the main author attribution techniques, and the state of the art of the application of these techniques have been summarized. In the third Section, the results of the application of ML algorithms to Maria Valtorta's texts have been synthesized, reporting in an appendix some insights. Finally, in the last Section, the main results have been discussed.

1. The state of the art in authorship attribution research

Research on text authorship identification relies on the structure of the text and the words used. Stylometry is part of this research area, and it is the study of distinct linguistic styles and individual writing practices with the aim of determining the authorship of a written text (Holmes 1998). A writing style represents the linguistic choices of a writer that persist throughout their work. Stylometric research is based on the hypothesis that each person has a unique and distinct writing style, called a "stylistic fingerprint" (Holmes 1994), which can be measured and learned by AI. An author's stylistic fingerprint indicates a set of characteristics frequently

used by that author, such as word length, sentence length, choice of specific words, and sentence structure. For these reasons, most features used for authorship identification are stylometric, especially in literary authorship (Holmes 1994; Juola 2006; Chaski 2001; Stamatatos et al. 2000; Anwar et al. 2019). In stylometric research, it is generally accepted that authors have unconscious writing habits (Chaski 1997). These habits become evident, for example, in the use of words and grammar. The more unconscious a process is, the less controllable it is. Therefore, words and grammar could be a reliable indication of an author's stylistic identity. These individual differences in language use are called idiolects. The unconscious use of syntax gives rise to the opportunity to perform authorship identification based on stylometric features.

One commonly used stylometric feature is based on so-called character n-grams. Experiments using character n-grams have been successful in determining text authorship (Clement and Sharp 2003). Additionally, structural information is also relevant for determining authorship, and successful classifications are reported when using the bi-gram of syntactic labels (Hirst and Feiguina 2007). Various factors can influence the performance of this task, such as the language of the messages used, the length of these messages, the number of authors and messages, the types of features, and the classification method. The number of features is often varied to determine the influence of certain types of features. In recent years, the domain of authorship analysis has embraced new research areas, typically with the emergence of ML techniques for text mining.

One of the recent and emerging trends in authorship analysis is the computational extraction of stylometric features from an author's text, instead of manually engineering stylometric features (Rosen-Zvi et al. 2004; Seroussi et al. 2011). The main goal of authorship identification is to decide the most probable author of a document from a list of known authors (Juola 2006). From an ML perspective, authorship identification can be perceived as a multi-class classification problem of a text with a specific label, where the role of the classes is played by competing authors (Sebastiani 2002). Sector studies in the field of authorship identification over the past two decades have revealed that it is a field of

great interest and has been mainly applied to the English language, with few other examples in other languages (Anwar et al. 2019).

A first approach to authorship identification was the use of univariate or multivariate measures that can reflect the style of a particular author. Measures such as word presence or frequencies of a specific word, average sentence length or average word length, vocabulary richness, have been proposed, although none of these univariate measures, taken individually, have proven to be adequate (Grieve 2007). The idea behind the multivariate approach, instead, is to take documents as points in the vector space and, using some suitable similarity measures, assign the document under inspection to that author whose documents are closest to it in the vector space (Argamon et al. 2009). Furthermore, other similarity parameters, based on different mathematical definitions of distance, such as Euclidean distance (Chaski 2001), Kullback-Leibler distance (Keřselj et al. 2003), or Hellinger distance (Raza et al. 2009), have been applied to various sets of features for text authorship identification.

A second approach involves statistical ML techniques. The individual author is identified by a category of texts, to which a specific value is assigned, and a classification model is created. ML techniques are further divided into two subgroups, one supervised and the other unsupervised. In supervised learning, a classifier is created using both the features and the value assigned to the category. Unsupervised models operate on unlabeled data (Blei et al. 2003). For authorship identification, supervised techniques include Support Vector Machine (SVM) (Argamon et al. 2009; Koppel et al. 2009; Stamatatos 2008), decision trees (Abbasi and Chen 2005), linear discriminant analysis (Chaski 2005), and neural networks (Tweedie et al. 1996; Zheng et al. 2006). SVM has outperformed other supervised techniques, such as linear discriminant analysis, in terms of accuracy (Anwar et al. 2019). Unsupervised classification techniques include principal component analysis (PCA) (Burrows 1987; Jamak et al. 2012), cluster analysis (Holmes 1992), distributed document representation (Le and Mikolov 2014), and Latent Dirichlet allocation (LDA) (Rosen-Zvi et al. 2004; Seroussi et al. 2011; Arun et al. 2009; Savoy 2013).

With the recent development and widespread use of ML models, however, authorship identification based on this latter approach has become the most promising solution. In this research area, vectorization techniques such as TF-IDF (Term Frequency-Inverse Document Frequency) are often used to transform texts into numerical formats that ML models can process (Juola 2006). Once identifying features have been extracted, traditional ML classification algorithms, such as decision trees, logistic regression, or SVM, are applied to train the model based on the known corpus. The model then learns to recognize the peculiarities of each author's writing style and can also be used to predict the author of new anonymous or contested texts (Brennan and Greenstadt 2009). Moreover, the use of advanced classification algorithms, such as SVMs, has significantly improved the accuracy of author attribution, allowing for more sophisticated and detailed analyses of texts. This model can learn from vast datasets, adapting to stylistic variations and recognizing complex patterns that would be difficult to detect through manual methods (Koppel et al. 2002).

2. Results of the analysis of Maria Valtorta's texts through ML techniques

In this section, we analyzed the texts of Maria Valtorta and the characters Jesus, the Holy Spirit (referred to hereafter as the Paraclete), and Azaria, using advanced ML techniques for textual classification. The goal is to explore the stylistic and linguistic differences among the texts attributed by Maria Valtorta to these characters as if they were different authors, employing binary and multiclass classification algorithms to identify distinctive patterns [Koppel et al. 2002]. We present in this section the main results obtained for the multiclass classification, directing the reader to the appendix A for technical details and to the Appendix B for results obtained through binary classification.

2.1. Division of the literary corpus into homogeneous blocks

The dataset used in this study consists of a total of 411 texts attributed to four alleged distinct authors, totaling 689,597 words: the writer Maria Valtorta (Valtorta 2009), Jesus (Valtorta 2015; 2006a; 2006b; 2006c; 2006d), Azaria (the guardian angel of Maria Valtorta) (Valtorta 2007), and the Holy Spirit (Paraclete) (Valtorta 2008)] all characters present in the works of Maria Valtorta. To ensure a representative database, 213 texts attributed to Jesus (Valtorta 2015; 2006a; 2006b; 2006c; 2006d) were selected, comprising 361,687 words; 90 texts attributed to Maria Valtorta, extracted from her Autobiography (Valtorta 2009) totaling 149,959 words; 60 texts attributed to Azaria (Valtorta 2007), totaling 99,520 words; and 48 texts attributed to the Paraclete (Valtorta 2008), totaling 78,431 words. To ensure comparable size among the analyzed documents of each class, the total texts attributed to the 4 characters were divided into blocks of texts consisting of an average of approximately 1660 words (about 2 pages). Table 1 summarizes these statistical data. This extensive collection of materials extracted from the writings of Maria Valtorta has allowed for a thorough comparison of the styles of the different alleged authors.

Table 1. Division of Maria Valtorta’s literary corpus into homogeneous text blocks by character category and length

Character	Number of texts	Total words	Average words per text
Jesus	213	361,687	1,698
Maria Valtorta	90	149,959	1,666
Azaria	60	99,520	1,659
Paraclete	48	78,431	1,634

2.2. Classification models

For each of the three pairs of characters (Maria Valtorta and Jesus, Maria Valtorta and Azaria, Maria Valtorta and the Paraclete), we implemented two different binary classification models: Support Vector Classifier

(SVC) (Cortes and Vapnik 1995) and Random Forest (RF) (Breiman 2001), selecting hyperparameters through careful tuning, as shown in the example provided in the Appendix B. This approach of binary comparisons allowed us to evaluate the effectiveness of each model in correctly identifying texts attributed to various authors.

The Support Vector Classifier (SVC), a specific type of the broader Support Vector Machine (SVM) algorithm, is primarily used for classification. It seeks to find the optimal hyperplane that separates data classes in the feature space, maximizing the margin between data points closest to the hyperplane, known as support vectors. This approach ensures good generalization, reducing the risk of overfitting. SVC is versatile, supporting both linear and non-linear kernels to handle both linearly separable and non-separable data (Cortes and Vapnik 1995).

The RF algorithm is an ensemble method based on decision trees, used for both classification and regression. It works by constructing a set of decision trees during the training phase and aggregating their predictions to produce a more accurate and stable final result. Each tree is trained on a random subset of data with randomly selected features, which helps reduce the risk of overfitting and increases the model's robustness. RF is appreciated for its ease of use, effectiveness in handling large datasets, and ability to provide insight into feature importance (Breiman 2001).

2.3. Results of multiclass analysis

After performing binary classification by pairwise comparison of the 4 characters, as described in the example provided in the Appendix B, the analysis was subsequently extended to multiclass classification, where we examined the algorithms' ability to distinguish between texts directly attributed to Maria Valtorta and those referring to the other 3 characters in a single classification model. In addition to multiclass SVC classification, we implemented supervised multiclass classification using Logistic Regression (LR) (Hosmer et al. 2013). Logistic Regression is a statistical method of supervised learning that models the probability of a categorical dependent variable based on one or more predictor variables, employing a logistic function to estimate these probabilities.

This approach is ideal for binary and multiclass classification tasks and is appreciated for its interpretive simplicity and computational efficiency (Hosmer et al. 2013).

The use of specific balancing techniques for multiclass classification further strengthened the effectiveness of our model, allowing for more precise capture of the stylistic and thematic nuances that distinguish texts attributed to various classes of presumed authors. The goal was to distinguish between texts directly attributed to Maria Valtorta and those referring to the other three alleged characters in her writings in a single classification model. In summary, we adopted the following methodology for multiclass analysis, discussed in detail in Appendix A:

- Data preprocessing, removing special characters, lemmatizing words, converting valid tokens to lowercase, and removing punctuation, numbers, and stop-words;
- Text vectorization, applying the TF-IDF technique to convert normalized texts into numerical representations;
- Dimensionality reduction, using the Truncated SVD approach to compress the data into a three-dimensional space to facilitate graphical visualization of class distribution;
- Dataset balancing, introducing class weighting through the `compute_class_weight` function of scikit-learn, to mitigate the problem of class imbalance in a multiclass context, where the goal was to treat minority classes fairly without reducing the dataset size through under-sampling.

In this multiclass context, class balancing using `compute_class_weight` played a crucial role in ensuring that the model assigned proportional importance to all classes during training, thus compensating for discrepancies in the number of examples per class. This approach allowed the models to learn more effectively from less represented classes, improving generalization and discriminative ability.

Figure 1 shows the 3D representation of the data through dimensionality reduction obtained from Truncated SVD of the TF-IDF matrix to demonstrate the distribution of texts attributed to Maria Valtorta, Jesus, Paraclete, and Azaria. It can be observed how the four text classes are markedly complementary to each other, implying a stylistic

difference present among them. The different colored points form homogeneous clusters in well-defined regions of the 3D plot, with few exceptions.

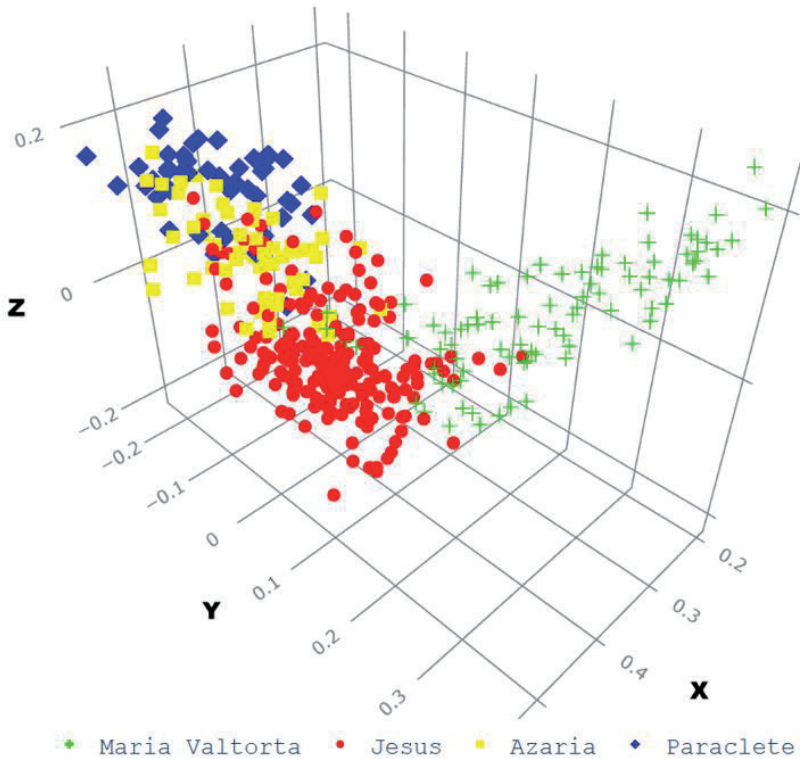


Figure 1. 3D representation of the data through dimensionality reduction obtained from Truncated SVD of the TF-IDF matrix to show the distribution of texts attributed to Maria Valtorta, Jesus, the Paraclete, and Azaria. Each point represents a text, with colors indicating authorial attribution. The complementary visualization in 3D space of the four colors highlights the stylistic differences present among the analyzed categories.

2.4. Cross-validation and evaluation metrics

Cross-validation is an evaluation technique aimed at testing the model's ability to generalize on an independent dataset. Specifically, k-fold cross-validation divides the dataset into k distinct subsets. The model is trained

on k-1 of these folds as training data and tested on the remaining fold to calculate performance. This process is repeated k times, with each fold used exactly once as the test set. Cross-validation accuracy, calculated as the average performance across all folds, provides a robust estimate of the model's efficiency (Kohavi 1995).

The main evaluation metrics used are accuracy, precision, recall, and F1 score. To fully understand these metrics, it is important to explicitly define the concepts of true positives (TP), false positives (FP), true negatives (TN), and false negatives (FN):

- True Positives (TP): number of positive instances correctly identified as positive by the model.
- False Positives (FP): number of negative instances incorrectly identified as positive by the model.
- True Negatives (TN): number of negative instances correctly identified as negative by the model.
- False Negatives (FN): number of positive instances incorrectly identified as negative by the model.

The formulas for the metrics are:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}. \quad (1)$$

Accuracy represents the percentage of correct predictions out of the total instances. It provides a general measure of the model's effectiveness.

$$Precision = \frac{TP}{TP+FP}. \quad (2)$$

Precision indicates the percentage of true positives among all instances identified as positive by the model. It measures the model's ability to not label a negative instance as positive.

$$Recall = \frac{TP}{TP+FN}. \quad (3)$$

Recall measures the percentage of true positives correctly identified out of the total actual positives. It evaluates the model's ability to find all positive instances.

$$F1\ Score = 2 * \frac{Precision * Recall}{Precision + Recall} \quad (4)$$

F1 score is the harmonic mean of precision and recall. It provides a single score that balances both metrics, useful when a comprehensive assessment of the balance between precision and recall is desired.

These metrics are crucial for understanding not only how well a model predicts in general (accuracy), but also how reliable its positive predictions are (precision) and how capable it is of identifying all positive cases (recall). The *F1 score* is particularly valuable when classes are imbalanced or when maintaining a balance between precision and recall is important (Powers 2011).

Table 2 summarizes the values of optimized hyperparameters and performances for the SVC and Logistic Regression (LR) models applied to the comparison between Maria Valtorta, Azaria, Jesus, and the Paraclete.

Table 2. Hyperparameters of multiclass classification models

Classifier	Hyperparameters	Accuracy in Cross- Validation %	Accuracy %
SVC	Kernel = 'sigmoid'; gamma = 2.85; class_weight = class_weight_dict; probability = True; C=4; coef0 = 0.55 random_state = 42	96.1	97.6
Logistic Regression (LR)	penalty = 'elasticnet' l1_ratio = 0.19; C = 2.6; solver = 'saga'; max_iter = 1000; class_weight = class_weight_dict; random_state = 42	94.4	96.0

In SVC, the kernel flag defines the type of kernel function selected for calculating the hyperplane, used to separate classes; gamma regulates the influence of each example; C is the regularization parameter that balances the classification error and model complexity; the 'coef0' parameter is an independent term in polynomial and sigmoidal kernels; the 'probability' flag enables probability estimation. LR utilizes the elasticnet-penalty

parameter, which combines two different norms for regularization, with the `l1_ratio` parameter calibrated to balance the two norms, and the `C` parameter for regularization. Both models employ: the `class_weight` function to handle imbalanced data classes; the `max_iter` parameter to manage the maximum number of iterations, and the `random_state` parameter for reproducibility.

Table 3 summarizes the precision, recall, and F1 score values for the SVC classification.

Table 3. Precision, recall, and F1 score values for the SVC classification

Character	Precision (%)	Recall (%)	F1 (%)
Jesus	98.4	96.9	97.6
Paraclete	93.8	100.0	96.8
Azaria	94.7	100.0	97.3
Maria Val-torta	100.0	96.3	98.1

Table 4 summarizes the precision, recall, and F1 score values for the LR classification.

Table 4. Precision, recall, and F1 score values for the LR classification

Character	Precision (%)	Recall (%)	F1 (%)
Jesus	92.8	100.0	96.2
Paraclete	100.0	93.3	96.6
Azaria	100.0	94.4	97.1
Maria Val-torta	100.0	88.9	94.1

Tables 3 and 4 show very high precision and recall values, very close to 100%, demonstrating that the four text classes present in the writings

of Maria Valtorta are indeed stylistically so well distinguishable that they could be attributed to 4 distinct authors.

2.5. Confusion matrix

The confusion matrix details the number of correct and incorrect predictions made by the model, allowing for a deeper analysis of specific performance in each class. It compares the actual labels of the texts with the model's predictions, where rows represent the true characters and columns represent the predicted characters. Values along the diagonal indicate correct classifications, showing how many texts were accurately attributed to the right character. Off-diagonal values represent misclassifications, indicating instances where texts were incorrectly assigned to a different character. Through this matrix, we can directly evaluate the model's ability to distinguish between texts of Jesus, Azaria, the Paraclete, and those of Maria Valtorta. Figures 2 and 3 depict the confusion matrices obtained, respectively, from the SVD and LR classifications.

Figures 2 and 3 show very few off-diagonal elements, demonstrating that the style of the first-person texts written by Maria Valtorta is always well distinguishable from those attributed to the three characters, Jesus, Azaria, and the Paraclete. Furthermore, the texts of the three characters are also clearly distinguishable from each other, as if they were texts from different authors.

It is important to note that in the evaluation of the multiclass classification, we used k-fold cross-validation to select and optimize the model, evaluating it on the entire available dataset to provide an accurate estimate of its average performance. Subsequently, the final model was trained using a 70% split for training and 30% for validation, providing an additional and independent verification of its ability to generalize on unseen data.

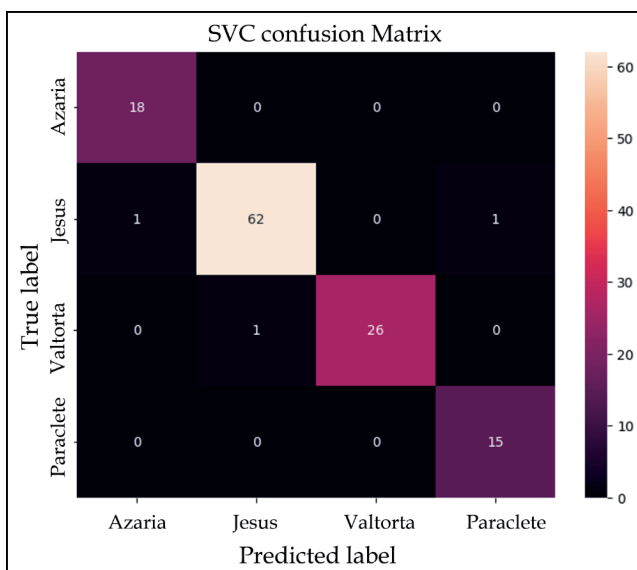


Figure 2. Confusion matrix for the SVC classification

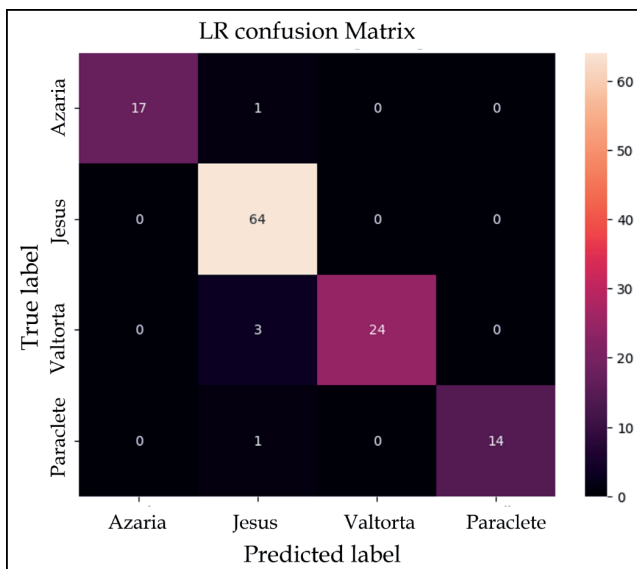


Figure 3. Confusion matrix for the LR classification

While the final model was trained on a fixed portion of the dataset, the k-fold cross-validation process allowed it to be tested on multiple configurations, ensuring a robust estimation of performance. This combination of cross-validation and final validation was adopted to ensure the maximum reliability of the results: cross-validation provided a comprehensive assessment using different subsets of the data, and the final evaluation on separate data confirmed that the performance was consistent and did not depend on the specific chosen split.

3. Discussion and conclusions

In this study, we subjected the writings of Maria Valtorta to binary and multiclass classification using ML algorithms to verify the potential presence of any significant stylometric differences between the texts written by the author in the first person and those attributed to three distinct characters: Jesus, her guardian angel named Azaria, and the Holy Spirit Paraclete. Indeed, Maria Valtorta claimed about receiving mystical visions of the life of Jesus, the Virgin Mary, and dictations, comments on biblical texts, and spiritual suggestions from Jesus, Azaria, and the Paraclete. For this purpose, we introduced four distinct categories of texts, each attributable to a specific character in her writings: Maria Valtorta, Jesus, the Paraclete, and Azaria, as if they were texts produced by different authors. The aim of the analysis was to determine whether the ML consistently identifies the same authorship for all four classes of texts or, conversely, confirms that we are dealing with a very different situation, as these texts would be written by four different authors, and with what probability this occurs. Classification metrics and confusion matrices demonstrated that the four classes of texts (Maria Valtorta, Jesus, Azaria, Paraclete) are stylistically so well separated that they could be also attributed to four different authors. This result confirms what has been already obtained in (Matricciani and De Caro 2018).

Considering all the results obtained – in this study and in (Matricciani and De Caro 2018), with two approaches fully complementary – we can conclude that the writing style of the characters Jesus, Azaria, Paraclete

is clearly distinguishable from that of Maria Valtorta with a level of accuracy usually obtained for writings of different authors. It should be argued that, in principle, different topics could influence the efficiency of algorithms used for authorship attribution. In the case of Maria Valtorta, however, four different authors emerged (nominally, Maria Valtorta, the Angel Azariah, Jesus and the Paraclete) and, excluding the autobiography, all the other texts are characterized by the same topic, linked to themes of religion and faith. Therefore, if the authorship attribution was effectively influenced by the topic, with our AI-based approach we should have identified only two distinct writing styles, the first of the author Maria Valtorta for the autobiography, the second related to a second author, once called Azariah, once called Jesus, and once called Paraclete by Maria Valtorta, but not four different writing styles, as already verified in (Matricciani and De Caro 2018). This finding demonstrates that the topic could have only a small influence (if any) in affecting the writing style of Maria Valtorta, as shown in the present study and already shown in (Matricciani and De Caro 2018).

In conclusion, the results of our stylometric analysis suggest that four distinct writing styles can be identified in the literary production of Maria Valtorta, with divergences comparable to those found between different authors. Considering that these texts—except for the autobiography—deal with homogeneous themes related to faith and religion, it is unlikely that such stylistic variations can be attributed solely to content. While authorship attribution cannot provide answers regarding the supernatural origin of the texts, it can serve as a preliminary tool to support the analysis of alleged mystical revelations. In this sense, the results obtained could contribute useful elements to the Dicastery for the Doctrine of the Faith, enabling a more informed and cautious evaluation of similar cases by integrating theological investigation with objective linguistic and scientific analysis, where possible, as widely demonstrated by numerous studies on the alleged private revelations of Maria Valtorta.

Appendix A: Pre-processing, TF-IDF Vectorizer, Dimensionality reduction, Dataset balancing

Pre-processing

The original texts of any author are not immediately suitable for use in machine learning models. Without proper preliminary text normalization, analysis through machine learning algorithms may yield suboptimal results. Therefore, the following preliminary actions were taken on the data:

- *Removal of special characters*: Elimination of symbols such as @, #, \$, %, &, * that do not contribute to the semantic content of the text and may interfere with analysis.
- *Lemmatization of words*: Process of reducing words to their base or lemma form. For example, words like “going”, “went”, “gone” are all reduced to the lemma “go”. This unifies different inflected forms of a word, improving data consistency.
- *Conversion of valid tokens to lowercase*: Transforming all words to lowercase letters to prevent “Dog” and “dog” from being treated as different terms.
- *Removal of punctuation, numbers, and stop-words*: Elimination of punctuation marks (.,;:!?), numerical digits, and very common words in the language (such as “the”, “and”, “but”) that do not provide useful information for analysis.

It is worth noting this aspect of the removal of all punctuation marks which, conversely, are fundamental in other approaches for analyzing the deep language structure of texts (Matricciani and De Caro 2018). This finding is important for two reasons. The first is that machine-learning-based authorship identification can be considered fully complementary to that based on previous approaches (Matricciani and De Caro 2018). The second is that, sometimes, mystical writers are almost illiterate. In these cases, punctuation marks are often missing and, therefore, they are introduced by the editor which first publish mystical writings. Evidently in these cases the approach discussed in (Matricciani and De Caro 2018) is not applicable, unlike approaches based on machine learning.

For text processing and normalization, we employed the ‘it_core_news_sm’ linguistic model from spaCy (Juola 2006), specifically trained to understand and analyze the Italian language. This allowed us to perform lemmatization operations, removal of non-printable characters, substitution of non-separable spaces with normal spaces, and identification and exclusion of punctuation, numbers, and stop-words while maintaining the semantic integrity of the original texts.

TF-IDF Vectorizer

After applying the data pre-processing techniques described previously, we performed text vectorization using TF-IDF and dimensionality reduction with Truncated-SVD. The TF-IDF vectorization technique (Term Frequency-Inverse Document Frequency) converts normalized texts into numerical representations, facilitating their use in machine learning algorithms.

TF-IDF allows for assessing the importance of each word in a document relative to a comprehensive set of documents. This balance is achieved by weighting the term’s frequency in the document (TF) with its rarity across the entire document collection (IDF). Specifically:

- *Term Frequency (TF)*: Calculated as the number of times a word appears in a document divided by the total number of words in the document. For example, if the word “learning” appears 3 times in a document of 100 words, the TF is 0.03.
- *Inverse Document Frequency (IDF)*: Measures the importance of a word by reducing the weight of common words. It is calculated as the logarithm of the total number of documents divided by the number of documents containing the word. For example, if “learning” appears in 10 out of 1,000 documents, the IDF is $\log(1000/10) = 2$.

The TF-IDF weight for a word is given by the product $TF * IDF$, ensuring that rare but significant terms have a higher weight compared to common terms.

For text vectorization, we adopted the unigram-based approach, focusing exclusively on individual words. This choice aims to maintain relevant content analysis capability while limiting the complexity of the model and ensuring greater computational efficiency.

Originally developed to address information retrieval problems, TF-IDF has become a fundamental tool in document classification and textual analysis due to its simplicity of implementation and ability to highlight words that significantly characterize a document within a large corpus (Salton and Buckley 1988).

Dimensionality reduction

In our study, we chose to use the Truncated Singular Value Decomposition (Truncated SVD) algorithm to reduce the dimensionality of textual data transformed via TF-IDF. After TF-IDF vectorization, we obtain a large sparse matrix where:

- *The rows represent the documents (texts):* each row corresponds to a text in the corpus.
- *The columns represent the unique words in the vocabulary:* the number of columns equals the total number of distinct terms in the vocabulary built from the corpus.
- *The matrix elements:* each element (i, j) of the matrix corresponds to the TF-IDF weight of the j -th term in the i -th document. If a term does not appear in a specific document, the value is zero.

This method, known for its effectiveness in decomposing matrices into principal components, allowed us to reduce complexity while retaining the most relevant information. By applying Truncated SVD to this matrix, we compressed the high-dimensional representations into a three-dimensional space, reducing the number of columns to three principal components. This process not only preserves key relationships between documents but also enables a three-dimensional representation of the data.

The dimensionality reduction facilitated the analysis of text clustering patterns, allowing for visually observing how documents attributed to different authors distinguish themselves in the reduced three-dimensional space. The resulting graphical representation offers an intuitive view of distinctions between various text classes, emphasizing the effectiveness of Truncated SVD in highlighting distinctive text features in an understandable manner (Landauer et al. 1998).

Dataset balancing

Additionally, we adopted specific strategies to balance the datasets, such as under-sampling for binary classification and class weighting for multiclass classification, thereby ensuring fair and representative analysis. To ensure that classification models treat all classes equally, regardless of their representation in the dataset, specific balancing strategies were adopted for different types of classification: the use of the `compute_class_weight` function from scikit-learn (Pedregosa et al. 2011) for multiclass classification and the implementation of under-sampling for binary classification. In multiclass classification, the `compute_class_weight` function automatically computes weights inversely proportional to the class frequencies in the dataset, thereby assigning greater weight to less represented classes during the training phase. This approach compensates for the effect of an imbalanced dataset, where some classes are significantly more numerous than others (Pedregosa et al. 2011).

For binary classification, under-sampling was employed to reduce the discrepancy between classes by decreasing the number of samples in overrepresented classes. This technique involves randomly selecting a subset of samples from the most numerous classes, aiming to balance the number of samples for each class. The goal of under-sampling is to decrease the risk of overfitting on samples from the majority classes and to improve model generalization on minority classes. The adoption of these specific strategies for each type of classification allowed for the development of both binary and multiclass models that treated all classes equally, regardless of their original size in the dataset. Through this dual strategy, we aimed to optimize the performance of classification models, ensuring a fair and balanced representation of all classes in the learning process.

Appendix B: binary classification tests

To deepen the comparative analysis between the texts attributed to Maria Valtorta and the three distinct characters (Jesus, Azaria, and the Paraclete), we also adopted a binary classification approach. For each of the three pairs of characters/authors (Valtorta – Jesus, Valtorta – Azaria, Valtorta – Paraclete), we performed: data pre-processing; text vectorization; dimensionality reduction; dataset balancing, as already discussed in section 3 for the multiclass approach. For each character pair, we used two classification models: SVC and RF, selecting hyperparameters through a careful tuning phase. We found that the 3D graphical representation obtained from dimensional reduction shows a clear distinction between the texts attributed to each of the three different characters (Jesus, Azaria, Paraclete) compared to those of Maria Valtorta, highlighting significant stylistic differences between all three pairs of characters/authors. The metrics confirm the high discriminative capacity of the models used. Below are the details of the tests performed for the Valtorta – Jesus’ pair. Similar results are obtained in binary classifications with the other two characters’ pair Valtorta – Azaria and Valtorta –Paraclete.

Figure A1 shows the 3D graphical representation obtained through dimensionality reduction with Truncated SVD of the TF-IDF matrix. This visualization allows observing the separation between the two categories of texts, highlighting the stylistic distinction between the texts attributed to MV and Jesus.

Table A1 summarizes the hyperparameters for the binary classification models RF and SVC. This table provides an immediate overview of each model’s effectiveness in correctly identifying texts attributed to Jesus or MV, with SVC showing a slight advantage in terms of cross-validation accuracy. In RF the “max_depth” parameter limits the maximum depth of the trees; “min_samples_leaf” is the minimum number of samples required to form a leaf; “min_samples_split” represents the minimum number of samples required to split an internal node; “n_estimators” indicates the number of trees in the forest. In SVC, the kernel defines the type of kernel function used for calculating the hyperplane used to separate the classes; gamma regulates the influence of each example;

C is the regularization parameter that balances the classification error and the model complexity; the `coef0` parameter is an independent term in polynomial and sigmoidal kernels; the `probability` flag enables probability estimation.

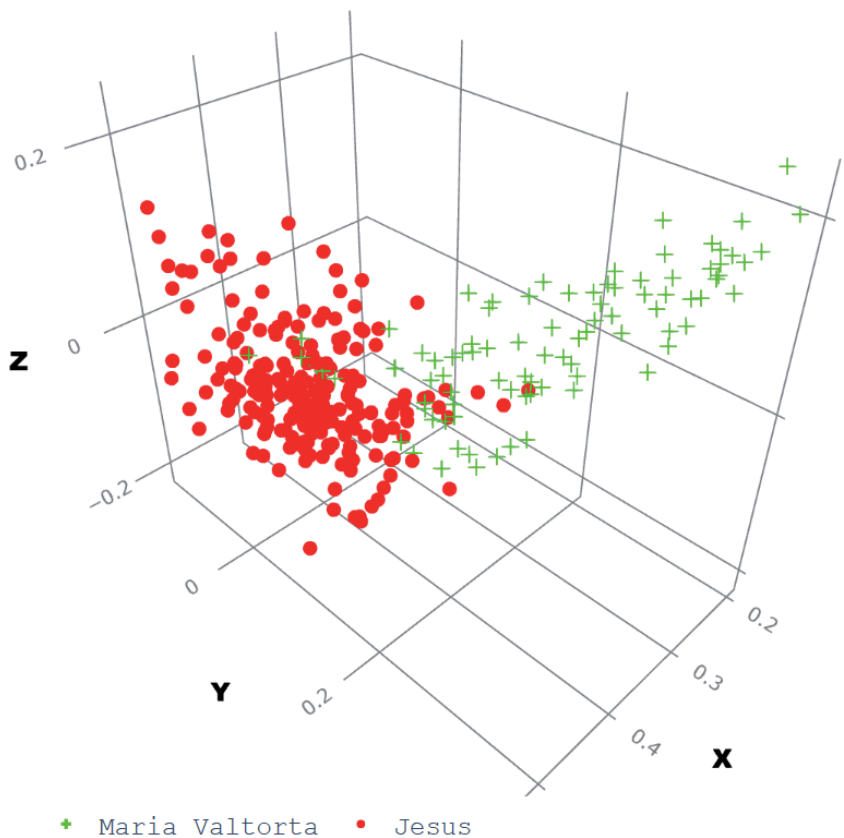


Figure A1. The 3D graphical representation, obtained through dimensionality reduction with Truncated SVD of the TF-IDF matrix, visualizes the spatial distribution of texts attributed to Jesus and MV. Each point in the graph represents a text, colored according to its authorial attribution. MV in the figure indicates Maria Valtorta's texts.

Table A1. Hyperparameters for RF and SVC classification models. Performance achieved in terms of both accuracy and cross-validation on the test set (last column) are shown

Classifier	Hyperparameters	Accuracy in Cross-Validation %	Accuracy %
RF	max_depth=4 min_samples_leaf=5 min_samples_split=2 n_estimators=180 random_state=42	98.3	98.9
SVC	kernel='sigmoid'; gamma=1.5; C=2; coef0=0.3; probability=True; random_state=42	99.4	98.9

Table A2 summarizes metrics' values obtained for the RF-supervised binary classification.

Table A2. Precision, recall, and F1 score values for the RF-supervised binary classification

Character	Precision (%)	Recall (%)	F1 (%)
Jesus	99.3	99.3	99.3
Maria Valtorta	96.3	96.3	96.3

Table A3 summarizes metrics' values obtained for the SVC-supervised binary classification.

Table A3. Precision, recall, and F1 score values for the SVC-supervised binary classification

Character	Precision (%)	Recall (%)	F1 (%)
Jesus	100.0	98.7	99.3
Maria Valtorta	93.1	100.0	96.4

Figure A2 shows the confusion matrix of the binary classification between Maria Valtorta and Jesus, supervised by the RF model. Finally, Figure A3 shows the confusion matrix of the binary classification between Maria Valtorta and Jesus supervised by the SVC model. Even the binary classification tests lead to very high values of accuracy, precision, and recall, very close to 100%. The confusion matrices, reported in Figures A2 and A3, show very few off-diagonal elements, demonstrating that the writing style of texts written in the first person by Maria Valtorta is always well distinguishable from that attributed to Jesus. Similar results are obtained for the pairs Maria Valtorta – Azaria and Maria Valtorta – Paraclete.

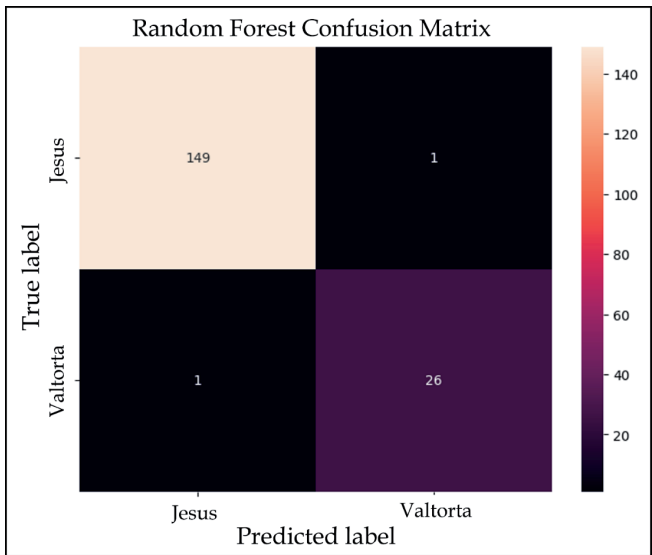


Figure A2. Confusion matrix of the binary classification between Maria Valtorta and Jesus supervised by the RF model

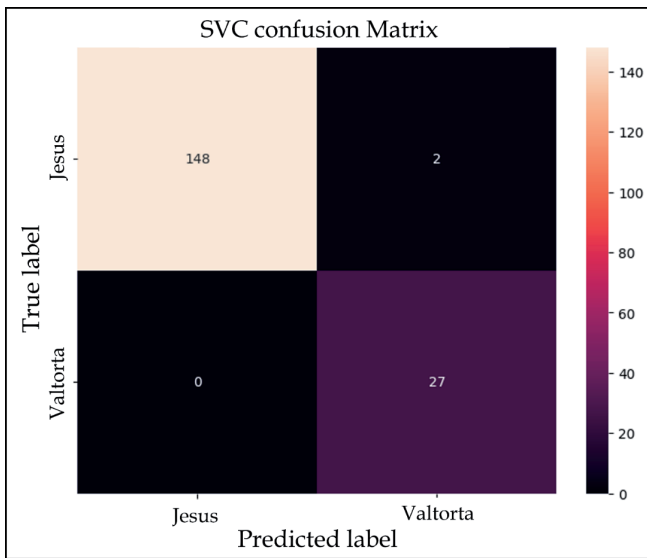


Figure A3. Confusion matrix of the binary classification between Maria Valtorta and Jesus supervised by the SVC model

References

- Abbasi, Ahmed, and Hsinchun Chen. 2005. "Applying authorship analysis to extremist-group Web forum messages." *IEEE Intelligent Systems* 20: 67–75. <https://doi.org/10.1109/MIS.2005.81>.
- Anwar, Waheed, Imran Sarwar Bajwa, and Shabana Ramzan. 2019. "Design and Implementation of a Machine Learning-Based Authorship Identification Model." *Scientific Programming*, vol. 2019, Article ID 9431073. <https://doi.org/10.1155/2019/9431073>.
- Argamon, Shlomo, Moshe Koppel, James W. Pennebaker, and Jonathan Schler. 2009. "Automatically profiling the author of an anonymous text." *Communications of the ACM* 52: 119–katasubramanian Suresh, M. Narasimha Murty, and C. E. Veni Madhavan. 2009. "Stopwords and stylometry: a latent Dirichlet allocation approach." *NIPS Work*, 1–4. https://users.umi.acs.umd.edu/~ying/nips_tm_workshop/12.pdf.

- Battisti, Gianfranco. 2020. "Geografie del sacro: la letteratura mistica come fonte di conoscenza." *Documenti Geografici*, 2: 1–22. http://dx.doi.org/10.19246/DOCUGEO2281-7549/201902_01.
- Blei, David M., Andrew Y. Ng, and Michael I. Jordan. 2003. "Latent dirichlet allocation." *Journal of Machine Learning Research* 3: 993–1022.
- Breiman, Leo. 2001. "Random Forests." *Machine Learning* 45: 5–32. <https://doi.org/10.1023/A:1010933404324>.
- Brennan, Michael R., and Rachel Greenstadt. 2009. "Practical Attacks Against Authorship Recognition Techniques." In *Proceedings of the Twenty-First Conference on Innovative Applications of Artificial Intelligence*, vol. 2, 60–65. Menlo Park, CA: AAAI Press.
- Burrows, John F. 1987. "Word-patterns and story-shapes: the statistical analysis of narrative style." *Literary and Linguistic Computing*, 2, 61–70.
- Can, Fazli, and James M. Patton. 2004. "Change of writing style with time." *Computers and the Humanities*, 38: 61–82.
- Chaski, Carole E. 1997. "Who wrote it? steps toward a science of authorship identification." *National Institute of Justice Journal*, 233: 15–22.
- Chaski, Carole E. 2001. "Empirical evaluations of language-based author identification techniques." *Forensic Linguistics*, 8: 1–65. <https://doi.org/10.1558/sll.2001.8.1.1>.
- Chaski, Carole E. 2005. "Who's at the keyboard? Authorship attribution in digital evidence investigations." *International Journal of Digital Evidence*, 4: 1–13.
- Clement, Ross, and David Sharp. 2003. "Ngram and bayesian classification of documents for topic and authorship." *Literary and linguistic computing*, 18: 423–47. <https://doi.org/10.1093/lc/18.4.423>.
- Cortes, Corinna, and Vladimir Vapnik. 1995. "Support-vector networks." *Machine Learning*, 20: 273–97. <https://doi.org/10.1007/BF00994018>.
- Crossley, Scott A. 2020. "Linguistic features in writing quality and development: An overview." *Journal of Writing Research*, 11: 415–43.
- De Caro, Liberato, Fernando La Greca, and Emilio Matricciani. 2020. "Saint Peter's First Burial Site According to Maria Valtorta's Mystical Writings, Checked Against the Archeology of Rome in the I Century." *J*, 3: 366–400. <https://www.mdpi.com/2571-8800/3/4/29>.
- De Caro, Liberato, Fernando La Greca, and Emilio Matricciani. 2021. "The Beginning of the Christian Era Revisited: New Findings." *Histories*, 1: 145–68. <https://www.mdpi.com/2409-9252/1/3/16>.
- Du, Xiaowei, and Yunmei Sun. 2022. "Linguistic features and psychological states: A machine-learning based approach." *Frontiers in Psychology*, 13: 823313.

- Faro, Giorgio. 2022. "L'ultima cena di Gesù era pasquale? Un'ipotesi da riconsiderare." *Annales Theologici*, 36: 47–79. <https://www.annalestheologici.it/article/view/3531>.
- Grieve, Jack. 2007. "Quantitative authorship attribution: an evaluation of techniques." *Literary and Linguistic Computing*, 22: 251–70. <https://doi.org/10.1093/llc/fqm020>.
- Hirst, Graeme, and Ol'ga Feiguina. 2007. "Bigrams of syntactic labels for authorship discrimination of short texts." *Literary and Linguistic Computing*, 22: 405–17. <https://doi.org/10.1093/llc/fqm023>.
- Hirst, Graeme, and Vanessa Wei Feng. 2012. "Changes in style in authors with Alzheimer's disease." *English Studies*, 93: 357–70.
- Holmes, David I. 1992. "A stylometric analysis of mormon scripture and related texts." *Journal of the Royal Statistical Society. Series A (Statistics in Society)*, 155: 91–120. <https://doi.org/10.2307/2982671>.
- Holmes, David I. 1994. "Authorship attribution." *Computers and the Humanities*, 28: 87–106. <https://doi.org/10.1007/BF01830689>
- Holmes, David I. 1998. "The evolution of stylometry in humanities scholarship." *Literary and Linguistic Computing*, 13: 111–17. <https://doi.org/10.1093/llc/13.3.111>.
- Hosmer, David W. Jr., Stanley Lemeshow, and Rodney X. Sturdivant. 2013. *Applied Logistic Regression*. Hoboken, NJ, USA: John Wiley & Sons.
- Jamak, Amra, Aleksandar Savatić, and Fazli Can. 2012. "Principal component analysis for authorship attribution." *Business Systems Research*, 3: 49–56. <https://doi.org/10.2478/v10305-012-0012-2>.
- Juola, Patrick. 2006. "Authorship attribution." *Foundations and Trends in Information Retrieval*, 1: 233–334. <https://doi.org/10.1561/1500000005>.
- Kešelj, Vlado, Fuchun Peng, Nick Cercone, and Calvin Thomas. 2003. "N-Gram-based Author profiles for authorship attribution." In *Proceedings of the conference Pacific Association for Computational Linguistics*, 255–64. Halifax, NS, Canada.
- Kohavi, Ron. 1995. "A study of cross-validation and bootstrap for accuracy estimation and model selection." In *IJCAI'95: Proceedings of the 14th international joint conference on Artificial intelligence*, 2, 1137–43.
- Koppel, Moshe, Shlomo Argamon, and Anat Rachel Shimoni. 2002. "Automatically categorizing written texts by author gender." *Literary and Linguistic Computing*, 17: 401–12. <https://doi.org/10.1093/llc/17.4.401>.
- Koppel, Moshe, Jonathan Schler, and Shlomo Argamon. 2009. "Computational methods in authorship attribution." *Journal of the American Society for Information Science and Technology*, 60: 9–26. <https://doi.org/10.1002/asi.20961>.

- La Greca, Fernando, and Liberato De Caro. L. 2017. "Nuovi studi sulla datazione della crocifissione nell'anno 34 e della nascita di Gesù il 25 dicembre dell'1 a.C.," *Annales Theologici*, 31, pp. 11–52. <https://www.annalestheologici.it/article/view/3416>.
- La Greca, Fernando, and Liberato De Caro. 2019. "La datazione della morte di Erode e l'inizio dell'era Cristiana." *Annales Theologici*, 33: 11–54. <https://www.annalestheologici.it/article/view/3457>.
- La Greca, Fernando, and Liberato De Caro. 2020. "Approfondimenti sulla nascita di Gesù nell'1 a.C. e sulla datazione della crocifissione nel 34." *Annales Theologici*, 34: 13–58. <https://www.annalestheologici.it/article/view/3480>
- La Greca, Fernando, Liberato De Caro, and Emilio Matricciani. 2024. "The "Galenic Question": A Solution Based on Historical Sources and a Mathematical Analysis of Texts." *Histories*, 14: 308–25. <https://doi.org/10.3390/histories4030015>.
- Landauer, Thomas K., Peter W. Foltz, and Danielle Laham. 1998. "An introduction to latent semantic analysis." *Discourse Processes*, 25: 259–84. <https://doi.org/10.1080/01638539809545028>.
- Le, Xuan, Ian Lancashire, Graeme Hirst, and Rina Jokel. 2011. "Longitudinal detection of dementia through lexical and syntactic changes in writing: A case study of three British novelists." *Literary and Linguistic Computing*, 26: 435–61.
- Le, Quoc V., and Tomas Mikolov. 2014. "Distributed representations of sentences and documents." In *Proceedings of International Conference on Machine Learning*, 1188–96. Beijing, China.
- Matricciani, Emilio, and Liberato De Caro. 2017a. "Finzione letteraria o antiche osservazioni astronomiche e meteorologiche nell'opera di Maria Valtorta?" *Scienza e Ricerche*, 44: 5–20. <https://www.calameo.com/books/003924817509d75f94026>.
- Matricciani, Emilio, and Liberato De Caro. 2017b. "Literary Fiction or Ancient Astronomical and Meteorological Observations in the Work of Maria Valtorta?" *Religions*, 8: 1–23. <https://www.mdpi.com/2077-1444/8/6/110>.
- Matricciani, Emilio, and Liberato De Caro. 2018. "A Mathematical Analysis of Maria Valtorta's Mystical Writings." *Religions*, 9: 1–23. <https://www.mdpi.com/2077-1444/9/11/373>.
- Matricciani, Emilio, and Liberato De Caro. 2020. "Jesus Christ's Speeches in Maria Valtorta's Mystical Writings: Setting, Topics, Duration and Deep-Language Mathematical Analysis." *J 3* (1): 100–23. <https://www.mdpi.com/2571-8800/3/1/10>.

- Matricciani, Emilio. 2022. "The Temporal Making of a Great Literary Corpus by a XX-Century Mystic: Statistics of Daily Words and Writing Time." *Open J. Statistics*, 12: 155–67.
- Pedregosa, Fabian, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, Jake Vanderplas, Alexandre Passos, David Cournapeau, Matthieu Brucher, Matthieu Perrot, and Édouard Duchesnay. 2011. "Scikit-learn: Machine Learning in Python." *Journal of Machine Learning Research*, 12: 2825–30.
- Pfenninger, Simone E., and Stefanie Polz. 2018. "Foreign language learning in the third age: A pilot feasibility study on cognitive, socio-affective and linguistic drivers and benefits in relation to previous bilingualism of the learner." *Journal of the European Second Language Association*, 2: 1–13.
- Powers, David M. W. 2011. "Evaluation: From Precision, Recall and F-Measure to ROC, Informedness, Markedness & Correlation." *Journal of Machine Learning Technologies*, 2: 37–63. <https://doi.org/10.48550/arXiv.2010.16061>.
- Ramnil, Hoshiladevi, Shireen Panchoo, and Sameerchand Pudaruth. 2016. "Authorship attribution using stylometry and machine learning techniques." In *Intelligent Systems Technologies and Applications*, edited by S. Berretti et al., Springer (Advances in Intelligent Systems and Computing, 384: 113–25). https://doi.org/10.1007/978-3-319-23036-8_10.
- Raza, Agha Ali, Awais Athar, and Sajid Nadeem. 2009. "N-Gram Based Authorship Attribution in Urdu Poetry." In *Proceedings of the Conference on Language & Technology 2009 (CLT09)*, Lahore, Pakistan, January 22–24, 2009, 88–93.
- Rosen-Zvi, Michal, Thomas L. Griffiths, Mark Steyvers, and Padhraic Smyth. 2004. "The author-topic model for authors and documents." In *Proceedings of 20th Conference on Uncertainty in Artificial Intelligence*, 487–94. Banff, Canada: AUAI Press.
- Salton, Gerard, and Christopher Buckley. 1988. "Term-Weighting Approaches in Automatic Text Retrieval." *Information Processing & Management* 24: 513–23. [https://doi.org/10.1016/0306-4573\(88\)90021-0](https://doi.org/10.1016/0306-4573(88)90021-0)
- Savoy, Jacques. 2013. "Authorship attribution based on a probabilistic topic model." *Information Processing & Management*, 49: 341–54. <https://doi.org/10.1016/j.ipm.2012.06.003>
- Sebastiani, Fabrizio. 2002. "Machine learning in automated text categorization." *ACM Computing Surveys*, 34: 1–47. <https://doi.org/10.1145/505282.505283>
- Seroussi, Yanir, Ingrid Zukerman, and Fabian Bohnert. 2011. "Authorship attribution with latent dirichlet allocation." In *Proceedings of Fifteenth Conference on Computational Natural Language Learning* edited by Sharon Goldwater

- and Christopher Manning, 181–89. Portland, OR, USA: Association for Computational Linguistics.
- Stamatatos, Efstathios. 2000. "A survey of modern authorship attribution methods." *Journal of the American Society for Information Science and Technology* 60: 538–56. <https://doi.org/10.1002/asi.21001>.
- Stamatatos, Efstathios, Nikos Fakotakis, and George Kokkinakis. 2000. "Text genre detection using common word frequencies." In *COLING 2000 Volume 2: The 18th International Conference on Computational Linguistics*, 808–814. Saarbrücken, Germany. <https://aclanthology.org/C00-2117>.
- Stamatatos, Efstathios. E. 2008. "Author identification: using text sampling to handle the class imbalance problem." *Information Processing & Management* 44: 790–99. <https://doi.org/10.1016/j.ipm.2007.05.012>.
- Tweedie, Fiona J., Sukhjot Singh, and David I. Holmes. 1996. "Neural network applications in stylometry: the Federalist Papers." *Computers and the Humanities* 30: 1–10. <https://doi.org/10.1007/BF00054024>.
- Valtorta, Maria. 2006a. *I Quaderni del 1943*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta, Maria. 2006b. *I Quaderni del 1944*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta, Maria. 2006c. *I Quaderni del 1944–1950*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta, Maria. 2006d. *I Quadernetti*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta Maria. 2007. *Il libro di Azaria*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta Maria. 2008. *Lezioni sull'Epistola di Paolo ai Romani*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta, Maria. 2009. *Autobiografia*. Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Valtorta Maria. 2015. *L'Evangelo come mi è stato rivelato*, Isola del Liri (Fr), Italy: Centro Editoriale Valtortiano.
- Zheng, Rong, Jiexun Li, Hsinchun Chen, and Zan Huang. 2006. "A framework for authorship identification of online messages: writing-style features and classification techniques." *Journal of the American Society for Information Science and Technology* 57: 378–93.