

Corporate risk stratification through an interpretable autoencoder-based model

Alessandro Giuliani ^{a,*}, Roberto Savona ^{b,1}, Salvatore Carta ^{a,1}, Gianmarco Addari ^{c,1},
Alessandro Sebastian Podda ^{a,1}

^a Department of Mathematics and Computer Science, University of Cagliari, Palazzo delle Scienze, Via Ospedale, 72, Cagliari, 09124, Italy

^b Department of Economics and Management, University of Brescia, Via San Faustino 74/B, Brescia, 25122, Italy

^c VisioScientiae S.r.l., Via San Tommaso D'Aquino, 20, Cagliari, 09134, Italy

ARTICLE INFO

Keywords:

Deep learning
Autoencoder
Balance sheets
Corporate risk
Financial sustainability

ABSTRACT

In this manuscript, we propose an innovative early warning Machine Learning-based model to identify potential threats to financial sustainability for non-financial companies. Unlike most state-of-the-art tools, whose outcomes are often difficult to understand even for experts, our model provides an easily interpretable visualization of balance sheets, projecting each company in a bi-dimensional space according to an autoencoder-based dimensionality reduction matched with a Nearest-Neighbor-based default density estimation. In the resulting space, the distress zones, where the default intensity is high, appear as homogeneous clusters directly identified. Our empirical experiments provide evidence of the interpretability, forecasting ability, and robustness of the bi-dimensional space.

1. Introduction

In essence, any model to predict distress risk has the goal to analyze a specific informational set *today* to forecast the firm's status *tomorrow*. Although the availability of high dimensional datasets and the advent of new computational algorithms have greatly improved the efficiency in processing such information, the COVID-19 and the Russia–Ukraine conflict raised new issues to deal with. We realized the importance of exploring the sensitivity to critical inputs (e.g., energy and commodity prices) and the mechanisms underlying the supply chain disruptions to better understand the propagation of the industry-specific shocks to the entire system. These new challenges added to common pitfalls the 2007–2009 Global Financial Crisis highlighted in default forecasting, namely delayed adaptation to changes in the state of the economy and limited ability to model complex non-linear interactions between economic, financial, and credit variables (Moscatelli et al., 2020). In such a more complex environment, Machine Learning (ML) techniques offer new ways to deal with unclear and intricate relationships between predictors and outcomes while providing high forecasting accuracy (Brown and Mues, 2012; Chakraborty and Joseph, 2017). However, economists and practitioners are often skeptical about this “black box” approach, which gives, in most cases, accurate outcomes (default prediction) but suffers from three critical points: (1) it lacks a clear economic

substrate; (2) results are often difficult to understand also for experts; (3) performances are unstable and dependent on the data they process.

This paper addresses the aforementioned issues by proposing a novel visual-based approach to detect potential distress and safe zones through dimensional data reduction. By processing firm-specific balance sheet data over time, we provide a mapping from high-to-low dimensional representation space such that the original balance sheet data are reduced to single points on a two-dimensional coordinate plane. We execute such a dimensional data reduction through an unsupervised Deep Learning algorithm (autoencoder), matched with a Nearest Neighbor-based default density estimation, thereby identifying distress zones, i.e., the regions where the risk of future distress is high. These zones appear as homogeneous clusters located in specific space partitions. Moving away from such a cluster corresponds to moving towards increasingly safe zones, where the risk of future distress is low.

Using four large databases containing balance sheet data of more than 120000 Small-Medium-Enterprises (SMEs) over the period 2017–2020, we provide evidence of the robust forecasting ability of our model by estimating popular ML classifiers: k-Nearest Neighbors (KNN), Naïve Bayes (NB), support vector machines (SVM), decision trees (DT) multilayer perceptron (MLP), random forest (RF), AdaBoost (ADA), and XGBoost (XGB).

* Corresponding author.

E-mail addresses: alessandro.giuliani@unica.it (A. Giuliani), roberto.savona@unibs.it (R. Savona), salvatore@unica.it (S. Carta), gianmarco.addari@visioscientiae.com (G. Addari), sebastianpodda@unica.it (A.S. Podda).

¹ These authors contributed equally to this work.

Our proposal aims to support different end-users interested in exploring the *financial sustainability* of non-financial firms, i.e., the complex relationships between firm value, business continuity, and risk dynamics of corporations, which is key to:

- enterprises, to validate budgets, test the resilience of their value chain, and compare their risk and performance with those of their competitors;
- banks, to estimate credit risk and connected capital requirements;
- investment funds, to assess the asset allocation composition in terms of risk exposure and check their investment strategy mandate compliance.

The objective is to better explain, measure, and forecast the corporate credit risk dynamics, being the main driver of financial sustainability for non-financial firms. We do this by providing a model that identifies distress zones for corporations in a straightforward way. We visualize complex outcomes through a map of risk, where each point is a single firm that moves and interacts with others, navigating across safe and distress zones.

Analytically, we compress a balance sheet of a given company in a given year within a bi-dimensional space, combining results over a specific time span and across companies. This procedure thus identifies the company trajectories towards specific risk and safe zones based on the projected statements. The reduced space we realize has high explanatory power and strong predictive ability of the past and future risk trajectories of the companies.

The remainder of this manuscript is as follows. Section 2 discusses the literature review. Section 3 presents the methodology, while Section 4 reports and discusses the results from our empirical experiment. Finally, Section 5 concludes.

2. Related work

The topic of corporate credit risk modeling, although the literature on it is extensive, dating back to Beaver (1966) and Altman (1968), is still challenging and has recently renewed academic interest, especially in assessing the creditworthiness of Small and Medium Enterprises (SMEs) (Modina et al., 2023). The first studies mainly focused on univariate analysis until the '60s, when the first multivariate study appeared (Bellovary et al., 2007). Ever since more than 500 works have been published, covering different methodological approaches, including recent advances pertaining to ML and AI (Shi and Li, 2019). The advent of such new techniques gained strong interest among academic circles and practitioners, as it seems these novel approaches outperform the classical purely statistical ones in financial distress and bankruptcy prediction (Shi et al., 2022). Their strengths lie in several key points: first, such models have great flexibility and scalability, providing high performances, especially when handling large datasets. Furthermore, the non-parametric nature makes them more appropriate for complex non-linear problems. Moreover, they are usually robust when dealing with noisy data. As a whole, the entire literature on default prediction can be related to three main approaches (Savona and Vezzoli, 2012):

- Intensity-based models, in which the default is assumed as a stochastic inaccessible process, usually modeled as a Poisson-like process whose intensity level drives the default time (see Duffie and Singleton, 2012).
- Structural-based models, in which the default is economically modeled as a triggering event based on the so-called distance-to-default, namely the standardized distance between the assets and liabilities of a firm (Merton, 1974). Basically, when the firm's assets fall below the level of the liabilities (default threshold), the company goes bankrupt.
- "Primitive" reduced-form models, which try to statistically predict a firm's defaults without formally connecting the methodology to the theory. These include (Breiman, 2001):

- parametric models, which assume a given stochastic data model generates the data. Early statistical approaches include: (i) Univariate (Beaver, 1966, 1968); (ii) Linear Discriminant Analysis (Altman, 1968), (iii) Logit (Ohlson, 1980) and Probit (Zmijewski, 1984) models, (iv) Multi-class Logit models (Lau, 1987).
- algorithmic models, which instead treat the data mechanism as unknown, no matter about *a-priori* theories on the data generating mechanism. These models are designed to handle nonlinear relationships among predictors while processing many (hundreds if not thousands) of variables without compromising model stability. This class includes (see Jones (2023) and the references herein for a comprehensive overview of the literature on machine learning methods as default prediction models): (i) First generation machine learning models, i.e. neural networks and recursive partitioning; (ii) Classification and regression trees (CARTs), namely binary recursive partitioning using and selecting the more important predictors, used one at the time (based on optimal splitting values); (iii) Advanced machine learning methods, such as Gradient Boosting models, AdaBoost, Random Forests, all using/combining single/weak tree/classifiers; (iv) Deep Learning, namely neural network consisting in many (usually more than three) layers for both inputs and the output; (v) Natural language processing, used to exploit the predictive value of signal embedded in unstructured text data, such as emails, social media posts, corporate annual reports.

By and large, the literature on corporate credit risk modeling explored all three models using accounting and market data. Accounting data relate to financial ratios/balance sheet items and are used as potential predictors of the probability of future distress, being assumed as fundamental variables reflecting the creditworthiness of the firm. As noted in Modina et al. (2023), models using the more relevant accounting variables remain the most widely used methodology for the prediction of default, especially for unlisted companies, although they exhibit some disadvantages (such as the fact that ratios might correlate with each other affecting the estimates). Market data relates to market prices and reflects the compensation required by investors to bear the credit risk of the firm. Moreover, other systematic/systemic variables can be used and relate to common risk factors (e.g., interest rates, equity market indices, GDP, or unemployment rate) assumed to impact the creditworthiness of the firm while not being firm-specific.

According to the aforementioned classification, our work is related to algorithmic model class (primitive reduced-form models) using accounting data. Specifically, we implement a novel ML methodology to predict distress risk based on balance sheet data of non-financial firms.

Accounting data, while being prone to some drawbacks such as opaqueness and consistency (when exploring firms using different methods of accounting or pertaining to different industrial sectors — see Ciampi, 2015), played a key role in early studies (Beaver, 1966; Altman, 1968) and maintain their central importance also in more recent works (Tian et al., 2015) to predict the expected performance of companies, particularly their creditworthiness (i.e., *credit scoring*), over short- and medium-term (e.g., Fuertes-Callén et al., 2022 exploited a set of statistical financial indicators for predicting the survival of Spanish startups by using 5-year balance sheet data). This is why our proposed work uses original balance sheet data (the so-called hard information) since our ambition is to be pioneering in the *exploitation* of historical accounting data to predict future business trends by means of ML approaches.

A vast number of studies have explored the possibility of predicting the risk of bankruptcy of companies by leveraging on the analysis of

accounting data through Machine/Deep Learning approaches (Nazareth and Ramana Reddy, 2023). Among them, Barboza et al. (2017) experimented with several classical ML algorithms for predicting bankruptcy risk from balance sheet data: they show that the *Random Forest* outperforms the other methods tested, and an average gain of ~ 10 percent in accuracy when using ML techniques with respect to traditional approaches. Bussmann et al. (2021) proposed an Explainable Artificial Intelligence (XAI) model to measure the credit risk of small and medium-sized businesses, highlighting that both risky and nonrisky borrowers can be grouped according to a set of similar financial characteristics, which can be used to explain and predict their financial sustainability. On the other hand, Son et al. (2019), by employing both financial statements and basic company information, proposed a method for resolving the skewness typical of financial data, improving the interpretability and accuracy of results by ~ 17 percent. García et al. (2019) proposed adopting ensemble methods to predict bankruptcy, showing how Bagging, Boosting, and Random Forest performed best. Moreover, Sadhwani et al. (2020) apply Deep Learning to investigate the impact of mortgage borrowers in several economic scenarios, finding that highly non-linear relationships correlate with borrowers' behaviors and risk factors. Deep Learning has also been used by Xu and He (2020), which relied on a deep belief network based on the Restricted Boltzmann Machine and classifier SOFTMAX for assessing online supply chain financial credit risk. The model showed an accuracy that is far beyond SVM and Logistic Regression.

In such a literature context, our work intends to open a novel line of research, exploiting interpretable AI methodologies and techniques to assess the financial sustainability of companies, with the goal of better representing their business model and forecasting their future performance.

3. Proposed methodology

As previously stated, the main research goals and innovations of the proposed methodology, detailed in the remainder of this section, consist of the following ones:

- A novel tool for identifying hidden patterns and information not recognizable with the state-of-the-art models.
- The identification of potential threats to financial sustainability in an early stage.
- An innovative data projection performed with a Neural Network-based dimensionality reduction and Nearest Neighbor-based density estimation.
- A visual outcome designed to be understandable and interpretable.
- Human-centered usability: the tool can be used without the need for domain knowledge and expertise.

3.1. System overview

We devised a model aimed at analyzing the balance sheets of companies and projecting them on a bi-dimensional space. The corresponding architecture, depicted in Fig. 1, comprises four main modules: first, the *Data Preprocessing* module cleans and encodes the data in a standardized representation. Afterward, the *Dimensionality Reduction* module transforms the data, projecting them in a reduced space. The module *Density Estimation* computes the failure densities, which are used, together with the reduced data, by the *Projection* module for plotting the corresponding data distribution in the 2D space, in which the distress zones will be highlighted. Hereinafter, we denote the corresponding plot as *Distress Map*.

3.1.1. Data preprocessing

The preprocessing module is aimed at cleaning and standardizing the data. First, irrelevant items are discarded. Indeed, as the bankruptcy process can usually take several years after it is declared distressed, during which the annual balance sheets are nevertheless drawn up, it is possible that, for a given year, although a company has a complete balance sheet, it is actually already bankrupt, and therefore could represent a noisy sample in the dataset. For such companies, we consider the balance sheet of the starting year of the distress process as a key indicator of a failed company, and we removed the balance sheet data of the subsequent years. Furthermore, samples with missing data are removed. Afterward, as the range of raw data values usually varies widely, each item is scaled in the range $[0,1]$ with an appropriate normalization algorithm following the typical requirement of many ML algorithms, including the Neural Networks, that require normalized data. We applied normalization rather than standardization, as financial data typically does not follow a Gaussian distribution (Gauss and Gottingensis, 1821). In particular, we adopt a classical MinMax scaler (Shalabi et al., 2006), which transforms the minimum–maximum values of a feature within the range $0-1$.

3.1.2. Dimensionality reduction

The main step of the proposed method, in terms of novelty and interpretability, consists of the dimensionality reduction module. As previously stated, this module aims to compress the information contained in the k -dimensional space ($\mathbb{R}^k$) of the balance sheet variables into a 2-dimensional space ($\mathbb{R}^2$). The resulting space allows us to represent all the information from the balance sheet of the companies in a visual way, as a point in a Cartesian plane (i.e., the *Distress Map*).

To comply with this step, the literature offers several established pathways. Among the most common ones, there are the *Principal Component Analysis* (PCA) (Karl Pearson, 1901) and the *t-Distributed Stochastic Neighbor Embedding* (t-SNE) (van der Maaten and Hinton, 2008). The PCA technique is a linear dimensionality reduction approach that seeks a linear transformation of the original attributes to obtain new variables, called *principal components*, that express the maximum variance in the data. On the other hand, the t-SNE method performs a non-linear transformation that seeks to preserve the similarity relationships between data points, reducing their dimensionality in a lower-dimensional space where these similarity relationships are preserved as much as possible. In the case of t-SNE, the non-linear transformation used is based on the minimization of an objective function that measures the divergence between the joint probability distribution of the points in the original space and the joint probability distribution of their corresponding points in the lower-dimensional space.

While potentially useful for our purpose, both approaches present some drawbacks that make them not fully suitable for achieving our goals.

In particular, the linear approach of the PCA may be unsuitable in capturing complex relationships between original variables or fail to account for outliers that are common in real-life data sets (Seheult and Green, 1989). Moreover, since balance sheet data are intrinsically correlated with causal relationships, a plain PCA approach seems not appropriate to adequately explain the business model ecosystem of the companies within the two-dimensional space, as the resulting components tend, by construction, to be significantly correlated with the most discriminant variables, then potentially resulting in a biased representation of the company business models. With regard to t-SNE, the main limitation (pointed out by its own author²) stems from the fact that this methodology learns a non-parametric mapping, i.e., it does not learn an explicit function that maps data from the input space to the map. Therefore, it is not possible to embed test points in an existing map without re-generating it. In our case, once the *Distress Map* is

² <https://lvdmaaten.github.io/tsne/>

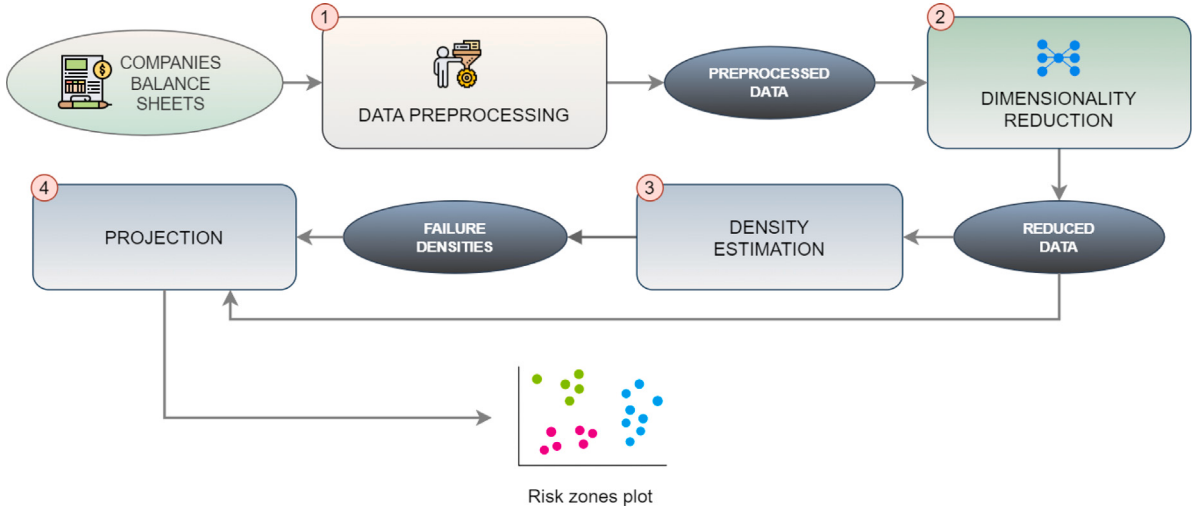


Fig. 1. General architecture of the proposed model.

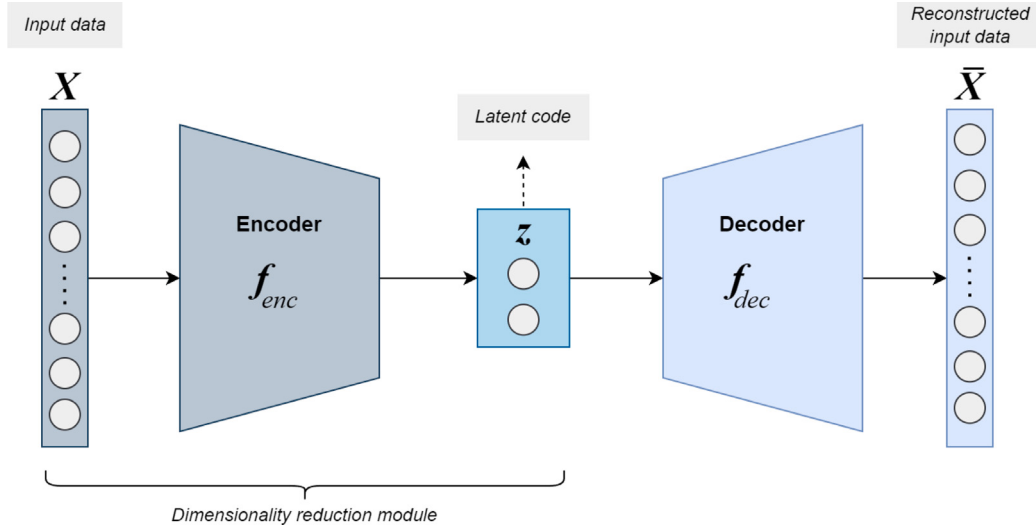


Fig. 2. Schema of the autoencoder.

generated from an existing set of companies, if that set of companies varies over time, the Distress Map appears to be not consistent.

Although some appropriate variations to the original methods can be applied to overcome these limitations, in this work, we propose a different approach, which natively combines the salient features of the two methods described, namely, the ability to extrapolate complex non-linear relationships from balance sheet data and the possibility of applying the generated mapping to new data.

Specifically, the approach we propose is the *autoencoder*, a neural network that learns a compressed representation of the input data by encoding it into a lower-dimensional latent space representation and then decoding it back into the original data. The adopted autoencoder architecture schema is shown in Fig. 2.

Such a network implements both an *encoder* function, denoted as f_{enc} , and a *decoder* function, indicated as f_{dec} . The encoder takes as input data a k -dimensional vector X , producing its compressed l -dimensional representation z (with $l < k$), also known as *latent code*:

$$z = f_{enc}(X)$$

Then, the decoder takes z as input, and reconstructs the output data $\hat{X}$, in the same space of X :

$$\hat{X} = f_{dec}(z)$$

The training process of the autoencoder thus aims to minimize the difference between the input data X and the reconstructed output data $\hat{X}$. To do so, we employ the standard *mean squared error* (MSE) loss function:

$$\mathcal{L} = \frac{1}{n} \sum_{i=1}^n \|X_i - \hat{X}_i\|^2$$

where n is the number of training examples, and X_i and $\hat{X}_i$ are the input and reconstructed output for the i th training example.

Since the purpose of this module, however, is to represent the balance sheet data in the $\mathbb{R}^2$ plane, the decoder is used only for the training phase, while the trained encoder will constitute the mapping function to represent the new data on the obtained (visually interpretable) space, the Distress Map.

3.1.3. Density estimation

The purpose of the Distress Map is not (only) to provide a straightforward visualization of the company statements but rather to provide crucial information about future and potentially risky trends. To this end, we devised an appropriate strategy to identify *distress zones*, i.e., regions with a high risk of future distress. To estimate the distress risk, after investigating many methodological alternatives, we

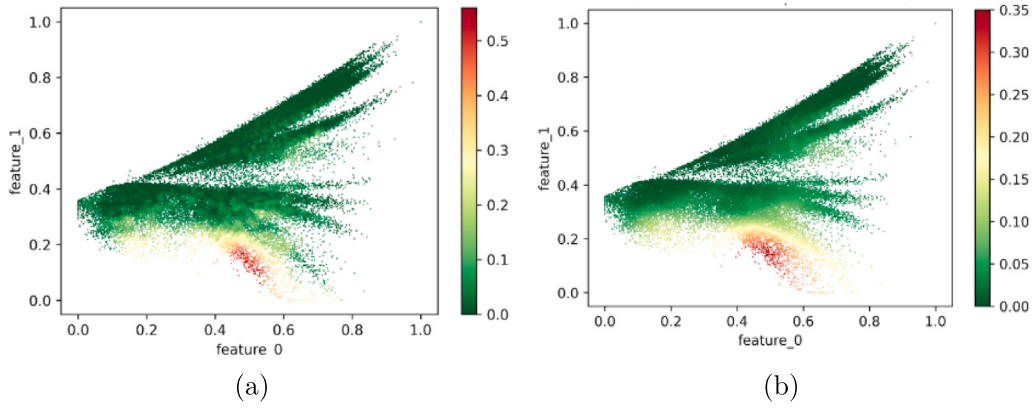


Fig. 3. Example of distress map obtained with $N = 100$ (a) and $N = 1000$ (b).

relied on a Nearest Neighbor-based default density estimation, which is the most effective in highlighting distinct distress zones. The Nearest Neighbor-based algorithm considers the rate of defaulted companies in the neighborhood of the point representing a given company. In detail, given a company i and its corresponding point c_i in the Distress Map at the year t , we consider the status of the N neighbors (including c_i itself) and compute the density δ_i as the ratio between the number of defaulted companies N_f at the year $t + 1$ and N , as reported in Formula (1):

$$\delta_i = \frac{N_f}{N} \quad (1)$$

High values of δ_i correspond to high risk of future distress.

The choice of the optimal N is strictly dependent on the dataset size and data distribution. Different datasets may require a significantly different value of N . Although we are aware of the need for a reliable policy for such a setting, this work is not focused on optimizing the defined Distress Map; rather, the aim is to assess its potential from a general perspective. However, a more in-depth investigation, including a sensitivity analysis to better define a strategy for selecting the optimal N , is planned as future work.

Nevertheless, for the sake of clarity, we provide an example of a Distress Map obtained in two different settings, i.e., with $N = 100$ (Fig. 3a) and $N = 100$ (Fig. 3b).

3.1.4. Projection

The density estimation permits labeling each point with the corresponding distress risk. Therefore, we plot each point with a proper color representing the distress risk in accordance with Formula (1). In particular, we apply a gradient color ranging from dark green (representing the lowest risk, i.e., 0) to red (representing the highest distress value.³). An example of a Distress Map is depicted in Fig. 4

The figure highlights how the distress zone (the red area) is notably defined. To summarize, the closer a company is to a red area, the higher the risk of future distress.

4. Results and discussion

This section reports the performed experiments, aiming to demonstrate the potential of our solution. As we propose a novel approach that presents a new perspective in analyzing and predicting balance sheets, to show the effectiveness of the proposal, we intend to give an answer to each of the following research questions:

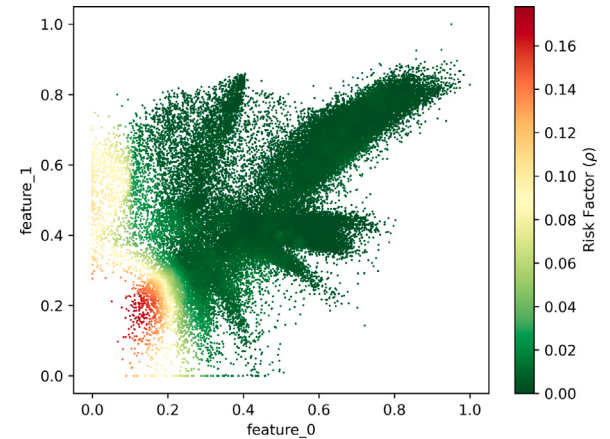


Fig. 4. Example of distress map.

- **Q1:** although the reduced space features lack a clear semantic meaning, can they appropriately express the balance sheet of a given company?
- **Q2:** is the model really effective in identifying and highlighting well-defined distress zones?
- **Q3:** how can we state that the autoencoder-based dimensionality reduction is reliable and robust to unknown data?
- **Q4:** how can we estimate the actual predictive power of the model?
- **Q5:** is the proposal effective on real-world data, which are typically extremely unbalanced?

Before reporting all the experiments, we briefly summarize how we address each question:

- **A1:** to assess the effectiveness of the proposed dimensionality reduction in representing the balance sheets, we perform an exploratory stage in which we analyze (i) the distribution of companies according to their business sizes and (ii) the correlation between the features of the reduced space and the original items;
- **A2:** we (i) train the autoencoder on all selected datasets, (ii) project the training data into the reduced space, (iii) perform the density estimation, and (iv) plot the corresponding Distress Maps, aiming to highlight remarkable distress zones;
- **A3:** we aim to assess the robustness of our tool by splitting each dataset into training and test sets, adopting a Monte Carlo validation approach. The training set is used as reported in point A2, whereas the test companies are projected with the trained model. The goal is to verify that the distress zones identified by

³ Its value depends on the data, the worst case is when all N neighbors will fail, i.e., the distress risk is 1.

the test data are highly similar to the ones highlighted by the training data;

- **A4:** although the Distress Map is based on past data, it can still provide accurate predictions for future company performance, even though we cannot know which companies will fail in advance. To this end, we train and test several classic classifiers, and we plot the Distress Maps obtained from their predictions. The goal is to verify that the more the classifier is effective, the more the predicted Distress Map is compliant with the actual one.
- **A5:** To assess the proposed method, we do not only consider literature datasets, but also we collected real-world datasets, as reported in Section 4.1.

4.1. Datasets

We consider two publicly available literature datasets, named, respectively, *Kaggle Financial Distress* and *Polish Bankruptcy*. Moreover, we adopt four real-world datasets, named *BVD2017*, *BVD2018*, *BVD2019*, and *BVD2020*, respectively. These datasets come from Aida, a popular database handled by the foremost publisher of financial information about companies and enterprises, and contain balance sheet data of non-financial companies in Italy, as we better explain below.

Kaggle financial distress (FD)

The dataset is publicly available,⁴ at the Kaggle portal⁵ a widespread online community platform allowing users involved or interested in Data Science and ML to collaborate. Among all resources, users and institutions may publish their datasets.

The FD dataset deals with predicting financial distress for a sample of companies. Each observation is represented by 83 features, i.e., financial and non-financial characteristics of companies, and one target variable, i.e., the financial distress (φ) the company reaches the following year. If φ , which is a continuous variable, is at most -0.5 , the company is considered *distressed*, a situation that, in most cases, leads the company to bankruptcy.

As our model requires a binary target, we considered a company *active* (class 0) if $\varphi \geq -0.5$. Otherwise, it is considered *distressed* (class 1). For completeness, let us point out that we discarded three meaningless categorical features, keeping 80 features for each sample. The final dataset is composed of 3536 active observations and 136 distressed companies.

Polish bankruptcy (PB)

PB dataset is a subset of the dataset released by Zięba et al. (2016), which contains data about Polish companies in five tasks corresponding to the bankruptcy prediction in the 1st, 2nd, 3rd, 4th, and 5th year.⁶ We adopt the subset containing information about bankruptcy in the following year. Samples are represented with 64 features, being classical financial indicators, and one label, i.e., a binary value that describes if the company failed the following year (value 1) or is still viable (value 0).

Real-world datasets

The previous datasets are built “ad-hoc” for research purposes, and although they are crucial for testing ML algorithms, they cannot reflect the real scenario. Indeed, in a given country or business sector, only a few companies usually fail in a year (typically, less than 1% of total companies fail yearly). With the goal of assessing the validity and potential of our tool on real-world data, we adopt several datasets

Table 1

Number of companies of each real-world dataset.

Dataset	Total	Active	Distressed
BVD2017	128 632	127 585	1047
BVD2018	129 748	128 874	847
BVD2019	129 982	129 361	621
BVD2020	123 161	122 909	252

collected from Aida,⁷ a database created by Bureau van Dijk,⁸ and contains financial information about Italian companies. Each sample is represented by 18 annual balance sheet items of a company, and the label is the binary value indicating if a company failed (class 1) or is active (class 0) the following year. In detail, we collected four different datasets covering the years 2017–2020, each dataset being related to a specific year and composed of the annual balance sheet of companies belonging to the *manufacturing* sector. Let us point out that such data are available from the Bureau Van Dijk portal under a proper license. Due to such restrictions, the generated datasets are not publicly available. Table 1 reports the details of each dataset (the suffix after “BVD” is the considered year).

4.2. Experimental settings

In this section, we briefly report the experimental settings. Let us remark that we carried out an empirical investigation on selecting a suitable set of parameter values for the autoencoder and the density estimation. In detail, we trained the autoencoder in a variety of settings, adjusting the number of layers (ranging from 1 to 3), the loss function (we tested *mean squared error* and *binary cross-entropy*), and the number of training epochs (ranging from 5 to 50). Notably, the behavior remained consistently strong across these variations, with minimum differences in performances, showcasing its versatility, as expected. Let us also remark that we select a suitable value of N for density estimation by preliminarily testing different values, i.e., 10, 100, and 1000.

Dataset Split. As previously reported (answer A3), we split each dataset into training and test sets. In detail, we used 70% and 30% of each dataset as the training and the test set, respectively.

Autoencoder. According to a preliminary hyperparameter tuning, we selected an autoencoder composed of 3 layers for both the encoder and decoder, with a *batch size* of 32, a *mean squared error* loss function, and an *adam* optimizer. For each experiment, we train the model with 10 epochs.

Validation. In this experiment scenario, we opted for a Monte Carlo validation approach, i.e., the dataset was repeatedly split into training and test sets with random shuffling across multiple iterations. Note that this is not unusual in the context of Machine Learning methods applied to financial data since such data often exhibits strong temporal dependencies and autocorrelation. Indeed, unlike traditional cross-validation, which enforces fixed folds, Monte Carlo validation can adapt to varying conditions, leading to more reliable performance estimates. Furthermore, we remark that the number of iterations for each experiment was not predetermined. Instead, a stopping criterion was employed, which halted the process once statistical confidence in the results was achieved.

Dimensionality Reduction. For each dataset, we adopted the corresponding training set to train the autoencoder. Subsequently, we exploited the code layer of each trained model to perform the dimensionality reduction. Finally, we projected the training and test data into the resulting bi-dimensional space for each dataset.

Density Estimation. We adopt the Nearest Neighbor-based algorithm, described in Section 3.1.3, setting up $N = 1000$. Let us remark that the density value δ varies in the range $[0,1]$.

⁴ <https://www.kaggle.com/datasets/shebrahimi/financial-distress>

⁵ <https://www.kaggle.com/>

⁶ <http://archive.ics.uci.edu/ml/datasets/polish+companies+bankruptcy+data>

⁷ <https://www.bvdinfo.com/it-it/le-nostre-soluzioni/dati/nazionali/aida>

⁸ <https://www.bvdinfo.com/>

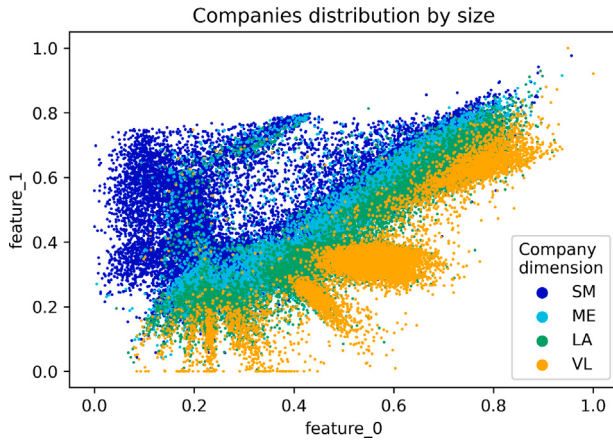


Fig. 5. Distribution of companies by business size (BVD2017 dataset).

4.3. Exploratory analysis

As the intrinsic meaning of the two dimensions in the reduced space may be unclear to a user, we performed a preliminary analysis to confirm that the Distress Map can adequately represent the balance sheets of a company in the bi-dimensional space. The real-world datasets provide information about the business size of a company, e.g., a company is annotated as *small* (SM), *medium* (ME), *large* (LA), or *very large* (VL). Fig. 5 reports the distribution of companies belonging to the BVD2017⁹ dataset, each company being a point in the reduced space colored according to its business size (reported in the legend). Such an analysis has not been possible for the literature datasets, as they have no information about the business size of the companies.

The figure highlights homogeneous clusters according to the business size, providing first evidence that the position in the reduced space may appropriately describe the balance sheet.

Moreover, we also analyzed the correlation between features to confirm our insight, computing the classical Pearson's correlation coefficient (r) (Pearson, 1895). The insight is that if the two features are somehow correlated with one or more balance sheet items, they may be representative of the balance sheet. Fig. 6 reports the heatmap associated with the correlation matrix comprising the pairwise correlations between the features of the BVD2017 dataset, including either the original balance sheet items (see Table A.5 in Appendix A for their meanings) and the reduced space features (represented with *feature_0* and *feature_1*). In particular, r ranges between -1 and 1 , and the stronger the correlation, the larger the absolute value of r .¹⁰

Fig. 6, on the one hand, highlights how *feature_0* is mainly correlated with *EBITDA* and *P_L*. *EBITDA* is a classical item used to assess the profitability of a company, as well as *P_L*, which is the final profit or loss after considering all costs and incomes. On the other hand, *feature_1* is correlated mainly with the variable *LTD*, i.e., the long-term debts. In particular, *feature_0* has a positive correlation with all the aforementioned items, whereas *feature_1* is negatively correlated with *LTD*.

Intuitively, *feature_0* represents the ability of a company to generate profit from its business activities, while *feature_1* is more related to the risk of becoming insolvent. Considering the Distress Map depicted in Fig. 5, we can assert that companies located in the upper right corner should have strong financial sustainability, as they should have high

financial performances and minimum (ideally zero) risk of becoming insolvent. Conversely, companies located in the lower left corner might have weak financial sustainability, as they would have low financial performances and high debts. Intuitively, the latter scenario might lead to high default risk. We investigate this feature in the following sections.

4.4. Density estimation and data projection

Figs. 5 and 6 show how the space is informative about the balance sheets. As said, the main evidence is that companies with higher financial risk tend to concentrate on specific areas (i.e., the distress zones). To check this, we relied on the density estimation, detailed in Section 4.2. In such a doing, for each point c_i representing a company i in the space, we obtain the estimated risk factor ρ_i associated with the corresponding point.

Distress Zones Identification. First, in order to validate the ability to identify distress zones, we plot the Distress Map realized using the training set of each dataset; see research answer A2. Our expectation is that it is possible to identify homogeneous distress zones for each dataset.

Distress Zones Assessment. As the *distress zones identification* is solely based on the data used for training the autoencoder, the main criticism one may state is that the proposed dimensionality reduction model may not handle unknown samples. As a generic autoencoder may deal with out-of-samples, i.e., a trained model may efficiently project new observations without being re-trained (Espadoto et al., 2019), we expect that the Distress Map obtained with the trained autoencoder is stable and reliable for analyzing unknown companies. To this end, the first challenge is to assess whether the distress zones identified by the test data are highly similar to the ones highlighted by the training data for each dataset.

Distress Map Comparisons. Figs. 7–12 report, for each dataset, the comparison between the original Distress Map (a) and the Distress Map obtained for the test set (b). As is clear, the Distress Maps are quite similar in each comparison. In more depth, the underlying experiment highlights that our approach identifies the same distress zones, thus proving that the reduced space is robust and reliable for analyzing new companies.

4.5. Predictions

To check the predictive ability of our Distress Map, we run an out-of-sample exercise contrasting our approach to different alternative classifiers. Indeed, as we already discussed, the Distress Map is in-sample, and then we can compare actual with predicted defaults. Being the developing of a visual tool conveying information on future distress risk our primary aim, we assess the predictive ability of the Distress Map by training several classic classifiers and plotting the Distress Maps obtained with the predictions of each corresponding test set. In light of our preliminary results, we expect that the Distress Map will show distress zones that are compliant with the original ones, and the more robust the classifier is, the more the predicted Distress Map will be similar to the actual one.

4.5.1. Classifiers evaluation

In ML, the classifier assessments can vary depending on the adopted algorithm. In this experiment, we compare the performance of several commonly used classifiers, i.e., k-Nearest Neighbors (KNN), Naïve Bayes classifier (NB), support vector machines (SVM), decision trees (DT) multilayer perceptron (MLP), random forest classifier (RF), Adaboost classifier (ADA), and XGBoost classifier (XGB). Moreover, as the effectiveness of predictions is significantly affected by strongly unbalanced datasets, as less than 1% of companies usually fail in a year, we adopt an oversampling technique on each training set to

⁹ For the sake of brevity, we report here the analysis of one dataset only, but similar assumptions may be made for all the selected real-world datasets.

¹⁰ Let us point out that we removed the variable *IE* (reported in Appendix A) from the correlation matrix, as each pairwise correlation resulted close to 0, meaning that such a particular variable is meaningless in our data.

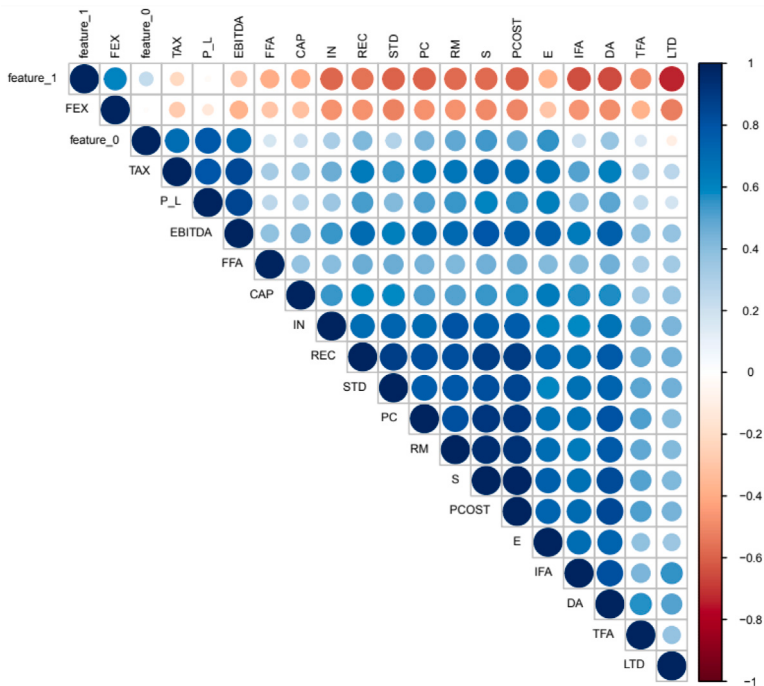


Fig. 6. Feature correlations (BVD2017 dataset).

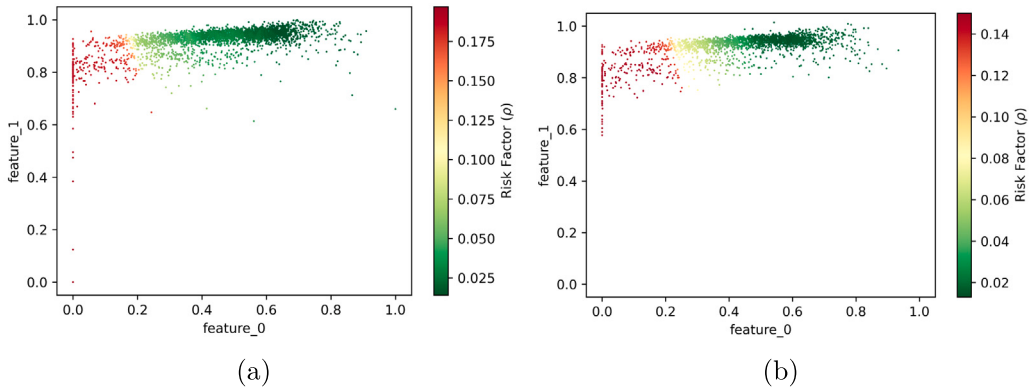


Fig. 7. Distress Map PB — Distress zones identification (a) and assessment (b).

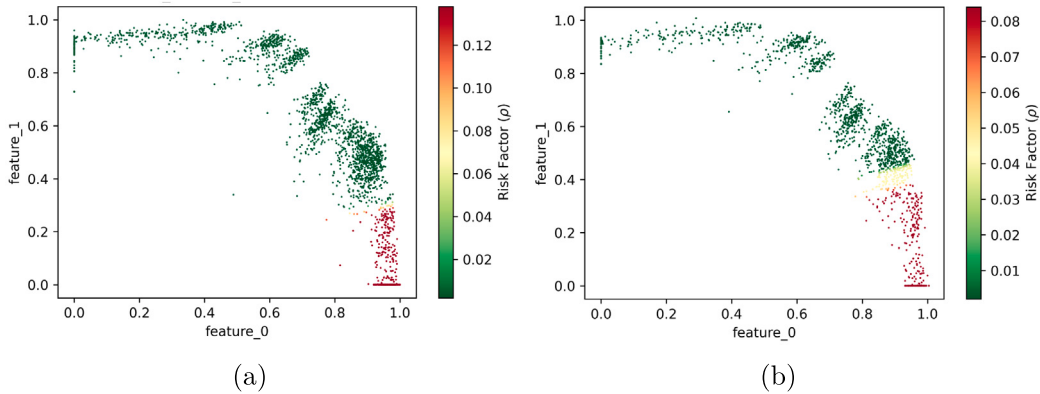


Fig. 8. Distress Map FD — Distress zones identification (a) and assessment (b).

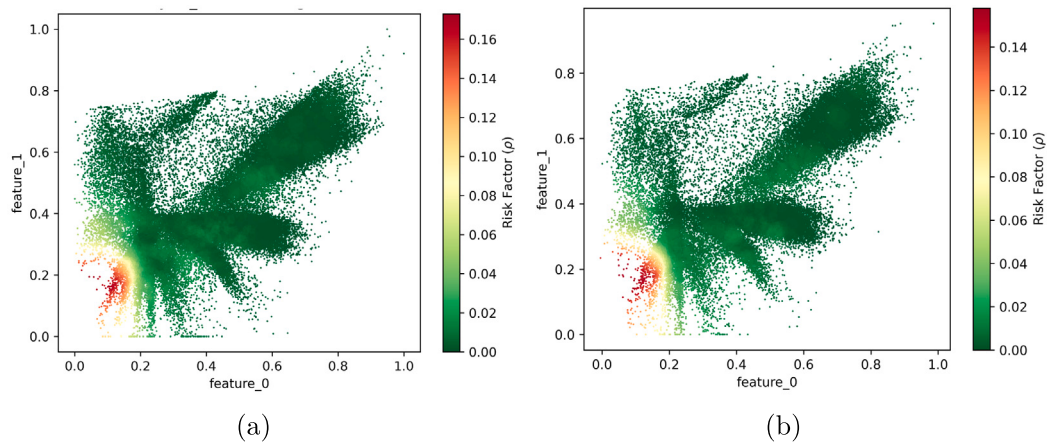


Fig. 9. Distress Map BVD2017 — Distress zones identification (a) and assessment (b).

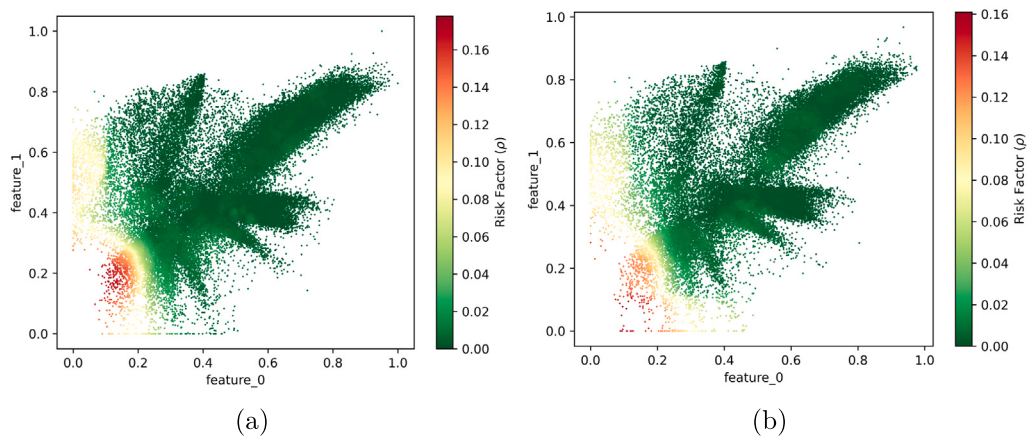


Fig. 10. Distress Map BVD2018 — Distress zones identification (a) and assessment (b).

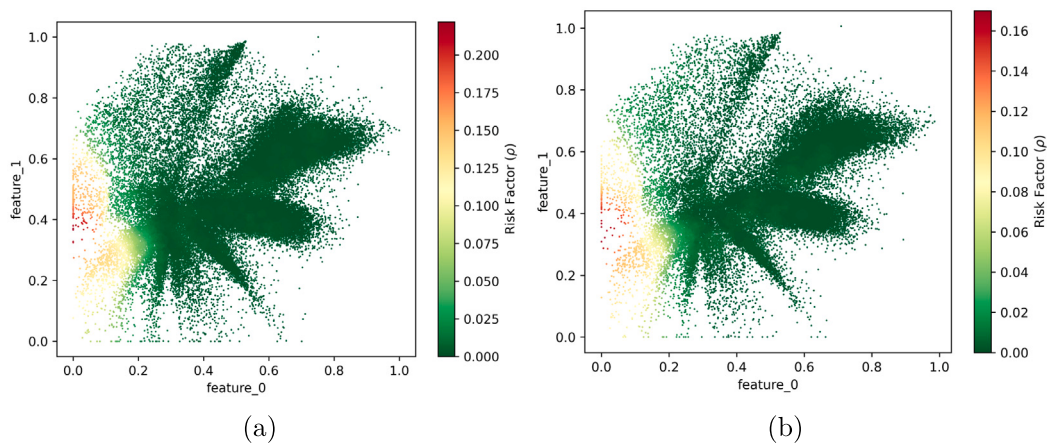


Fig. 11. Distress Map BVD2019 — Distress zones identification (a) and assessment (b).

address this issue. In detail, we adopt the Adaptive Synthetic Sampling approach (ADASYN) (He et al., 2008) to generate new instances in the minority class to balance the data distribution and perform more effective classifier training.

This experiment aims to determine which classifier is the most effective for predicting whether a company will fail. The results of this experiment will provide valuable insights into the impact of predicted data on tool reliability.

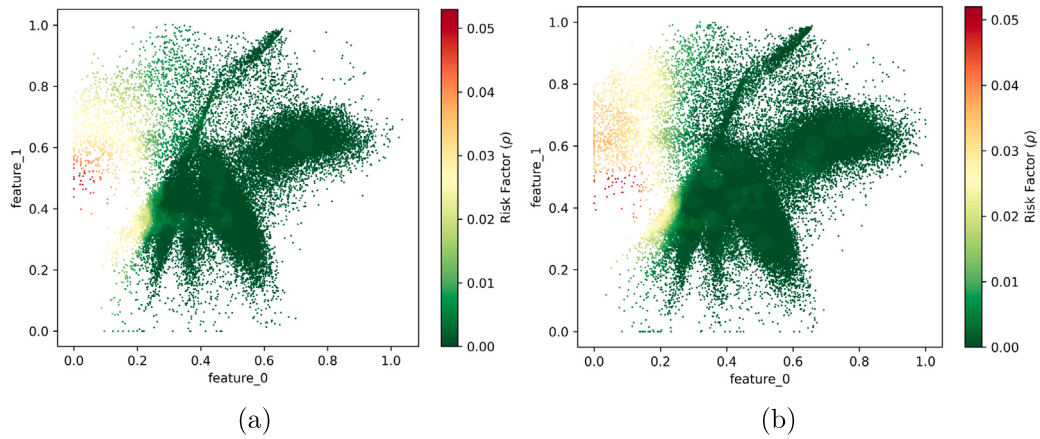


Fig. 12. Distress Map BVD2020 — Distress zones identification (a) and assessment (b).

Table 2 reports the performance of the selected classifiers on each dataset in terms of accuracy. The results highlight how the XGBoost classifier is the best algorithm for each dataset, whereas the Naïve Bayes classifier, especially for real-world datasets, is the less effective.

4.5.2. Predicted data projections

The last part of our experiments aims to plot the Distress Maps conveyed for a given dataset by computing the density estimation of the predictions obtained for each classifier. Hereinafter, we name a Distress Map resulting from a given prediction as *predicted Distress Map*. The goal is to verify whether a highly effective classifier leads to a predicted Distress Map highly similar to the actual one. Fig. 13 depicts the Distress Map corresponding to the actual Distress Map (plot Fig. 13a) and the Distress Maps obtained for each classifier (plots Figs. 13b–13i) when considering the BVD2017 dataset. In this section, we report the comparisons for one dataset for brevity, but similar results have been obtained for all other datasets. However, for completeness, results of the experiments on datasets BVD2018, BVD2019, and BVD2020 are reported in the Appendix B (Figs. B.14, B.15, and B.16, respectively).

The figure shows that the plot most similar to the actual Distress Map is the one depicted in Fig. 13i, i.e., the Distress Map associated with the XGBoost classifier. Conversely, the least similar appears to be the Distress Map shown in Fig. 13b, i.e., the Distress Map associated with the Naïve Bayes classifier. Tangentially, we also provide a quantitative assessment for Distress Map comparisons by computing the following similarity metric (*sim*):

$$\text{sim}(\Gamma_i, \hat{\Gamma}_i) = 1 - \epsilon(\Gamma_i, \hat{\Gamma}_i) \quad (2)$$

where

$$\epsilon(\Gamma_i, \hat{\Gamma}_i) = \frac{1}{N} \sum_{j=1}^N |\delta_j - \hat{\delta}_j| \quad (3)$$

In Formulas (2) and (3), Γ_i and $\hat{\Gamma}_i$ are, respectively, the actual and the predicted Distress Map for a given dataset i with N companies. For each point corresponding to a company j , δ_j and $\hat{\delta}_j$ are the actual and predicted densities. Note that, as the densities range from 0 to 1, also *sim* varies in the same range. In particular, if $\text{sim}(\Gamma_i, \hat{\Gamma}_i) = 1$, Γ_i and $\hat{\Gamma}_i$ are exactly the same Distress Maps.

Table 3 reports the similarities between the actual and each predicted Distress Map.

The table confirms what we inferred by the visual analysis, proving that the most similar predicted Distress Map is associated with the most effective classifier, i.e., the XGBoost, in all our experiments.

Moreover, we also report the corresponding standard deviation σ of the error ϵ in Table 4. Such a metric measures the dispersion of data points. Specifically, a substantial value of σ indicates that the errors can spread far from the mean error, whereas a small σ denotes

Table 2
Classifiers accuracy.

Classifier	Dataset					
	FD	PB	BVD2017	BVD2018	BVD2019	BVD2020
KNN	0.975	0.737	0.885	0.914	0.936	0.968
NB	0.941	0.858	0.689	0.763	0.847	0.811
SVM	0.937	0.870	0.823	0.855	0.900	0.921
DT	0.985	0.873	0.782	0.794	0.889	0.937
MLP	0.991	0.844	0.801	0.895	0.910	0.944
RF	0.965	0.807	0.854	0.847	0.895	0.900
ADA	0.993	0.838	0.834	0.862	0.900	0.938
XGB	0.993	0.935	0.973	0.980	0.988	0.995

Table 3
Similarities between actual and predicted Distress Maps.

Classifier	Dataset					
	FD	PB	BVD2017	BVD2018	BVD2019	BVD2020
KNN	0.988	0.814	0.895	0.922	0.943	0.974
NB	0.949	0.948	0.681	0.763	0.846	0.814
SVM	0.944	0.962	0.834	0.867	0.909	0.933
DT	0.991	0.941	0.791	0.797	0.892	0.945
MLP	0.975	0.888	0.862	0.855	0.905	0.912
RF	0.995	0.915	0.802	0.903	0.916	0.955
ADA	0.982	0.891	0.84	0.869	0.908	0.946
XGB	0.998	0.996	0.986	0.991	0.995	0.998

Table 4
Standard deviation of predicted Distress Maps.

Classifier	Dataset					
	FD	PB	BVD2017	BVD2018	BVD2019	BVD2020
KNN	0.001	0.072	0.093	0.099	0.092	0.055
NB	0.012	0.062	0.349	0.329	0.307	0.312
SVM	0.011	0.04	0.211	0.203	0.212	0.173
DT	0.004	0.035	0.193	0.186	0.212	0.107
MLP	0.002	0.08	0.209	0.221	0.213	0.194
RF	0.002	0.079	0.222	0.193	0.205	0.141
ADA	0.029	0.05	0.151	0.152	0.143	0.1
XGB	0.001	0.002	0.019	0.023	0.015	0.006

low dispersion around the mean. In other words, σ is a simple yet informative robustness metric of a predicted Distress Map.

As we note from the results reported in the table, the XGBoost-based predicted Distress Map is the most robust alternative relative to all other competitors.

To summarize, the XGBoost permits to obtain a more precise and robust predicted Distress Map. As the XGBoost is always, in our experiments, the most accurate classifier, this behavior confirms the insight that the better a classifier performs, the more the predicted Distress

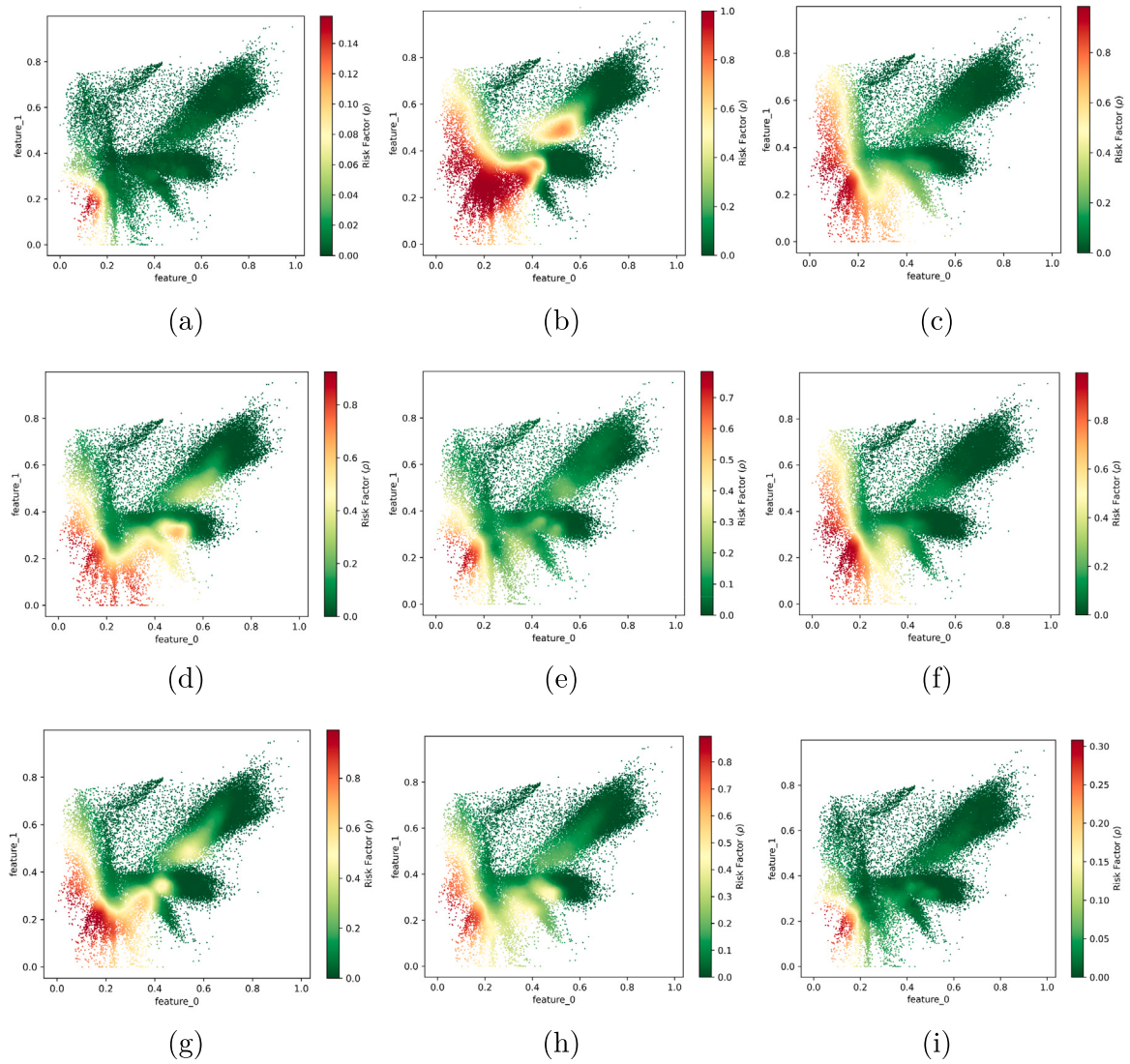


Fig. 13. Predicted Distress Maps for BVD2017 dataset: (a) Actual Distress Map, (b) Naive Bayes, (c) SVM, (d) Decision Tree, (e) KNN, (f) MLP, (g) Random Forest, (h) ADA Boost, (i) XGBoost.

Map is reliable for estimating the balance sheet of a company. As a whole, all these outcomes provide clear evidence of the predictive ability power of our proposed approach.

4.6. A comparison with key state-of-the-art ML default prediction models

Although our analysis (and related findings) has been executed to provide evidence on the explainability of the distress map as an effective and interpretable tool for financial risk analysis based on accounting data, rather than proving the risk prediction ability of our method only, in this section, we briefly discuss and compare our results with some key studies exploring the role of machine learning methods when using accounting indicators together with additional data (credit behavioral indicators, market data, soft information-based data).

The first is [Bacham and Zhao \(2017\)](#), which provides evidence that using a broader set of variables to predict defaults leads to a higher accuracy ratio, regardless of the methodology (machine learning vs. generalized additive model (GAM) framework¹¹). Specifically,

¹¹ In which nonlinear transformations of each variable are executed by assigning weights next combined into a single score.

when adding loan behavioral information to financial indicators, they improve the accuracy ratio by 8 to 10 percentage points.

Another key study is [Barboza et al. \(2017\)](#), in which the authors find a substantial improvement in prediction accuracy through machine learning techniques when using original Altman's Z-score variables together with complementary financial indicators. In more depth, random forest led to 87% accuracy, whereas logistic regression and Linear Discriminant Analysis led to 69% and 50% accuracy, respectively.

Finally, we compared with [Moscatelli et al. \(2020\)](#), in which the performance of a set of machine learning models in predicting default risk are explored relative to standard statistical methods (e.g., Logit) using a large dataset covering financial and credit behavioral indicators for Italian non-financial manufacturing firms. They show a predictive ability ranging from 72% (Linear Discriminant Analysis) to 77.3% (Gradient Boosted Trees) when using only accounting data, which increases on average for about ten percentage points when credit behavioral indicators are included.

On the other hand, our approach (see [Table 2](#)) exhibits a default predictive ability higher than 90% on average (using different classifiers over all six datasets used in the empirical experiment), with XGBoost the top performer. Moreover, and assessing a more suitable accuracy measure for our approach based on the similarity between the predicted and actual distress map, the forecasting ability of our approach

exceeds 99%. Interestingly, note that these results are obtained using accounting data only. In other terms, while we process a limited set of information relative to other studies using financial indicators and additional data, our proposal significantly outperforms the standard statistical methods as well as the machine learning models used in the literature (Bacham and Zhao, 2017; Barboza et al., 2017; Moscatelli et al., 2020). This finding seems to suggest that what matters more in credit risk prediction might be the way (method) with which we process the data, instead of the size of the informational set *per se*: balance sheet data, while lagged by construction with limited time frequency and *minimal* (at first glance), when properly treated could offer more information than one might reasonably expect.

5. Conclusions

In this manuscript, we propose an innovative visual-based approach for assessing the balance sheet of a company by identifying potential distress and safe zones. The main contribution of our proposal is a visual model that can address the most critical issues of the current state-of-the-art Machine Learning algorithm for analyzing and assessing balance sheets. Actually, the outcome of such models is usually difficult to understand even for experts, and the performance of the ML algorithms is often unstable and depends on the data used in the analysis. Our tool provides an easily interpretable, understandable, and robust visual outcome to support different end-users (e.g., companies, banks, investors) interested in exploring the risk dynamics of corporations.

In detail, the tool first analyzes the firm-specific balance sheet data and provides a mapping from high-to-low dimensional representation space through a Deep Learning-based dimensionality reduction approach; specifically, we use the encoder layer of a proper autoencoder to define the reduced space. Afterward, a Nearest Neighbor-based is defined for computing the default density estimation to identify the distress zones, i.e., the regions where the risk of future distress is high. Such regions appear as homogeneous clusters located in distinct space partitions. Moving away from such zones, the risk of future distress becomes low.

We performed several experiments on different datasets to assess the robustness of our approach. First, to check the usefulness of the autoencoder-based dimensionality reduction for representing the balance sheets, we performed a preliminary experiment in which we analyzed the distribution of companies and the correlation between the features of the reduced space and the original balance sheet items. We then plotted all the corresponding Distress Maps, each one showing remarkable distress zones. Next, we compared actual with predicted Distress Maps through a horse race of state-of-the-art ML classifiers. We found that the more the classifier is effective, the more the predicted Distress Map complies with the actual one, thereby confirming that the Distress Map also has significant predictive power.

Although our proposal is still in the preliminary stage, our results encourage further investigations. To this end, we plan to improve the approach in order to solve some limitations of our methodology, thereby providing a more robust and reliable tool. In particular, the crucial point is that the reduced space is obtained by only considering the balance sheet of a given company in a single year. A more reliable analysis should consider the time series together with the cross sections of the balance sheet data, as well as global economic/financial risk factors (e.g., interest rates, inflation, GDP, unemployment, etc.).

Moreover, the industrial sector in which a company operates is an informative factor alike. Future work will explore companies from all industrial sectors, not just manufacturing, as we did in this work.

CRedit authorship contribution statement

Alessandro Giuliani: Writing – original draft, Software, Methodology, Investigation, Conceptualization. **Roberto Savona:** Writing – original draft, Supervision, Conceptualization. **Salvatore Carta:** Supervision, Conceptualization. **Gianmarco Addari:** Writing – review & editing, Software. **Alessandro Sebastian Podda:** Writing – original draft, Visualization, Methodology, Investigation, Conceptualization.

Declaration of competing interest

The authors declare no conflict of interest in this manuscript. All authors confirm that this work has not been published or submitted elsewhere simultaneously.

Appendix A. Real-world datasets: balance sheet items description

Table A.5

List of all balance sheet items.

Feature Acronym	Description
IFA	<i>Intangible Fixed Assets:</i> assets not having a physical form but being valuable for their potential to generate future revenue or provide a competitive advantage, e.g., patents, copyrights, trademarks, software, brand names, customer lists, or goodwill
TFA	<i>Tangible Fixed Assets:</i> assets that have a physical value, e.g., business premises, equipment, inventory, or machinery
FFA	<i>Financial Fixed Assets:</i> assets comprised of money, a contractual right to receive money or other financial assets from another party, securities issued by another enterprise
IN	<i>Inventories:</i> all the raw materials, work in progress, and goods available for sale that a company owns
REC	<i>Receivables:</i> the amount of money due to a company to provide services. These are open invoices that the customers have to pay for a service.
CAP	<i>Share Capital:</i> the money a company raises by issuing common or preferred stock
E	<i>Net Worth:</i> balance sheet quantity obtained from the difference between assets and liabilities
STD	<i>Short-Term Debts:</i> the short-term obligations that are expected to be paid within one year, such as accounts payable
LTD	<i>Long-Term Debts:</i> debts a company owes third-party creditors that are payable beyond 12 months.
S	<i>Sales:</i> the money generated from normal business operations, calculated as the average sales price times the number of units sold
PCOST	<i>Production Costs:</i> the direct and indirect costs businesses face from manufacturing a product or providing a service
FEX	<i>Financial Expenses:</i> all positive and negative components of the economic result for the year related to the financial area of business management
TAX	<i>Taxes:</i> total Current, deferred, and prepaid income taxes.
P_L	<i>Profit & Loss:</i> the difference between a company's revenues and costs. If this difference is positive, it is called <i>profit</i> ; otherwise, it is called <i>loss</i> .
EBITDA	<i>Earning Before Interest, Tax, Depreciation & Amortization:</i> a further widely used measure of corporate profitability.
DM	<i>Direct (Raw) Material Cost:</i> the value of raw materials owned by the company.
PC	<i>Personnel Cost.</i>
DA	<i>Depreciation & Amortization:</i> non-cash expenses representing the cost of using capital assets (depreciation) and intangible assets (amortization).

Appendix B. Predicted distress maps projection

This section reports the results of the experiments on datasets *BV D2018*, *BV D2019*, and *BV D2020* (Figs. B.14, B.15, and B.16, respectively).

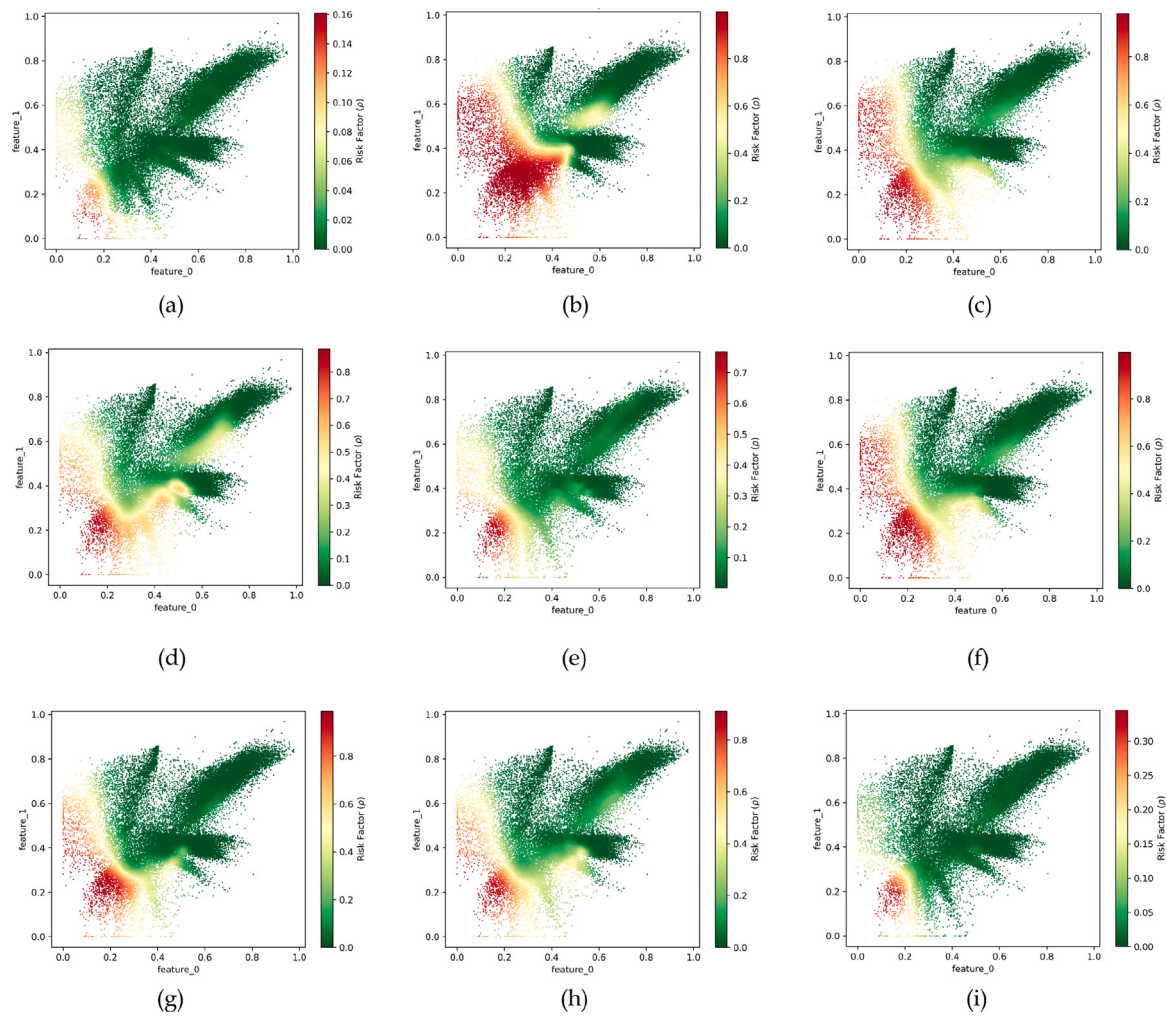


Fig. B.14. Predicted distress maps for BVD2018 dataset: (a) Actual Distress Map, (b) Naive Bayes, (c) SVM, (d) Decision Tree, (e) KNN, (f) MLP, (g) Random Forest, (h) ADA Boost, (i) XGBoost.

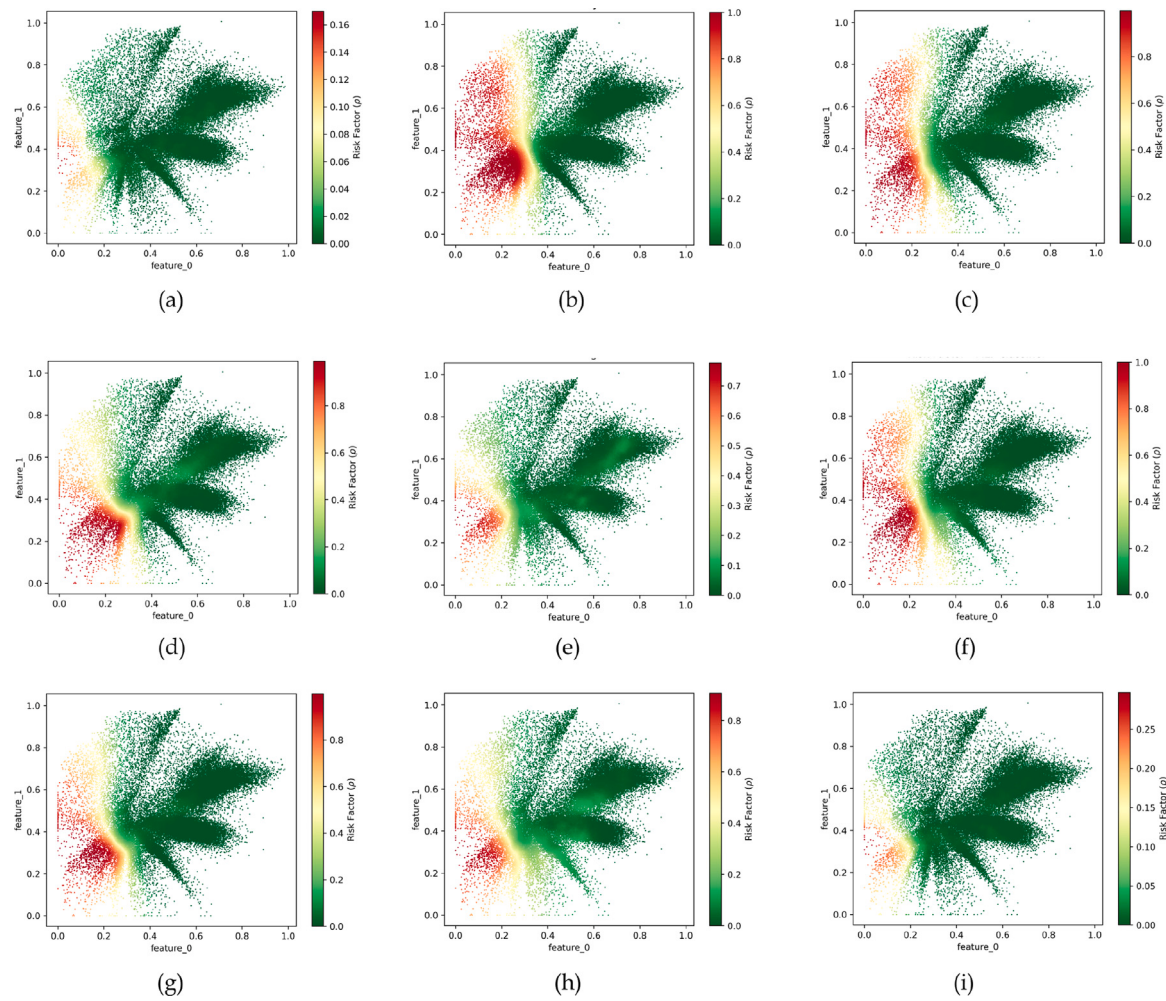


Fig. B.15. Predicted distress maps for BVD2019 dataset: (a) Actual Distress Map, (b) Naive Bayes, (c) SVM, (d) Decision Tree, (e) KNN, (f) MLP, (g) Random Forest, (h) ADA Boost, (i) XGBoost.

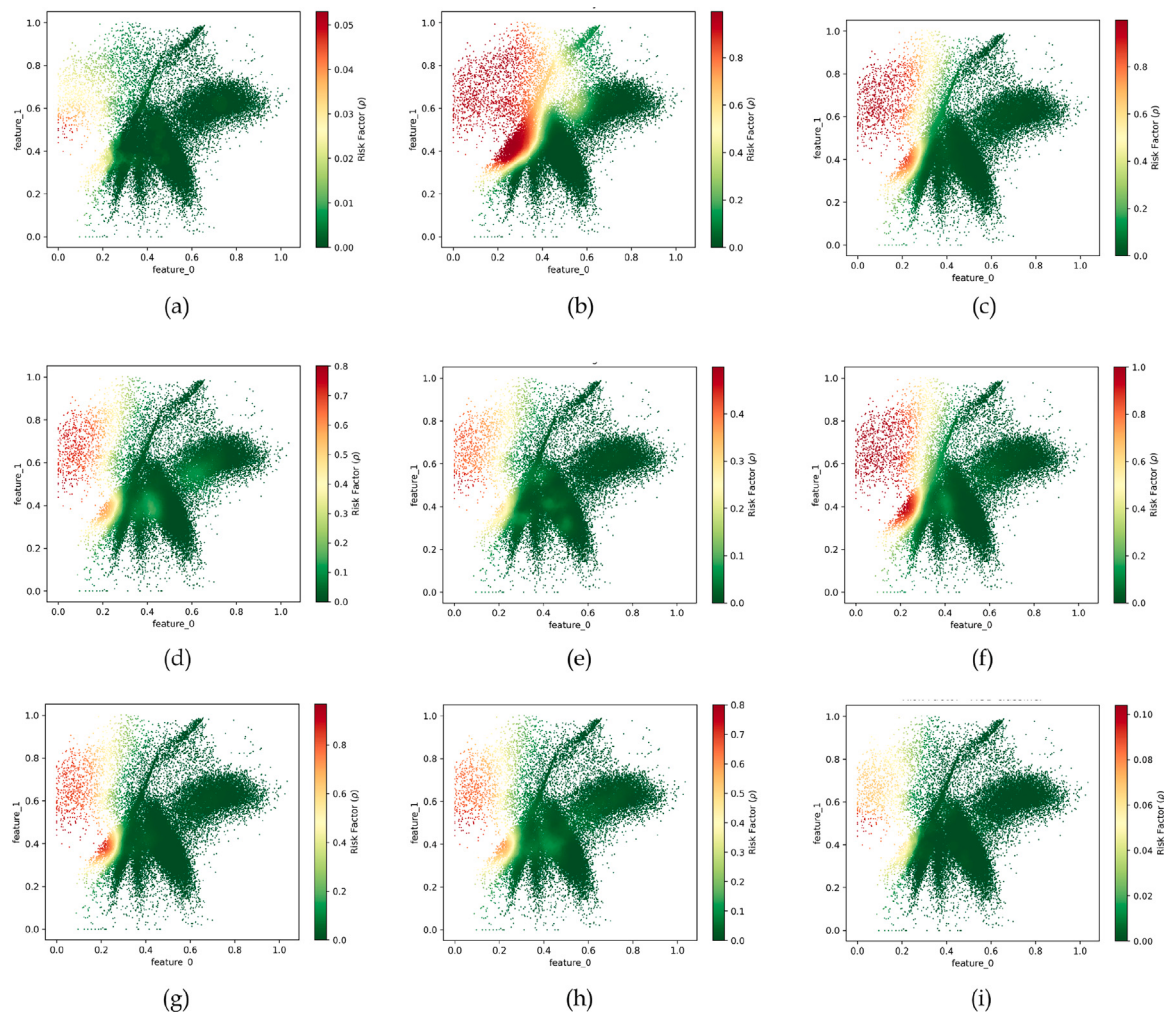


Fig. B.16. Predicted distress maps for BVD2020 dataset: (a) Actual Distress Map, (b) Naive Bayes, (c) SVM, (d) Decision Tree, (e) KNN, (f) MLP, (g) Random Forest, (h) ADA Boost, (i) XGBoost.

Data availability

The authors do not have permission to share data.

References

- Altman, E.I., 1968. Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *J. Finance* 23 (4), 589–609. URL: <https://www.jstor.org/stable/2978933>.
- Bacham, D., Zhao, J., 2017. Machine learning: challenges, lessons, and opportunities in credit risk modeling. *Moody's Anal. Risk Perspect.* 9, 30–35.
- Barboza, F., Kimura, H., Altman, E., 2017. Machine learning models and bankruptcy prediction. *Expert Syst. Appl.* 83, 405–417. <http://dx.doi.org/10.1016/j.eswa.2017.04.006>, URL: <https://www.sciencedirect.com/science/article/pii/S0957417417302415>.
- Beaver, W.H., 1966. Financial ratios as predictors of failure. *J. Account. Res.* 4, 71–111, URL: <http://www.jstor.org/stable/2490171>.
- Beaver, W.H., 1968. Market prices, financial ratios, and the prediction of failure. *J. Account. Res.* 6 (2), 179–192, URL: <http://www.jstor.org/stable/2490233>.
- Bellovary, J., Giacomino, D., Akers, M., 2007. A review of bankruptcy prediction studies: 1930 to present. *J. Financ. Educ.* 33 (WINTER), 1–42, URL: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-75949099657&partnerID=40&md5=39992813703986539653c4be3ff65a5f>.
- Breiman, L., 2001. Statistical modeling: The two cultures. *Statist. Sci.* 16 (3), 199–215. <http://dx.doi.org/10.1214/ss/1009213726>, URL: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-0000245743&doi=10.1214%2fs%2f1009213726&partnerID=40&md5=e57ba46e0cc6eb0ac6513fe9f6aee7d5>, Cited by: 3022; All Open Access, Bronze Open Access.
- Brown, I., Mues, C., 2012. An experimental comparison of classification algorithms for imbalanced credit scoring data sets. *Expert Syst. Appl.* 39 (3), 3446–3453. <http://dx.doi.org/10.1016/j.eswa.2011.09.033>, URL: <https://www.sciencedirect.com/science/article/pii/S095741741101342X>.
- Bussmann, N., Giudici, P., Marinelli, D., Papenbrock, J., 2021. Explainable machine learning in credit risk management. *Comput. Econ.* 57, 203–216, URL: <https://doi.org/10.1007/s10614-020-10042-0>.
- Chakraborty, C., Joseph, A., 2017. Machine Learning at Central Banks. Bank of England Working Papers 674, Bank of England, URL: <https://ideas.repec.org/p/boe/boewp/0674.html>.
- Ciampi, F., 2015. Corporate governance characteristics and default prediction modeling for small enterprises. An empirical analysis of Italian firms. *J. Bus. Res.* 68 (5), 1012–1025. <http://dx.doi.org/10.1016/j.jbusres.2014.10.003>, URL: <https://www.sciencedirect.com/science/article/pii/S0148296314003221>.
- Duffie, D., Singleton, K.J., 2012. Credit Risk: Pricing, Measurement, and Management. Princeton University Press, pp. 1–396, URL: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84923953954&partnerID=40&md5=23f8f73fce65d5cb62c7cb34fd288774>.
- Espadoto, M., Martins, R.M., Kerren, A., Hirata, N.S.T., Telea, A.C., 2019. Toward a quantitative survey of dimension reduction techniques. *IEEE Trans. Vis. Comput. Graphics* 27, 2153–2173, URL: <https://doi.org/10.1109/tvcg.2019.2944182>.
- Fuertes-Callén, Y., Cuellar-Fernández, B., Serrano-Cinca, C., 2022. Predicting startup survival using first years financial statements. *J. Small Bus. Manag.* 60 (6), 1314–1350. <http://dx.doi.org/10.1080/00472778.2020.1750302>.
- García, V., Marqués, A.I., Sánchez, J.S., 2019. Exploring the synergetic effects of sample types on the performance of ensembles for credit risk and corporate bankruptcy prediction. *Inf. Fusion* 47, 88–101. <http://dx.doi.org/10.1016/j.inffus.2018.07.004>, URL: <https://www.sciencedirect.com/science/article/pii/S1566253517308011>.
- Gauss, C., Gottingensis, S.R.S., 1821. *Theoria Combinationis Observationum: Erroribus Minimis Obnoxiae*. URL: <https://books.google.it/books?id=4BrCswEACA AJ>.

- He, H., Bai, Y., Garcia, E.A., Li, S., 2008. ADASYN: Adaptive synthetic sampling approach for imbalanced learning. In: 2008 IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence). pp. 1322–1328, URL: <https://doi.org/10.1109/IJCNN.2008.4633969>.
- Jones, S., 2023. A literature survey of corporate failure prediction models. *J. Account. Lit.* 45 (2), 364–405, URL: <https://www.emerald.com/insight/content/doi/10.1108/JAL-08-2022-0086/full/html>.
- Karl Pearson, F., 1901. LIII. On lines and planes of closest fit to systems of points in space. *Lond. Edinb. Dublin Philos. Mag. J. Sci.* 2 (11), 559–572. <http://dx.doi.org/10.1080/14786440109462720>.
- Lau, A.H.-L., 1987. A five-state financial distress prediction model. *J. Account. Res.* 25 (1), 127–138, URL: <http://www.jstor.org/stable/2491262>.
- van der Maaten, L., Hinton, G., 2008. Visualizing data using t-SNE. *J. Mach. Learn. Res.* 9 (86), 2579–2605, URL: <http://jmlr.org/papers/v9/vandermaaten08a.html>.
- Merton, R.C., 1974. On the pricing of corporate debt: The risk structure of interest rates. *J. Finance* 29 (2), 449–470, URL: <http://www.jstor.org/stable/2978814>.
- Modina, M., Pietrovito, F., Gallucci, C., Formisano, V., 2023. Predicting SMEs' default risk: Evidence from bank-firm relationship data. *Q. Rev. Econ. Finance* 89, 254–268. <http://dx.doi.org/10.1016/j.qref.2023.04.008>, URL: <https://www.sciencedirect.com/science/article/pii/S1062976923000558>.
- Moscatelli, M., Parlapiano, F., Narizzano, S., Viggiano, G., 2020. Corporate default forecasting with machine learning. *Expert Syst. Appl.* 161, 113567. <http://dx.doi.org/10.1016/j.eswa.2020.113567>, URL: <https://www.sciencedirect.com/science/article/pii/S0957417420303912>.
- Nazareth, N., Ramana Reddy, Y.V., 2023. Financial applications of machine learning: A literature review. *Expert Syst. Appl.* 219, 119640. <http://dx.doi.org/10.1016/j.eswa.2023.119640>, URL: <https://www.sciencedirect.com/science/article/pii/S0957417423001410>.
- Ohlson, J.A., 1980. Financial ratios and the probabilistic prediction of bankruptcy. *J. Account. Res.* 18 (1), 109–131, URL: <http://www.jstor.org/stable/2490395>.
- Pearson, K., 1895. Note on regression and inheritance in the case of two parents. *Proc. Royal Soc. Lond.* 58, 240–242, URL: <http://www.jstor.org/stable/115794>.
- Sadhwani, A., Giesecke, K., Sirignano, J., 2020. Deep Learning for Mortgage Risk*. *J. Financ. Econom.* 19 (2), 313–368. <http://dx.doi.org/10.1093/jfinec/nbaa025>, URL: <https://doi.org/10.1093/jfinec/nbaa025>.
- Savona, R., Vezzoli, M., 2012. Multidimensional distance-to-collapse point and sovereign default prediction. *Int. J. Intell. Syst. Account. Financ. Manage.* 19 (4), 205–228, [arXiv:https://onlinelibrary.wiley.com/doi/pdf/10.1002/isaf.1332](https://onlinelibrary.wiley.com/doi/pdf/10.1002/isaf.1332), URL: <https://onlinelibrary.wiley.com/doi/abs/10.1002/isaf.1332>.
- Seheult, A., Green, P., 1989. Robust Regression and Outlier Detection. Wiley Online Library, URL: <https://onlinelibrary.wiley.com/doi/book/10.1002/0471725382>.
- Shalabi, L.A., Shaaban, Z., Kasasbeh, B., 2006. Data mining: A preprocessing engine. *J. Comput. Sci.* 2 (9), 735–739. <http://dx.doi.org/10.3844/jcssp.2006.735.739>, URL: <https://thesaispub.com/abstract/jcssp.2006.735.739>.
- Shi, Y., Li, X., 2019. An overview of bankruptcy prediction models for corporate firms: A systematic literature review. *Intangible Cap.* 15 (2), 114–127. <http://dx.doi.org/10.3926/ic.1354>, URL: <https://www.intangiblecapital.org/index.php/ic/article/view/1354>.
- Shi, S., Tse, R., Luo, W., D'Addona, S., Pau, G., 2022. Machine learning-driven credit risk: A systemic review. *Neural Comput. Appl.* 34 (17), 14327–14339. <http://dx.doi.org/10.1007/s00521-022-07472-2>.
- Son, H., Hyun, C., Phan, D., Hwang, H., 2019. Data analytic approach for bankruptcy prediction. *Expert Syst. Appl.* 138, 112816. <http://dx.doi.org/10.1016/j.eswa.2019.07.033>, URL: <https://www.sciencedirect.com/science/article/pii/S0957417419305123>.
- Tian, S., Yu, Y., Guo, H., 2015. Variable selection and corporate bankruptcy forecasts. *J. Bank. Financ.* 52, 89–100, URL: <https://doi.org/10.1016/j.jbankfin.2014.12.003>.
- Xu, R., He, M., 2020. Application of deep learning neural network in online supply chain financial credit risk assessment. In: 2020 International Conference on Computer Information and Big Data Applications. CIBDA, pp. 224–232. <http://dx.doi.org/10.1109/CIBDA50819.2020.00058>.
- Zięba, M., Tomczak, S.K., Tomczak, J.M., 2016. Ensemble boosted trees with synthetic features generation in application to bankruptcy prediction. *Expert Syst. Appl.* 58, 93–101. <http://dx.doi.org/10.1016/j.eswa.2016.04.001>, URL: <https://www.sciencedirect.com/science/article/pii/S0957417416301592>.
- Zmijewski, M.E., 1984. Methodological issues related to the estimation of financial distress prediction models. *J. Account. Res.* 22, 59–82, URL: <http://www.jstor.org/stable/2490859>.