

Article

Machine Learning-Based Data Quality Assessment for the Textile and Clothing Digital Product Passport

Estrela Ferreira Cruz ^{1,2,*} , Pedro Silva ¹ , Sérgio Serra ¹ , Rodrigo Rodrigues ¹ , Marcelo Alves ³, João Oliveira ⁴ and António M. Rosado da Cruz ^{1,*} 

¹ ADiT-Lab-Instituto Politécnico de Viana do Castelo, 4900-348 Viana do Castelo, Portugal

² Algoritmi Research Centre, Escola de Engenharia, Universidade do Minho, 4800-058 Guimarães, Portugal

³ INFOS-Informática e Serviços S.A., Rua Veloso Salgado 971 1011, 4450-801 Leça da Palmeira, Portugal; marcelo.alves@infos.pt

⁴ CITEVE-Centro Tecnológico das Indústrias Têxtil e do Vestuário, Rua Fernando Mesquita 2785, 4760-034 Vila Nova de Famalicão, Portugal; joliveira@citeve.pt

* Correspondence: estrela.cruz@estg.ipvc.pt (E.F.C.); miguel.cruz@estg.ipvc.pt (A.M.R.d.C.)

Abstract

Transparency in business practices is essential for sustainability, ensuring that resources are used responsibly and that environmental and social impacts are properly measured and monitored, allowing the end consumer to make informed purchasing decisions without feeling cheated. The Digital Product Passport (DPP) promotes transparency by providing detailed information about a product's origin, composition, and life-cycle activities, enabling more sustainable and responsible choices. The implementation of the DPP for textile and clothing items faces many challenges due to the large number and diversity of companies involved in the value chain of these products, combined with the large amount and variability of information that needs to be collected. Therefore, the integration and standardization of data from these companies is one of the largest present challenges. In this article, we study the use of Machine Learning (ML) algorithms for validating, in a homogeneous way, the quality of the data submitted by each company for the implementation of the DPP. We have studied four solutions that, using datasets organized in different ways and using different ML algorithms, enable selecting the solution that best suits each particular situation.

Keywords: circular economy; data anomaly detection; data quality assessment; digital product passport; machine learning; sustainability; textile and clothing value chain; traceability



Academic Editor: Douglas O'Shaughnessy

Received: 27 July 2025

Revised: 7 September 2025

Accepted: 16 September 2025

Published: 20 September 2025

Citation: Cruz, E.F.; Silva, P.; Serra, S.; Rodrigues, R.; Alves, M.; Oliveira, J.; Rosado da Cruz, A.M. Machine Learning-Based Data Quality Assessment for the Textile and Clothing Digital Product Passport. *Appl. Sci.* **2025**, *15*, 10259. <https://doi.org/10.3390/app151810259>

Correction Statement: This article has been republished with a minor change. The change does not affect the scientific content of the article and further details are available within the backmatter of the website version of this article.

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

The Textile and Clothing (T&C) industrial sector has come under considerable scrutiny over the years. This results from the poor working conditions observed in many factories, the child labor it has been accused of in the past, and the realization of its enormous environmental footprint. The textile industry has been accused of excessive water consumption for the production of raw materials associated with excessive soil exploitation, as well as the pollution of rivers by chemicals used in its manufacturing processes. Currently, the main problem is the exaggerated increase in textile waste, a consequence of excessive consumption encouraged by fast fashion, which stimulates major consumption and the rapid disposal of clothes [1]. Furthermore, the T&C sector is one of the major contributors to the increase in greenhouse gases (GHGs) in the atmosphere [1]. The textile industry needs

to be adapted and “reinvented” to reduce its environmental and social impacts, becoming more sustainable. Furthermore, it is necessary to increase transparency at every stage of the life-cycle of these products by implementing the Digital Product Passport (DPP).

According to the European Union (EU), the “DPP is the combination of an identifier, the granularity of which can vary throughout the life-cycle (from a batch to a single product), and data characterizing the product, processes and stakeholders, collected and used by all stakeholders involved in the circularity process” [2]. To implement the DPP in the T&C sector, it is necessary to collect information about every stage and all the participants in the value chain and integrate this information. There are several categories of stakeholders in the T&C DPP, namely supply chain companies, brands, authorities, certification and assessment companies, retailers, marketing companies, consumers, and circularity operators [2].

The EU has already created rules that companies wishing to trade in Europe must follow. The DPP must contain information about composition, product description (size, weight, etc.), the supply chain, transportation, social and environmental impacts, circularity, the brand, usage and maintenance of garments (customer feedback), tracking and tracing after sales, etc. [2]. In other words, to implement the DPP, it is necessary to collect a wide variety of information involving a large number of companies and even individual consumers. Thus, the collection of this information may be subject to errors. Therefore, it is necessary to create a way to uniformly evaluate this data before it is stored and integrated into the DPP.

The strong recent investments in the area of Artificial intelligence (AI), and more particularly in its sub-field of Machine Learning (ML), have led to great and surprising advances, with a very positive impact on many areas of society. ML, with its ability to learn and make decisions based on data from previous experience, has been adopted in several business areas that touch our daily lives, including image recognition, Natural Language Processing, fraud analysis, and many others. The detection of non-standard or anomalous values is one of the areas in which ML demonstrates good results [3–5].

In the literature review presented in Section 3, a clear gap in data quality assurance for the T&C DPP has been identified. In this article, the main objective is to study and compare ML-based approaches for data anomaly detection using real data from the T&C value chain and integrate the results into an API that will be used to validate the data before integration into the DPP platform. We will emphasize information about the value chain and the collection of data that, in addition to enabling a product’s traceability, enables calculating the environmental and social impacts of the product. Therefore, it is necessary to receive information from and about all the companies involved in the manufacturing of a given garment, including transportation between the various companies. There are many companies involved, potentially with very different levels of digital maturity, which is why it is necessary to standardize the data collected and the validation process that it must be subjected to.

The data anomaly detection solutions studied in this article, which use ML algorithms, may be employed to validate the data collected by the various business partners involved in the T&C value chain for the implementation of the DPP. More precisely, we started by collecting data on sustainability indicators from several companies involved in different production activities. Next, we worked with this data to prepare different datasets so that different types of ML algorithms (supervised learning and unsupervised learning) could be trained. Finally, we compared the approaches and analyzed which one best suits each situation. Again, the main goal is to study and compare ML-based approaches for data anomaly detection that can be used in the future for validating data before integration into the DPP platform. The goal is not to propose a final ready-to-use implementation but to

study the applicability of ML-based anomaly detection approaches to validate real textile data from the T&C value chain and provide access through an API that may be used for data quality assessment before integration with the DPP.

The remainder of this paper is structured as follows: In Section 2, some concepts related to AI and ML, including some ML algorithms, are presented. In Section 3, the related works previously developed in the area are analyzed, including previous works developed by the same authors. In Section 4, the methodology used for this research is presented. Section 5 presents the dataset (Section 5.1) used to train the algorithms, as well as the studied solutions (Section 5.2). For each solution, its objective and design are presented, together with a discussion of the results obtained. Section 5.3 explains how the validation solutions may be made available to external applications through an integration Application Programming Interface (API). In Section 6, a discussion regarding the results is presented. Finally, Section 7 presents the conclusions and outlines future work.

2. Machine Learning Models for Data Validation

Artificial Intelligence is a field of science that focuses on creating machines that, either rule-based or data-driven, can act by taking actions that normally require human intelligence. AI already has a long history. Since the creation of the Turing Machine by Alan Turing in 1936, which demonstrated that any computational problem could be solved by a machine with well-defined rules, AI has undergone great advances. In recent years, and especially in its sub-field of Machine Learning, the developments have been extraordinary. Nowadays, AI (particularly ML) is present in all areas of society, from finance [6] to medical diagnostics [7], image detection and facial recognition [8], cyber security [9], fraud detection [10], Natural Language Processing (ChatGPT, Siri, etc.) [11], machine translation (Google Translate, etc.), industrial automation [12], autonomous vehicles [13], etc.

ML is a branch of AI that enables machines to make decisions based on prior experiences. It achieves this by utilizing algorithms and models that facilitate learning from data stored in datasets and subsequently making intelligent decisions autonomously. ML algorithms employ statistical techniques to identify patterns and anomalies within often large datasets, allowing machines to learn from the data. The dataset size can be so extensive that it becomes challenging for humans to process all the information, which often enables ML to offer superior solutions to problems compared to humans [14].

Since models learn from data, the data used is very important and needs to be properly prepared. There are many algorithms that can be used. Based on how they learn from data, algorithms can be categorized into four main types: supervised learning, unsupervised learning, reinforcement learning and deep learning [15–17].

- **Supervised Learning**—This category includes algorithms that learn from labeled data. It comprises classification algorithms, which categorize data into predefined classes, and may be binary (e.g., yes/no; correct/incorrect) or multiclass; and regression algorithms, which predict continuous values, such as the amount of energy consumed. The most common classification algorithms are Naive Bayes (NB), Support Vector Machine (SVM), Decision Tree (DT), Logistic Regression, K-Nearest Neighbor (KNN), Random Forest (RF), and Multilayer Perceptron. The most common regression algorithms are SVM, DT, Polynomial Regression, Linear Regression, KNN, and RF.
- **Unsupervised Learning**—These algorithms learn from unlabeled data and are useful for discovering hidden patterns or structures. The common tasks include clustering, frequent pattern mining, and dimensionality reduction. This group includes algorithms such as K-means, which groups data into K clusters; Hierarchical Clustering, which creates a hierarchical structure of clusters; DBSCAN, which identifies dense regions of data items and is well-suited for unstructured data; and Isolation Forest,

a method that explicitly isolates anomalies and is one of the most popular anomaly detection methods [18,19].

- **Reinforcement Learning**—These algorithms learn optimal behaviors through interactions with an environment, using rewards and penalties as feedback. They are commonly applied in areas such as robotics and game playing. The popular algorithms in this category are Q-Learning, a table-based algorithm for finding the best action, Deep Q-Network (DQN), and Proximal Policy Optimization (PPO).
- **Deep Learning**—These models use artificial neural networks with multiple layers to learn complex representations of data. So, for these algorithms to have good results, very large datasets are required. In this group, there are algorithms like Artificial Neural Network (ANN), Convolutional Neural Network (CNN), Recurrent Neural Network (RNN), and Pre-Trained Language Models (PLMs) such as Bidirectional Encoder Representations from Transformers (BERT) and Generative Pre-Trained Transformer (GPT).

Between the supervised and unsupervised learning categories, a semi-supervised learning category can be considered, which uses a small amount of labeled data together with a large amount of unlabeled data to train ML models. The goal is to improve learning accuracy when labeled data is scarce [20].

The best algorithm to be used in any specific case largely depends on the objectives and the training dataset.

ML models can be evaluated in several aspects. For classification models, the typical evaluation metrics include the following [21]:

- **Accuracy**—measures the proportion of correct predictions made by a model relative to the total.
- **Precision**—measures how many of the items classified as positive are actually positive.
- **Recall**—measures how many of the positive items were correctly identified.
- **F1-Score**—an evaluation metric that combines precision and recall ($2 \times (\text{precision} \times \text{recall}) / (\text{precision} + \text{recall})$).

For regression models, the typical evaluation metrics include the following:

- **Mean Absolute Error (MAE)**—average absolute difference between predicted and actual values.
- **Mean Squared Error (MSE)**—average of squared differences.
- **Root Mean Squared Error (RMSE)**—square root of MSE.
- **R^2 Score**—also known as Coefficient of Determination. It measures the proportion of variance in the target explained by the model.

3. Literature Review

This section addresses the related work on data validation approaches in the T&C value chain, the complexity of the T&C value chain, and also other works in the area of data anomaly detection.

3.1. Textile and Clothing Value Chain

As research and innovation advance, the variety of raw materials continues to grow. In addition to traditional natural fibers such as wool, linen, silk, and cotton, there are now synthetic materials derived from petroleum, including polyester, acrylic, nylon, plastic, elastane, and lycra. Emerging alternatives also include plant-based materials like jute, bamboo, hemp, and banana leaves. More recently, with the growing emphasis on the circular economy, we can find raw materials such as recycled cotton, recycled polyester, and recycled polyester resin. Other materials are also commonly used in textile applications,

such as coconut, wood, metal, seashell, leather, and vegetable ivory [22]. Within each type of raw material, further sub-classifications may exist, such as organic silk, organic cotton, organic hemp, Merino wool. The treatment and processing of these diverse raw materials typically involve different specialized companies. As a result, the number of companies involved in the T&C value chain is increasing, becoming more diversified and spread throughout the world [23]. Considering, for example, the production of a shirt, with a composition of 50% silk + 50% cotton, with buttons, this will involve a sequence of several value-chain activities: silk production, cotton production, transportation, silk spinning, cotton spinning, dyeing, weaving, and garment manufacturing (including cutting, sewing, and assembly). In addition, the buttons must also be produced, transported, and applied.

The size and complexity of the T&C value chain is very demanding on the quality of intermediate and final products [24–26], and also on the quality of the industrial processes and of the waste produced [27]. Thus, in [24], the authors proposed a methodology that leverages Decision Trees and Naïve Bayes classifiers to predict defect types based on input material characteristics and technical sewing parameters. Since quality influences both the product's lifespan and its suitability for recycling, the authors in [25] presented a study aimed at developing a method for sorting and assessing the quality of textiles in household waste, with the evaluation based on product type, manufacturing method, and fiber composition. In [26], the authors presented a state-of-the-art review on anomaly detection in wool carpets using unsupervised detection models. In [27], the authors assessed the quality of the waste produced and examined its potential for use as raw material in the production of new fibers. While several studies have focused on validating product quality, research addressing the quality of the data collected remains scarce.

3.2. Solutions for Data Quality Assessment in the Scope of the DPP

The DPP for T&C products is currently being widely discussed as a way to increase product transparency and traceability [28].

In the literature, discussions on quality or anomaly detection in relation to textiles and clothing typically focus on aspects such as product quality, as mentioned in the previous subsection, rather than on the quality of the data collected. Consequently, research is lacking regarding data quality assessment or data anomaly detection in terms of the metrics gathered for the DPP for T&C products. The few previously proposed solutions for validating data collected for implementing the DPP are presented next.

The DPP is part of the European Green Deal and the European Circular Economy Action Plan. The study presented in [29] analyzed the possibilities and limitations in implementing the DPP and in its use for other fields of application.

With the aim of improving long-term sustainability, the authors in [30] proposed a model to extend the useful lifespan of a product and close the product life-cycle loop, contributing to the circular economy. The cyclical model includes reverse logistics of components and raw materials. The model also includes some information on how to manage data at the end of each life cycle.

In [31], the authors conducted a study on the potential benefits of implementing the DPP in the manufacturing industry and in particular in the electronics industry. The authors concluded that the implementation of the DPP and the collection of information along the value chain and manufacturing processes contribute to increased transparency in the manufacturing industry and increased sustainability. Furthermore, the information collected can be very beneficial for resource management and subsequent decision-making [31].

3.3. Solutions for Data Quality Assessment Using ML

With the constant increase in the variety and volume of data to be analyzed, the traditional strategies for detecting data anomalies and outliers are no longer efficient. This necessitates new strategies, such as the use of ML algorithms for anomaly detection and assessing data quality.

In [32], the authors studied existing approaches for network anomaly detection. The authors concluded that the existing approaches are not effective, mainly due to the large volumes of data involved, collected in real time through connected devices. To improve efficacy, the authors studied the use of ML algorithms for anomaly detection within a large volume of data collected in real time.

In [5], the authors performed a systematic literature review on ML models used in anomaly detection and their applications. The study analyzed ML models from four perspectives: the types of ML techniques used, the applications of anomaly detection, the performance metrics for ML models, and the type of anomaly detection (supervised, semi-supervised, or unsupervised). The authors concluded that unsupervised anomaly detection has been adopted by more researchers than the other classification anomaly detection systems. Semi-supervised is practically not used. The authors also concluded that a large portion of the research projects employed a combination of several ML techniques to obtain better results and that SVM is the most frequently used algorithm, either alone or combined with others [5].

Some approaches are more specific and study the use of ML in detecting anomalies in specific areas, such as data collected by IoT devices [3,33], water quality assessment [34], and other areas.

Nowadays, a great deal of information is collected through the use of IoT devices. These devices continuously collect large volumes of data, which can make detecting anomalous data challenging. In [33], the authors reviewed the literature on the main problems and challenges in IoT data and regarding the ML techniques used to solve such problems. SVM, KNN, Naive Bayes and Decision Tree were the supervised classification algorithms mentioned most often [33]. The unsupervised Machine Learning algorithms most frequently used were K-means clustering and the Gaussian mixture model (GMM).

Uddin et al. [35] also employed ML algorithms to assess water quality. The authors used ML algorithms to identify the best classifier to predict water quality classes. The ML algorithms used were SVM, Naïve Bayes (NB), RF, KNN, and gradient boosting (XGBoost) [35].

Zhu et al. carried out a review on the use of ML in water quality assessment [34]. The increasing volume and diversity of data collected on water quality has made traditional data analysis more difficult and time-consuming, which has driven the use of ML for water quality analysis. ML algorithms are used in monitoring, simulation, evaluation, and optimization of various water treatment and management systems and in different aquatic environments (sewage, drinking water, marine, etc.). The authors evaluated the performance of 45 ML algorithms, including SVM, RF, ANN, DT, Principal Component Analysis (PCA), XGBoost, etc. [34]. The authors concluded that conditions in real systems can be extremely complex, which makes the widespread application of Machine Learning approaches very difficult. The authors suggested improving the algorithms and models so that they are more universal according to the needs of each system [34].

Eduardo Nunes, in 2022, presented a literature review on the use of ML techniques in anomaly detection for the specific case of “Smart Shirt”, that is, shirts that have the ability to monitor and collect signals emitted by our body. These “Smart Shirts” are typically used for medical purposes. The author, among other things, studied the ML techniques that

are being used for anomaly detection and concluded that the SVM algorithm is the most widely used, followed by K-Nearest Neighbor (KNN) and Naive Bayes (NB) [3].

3.4. Previous Work by the Same Authors

In a previous study [36] conducted by the same authors, within the same project, a non-AI rule-based approach was described for validating data collected by each business partner in the T&C value chain prior to its integration into a traceability platform for the DPP.

The approach introduced four formulas for validating sustainability indicator values [36]. These formulas assess whether a given value lies within a specified range, defined by minimum and maximum thresholds, and evaluate its proximity to a central or average value. This central value, reflecting the central tendency of the indicator's distribution, can be measured using the arithmetic mean, median, or mode. The arithmetic mean represents the average of all values, the median is the middle value when ordered, and the mode is the most frequently occurring value. The values closer to the central tendency are assigned higher validity, whereas those further away are assigned lower validity. The rate at which validity decreases with increasing deviation from the mean can vary depending on the metric. In cases where a value falls outside the defined minimum or maximum, it may be classified as invalid or still acceptable depending on the metric. The parameters for these validations can be adjusted for each specific metric [36].

The study also suggested that future work could explore the application of Machine Learning algorithms to enable more dynamic and accurate validation. This suggestion was implemented and presented in [37]. In [37], the ML algorithms RF, DT, SVM, and KNN were studied as approaches to enable data anomaly detection in a consistent and homogeneous manner throughout the value chain.

The approach presented here builds upon the work in [37] by reorganizing the datasets, exploring new Machine Learning algorithms, and comparing the resulting performance.

4. Methodology

As seen before, although product quality validation is firmly established in the literature, the validation of collected data remains critically overlooked and underexplored. For this research, the Design Science Research (DSR) method was used as it adapts to the needs of practical and useful research in the scientific areas of Information Systems/Information Technology (IS/IT) [38]. The DSR method drives research forward through a series of cycles in which an artifact is produced and evaluated. The evaluation of the produced artifact can identify strengths and weaknesses for improvement, as well as generate new ideas that can be implemented in the next cycle [38,39]. The DSR method involves six main steps [38,40]:

1. **Problem Identification and Motivation**—As discussed in Section 1, the volume and complexity of data collected for the implementation of the DPP demand the development of new strategies for effective data validation. To address this, it is essential to establish a standardized and flexible mechanism for evaluating data consistently across all companies involved in the T&C value chain.
2. **Definition of Objectives**—The objective is to develop an artifact that leverages ML techniques to validate incoming data for the implementation of the DPP while ensuring seamless integration into the broader DPP platform.
3. **Artifact Design and Development**—The created artifacts include a set of Machine Learning models for anomaly detection in data to be integrated in the DPP and an API that provides a set of services based on those models designed to be easily accessible and usable by all participants in the value chain for validating the data collected for DPP implementation. The created ML models are presented in Section 5.2.

4. **Demonstration**—The demonstration involves creating and preparing appropriate datasets for each proposed approach, with which the selected ML algorithms are trained and tested. For this, 80% of the data has been used for training and 20% for testing purposes. When the results obtained had low performance values, or to confirm the first results, we also used K-fold cross-validation in an attempt to circumvent the fact that the used datasets were small. The results are evaluated, as detailed in Section 5.2.
5. **Evaluation**—The created models are evaluated iteratively through quantitative analysis of the results, as presented in Section 5.2. Ultimately, the validation services of the API are tested, with the first data items being integrated into the general DPP platform.
6. **Communication**—Communication and dissemination of results occur through the publication of scientific articles. The findings from the first iteration were presented in [36], the second iteration in [37], and the final results are presented in the current article.

The created artifacts (ML models and API) have several limitations, which are addressed in Section 6, but they are still useful and innovative and may serve as a basis for a data anomaly detection solution when integrated with the DPP data from companies involved in the T&C value chain.

5. Results

This section presents the research carried out towards a solution for anomaly detection in textile data before integrating it into the DPP. The section begins by introducing real data obtained from companies to construct the datasets for training the ML models, followed by four approaches that employ different Machine Learning algorithms for data anomaly detection. Finally, the API for data validation is presented.

5.1. T&C Value Chain Dataset Preparation

Preparing a dataset is one of the most critical stages in a Machine Learning (ML) project. The goal of data preparation is to format and structure the data in a way that aligns with the requirements of the Machine Learning algorithms to be used [41]. As outlined in [41], dataset preparation typically follows a series of steps, beginning with data collection and followed by data cleaning, transformation, and reduction:

- **Data Collection**—This is the initial phase of any Machine Learning pipeline. It involves identifying, selecting, and acquiring the appropriate data needed for the specific algorithm and intended outcomes.
- **Data Cleaning**—This step involves identifying and correcting errors, such as missing values, noisy data, and anomalies, to ensure the dataset is accurate and reliable.
- **Data Transformation**—This process involves converting raw data into a suitable format for modeling. Common tasks include normalization (scaling numerical data to a specific range) and discretization (converting continuous variables into discrete buckets or categories). These transformations are essential for improving algorithm performance and interpretability.
- **Data Reduction**—This process is used in situations where datasets are too large or complex. In these cases, reducing data may be of aid to deal with issues like overfitting, computational inefficiency, and difficulty in interpreting models with numerous features and is crucial to enhance the efficiency and effectiveness of Machine Learning analysis.

These steps, mentioned in [41], have been followed to create and prepare the datasets used in the study presented here. Depending on the ML approach followed, the dataset is prepared in different ways; however, the steps presented below are common to all datasets.

5.1.1. Data Collection

The data used for this study has been provided by Portugal's Textile and Clothing Technological Center (CITEVE) and consists of real data collected from various companies representing different types of production activities involved in the textile and clothing sector, such as spinning, weaving, dyeing, and manufacturing, among others.

The data was provided through multiple Excel files, each composed of tables for a product batch/lot resulting from a production activity, where the batch identification and size were listed, along with environmental impact metrics and their respective values.

To use this information in the ML training phase, it was necessary to consolidate all the data in a structured format. Thus, a single table has been created where each row contains information about the batch size, production activity, responsible firm or organization, materials used to produce it (e.g., cotton, linen, or polyester) and the percentages of each, and the corresponding metrics information: name, unit of measurement, and value.

The information provided was simple to understand, so there was no need to group attributes to reduce their volume. Furthermore, each of the various parameters was labeled, aiding in their comprehension. Although the data supplied was of good quality and featured real information from different types of companies, the total number of records is small. Due to its complexity and reliability, it was not possible to use data augmentation methods to address the issue of a small number of records. At the end of this stage, the dataset consisted of 223 rows.

Each row contained information about the size of the product batch, the production activity that created it, the materials used for its creation, and their percentages and the name of an attribute, the unit of measurement, and their value.

In total, we received information about 28 sustainability indicators that can be used to measure the environmental impact of the batch produced by the identified production activities, i.e., regarding the transformation of the input material/product (or products) into the output product.

For a given batch, data was collected on environmental sustainability indicators related to the resources consumed and outputs produced during a specific production activity, aimed at transforming and manufacturing a product of a defined weight. The following environmental impact metrics were analyzed regarding consumption: consumption of water, non-renewable electricity, renewable electricity, self-consumption of renewable energy, biomass, coal, total electricity, natural gas, propane gas, fuel oil, water, gasoline, diesel, purchased steam, chemical products, etc. The following environmental impact metrics were analyzed regarding production: quantity of SVHC (Substance of Very High Concern) in products, quantity of non-hazardous chemicals, quantity of recovered chemicals, quantity of non-hazardous waste, quantity of solid waste, quantity of textile waste, quantity of recovered waste, quantity of recovered textiles, quantity of recycled water, and volume of liquid effluent, among others. Some environmental impact metrics make more sense for some production activities and less, or no sense at all, for others.

After collecting the data, initial data processing was carried out. These treatments were necessary for all the approaches that will be presented in Section 5.2. However, some of the approaches presented there also required additional treatment.

5.1.2. Data Cleaning

In the case of the received data, there were no missing values or noisy data, so it was not necessary to implement replacement techniques. There was also no need for outlier treatment since all data represented real data obtained from industrial activities.

5.1.3. Data Transformation

The original data received underwent some transformations. The categorical data were encoded using label encoding, and all values were converted to numerical format to ensure greater compatibility with ML models. Proper normalization was also applied to the numerical data to ensure consistency between the different value scales to avoid loss of information.

5.1.4. Data Reduction

In the case presented here, the amount of data obtained was not too large, so this step was not performed.

After this first data processing, the dataset consisted of 223 lines. Each row contained information on the product identification and size of the product batch produced, the production activity in which the batch was produced, the materials used in its creation and corresponding percentages, the name of the metric to be evaluated, and its measurement unit and its value.

Given that we are evaluating a large number of production activities and a large number of environmental impact indicators, it is natural that the most appropriate value validation solution varies from one indicator to another. It may depend on the production activity and the indicator being evaluated.

5.2. Proposed ML-Based Solutions for Data Validation in the Context of the T&C DPP

Recognizing that data quality issues can manifest in various forms, in this section, we study four ML-based approaches for data anomaly detection and discuss their results. This multi-faceted strategy allows us to address the problem from different angles, from detecting unusual individual data points to assessing the collective consistency of an entire record. The rationale for exploring different model types for anomaly detection, namely unsupervised, supervised classification, and supervised regression, is to provide a comprehensive and flexible validation toolkit.

Using the same set of real-world data collected from companies in the T&C sector (as detailed in Section 5.1), each approach involves distinct data preparation strategies and uses different ML algorithms for data validation and anomaly detection. An anomaly is defined as an observation that significantly deviates from the expected behavior or pattern within a dataset. Also referred to as an outlier or deviation, an anomaly represents an unusual or rare value among a collection of observed data points. Anomalies can generally be classified into three main categories [5,32]:

- Point Anomaly: A single value or data point that is anomalous compared to the rest of the data.
- Contextual Anomaly: A data point that is considered anomalous only within a specific context (e.g., time and location).
- Collective Anomaly: A group of related data points that together form an anomalous pattern, even if individual points may not appear abnormal on their own.

The four proposed approaches are presented next:

- Approach 1—The first approach uses unsupervised anomaly detection models combined with predefined threshold values to individually assess the reliability of each environmental impact metric.
- Approach 2—The second approach uses supervised learning models to both predict values for individual metrics (regression) and to classify whether a received value is correct or not (classification).
- Approach 3—The third approach focuses on analyzing the relationships among the different environmental impact metrics. All the metric values are considered together

within a single record (i.e., one record includes all the metrics), and the goal is to determine whether the values within a given record are consistent with each other. Supervised learning algorithms are used to perform this consistency check.

- Approach 4—In the fourth approach, the goal is to predict the value of a specific metric based on the remaining metrics in the record. The objective is to assess whether the value received for a particular metric in a new record is plausible given the values of the other metrics. Supervised models are used for this purpose.

As the dataset is relatively small, neural network-based algorithms were not considered due to their typically high data requirements.

5.2.1. Approach 1—Unsupervised Anomaly Detection Based on Individual Metrics

This approach focuses on validating the data of a certain batch according to an attribute (metric) using ML anomaly detection algorithms. Each record includes the batch size, production activity, and the name and value of the attribute to be validated. Given the objective of assessing and ensuring the quality of the data, anomaly detection provides an effective way to identify inconsistencies and deviations that may indicate errors or irregularities.

Data Preparation

This approach uses the dataset presented previously in Section 5.1. The dataset consisted of 223 rows, each with 11 columns.

Only the organization's column was removed from the initial dataset. Of these eleven columns, one contains the size of the lot, one specifies the production activity, two indicate the types of materials used in the lot, and two record the corresponding material percentages. Another column contains the name of the metric, followed by its unit of measurement and its recorded value. The last two columns contain information about the data's validity. One includes a validity score (ranging from 0 to 100), and the other assigns a validity category (valid, suspect, or invalid). For ML purposes, the dataset was preprocessed to handle missing and noisy data, and all non-numeric values were converted to numeric representation.

After analyzing the created dataset, a key issue was identified concerning the validity category of the data: there were no data records categorized as "suspect". This posed a challenge for training the models to recognize suspect data, which is important for real-world applications where such cases may occur. To solve this issue, additional data was generated based on the existing records and using the API previously developed and presented in [36]. The API uses formulas to evaluate the validity of the data. By using existing data and modifying the values of the attributes sufficiently for the API's formulas to classify them as "suspect", it was possible to create new data entries. To preserve the integrity and trustworthiness of the dataset, only 19 new rows were created using this method. At the end, the dataset had 242 records.

ML Algorithms

The models chosen for this approach were selected for their ability to detect anomalies in both large and small datasets. This made them suitable for the dataset used here and for potential future datasets. Other selection criteria included scalability, widespread adoption, distinct anomaly detection mechanisms, and their established use in fields such as environmental science and industrial systems [4,42]. The approach presented in this subsection uses the algorithms Isolation Forest (IF), Local Outlier Factor (LOF), One-Class SVM (OCSVM), Gaussian Mixture Model (GMM), and Robust Covariance (RC), the latter employing Elliptic Envelope (EE) for covariance estimation.

The IF (Isolation Forest) is an anomaly detection algorithm that identifies outliers by isolating data points that differ significantly from the rest. It operates by randomly selecting features and recursively splitting the data. Anomalies are easily isolated because they are fewer in number and more distinct from the bulk of the data [18].

The LOF operates by contrasting the local density of a point with the local density of k of its surrounding neighbors. If this point has significantly lower density than its surrounding neighbors, it is considered an outlier [43].

The OCSVM works by creating a boundary that encloses the majority of the data points while maximizing the margin around them. Any data point that falls outside this boundary is considered an anomaly [44].

GMM models the data as a combination of multiple Gaussian distributions, assigning each data point a likelihood of belonging to each component distribution. Outliers are detected by identifying points with a low probability of belonging to any of the distributions [45].

The RC method using Elliptic Envelope assumes that the data follows a Gaussian distribution, forming an elliptical shape. Data points that fall outside this elliptical boundary are classified as outliers [46].

Because these are unsupervised models, their effectiveness does not rely on traditional classification metrics. Instead, they are often interpreted through visualizations such as plots. In this work, however, anomaly detection is treated as a classification task. Each record was classified as “valid”, “invalid”, or “suspect” based on its anomaly score, as predicted by the models, and according to the thresholds defined for classification.

These thresholds were determined based on the proportion of data points in each validity category (29.3% invalid, 7.9% suspect, and 62.8% valid; see Figure 1). This technique was feasible because the dataset had each row labeled according to validity (valid, suspect, or invalid).

Distribution of Value Records Validity Categories

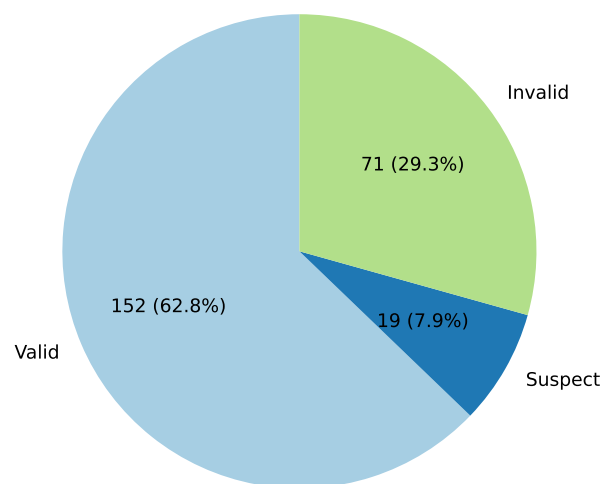


Figure 1. Distribution of data point classifications based on data validity category.

For each model, an anomaly score s_i has been computed for each record i . The anomaly score corresponds to

- The log-likelihood of each record for GMM;
- The negative of `negative_outlier_factor`, for LOF;
- The output of `decision_function` for other algorithms.

From the set of all scores

$$S = \{s_1, s_2, \dots, s_n\},$$

two percentile-based thresholds are calculated:

$$T_{\text{invalid}} = Q_{p_1}(S), \quad T_{\text{valid}} = Q_{p_2}(S),$$

where $Q_p(S)$ is the empirical p -th percentile of S . In the dataset used for this study, $p_1 = 29.3\%$ and $p_2 = 37.2\%$, chosen to match the observed proportions of invalid (29.3%), suspect (7.9%), and valid (62.8%) records (see Figure 1).

The validity category is then assigned as

$$C_i = \begin{cases} \text{Invalid,} & \text{if } s_i \leq T_{\text{invalid}}, \\ \text{Valid,} & \text{if } s_i \geq T_{\text{valid}}, \\ \text{Suspect,} & \text{otherwise.} \end{cases}$$

These threshold values were selected for this initial experiment because the dataset contained a small number of records (242). In the future, with a larger dataset, thresholds can be optimized in a more robust data-driven way, for example by selecting T_{invalid} and T_{valid} to maximize a chosen evaluation metric, such as macro F1-score, through K-fold cross-validation over the labeled training set. This will be able to yield category boundaries more closely aligned with the model's performance and the underlying score distribution.

The hyperparameters used in this approach were mostly the default ones for the used library, namely Scikit-learn (<https://scikit-learn.org/> accessed on 3 June 2025). Table 1 summarizes the hyperparameters used in each model.

Table 1. Hyperparameters and values used for the anomaly detection models in Approach 1.

Model	Hyperparameters and Values
Isolation Forest	contamination = ‘auto’, random_state = 42
Local Outlier Factor	novelty = True (defaults: n_neighbors = 20, algorithm = ‘auto’, leaf_size = 30, metric = ‘minkowski’, p = 2, contamination = 0.1)
One-Class SVM	kernel = ‘rbf’, gamma = ‘scale’ (default: nu = 0.5)
Gaussian Mixture	n_components = 10, random_state = 42 (defaults: covariance_type = ‘full’, max_iter = 100)
Elliptic Envelope/Robust Covariance	support_fraction = 1 (default: assume_centered = False)

For LOF, setting novelty = true enables predicting on new unseen data, which is necessary for anomaly detection outside the training set. For Isolation Forest, ‘auto’ for contamination lets the model estimate the proportion of outliers automatically.

This unsupervised approach incorporates supervised evaluation, employing labeled validity categories to assess model performance using traditional performance metrics. In this case, the accuracy, precision, recall, and F1-score of each model were measured and confusion matrices were generated to assess their performance.

Results

Figure 2 presents a visual representation of the output of the five anomaly detection models. In the plots, each point represents a batch, with the x-axis indicating the row index and the y-axis representing the quantity of products in each batch. The points are color-coded according to their validity category.

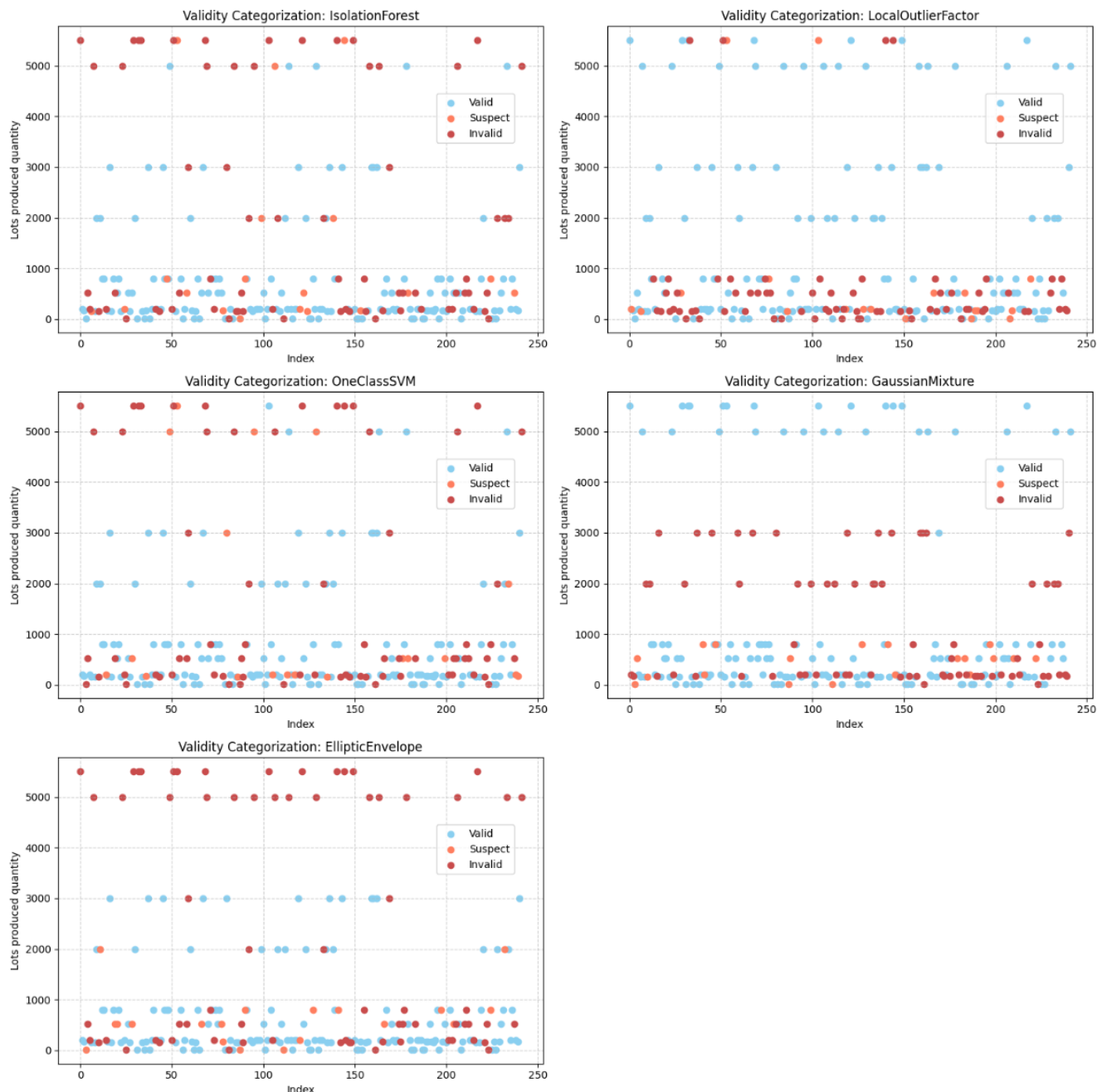


Figure 2. Comparative visualization of data validity categorization using multiple unsupervised anomaly detection models.

By analyzing Figure 2, a noticeable pattern emerges. Indeed, suspect and invalid data points seem to be more frequently associated with lower product quantities, while there is significant variation in the detection results across the different models, indicating inconsistencies in how anomalies are identified. This diversity underscores the need for additional evaluation methods to determine which models are most suitable for the specific characteristics and requirements of this use case.

As previously mentioned, we evaluated the performance of each anomaly detection model using traditional classification metrics. Table 2 summarizes the performance of each model. According to the results, the best-performing model was LOF with an accuracy of 52.48%, followed by OCSVM with 46.69%, IF with 45.04%, GMM with 44.63%, and RC using EE with 41.74%.

Table 2. Performance of anomaly detection ML models.

ML Models	PM	Results
IsolationForest	Accuracy	45.04%
	F1-score	29.69%
	Precision	29.69%
	Recall	29.69%
LocalOutlierFactor	Accuracy	52.48%
	F1-score	34.61%
	Precision	34.61%
	Recall	34.61%
OneClassSVM	Accuracy	46.69%
	F1-score	30.82%
	Precision	30.82%
	Recall	30.82%
GaussianMixture	Accuracy	44.63%
	F1-score	29.22%
	Precision	29.22%
	Recall	29.22%
RC with EllipticEnvelope	Accuracy	41.74%
	F1-score	25.90%
	Precision	25.90%
	Recall	25.90%

These findings suggest that, among the models tested, LOF was relatively more effective at identifying anomalous data. However, even its performance was modest, indicating limitations across all the models. This could be due to the limited size and nature of the dataset as anomaly detection models generally require large volumes of data to perform well. Additionally, more advanced data processing techniques may be needed beyond basic data normalization. Another important factor to consider is the lack of parameter tuning as the standard default values were used for all the models. These observations indicate that there is still significant room for improvement in both the data preparation and modeling stages in order to enhance the models' performance.

To enhance model performance and better understand the impact of dimensionality reduction, PCA was applied during the data transformation stage. PCA is a dimensionality reduction technique, which transforms the original features into a new set of uncorrelated linear combinations, known as principal components, which capture the maximum variance present in the data while preserving the most significant patterns and structures. This reduces the amount of redundant or less-informative features, improving the performance of unsupervised algorithms and their ability to detect anomalies [47].

Before applying PCA, it is essential to determine the optimal number of components to reduce dimensionality while retaining most of the dataset's variance. Figure 3 illustrates the cumulative explained variance as a function of the number of principal components. The red dashed line indicates the 95% variance threshold, while the green line marks the minimum number of components required to capture at least 95% of the total variance. This

ensures that dimensionality is effectively reduced without significant loss of information. In this case, six principal components were selected.

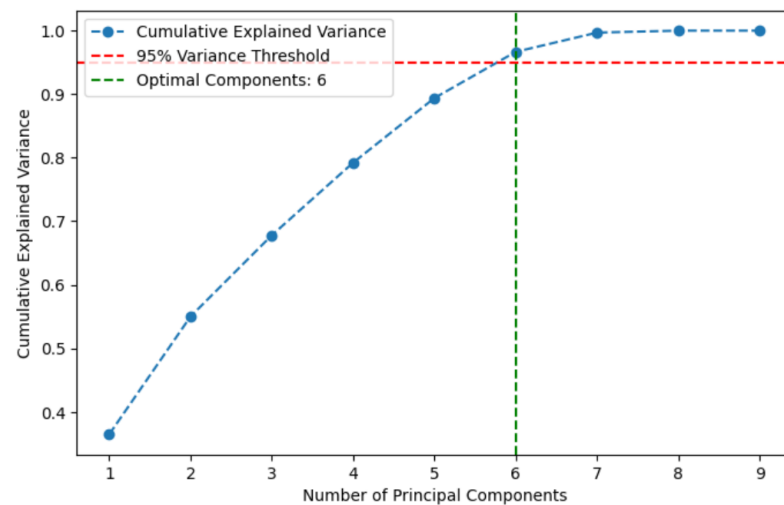


Figure 3. Cumulative explained variance as a function of the number of principal components.

Table 3 presents the updated performance results of the anomaly detection models after applying PCA. As shown, LOF continues to be the best model in terms of accuracy (50.00%). However, both OCSVM and IF demonstrate comparable accuracy (49.59%) and outperform LOF in other key metrics, including F1-score, precision, and recall. On the other hand, GMM and RC with EE continue to yield the weakest results among the five models, with GMM showing the lowest overall performance.

Table 3. Performance of anomaly detection ML models with PCA.

ML Models	PM	Results
IsolationForest	Accuracy	49.59%
	F1-score	34.39%
	Precision	34.39%
	Recall	34.39%
LocalOutlierFactor	Accuracy	50.00%
	F1-score	32.54%
	Precision	32.54%
	Recall	32.54%
OneClassSVM	Accuracy	49.59%
	F1-score	36.43%
	Precision	36.43%
	Recall	36.43%
GaussianMixture	Accuracy	45.87%
	F1-score	27.84%
	Precision	27.84%
	Recall	27.84%
RC with EllipticEnvelope	Accuracy	46.69%
	F1-score	32.85%
	Precision	32.85%
	Recall	32.85%

Compared to the results in Table 2, there was an overall improvement in the performance of the models with the use of PCA, with the exception of LOF, which demonstrated

worse performance, indicating that this model works better with high-dimensional data, while the others do not.

This approach still requires further refinement, such as parameter tuning, to maximize model performance. However, this can be time-consuming, especially when working with larger or continuously updated datasets. Another problem is the use of fixed validity thresholds, which may need to be adapted for different contexts to avoid misclassifications.

The results presented above have been obtained with a single dataset split into 80% for training and 20% for testing model performance. This raises a problem related to the credibility of the results obtained. To address this issue and better understand the models' performance, we have also used stratified K-fold validation with 100 splits over the 242 rows of the used dataset. Table 4 shows the mean and standard deviation of the obtained models' performance metrics.

Table 4. Mean and standard deviation of the studied anomaly detection models' performance metrics after using stratified K-fold validation.

ML Models	PM	Mean	Std. Dev.
IsolationForest	Accuracy	50.7%	±10.8
	F1-score	31.4%	±7.6
	Precision	42.8%	±13.2
	Recall	25.3%	±5.4
LocalOutlierFactor	Accuracy	50.3%	±10.9
	F1-score	31.2%	±7.7
	Precision	42.7%	±13.5
	Recall	25.2%	±5.4
OneClassSVM	Accuracy	51.0%	±10.8
	F1-score	31.5%	±7.4
	Precision	43.0%	±13.0
	Recall	25.5%	±5.4
GaussianMixture	Accuracy	51.7%	±10.7
	F1-score	31.9%	±7.1
	Precision	43.3%	±12.5
	Recall	25.8%	±5.4
RC with EllipticEnvelope	Accuracy	50.7%	±10.8
	F1-score	31.4%	±7.6
	Precision	42.8%	±13.2
	Recall	25.3%	±5.4

The results obtained reveal that this approach, with anomaly detection based on individual metrics, has low performance and a high degree of variability. The accuracy around 50% for all the models indicates that the dataset is imbalanced. Precision around 43% allows us to conclude that, when the model predicts positive, it is correct 43% of the time. However, the very low recall values mean that there are many false negatives. All these are indicators of an imbalanced dataset, which may also relate to the fact that it is a rather small dataset.

Despite the lack of accuracy, these models were capable of detecting deviations without the need for labeled data, which can be helpful in scenarios where labeled datasets are difficult to obtain. In the end, the LOF, OCSVM, and IF models were chosen to be used in a validation API (refer to Section 5.3) as they were the best-performing models.

5.2.2. Approach 2—Supervised Classification and Regression Models for Anomaly Detection Based on Individual Metrics

This approach employs both regression and classification supervised learning models to predict the degree and category of validity for the value received for a specific environmental metric related to the production of a textile product batch. In other words, it evaluates the reliability of a given metric within the context of transforming one or more raw materials into a final product. For example, the model may be used to predict the validity (both in degree and category) of the reported amount of electricity consumed in producing a 900 kg batch of fabric, where cotton and silk yarn are used as raw materials.

Data Preparation

The dataset used for this approach is the same as the one described in Section 5.2.1, with the structure and content previously detailed in Section 5.1 and additional preparation steps explained in Section 5.2.1. As mentioned previously, the dataset contains 242 rows and 11 columns.

ML Algorithms

In this approach, the algorithms RF, DT, SVM, NB, and KNN are used to classify data based on their validity category. To predict the degree of data reliability (i.e., a continuous score), the models used were DT, RF, and eXtreme Gradient Boosting (XGBoost). These algorithms were selected due to their effectiveness on small datasets, as well as their popularity and scalability in practical applications. XGBoost is a widely used ML algorithm known for its high efficiency, speed, accuracy, and scalability in real-world scenarios [48]. It is an ensemble learning method based on gradient boosting, which builds a series of Decision Trees, where each new tree tries to correct the errors made by the previous ones.

All the models were developed with Python (version 3.9.22) and Scikit-learn (version 1.6.1) and XGBoost (version 2.1.4) libraries, using the hyperparameters shown in Table 5.

Table 5. Machine learning models, hyperparameters, and target type (regression vs. classification) for Approach 2.

Model	Hyperparameters and Values	Target Type
RandomForestClassifier	n_estimators = 100, random_state = 42	Classification
DecisionTreeClassifier	max_depth = 10, min_samples_split = 5, random_state = 42	Classification
OneClassSVM (Classifier)	kernel = 'rbf', gamma = 'scale'	Classification
Naive Bayes (GaussianNB)	Default parameters	Classification
KNeighborsClassifier	n_neighbors = 5	Classification
DecisionTreeRegressor	max_depth = 10, min_samples_split = 5, random_state = 42	Regression
RandomForestRegressor	n_estimators = 100, random_state = 42	Regression
XGBoost Regressor	objective = "reg:squarederror", n_estimators = 100, learning_rate = 0.1, max_depth = 3, random_state = 42	Regression

To measure the performance of these models, standard performance metrics, such as accuracy, F1-score, recall, and precision, were used to measure the effectiveness of the classification algorithms. For the regression models, we selected the R^2 score as the sole

performance metric as it provides a clear indication of how well the predicted values approximate the actual degree of validity.

Results

Table 6 presents the results of the classification models using the mentioned dataset. As can be seen in the table, the best-performing model was RF with 83.67% accuracy, followed by DT and SVM with both 77.55%, then KNN with 75.51%, and finally NB with 34.69%. These results show that the models were able to correctly classify most of the data; however, they also show that there is still room for improvement since they are far from being excellent.

Table 6. Performance of ML classification models.

ML Models	PM	Results
RF	Accuracy	83.67%
	Precision	76.80%
	Recall	83.67%
	F1-score	80.02%
DT	Accuracy	77.55%
	Precision	76.83%
	Recall	77.55%
	F1-score	77.00%
SVM	Accuracy	77.55%
	Precision	69.96%
	Recall	77.55%
	F1-score	73.52%
KNN	Accuracy	75.51%
	Precision	67.23%
	Recall	75.51%
	F1-score	70.95%
NB	Accuracy	34.69%
	Precision	78.78%
	Recall	34.69%
	F1-score	31.98%

One important thing to note for NB is that its accuracy is lower than its precision, which means that the model has overall poor predictive performance, but, when it does predict the most dominant class (“valid”), it is often correct. This might be because of class imbalance since there are more valid data values than the other types, which can be solved by training the model with more diverse data.

Table 7 presents the results of the regression models when assessing the degree of validity of the textile data. As can be seen in Table 7, RF is the best-performing model, with an R^2 score of 77.18%, followed by XGBoost with 76.89% and DT with 62.39%. These results are good; however, much like the classification models, there is still room for improvement, such as using more data in the training of these models.

When comparing the performance of the classification and regression models, it is possible to see that RF is the most suitable model for our data for both tasks as it was the best model overall. DT, on the other hand, is more suitable for categorizing the validity of the data rather than predicting their degree of validity. This may be due to the structure of the dataset, which may contain discrete or categorical patterns that are easier to separate into classes than to model as a continuous output.

Table 7. Performance of ML regression models.

ML Models	PM	Results
RF	R^2 Score	77.2%
XGBoost	R^2 Score	76.9%
DT	R^2 Score	62.4%

The results above, both for classification models and for regression models, have been obtained with a single dataset split of 80%/20% for testing model performance. For confirming the reliability of the results obtained, we have also used stratified K-fold validation with 100 splits for further validating the classification models in Approach 2, and K-fold validation with 10 splits for regression models.

Table 8 presents the results of supervised classification models evaluated using stratified K-fold cross-validation.

Table 8. Performance metrics of ML classification models after applying stratified K-fold validation with 100 splits.

ML Models	PM	Mean	Std. Dev.
RF	Accuracy	95.3%	± 11.6
	F1-score	91.6%	± 20.9
	Precision	90.7%	± 23.2
	Recall	93.0%	± 17.4
DT	Accuracy	92.7%	± 17.1
	F1-score	89.1%	± 24.4
	Precision	89.6%	± 23.5
	Recall	89.3%	± 24.3
SVM	Accuracy	95.0%	± 12.9
	F1-score	91.5%	± 21.3
	Precision	90.8%	± 23.1
	Recall	92.6%	± 18.7
KNN	Accuracy	96.7%	± 10.1
	F1-score	94.0%	± 18.1
	Precision	93.3%	± 20.1
	Recall	95.0%	± 15.1
NB	Accuracy	96.0%	± 10.1
	F1-score	92.8%	± 19.6
	Precision	92.0%	± 21.8
	Recall	94.0%	± 16.3

Overall, the classifiers demonstrated strong predictive ability, with mean accuracies above 92% and F1-scores exceeding 89%. Precision and recall were also consistently high, indicating that the models were capable of correctly identifying anomalies while maintaining a low rate of false positives. Despite these strong averages, relatively large standard deviations were observed across all the metrics (ranging from $\pm 10\%$ to $\pm 24\%$). This variability suggests that performance is somewhat sensitive to data splits, which may reflect the limited dataset size or inherent heterogeneity in the samples. Nevertheless, the consistently high mean values across classifiers provide evidence that supervised classification constitutes an effective approach for anomaly detection based on individual metrics.

In contrast, the regression-based approaches performed poorly. As shown in Table 9, the mean R^2 score values were near zero or negative for all the models, with particularly poor performance for Decision Trees (-97.4% and ± 90.1). Random Forest achieved a

slightly positive mean R^2 (3.5%) but with high variability (± 29.6), indicating instability and inconsistent generalization. XGBoost also underperformed, with a negative mean R^2 (-16.4%) and large variance (± 39.7). These results suggest that regression models based on individual metrics are not suitable for anomaly detection of textile industry indicators as their predictions frequently fail to outperform a naive mean-based baseline. The discrepancy between the strong performance of classification models and the weak results from regression highlights the importance of framing anomaly detection as a classification problem rather than as a regression task when using individual metrics.

Table 9. Performance metrics of ML regression models after applying K-fold validation with 10 splits.

ML Models	PM	Mean	Std. Dev.
RF	R^2 Score	3.5%	± 29.6
XGBoost	R^2 Score	-16.4%	± 39.7
DT	R^2 Score	-97.4%	± 90.1

Based on the results, the models selected for implementation in the proposed ML API were the classification models RF, DT, and SVM.

5.2.3. Approach 3—Supervised Models for Anomaly Detection of Records with Collective Related Attributes

The goal of this approach is to develop ML models capable of assessing the reliability of sustainability indicator values for a given batch while also verifying the proportionality of these values in relation to other indicators. To achieve this, the dataset described in Section 5.1 has been reused, but it has now been reorganized to enable meaningful comparisons between sustainability indicators (metrics).

The restructured dataset is designed to help the models learn relationships between the various metrics within each different production activity. For instance, when evaluating total water consumption during a production activity for producing a given product batch, the model also considers whether this value aligns with the volume of liquid effluents generated during the same process. By incorporating such relationships, the models can more effectively assess data reliability, leading to more robust and accurate validation of textile sustainability data.

Data Preparation

The data was structured so that each row aggregates all sustainability indicators related to a production activity for producing a single batch. Each row now contains detailed information about a production activity for a given batch, including the batch size, production activity, types of raw materials used along with their respective percentages, and the values of 28 environmental impact metrics (sustainability indicators). Additionally, there are columns indicating the data validity category (invalid, suspect, or valid) and the degree of validity of the information. Depending on the type of production activity, some metrics may have zero value. This happens when the metric is not measured for that production activity. With the aggregation of several records from the initial dataset within one record for the same production activity and batch reference, the total number of records is now down to 13 rows with real (valid) data, and the number of columns has risen to 37. As this meant a very small number of rows, we have created new valid rows by applying different factors to the existing valid rows, and have created some suspect and invalid rows, with a validity degree calculated the by formulas presented in [36].

After all these preparation steps, the final dataset had forty-one rows (twenty-six valid rows (63.4%), six suspect (14.6%), and nine invalid (22.0%) and thirty-seven columns: a

column with the lot produced quantity, another with the production activity name, two columns with the types of materials used, and another two representing their respective percentages, and twenty-eight columns corresponding to the twenty-eight environmental impact metrics. Columns with the organization and record validity degree and corresponding validity category complete the thirty-seven columns. The organization is dropped before running the algorithms, reducing the number of columns to thirty-six.

After the preparation of the dataset, the correlation matrix was generated to show how attribute values vary in relation to each other. The correlation matrix for the created dataset is shown in Figure 4, and it represents the degree of relationship between the environmental impact metrics of the dataset. Possible values for each row and column of the correlation matrix range from -1 to $+1$. The value “ -1 ” means that there is a negative correlation; that is, the value of one attribute increases while the other decreases. On the opposite side, the value “ $+1$ ” means that there is a positive correlation; that is, the attribute values grow together. The value 0 means that there is no relationship between the attributes. The correlation matrix also demonstrates color tones that vary between blue, representing strong negative correlation, and red, representing strong positive correlation.

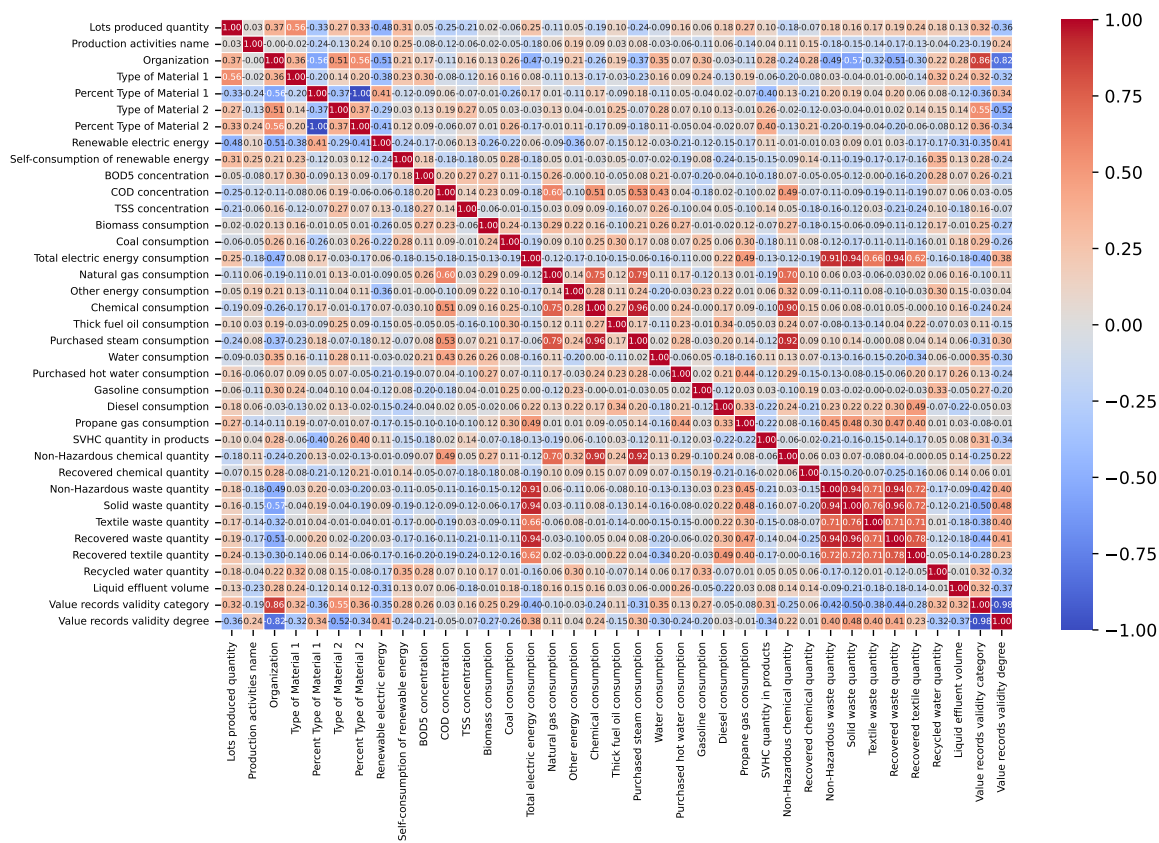


Figure 4. Correlation matrix—relation between the environmental impact metrics.

As can be seen in the correlation matrix, some metrics exhibit strong relationships with others (with correlation values close to 1). For example, the metrics “Non-Hazardous Waste Quantity”, “Solid Waste Quantity”, “Textile Waste Quantity”, “Recovered Waste Quantity”, and “Recovered Textile Quantity” all show a strong positive correlation with “Total Electric Energy Consumption”. Also, “Purchased Steam Consumption” demonstrates a strong positive correlation with both “Chemical Consumption” and “Non-Hazardous Chemical Quantity” produced.

ML Algorithms

Having each record labeled as invalid, suspect, or valid, the dataset is prepared for use with supervised learning algorithms. This approach explores the application of commonly used algorithms for anomaly detection, namely RF, DT, KNN, SVM, and NB [3]. The study is divided into classification and regression models, each analyzed separately to evaluate their effectiveness in detecting anomalies in the data.

In this approach, the algorithms and hyperparameters used to train the models are the same as those employed in Approach 2.

Results

As before, classification models have been evaluated according to standard performance metrics for classification models, namely accuracy, precision, recall, and F1-score. Table 10 presents the results with regard to accuracy, F1-score, and precision of each model based on the created dataset.

Table 10. Performance of ML classification models.

ML Models	PM	Results
RF	Accuracy	96%
	F1-score	98%
	Precision	97%
DT	Accuracy	95%
	F1-score	98%
	Precision	96%
KNN	Accuracy	70%
	F1-score	71%
	Precision	76%
SVM	Accuracy	63%
	F1-score	66%
	Precision	79%
NB	Accuracy	67%
	F1-score	65%
	Recall	67%

Regression is a supervised learning technique, used to make predictions, that models a relationship between independent and dependent variables [49]. As previously, R^2 score has been used for assessing the performance of regression models. The results of the regression ML models trained and evaluated via this approach are represented in Table 11. The table presents the models sorted in descending order based on their R^2 scores. As shown, the DT model achieved the highest performance among all the evaluated models.

Table 11. Performance of ML regression models.

ML Models	PM	Results
DT	R^2 Score	85%
RF	R^2 Score	56%
XgBoost	R^2 Score	52%

As before, given the very small dataset, K-fold validation has also been used to better understand the performance of the models.

Table 12 shows the performance metrics of the ML classification models studied in Approach 3 after using stratified K-fold validation with 10 splits. Moreover, in Table 13, the R^2 performance indicators of the regression models of this approach, obtained using K-fold validation with 10 splits, are shown.

Table 12. Performance of ML classification models of Approach 3 after using stratified K-fold validation with 10 splits.

ML Models	PM	Mean	Std. Dev.
RF	Accuracy	88.0%	± 12.1
	F1-score	84.7%	± 15.6
	Precision	82.8%	± 17.7
DT	Accuracy	83.0%	± 15.8
	F1-score	79.4%	± 17.9
	Precision	77.2	± 19.7
KNN	Accuracy	80.5%	± 14.9
	F1-score	73.1%	± 20.1
	Precision	68.4%	± 23.2
SVM	Accuracy	88.0%	± 12.1
	F1-score	84.5%	± 15.7
	Precision	82.3%	± 18.1
NB	Accuracy	63.5%	± 19.8
	F1-score	64.3%	± 20.1
	Recall	70.1%	± 24.5

Table 13. Performance of ML regression models of Approach 3 after using K-fold validation with 10 splits.

ML Models	PM	Mean	Std. Dev.
DT	R^2 Score	53.1%	± 43.1
RF	R^2 Score	53.1%	± 40.9
XgBoost	R^2 Score	58.7%	± 40.8

In summary, this approach evaluates two groups of models with a dataset with multiple related features in each data record: regression and classification. The results indicate that the used classification models outperform the regression models. Random Forest and SVM are the most reliable classifiers (*Accuracy* $\approx 88\%$; strong *F1* $\approx 84\text{--}85\%$) and show moderate variability (Std. Dev. ~ 12). Decision Tree also performs decently (83% accuracy) but is less stable (Std. Dev. 15.8). KNN is weaker, with lower accuracy (80.5%) and weaker precision. Naive Bayes performs poorly (accuracy 63.5%) and has large variability (Std. Dev. 19.8).

Due to the relatively small dataset, regression models struggle to predict continuous values with high accuracy. XGBoost is the strongest regressor ($R^2 = 58.7\%$) but still unstable. All the regressor models show very high standard deviations of R^2 (above 40). This means that model performance changes considerably depending on the fold (dataset is small and sensitive to sampling). Compared to the 80/20 split, the R^2 of DT dropped dramatically (from 85% to $\sim 53\%$). This confirms that the 80/20 results were optimistic and overfitted. Data likely needs more samples and better preprocessing.

In contrast, classification models, working with discrete categories (invalid, suspect, and valid), tend to generalize better with limited data, making them more robust and effective in this context.

5.2.4. Approach 4: Supervised Models for Anomaly Detection of Individual Metrics Based on Collective Related Attributes

This approach explores two distinct supervised learning strategies aimed at predicting the values of environmental metrics. The primary goal is to estimate the value of a specific metric, associated with a given production activity and a produced batch, by leveraging information from the remaining metrics and batch-specific attributes, such as raw material composition, batch size, water consumption, and liquid effluent volume. For example, validating the total water consumption for a batch may involve using features like the batch size, types of materials used, production activity, and other related metrics, such as the volume of liquid effluents. This method captures the interdependencies between various indicators and production characteristics to accurately validate the target metric.

Scikit-learn and XGBoost Python libraries have been used for the models in this approach, mostly with the libraries' default hyperparameters, as can be seen in Table 14.

Table 14. Hyperparameters and values for the models used in Approach 4.

Model	Hyperparameters and Values
RandomForestRegressor	n_estimators = 100, random_state = 42, max_depth = None, min_samples_split = 2, other defaults
DecisionTreeRegressor	random_state = 42, max_depth = None, criterion = 'squared_error', other defaults
KNeighborsRegressor	n_neighbors = 5, weights = 'uniform', metric = 'minkowski', p = 2
XGBRegressor	n_estimators = 100, random_state = 42, verbosity = 0, max_depth = 6, learning_rate = 0.3, other defaults

Approach 4.1: Dedicated Model per Metric

The first strategy involves developing twenty-eight distinct Machine Learning (ML) regression models, one for each of the twenty-eight metrics being validated.

The dataset utilized was identical to the one created in Section 5.2.3, inheriting the same data preparation pipeline. The key distinction lies in the model training setup: for each of the 28 metrics that can be validated, a dedicated model has been trained. In each case, the target metric was removed from the input feature set and designated as the output variable, while the remaining 27 metrics plus the batch characteristics served as input features.

This design allows each model to specialize in predicting its assigned metric, potentially leading to higher accuracy as each model is tuned to the specific nuances of its target variable within the context of the other batch data.

ML Model Evaluation

Following data preparation, appropriate regression models were selected. Given the relatively small dataset size, neural network architectures were not considered.

The same set of regression algorithms evaluated in Section 5.2.3 were employed here. In total, one-hundred-twelve models were trained (one for each of the twenty-eight metrics, potentially across four different algorithms). Model performance was evaluated using the Coefficient of Determination (R^2 score). After evaluating all models, the top five performers based on R^2 score were identified, as shown in Table 15.

Table 15. Performance of the top 5 ML regression models from Section 5.2.4—Approach 4.1 (dedicated model per metric).

ML Algorithm	Target Metric	Perf. Metric	Result
Decision Tree (DT)	Total electricity consumption	R^2 Score	99%
Random Forest (RF)	Amount of solid waste	R^2 Score	93%
XGBoost	Consumption of chemical products	R^2 Score	91%
XGBoost	Quantity of waste recovered	R^2 Score	91%
Random Forest (RF)	Purchased steam consumption	R^2 Score	86%

The highest performance was observed with a Decision Tree model for predicting/validating ‘Total electricity consumption’ ($R^2 = 99\%$), indicating excellent predictive capability for this metric. Random Forest and XGBoost models also showed strong performance for metrics such as ‘Amount of solid waste’, ‘Consumption of chemical products’, and ‘Quantity of waste recovered’. These results suggest that tree-based methods are particularly effective for this specialized modeling task.

Approach 4.2: Single Unified Model

This second strategy aimed to develop a single ML model capable of predicting any of the 28 metrics rather than training metric-specific models.

The same dataset and preprocessing steps from Section 5.2.3 were used as before. The key transformation in this approach involved restructuring the dataset: each original record (representing a single production activity and batch) was expanded into twenty-eight new records, one for each metric considered as a potential prediction target. A new categorical feature, ‘target_name’, was introduced to each expanded record to indicate which metric serves as the target variable for that instance. The value of the specified target metric became the output variable of that row, while the remaining twenty-seven metrics, along with batch-specific attributes (e.g., production activity, batch size, material types, and their percentages) and the ‘target_name’ itself formed the input features.

This transformation resulted in a dataset with 392 rows (based on 14 original batches, $14 \times 28 = 392$) and 35 columns. Among these, twenty-eight columns represented metric values (with one serving as the output and the remaining twenty-seven as inputs depending on the ‘target_name’). The rest included features for production activity, batch size, material types and proportions, and the ‘target_name’ identifier. In total, each row contained thirty-four input features and one target variable, the value of the designated target metric.

ML Model Evaluation

Model selection and evaluation followed the same methodology as in Section 5.2.4—Approach 4.1. Table 16 presents the performance of the best models identified for this unified approach.

Table 16. Performance of the top 4 ML regression models from Section 5.2.4—Approach 4.2 (single unified model).

ML Algorithm	Perf. Metric	Result
Decision Tree (DT)	R^2 Score	98.6%
Random Forest (RF)	R^2 Score	85.4%
XGBoost	R^2 Score	93.5%
K-Nearest Neighbor (KNN)	R^2 Score	45.3%

As shown in Table 16, DT and tree-based ensemble models (Random Forest and XGBoost) demonstrated very strong performance, achieving high overall R^2 scores. In contrast, the K-Nearest Neighbor (KNN) algorithm performed poorly ($R^2 = 46\%$). This discrepancy is likely due to the challenge faced by distance-based methods like KNN in generalizing effectively across the heterogeneous nature of the target variable, which encompasses 28 different metrics with potentially diverse scales and distributions, all handled within a single model structure. The added 'target_name' feature aids tree-based models in partitioning the data appropriately, a task potentially harder for KNN in this context. For confirming these results, K-fold cross-validation has once again been used. The obtained results can be seen in Table 17.

Table 17. Performance of ML classification models of Approach 4.2 after using K-fold cross-validation with 5 splits.

ML Models	Perf. Metric	Mean	Std. Dev.
DT	R^2 Score	92.8%	± 8.15
	MSE	3.92×10^5	—
	MAE	57.9	—
RF	R^2 Score	88.0%	± 9.98
	MSE	5.98×10^5	—
	MAE	80.9	—
XGBoost	R^2 Score	91.9%	± 8.09
	MSE	4.46×10^5	—
	MAE	62.3	—
KNN	R^2 Score	76.4%	± 16.84
	MSE	8.76×10^5	—
	MAE	120.5	—

These results somewhat confirm what we had already concluded. Decision Tree and XGBoost are clearly the best models for predicting any of the 28 considered textile metrics based on the collective related metrics for the same lot. They achieve a very high R^2 ($\sim 92\text{--}93\%$) and present the lowest MSE and MAE, having the most accurate predictions. Also, DT and XGBoost have the lowest R^2 standard deviation, being fairly consistent across folds.

RF also performs well ($R^2 \sim 88\%$), but not as strong as DT and XGBoost as its MAE is much higher (~ 81 vs. $58\text{--}62$), which means that predictions are less precise. Its R^2 standard deviation is also higher (~ 10).

KNN shows moderate performance ($R^2 \sim 0.76$) but high variance, meaning instability.

Comparison Between the Two Approaches

This section compares the two strategies presented in this section, Approach 4.1 (28 models, denoted 28/28) and Approach 4.2 (single unified model, denoted 1/28). The comparison considers predictive performance, maintenance effort, and scalability, as summarized in Table 18.

Approach 4.1 (28/28) offers the potential for superior accuracy on a per-metric basis due to model specialization. However, this comes at the cost of significantly higher complexity in terms of development, deployment, and ongoing maintenance (managing 28 individual models).

Approach 4.2 (1/28) presents a much simpler and more scalable architecture. While its overall performance is excellent, the accuracy for any single specific metric might be lower than that achievable with a dedicated model from Approach 4.1. The single model

must generalize across all metrics, which might compromise performance for metrics with unique patterns or weaker correlations with the input features.

Table 18. Comparison between the dedicated models (28/28) and single unified model (1/28) approaches.

Criteria	Approach 4.1 (28/28)	Approach 4.2 (1/28)
Overall R^2 score	High for best models (metric-specific)	High overall (e.g., 98%), but potentially variable per metric
Metric-Specific R^2	Potentially higher per metric	Can be lower for specific metrics vs. dedicated models
Maintenance	Complex (manage 28 models)	Simpler (manage 1 model)
Scalability	More difficult (adding metrics requires new models)	Easier (handles new targets via 'target_name' if structure is similar)
Recommendation	When highest precision per metric is critical	When scalability, maintainability, and good overall performance are key

Considering the operational context, where scalability and ease of maintenance are often crucial, Approach 4.2 was deemed more appropriate. The high overall R^2 achieved suggests strong predictive power across the board, and the operational benefits of managing a single model outweigh the potential for slightly lower accuracy on some individual metrics compared to the more complex 28-model system. Therefore, the unified model strategy (Section 5.2.4—Approach 4.2) was selected for further development and deployment regarding the validation API.

5.3. Integration API

For the DPP to serve as a trustworthy source of detailed information about a product's origin, composition, and life-cycle activities, promoting transparency and supporting more sustainable and responsible decision-making, the data used to build the DPP must be validated prior to its integration into the DPP platform. This validation is particularly critical in the complex landscape of the T&C industry, where numerous organizations generate and manage data through different software systems before providing some of that data as indicators to be collected and incorporated into the DPP. To this purpose, an API has been developed to facilitate the validation of such indicators or metrics. The API provides endpoints that leverage a range of Machine Learning algorithms, selected based on the research and evaluation presented in Sections 5.2.1–5.2.4, to accomplish the data validation process.

Within the context of the DPP, where both data integrity and interoperability are essential, this API-driven solution becomes indispensable. Data must be collected from various points across the T&C value chain, including Enterprise Resource Planning (ERP) systems, Internet of Things (IoT) devices, and other data collection platforms used by participating companies. The developed API enables seamless submission of product and sustainability data to the validation service, ensuring that only verified and reliable information is integrated into the DPP.

The API data validation services evaluate the submitted data for accuracy, consistency, and compliance with predefined standards, returning a validation status, ensuring that only high-quality trustworthy data is integrated into the Digital Product Passport.

Along with the data validation services, the API includes services for managing the different types of data entities, such as production activities, raw materials, and environmental indicators, which were addressed in [36,37]. This standardized interface enables consistent data validation practices across the textile and clothing (TC) value chain, thereby improving the reliability and usefulness of the DPP.

To implement and manage the proposed solutions, a software framework was developed using Spring Boot in Java. The system was designed as a RESTful Web API, ensuring

seamless integration and scalability. Additionally, authentication and authorization mechanisms are managed via Keycloak, ensuring secure access control. To enhance operational monitoring and diagnostics, Grafana Loki and Tempo are employed for centralized log management and distributed tracing visualization. Prometheus is used for real-time metric collection and analysis, while Grafana serves as a comprehensive visualization and dashboard tool.

The validation services provided by the API include endpoints for using the ML-based validation services. These are presented in Table 19. As can be seen in the table, there are three endpoint services for anomaly detection using the IF, OCSVM, and LOF models selected in Section 5.2.1 to be used in the validation API.

From Approach 2 (Section 5.2.2), the models selected for implementation in API were RF, DT, and SVM for classification tasks.

From Approach 3 (Section 5.2.3), the services available on the API apply classification models based on DT, RF, and SVM, and regression models based on DT, RF, and XGBoost.

From Approach 4, presented in Section 5.2.4, the unified model strategy (Approach 4.2) was the one selected for the API. The services available use DT, RF, and XGBoost.

Table 19. ML API endpoints available.

Dataset	Endpoint	Type	Description
Approach 1	POST /approach1/classification/iso_forest	Classification	Models trained to detect anomalous deviations in metrics, classifying them as valid, invalid, or suspect, based on normal behavior patterns.
	POST /approach1/classification/ocsvm		
	POST /approach1/classification/lof		
Approach 2	POST /approach2/classification/random_forest	Classification	The classification models were trained to predict the validity of a metric (valid, invalid, or suspect).
	POST /approach2/classification/decision_tree		
	POST /approach2/classification/naive_bayes		
Approach 3	POST /approach3/classification/dt	Classification	Classification models trained with dataset 3 to assess the validity of records based on their characteristics, classifying them as valid, invalid, or suspect.
	POST /approach3/classification/rf		
	POST /approach3/classification/svm		
	POST /dataset3/regression/dt	Regression	Regression models to predict a continuous value associated with a record quality degree. The decision follows the same logic: < 25 invalid, between 25 and 75 suspect, > 75 valid.
	POST /approach3/regression/rf		
	POST /approach3/regression/xgb		
Approach 4	POST /approach4/regression/dt	Regression	Regression models to predict environmental metric values.
	POST /approach4/regression/rf		
	POST /approach4/regression/xgb		

API Integration Case Validation

The validation API for textile environmental indicators is currently being used by INFOS (<https://infos.pt/en/projetos-de-inovacao/> accessed on 25 August 2025) in the implementation of a Digital Product Passport for the textile sector. The current phase includes major contributions, such as the selection of more real data records for building larger datasets, with a wider spectrum of production activities and types of textile industries, which will enable new models to be trained better with the same algorithms. Another major contribution is the building of knowledge about which validation approach is better for each environmental indicator registered in the DPP.

6. Discussion

This section analyzes the approaches presented in the previous section. Table 20 presents a comparative summary of the four proposed Machine Learning approaches for assessing the quality of textile data.

Table 20. Comparison of the four proposed ML-based approaches for textile data validation.

Approach	Model Type	Dataset Type	Strengths	Limitations	Best Use Case
1. Anomaly Prediction of Individual Metrics	Unsupervised	Single metric	- Suitable when labeled data is limited - Can learn from labeled and unlabeled data - Effective in anomaly detection with minimal supervision - Easy implementation and scalability	- Performance depends on the quality and distribution of unlabeled data - Harder to train and fine-tune compared to fully supervised methods - Difficult to find the right threshold values - Lower accuracy when large labeled datasets are available	Applicable when labeled data is scarce and metrics need to be validated individually
2. Supervised (Classification and Regression) Prediction of Individual Metrics	Supervised	Single metric	- High accuracy with sufficient labeled data - Allows for specialized models tailored to each individual metric - Easier to interpret and validate metric by metric	- Ignores relationships between different metrics - Requires substantial labeled data for each metric - Less scalable with many diverse metrics	Applicable when each metric can be treated independently and labeled data is sufficient
3. Supervised Prediction based on collective attributes	Supervised	All metrics together	- Captures correlations and dependencies between different metrics - Robust for detecting inconsistencies across the full dataset - Effective for identifying systemic data quality issues - Useful for identifying systemic inconsistencies	- Less interpretable with many metrics - Requires clean and normalized input data - May overlook isolated issues in individual metrics - Overfitting likely when metrics outnumber samples	Applicable when validating global consistency of records across all metrics
4.1 Supervised Prediction of Individual Metrics based on collective attributes (28 models for 28 metrics)	Supervised	All metrics together	- Higher accuracy potential per metric due to task-specific tuning - Allows for custom hyperparameter optimization per model	- Requires maintaining and deploying 28 separate models, increasing system complexity - More computationally expensive to train and manage - Performance depends on the existence of real correlations between the metrics	Applicable when validating specific metric values based on the context provided by other metrics
4.2 Supervised Prediction of Individual Metrics based on collective attributes (1 model for 28 metrics)	Supervised	All metrics together	- Single model architecture, simplifying deployment and scaling - Potential for generalization to new or unseen metrics (if similar)	- May lead to lower accuracy for some metrics compared to specialized models - Performance depends on the existence of real correlations between the metrics	Applicable when validating specific metric values based on the context provided by other metrics

The table highlights each approach's model type, dataset structure and characteristics, the key advantages and limitations, and the most appropriate scenarios for their application.

Approach 1 (anomaly detection based on individual metrics) is an unsupervised approach that may achieve good results when labeled data is scarce. It can detect anomalies using patterns in unlabeled data and is easy to scale and implement. However, it suffers from lower accuracy in well-labeled contexts, and performance heavily depends on proper threshold selection and data distribution. This approach is best suited for early-stage deployments or cases with limited annotation.

Although the results are not the best, this approach shows that unsupervised anomaly detection algorithms may be a viable approach to detect anomalies in textile data.

Nevertheless, several limitations must be acknowledged. The small dataset constrains generalizability, the thresholds were set without extensive model tuning, and class imbalance, particularly regarding the limited "suspect" cases, may have influenced its performance. These factors, along with the absence of variability measures in the reported results, should be considered in future work, specifically by expanding the datasets with more records from diverse companies in the textile value chain, introducing light or domain-informed parameter tuning, and applying stratified K-fold cross-validation to improve the reliability and robustness of the model evaluation.

Approach 2 (supervised classification and regression models for anomaly detection based on individual metrics) applies classification and regression models independently to each metric. It provides high accuracy when sufficient labeled data is available and is easier to interpret on a metric-by-metric basis. However, it fails to account for inter-metric relationships and may not scale well with many diverse metrics.

In Approach 3 (supervised models for anomaly detection of records with collective related attributes), instead of assessing individual metrics, the models assess the consistency of an entire record by considering all the metrics together. This allows for detecting systemic errors and exploiting correlations between metrics. While powerful for spotting broad inconsistencies, it is less interpretable and may struggle with isolated metric-level anomalies, especially in small datasets with many features.

Approach 4 (supervised models for anomaly detection of individual metrics based on collective related attributes) predicts each individual metric using the rest of the metrics and batch characteristics as context. It combines the benefits of both localized and contextual validation. Approach 4.1 allows for fine-tuning per metric while leveraging interdependencies. Despite its higher implementation and computational complexity, it offers strong generalization and accuracy potential across the metrics. Approach 4.1 is computationally demanding, and the models' scalability, training latency, and inference throughput need to be critically assessed in future work. Future work should also better analyze the computationally demanding algorithms of Approach 4.1, for which K-fold validation has not been used. The possible advantage of using a specialized regression model for each metric that is able to make use of the context provided by the other metrics (Approach 4.1) must be further analyzed when compared with a unique regression model for all the metrics based on the global metrics context (Approach 4.2).

As mentioned before, the validation API for textile environmental indicators is currently being used by INFOS in the development of a DPP for textile products, enabling the gathering of real data for building a larger dataset. With the collection of new data and the construction of a new dataset from several companies in the value chain, including from companies in different countries, future work will be able to retrain all the models of the proposed approaches and solve most of the problems related to dataset scale and geographical bias, enhancing the generalizability of the findings.

Further research should also be conducted in the future regarding the exploration of data augmentation methods (e.g., SMOTE variants and GANs) to mitigate class imbalance and scale methods for larger datasets.

7. Conclusions

The implementation of the DPP is expected to contribute many benefits in terms of sustainability and the circular economy by providing transparency and traceability in the value chain and increasing social and environmental responsibility in the T&C industrial sector.

For companies to achieve returns on their investments and be able to compete fairly with each other, this implementation must be required at the global level. However, beyond the initial investment required, the rollout of the DPP presents several challenges. At the top are the difficulties involved in collecting and integrating data across companies with widely varying levels of digital maturity, as well as the complexity of obtaining certain sustainability indicators.

The T&C value chain generates a vast volume of data due to the large number of production activities and the wide range of organizations involved in data collection. This data is sourced from diverse IoT devices and various software platforms (e.g., ERPs and MRPs), leading to significant variability in format and quality. To address this challenge, this article proposes an API that offers a comprehensive suite of data validation services. The API serves as a unifying layer to validate and standardize data before it is stored in the platform that supports the DPP.

Given the dynamic nature of the T&C value chain and the high volume and velocity of incoming data, traditional data validation methods may fall short in meeting performance requirements. Therefore, this work explores a set of Machine Learning-based approaches to enable more efficient and scalable data validation procedures within the DPP implementation. By leveraging pre-trained ML models, the system can significantly enhance validation performance, especially in scenarios involving high data throughput and rapid data ingestion.

The API services are prepared to validate data received from new companies and can be easily adapted to receive data from new raw materials, new production activities, and new types of indicators to be measured and validated.

The results presented in this study are constrained by the limited size of the training dataset. It is anticipated that, with a larger volume of data, more insightful and robust conclusions could be drawn. This article presents the idea of using Machine Learning to validate the data used to create the DPP, acknowledging that the current dataset is limited and potentially biased as it originates from a small set of companies. To improve the results, the models will need to be retrained. As real data will be continuously collected and the volume and diversity of the dataset will continuously increase, the models should be periodically retrained. This process allows them to adapt to new trends, identify emerging patterns, enhance model accuracy, and address potential biases.

A key limitation of this study is the absence of systematic hyperparameter tuning, especially for the unsupervised anomaly detection models in Approach 1. In this work, we relied on sensible defaults and domain-informed thresholds to avoid overfitting. While this provides a useful baseline, it may underestimate the potential of these methods.

In future work, we plan to replicate this study using an expanded dataset to validate the findings reported here. With access to more data, we also aim to explore the use of deep learning models, such as ANNs and RNNs, which typically require larger datasets to perform effectively. Additionally, we intend to implement lightweight and reproducible tuning

pipelines, such as through stratified K-fold cross-validation with repeated folds, constrained grid or random search over high-impact parameters, and threshold sensitivity analysis.

Future work should also address the integration and testing of the validation API in a real-time pipeline that integrates directly with textile industry traceability and DPP platforms, enabling on-the-fly anomaly detection during data ingestion.

Author Contributions: This research article is a result of research and findings by the authors, with the individual contributions as follows: conceptualization, E.F.C. and A.M.R.d.C.; methodology, E.F.C. and A.M.R.d.C.; software, P.S., S.S., R.R. and A.M.R.d.C.; validation, P.S., S.S., R.R., A.M.R.d.C. and M.A.; investigation, E.F.C., A.M.R.d.C., P.S., S.S. and R.R.; resources, J.O. and E.F.C.; data curation, E.F.C., A.M.R.d.C., P.S., S.S. and R.R.; formal analysis, E.F.C., A.M.R.d.C., P.S., S.S. and R.R.; writing—original draft preparation, E.F.C., A.M.R.d.C., P.S., S.S., R.R., M.A. and J.O.; writing—review and editing, E.F.C. and A.M.R.d.C.; supervision, E.F.C.; project administration, E.F.C. All authors have read and agreed to the published version of the manuscript.

Funding: This research has been developed in the context of Project “BE@T: Bioeconomia Sustentável fileira Têxtil e Vestuário-Medida 1”, funded by “Plano de Recuperação e Resiliência” (PRR), through measure TC-C12-i01 of the Portuguese Environmental Fund (“Fundo Ambiental”).

Data Availability Statement: The original contributions presented in this study are included in the article. Further inquiries can be directed to the corresponding authors.

Conflicts of Interest: The authors declare no conflicts of interest.

Abbreviations

The following abbreviations are used in this manuscript:

AI	Artificial Intelligence
ANN	Artificial Neural Network
API	Application Programming Interface
BERT	Bidirectional Encoder Representations from Transformers
CNN	Convolutional Neural Network
DPP	Digital Product Passport
DSR	Design Science Research
DT	Decision Tree
EE	Elliptic Envelope
ERP	Enterprise Resource Planning
EU	European Union
GMM	Gaussian Mixture Model
GPT	Generative Pre-Trained Transformer
IF	Isolation Forest
IoT	Internet of Things
KNN	K-Nearest Neighbor
LOF	Local Outlier Factor
ML	Machine Learning
NB	Naive Bayes
OCSVM	One-Class SVM
PCA	Principal Component Analysis
PLM	Pre-Trained Language Model
PM	Performance Metric
RF	Random Forest
RNN	Recurrent Neural Network
SVM	Support Vector Machine
T&C	Textile and Clothing
XGBoost	eXtreme Gradient Boosting

References

- Alves, L.; Cruz, E.F.; da Cruz, A.M.R. Tracing Sustainability Indicators in the Textile and Clothing Value Chain using Blockchain Technology. In Proceedings of the 2022 17th Iberian Conference on Information Systems and Technologies (CISTI), Madrid, Spain, 22–25 June 2022; pp. 1–7. [\[CrossRef\]](#)
- Union, E. *Digital Product Passport for the Textile Sector*; EPRS | European Parliamentary Research Service: Brussels, Belgium, 2024. [\[CrossRef\]](#)
- Nunes, E.C. Machine Learning based Anomaly Detection for Smart Shirt: A Systematic Review. *arXiv* **2022**. [\[CrossRef\]](#)
- Kraljevski, I.; Ju, Y.C.; Ivanov, D.; Tschöpe, C.; Wolff, M. How to Do Machine Learning with Small Data?—A Review from an Industrial Perspective. *arXiv* **2023**. [\[CrossRef\]](#)
- Nassif, A.B.; Talib, M.A.; Nasir, Q.; Dakalbab, F.M. Machine Learning for Anomaly Detection: A Systematic Review. *IEEE Access* **2021**, *9*, 78658–78700. [\[CrossRef\]](#)
- Sahu, S.K.; Mokhade, A.; Bokde, N.D. An Overview of Machine Learning, Deep Learning, and Reinforcement Learning-Based Techniques in Quantitative Finance: Recent Progress and Challenges. *Appl. Sci.* **2023**, *13*, 1956. [\[CrossRef\]](#)
- Lohani, B.P.; Thirunavukkarasan, M. A Review: Application of Machine Learning Algorithm in Medical Diagnosis. In Proceedings of the 2021 International Conference on Technological Advancements and Innovations (ICTAI), Tashkent, Uzbekistan, 10–12 November 2021; pp. 378–381. [\[CrossRef\]](#)
- Sharma, S.; Bhatt, M.; Sharma, P. Face Recognition System Using Machine Learning Algorithm. In Proceedings of the 2020 5th International Conference on Communication and Electronics Systems (ICCES), Coimbatore, India, 10–12 June 2020; pp. 1162–1168. [\[CrossRef\]](#)
- Ozkan-Okay, M.; Akin, E.; Aslan, O.; Kosunalp, S.; Iliev, T.; Stoyanov, I.; Beloev, I. A Comprehensive Survey: Evaluating the Efficiency of Artificial Intelligence and Machine Learning Techniques on Cyber Security Solutions. *IEEE Access* **2024**, *12*, 12229–12256. [\[CrossRef\]](#)
- Alarfaj, F.K.; Malik, I.; Khan, H.U.; Almusallam, N.; Ramzan, M.; Ahmed, M. Credit Card Fraud Detection Using State-of-the-Art Machine Learning and Deep Learning Algorithms. *IEEE Access* **2022**, *10*, 39700–39715. [\[CrossRef\]](#)
- Min, B.; Ross, H.; Sulem, E.; Veyseh, A.P.B.; Nguyen, T.H.; Sainz, O.; Agirre, E.; Heintz, I.; Roth, D. Recent Advances in Natural Language Processing via Large Pre-trained Language Models: A Survey. *ACM Comput. Surv.* **2023**, *56*, 1–40. [\[CrossRef\]](#)
- Maschler, B.; Weyrich, M. Deep Transfer Learning for Industrial Automation: A Review and Discussion of New Techniques for Data-Driven Machine Learning. *IEEE Ind. Electron. Mag.* **2021**, *15*, 65–75. [\[CrossRef\]](#)
- Parekh, D.; Poddar, N.; Rajpurkar, A.; Chahal, M.; Kumar, N.; Joshi, G.P.; Cho, W. A Review on Autonomous Vehicles: Progress, Methods and Challenges. *Electronics* **2022**, *11*, 2162. [\[CrossRef\]](#)
- Jiang, Y.; Li, X.; Luo, H.; Yin, S.; Kaynak, O. Quo vadis artificial intelligence? *Discov. Artif. Intell.* **2022**, *2*, 4. [\[CrossRef\]](#)
- Badillo, S.; Banfai, B.; Birzele, F.; Davydov, I.I.; Hutchinson, L.; Kam-Thong, T.; Siebourg-Polster, J.; Steiert, B.; Zhang, J.D. An Introduction to Machine Learning. *Clin. Pharmacol. Ther.* **2020**, *107*, 871–885. [\[CrossRef\]](#)
- Pouyanfar, S.; Sadiq, S.; Yan, Y.; Tian, H.; Tao, Y.; Reyes, M.P.; Shyu, M.L.; Chen, S.C.; Iyengar, S.S. A Survey on Deep Learning: Algorithms, Techniques, and Applications. *ACM Comput. Surv.* **2018**, *51*, 1–36. [\[CrossRef\]](#)
- Rosado da Cruz, A.; Cruz, E.F. Machine Learning Techniques for Requirements Engineering: A Comprehensive Literature Review. *Software* **2025**, *4*, 14. [\[CrossRef\]](#)
- Xu, H.; Pang, G.; Wang, Y.; Wang, Y. Deep Isolation Forest for Anomaly Detection. *IEEE Trans. Knowl. Data Eng.* **2023**, *35*, 12591–12604. [\[CrossRef\]](#)
- Liu, F.T.; Ting, K.M.; Zhou, Z.H. Isolation forest. In Proceedings of the 2008 Eighth IEEE International Conference on Data Mining, Pisa, Italy, 15–19 December 2008; pp. 413–422.
- van Engelen, J.E.; Hoos, H.H. A survey on semi-supervised learning. *Mach. Learn.* **2020**, *109*, 373–440. [\[CrossRef\]](#)
- Chala, A.T.; Ray, R. Assessing the Performance of Machine Learning Algorithms for Soil Classification Using Cone Penetration Test Data. *Appl. Sci.* **2023**, *13*, 5758. [\[CrossRef\]](#)
- Alves, L.; Sá, M.; Cruz, E.F.; Alves, T.; Alves, M.; Oliveira, J.; Santos, M.; Rosado da Cruz, A.M. A Traceability Platform for Monitoring Environmental and Social Sustainability in the Textile and Clothing Value Chain: Towards a Digital Passport for Textiles and Clothing. *Sustainability* **2024**, *16*, 82. [\[CrossRef\]](#)
- De Felice, F.; Fareed, A.G.; Zahid, A.; Nenni, M.E.; Petrillo, A. Circular economy practices in the textile industry for sustainable future: A systematic literature review. *J. Clean. Prod.* **2025**, *486*, 144547. [\[CrossRef\]](#)
- Le, S.T.Q.; Bui, H.M.; Trinh, K.H.T. Enhancing clothing industry quality: Predicting knitwear defects with data mining. *Res. J. Text. Appar.* **2025**. [\[CrossRef\]](#)
- Nørup, N.; Pihl, K.; Damgaard, A.; Scheutz, C. Development and testing of a sorting and quality assessment method for textile waste. *Waste Manag.* **2018**, *79*, 8–21. [\[CrossRef\]](#) [\[PubMed\]](#)
- Forsberg, B.; Williams, D.H.; MacDonald, P.B.; Chen, T.; Hulse, D.K. Textile Anomaly Detection: Evaluation of the State-of-the-Art for Automated Quality Inspection of Carpet. *arXiv* **2024**. [\[CrossRef\]](#)

27. Sanches, R.A.; Mendes, F.D.; de Campos Araújo, M.; Issac, M.G. Textile Waste Is the Raw Material for New Fashion Products. In *Perspectives on Design III: Research, Education and Practice*; Springer Nature: Cham, Switzerland, 2024; pp. 297–310. [\[CrossRef\]](#)
28. Wenning, R.; Papadakos, P.; Bernier, C. *DPP System Architecture (V1.9)*; Technical Report; CIRPASS Consortium: Diegem, Belgium, 2024.
29. Durand, A.; Goetz, T.; Hettesheimer, T.; Tholen, L.; Hirzel, S.; Adisorn, T. Enhancing evaluations of future energy-related product policies with the Digital Product Passport. In *Proceedings of the Energy Evaluation Europe 2022 Conference*, Paris-Saclay University, Saint-Denis, France, 28–30 June 2022; pp. 28–30.
30. Lindström, J.; Kyösti, P.; Psarommatis, F.; Andersson, K.; Starck Enman, K. Extending Product Lifecycles—An Initial Model with New and Emerging Existential Design Aspects Required for Long and Extendable Lifecycles. *Appl. Sci.* **2024**, *14*, 5812. [\[CrossRef\]](#)
31. Psarommatis, F.; May, G. Digital Product Passport: A Pathway to Circularity and Sustainability in Modern Manufacturing. *Sustainability* **2024**, *16*, 396. [\[CrossRef\]](#)
32. Ariyaluran Habeeb, R.A.; Nasaruddin, F.; Gani, A.; Targio Hashem, I.A.; Ahmed, E.; Imran, M. Real-time big data processing for anomaly detection: A Survey. *Int. J. Inf. Manag.* **2019**, *45*, 289–307. [\[CrossRef\]](#)
33. Al-amri, R.; Murugesan, R.K.; Man, M.; Abdulateef, A.F.; Al-Sharafi, M.A.; Alkahtani, A.A. A Review of Machine Learning and Deep Learning Techniques for Anomaly Detection in IoT Data. *Appl. Sci.* **2021**, *11*, 5320. [\[CrossRef\]](#)
34. Zhu, M.; Wang, J.; Yang, X.; Zhang, Y.; Zhang, L.; Ren, H.; Wu, B.; Ye, L. A review of the application of machine learning in water quality evaluation. *Eco-Environ. Health* **2022**, *1*, 107–116. [\[CrossRef\]](#)
35. Uddin, M.G.; Nash, S.; Rahman, A.; Olbert, A.I. Performance analysis of the water quality index model for predicting water state using machine learning techniques. *Process Saf. Environ. Prot.* **2023**, *169*, 808–828. [\[CrossRef\]](#)
36. Rosado da Cruz, A.M.; Silva, P.; Serra, S.; Rodrigues, R.; Pinto, P.; Ferreira Cruz, E. Data Quality Assessment for the Textile and Clothing Value-Chain Digital Product Passport. In *Proceedings of the 26th International Conference on Enterprise Information Systems—Volume 2: ICEIS*, Angers, France, 28–30 April 2024; INSTICC, SciTePress: Setúbal, Portugal, 2024; pp. 288–295. [\[CrossRef\]](#)
37. Ferreira Cruz, E.; Rodrigues, R.; Serra, S.; Silva, P.; Rosado da Cruz, A.M. Using Machine Learning for Data Quality Assessment for the Textile and Clothing Digital Product Passport. In *Proceedings of the 11th IFAC Conference on Manufacturing Modelling, Management and Control (IFAC MIM)*, Trondheim, Norway, 30 June–3 July 2025.
38. Cruz, E.F.; Cruz, A.M.R.d. Design Science Research for IS/IT Projects: Focus on Digital Transformation. In *Proceedings of the 2020 15th Iberian Conference on Information Systems and Technologies (CISTI)*, Sevilla, Spain, 24–27 June 2020; pp. 1–6. [\[CrossRef\]](#)
39. Hevner, A.R.; March, S.T.; Park, J.; Ram, S. Design Science in Information Systems Research. *MIS Q.* **2004**, *28*, 75–105. [\[CrossRef\]](#)
40. Pfeffers, K.; Tuunanen, T.; Gengler, C.E.; Rossi, M.; Hui, W.; Virtanen, V.; Bragge, J. The design science research process: A model for producing and presenting information systems research. In *Proceedings of the First International Conference on Design Science Research in Information Systems and Technology (DESIST 2006)*, Claremont, CA, USA, 24–25 February 2006; pp. 83–106.
41. Ndung’u, R.N. Data Preparation for Machine Learning Modelling. *Int. J. Comput. Appl. Technol. Res.* **2022**, *11*, 231–235. [\[CrossRef\]](#)
42. Agyemang, E.F. Anomaly detection using unsupervised machine learning algorithms: A simulation study. *Sci. Afr.* **2024**, *26*, e02386. [\[CrossRef\]](#)
43. Johannesen, N.J.; Kolhe, M.L.; Goodwin, M. Vertical approach anomaly detection using local outlier factor. In *Power Systems Cybersecurity: Methods, Concepts, and Best Practices*; Springer: Berlin/Heidelberg, Germany, 2023; pp. 297–310.
44. Barbado, A.; Corcho, Ó.; Benjamins, R. Rule extraction in unsupervised anomaly detection for model explainability: Application to OneClass SVM. *Expert Syst. Appl.* **2022**, *189*, 116100. [\[CrossRef\]](#)
45. Setiadi, D.R.I.M.; Muslikh, A.R.; Iriananda, S.W.; Warto, W.; Gondohanindijo, J.; Ojugo, A.A. Outlier Detection Using Gaussian Mixture Model Clustering to Optimize XGBoost for Credit Approval Prediction. *J. Comput. Theor. Appl.* **2024**, *2*, 244–255. [\[CrossRef\]](#)
46. Airlangga, G. Analysis of Machine Learning Algorithms for Seismic Anomaly Detection in Indonesia: Unveiling Patterns in the Pacific Ring of Fire. *J. Lebesgue J. Ilm. Pendidik. Mat. Mat. Dan Stat.* **2024**, *5*, 37–48. [\[CrossRef\]](#)
47. Greenacre, M.; Groenen, P.J.; Hastie, T.; d’Enza, A.I.; Markos, A.; Tuzhilina, E. Principal component analysis. *Nat. Rev. Methods Prim.* **2022**, *2*, 100. [\[CrossRef\]](#)
48. Chen, T.; Guestrin, C. XGBoost: A Scalable Tree Boosting System. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York, NY, USA, 13–17 August 2016; pp. 785–794. [\[CrossRef\]](#)
49. Tatachar, A.V. Comparative assessment of regression models based on model evaluation metrics. *Int. Res. J. Eng. Technol. (IRJET)* **2021**, *8*, 853–860.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.