



HOKKAIDO UNIVERSITY

Title	Bankruptcy prediction model based on business risk reports : Use of natural language processing techniques
Author(s)	Rasolomanana, Onjaniaina Mianin' Harizo
Citation	經濟學研究, 72(1), 47-56
Issue Date	2022-06-09
Doc URL	https://hdl.handle.net/2115/85951
Type	departmental bulletin paper
File Information	40_ES_72(1)_047.pdf



Bankruptcy prediction model based on business risk reports: Use of natural language processing techniques

Rasolomanana Onjaniaina Mianin'Harizo

1. Introduction

Many studies have proposed bankruptcy prediction models using statistical (Altman, 1968; Altman et al., 1977; Ohlson, 1980; Shirata, 2003) or machine learning models (Rasolomanana, 2021; Shin et al., 2005; Shirata, 2019; Tam & Kiang, 1992). Although some researchers attempted to incorporate non-financial information, such as industry or number of employees, previous studies have mostly focused on quantitative information, namely financial ratios, to analyze the financial situation of companies. Calderon & Cheh (2002) summarized thoroughly the inputs and outputs used in previous studies. In practice, qualitative information is as much important and insightful for such analysis. Thus, it is necessary to implement a model using qualitative information as input variables (Chung, 2014).

Qualitative information can be internal disclosures, such as managerial reports, and external, such as macroeconomic data from news articles. Previous studies have shown that reports relative to auditor's opinion can be useful when it comes to bankruptcy prediction (Matin et al., 2019). In the case of Japanese companies, audit reports do not include information related to going concern assumption. Managers formulate this assumption, along other risks information, and are reported in the business risks section of the securities report. Therefore, this study presents the hypothesis that the risk of bankruptcy within the next year can be evaluated using information relative to business risk, from training a machine learning model.

For computers to learn features in textual data, it is first necessary to process the texts and extract their features, converting the textual data into their representative numerical values. Natural Language Processing (NLP) is a technology that allows computers to derive insights from human languages. Most NLP applications rely on machine learning methods to mine documents. Support Vector Machines (SVM) have been amongst the most popular techniques (Kim et al., 2005; Shin et al., 2005), followed by Naïve Bayes and neural networks (Amani & Fadlalla, 2017). The present study will be focused on the latter. Neural networks are powerful learning models that have gained in popularity as the quantity of data increases constantly making it more complex to process (Goldberg, 2016). The performance of the neural network model will then be compared to the performances of SVM and Naïve Bayes.

The present research answers the questions: Is there a positive-negative polarity in business risk reports? Is the sentiment in the text extracted from NLP techniques linked to the financial situation, thus, the risk of bankruptcy? Which machine learning method performs best on data from Japanese listed companies?

We propose a bankruptcy prediction model for Japanese listed companies using qualitative information from managerial reports related to business risks. The performance of the model will be evaluated based on how accurately it classifies the texts as bankrupted or non-bankrupted. The objectives of this study are, therefore, to (1) find out if business risk reports are useful in predicting bankruptcy, and (2) to build a machine learning model that can classify texts.

2. Related works

2.1. Studies related to bankruptcy prediction using NLP and machine learning

Financial statements have long been the number one source of information to assess the financial situation of a company. However, companies are required to disclose different reports regularly, which constitutes another source of information. Texts are unstructured data that contain insights. Although some prior works have developed machine learning models using texts as input variables, it is still not has been discussed actively so far, making it underexploited. Human language is very complex and can sometimes be ambiguous even to humans, let alone to computers. NLP is, in fact, a rigorous series of many tough tasks, but it is a powerful tool for interpreting textual data.

Using the techniques of NLP, some studies (see Table 1) have shown that information from financial news brings insights not found in financial quantitative variables (Cerchiello et al., 2017, 2018). (Matin et al., 2019; Muñoz-Izquierdo et al., 2020) found out that statements from auditors and managements also contribute to the prediction of distress. Other works showed text segments in business management reports, such as Management Discussion and Analysis (MD&A), help detect financial distress (Ahmadi et al., 2019; Mai et al., 2019).

The main task of sentiment analysis is to identify inner expression in the text (Ragini et al., 2018). In the present study, we demonstrate that texts segments related to the business risks contain decisive

Table 1: Literature related to bankruptcy prediction using NLP and machine learning

Article	Data		Text processing method	Results
	Textual data	Origin		
Cerchiello et al. (2017, 2018)	News article	European banks	Semantic vectors	Relative usefulness= 13%
Matin et al. (2019)	Auditors' reports and managements' statements	Denmark	Word embeddings and Convolutional Neural network	AUC and log score • NN_{aud}: 0.844 and 0.1064 • NN_{man}: 0.836 and 0.1078 • $NN_{aud+man}$: 0.843 and 0.1070
Ahmadi et al. (2019)	Annual business reports	Germany	DSCNN	Accuracy • SVM = 0.7679 • CNN = 0.5712 • LTSM = 0.6549 • DSCNN = 0.8414
Mai et al. (2019)	MD&A	US	Word embeddings and Convolutional Neural network	Accuracy and AUC • DL-embedding = 0.568 and 0.784 • DL-CNN = 0.428 and 0.714 • Logistic Regression = 0.434 and 0.717 • SVM = 0.422 and 0.71 • Random Forest = 0.733 and 0.716

Source: prepared by the author

Note: The relative usefulness is a measure of the relative performance gain of the model compared to a perfect model (Sarlin, 2013). If relative usefulness is equal to 1, the model loss is equal to 0, meaning that the model is perfect. The Area Under the receiver operating characteristics Curve (AUC) is a metric that tells how much the model can distinguish between classes. Dependency Sensitive Convolutional Neural network (DSCNN) consists of a convolutional layer built on top of two LSTM networks, because after filtering, the texts are still too long for one-layer- CNN, making it difficult to capture dependencies.

information that defines the sentiment of the text. Therefore, two sentiments are considered:

- Negative sentiment, meaning that the risk of bankruptcy is high, and
- Positive sentiment, meaning that bankruptcy is unlikely to happen within the next year.

In Japanese securities report, the risk information section is separated from the Japanese equivalent of the MD&A section, implying that both carry complementary but different information. Also, unlike previous research, sentiment analysis based on TF-IDF (Term Frequency - Inverse Document Frequency) scores is performed in this work. This is because instead of understanding the relationship between the words, we want to capture the lexical features that are descriptive of bankruptcy or non-bankruptcy; and TF-IDF scores measure the importance of each word in the corpus.

2.2. Machine learning techniques in NLP

Many machine learning techniques are used along NLP tools, but the most used in previous studies dealing with bankruptcy prediction are the SVM, Neural networks and Naïve Bayes (Qu et al., 2019).

Support Vector Machines, introduced by Cortes et al. (1995), are widely used supervised machine learning models for binary classification problems. They are easily interpretable and can work well even with a small data set. Its kernel function assigns a hyperplane that best divides a dataset into two classes, by transforming the data into a high-dimensional one. This is particularly useful for nonlinear data sets.

Neural networks and deep learning are powerful tools that learn by mistakes. They have become popular in both academic and practical applications. The classical Neural network is a group of multiple neurons organized in layers. It can learn linear and non-linear functions, making it a proper choice when the relationship between input and output is complex. Another typical model is the Recurrent Neural network (RNN) which is commonly used in stock price prediction since the RNN is suitable for time series analysis, where sequence is key. It is also used in text classification. Another major model in deep learning is the Convolutional Neural network (CNN), generally used for image recognition, as it was initially designed for 3-dimensional data, later successfully experimented on sequential data as well.

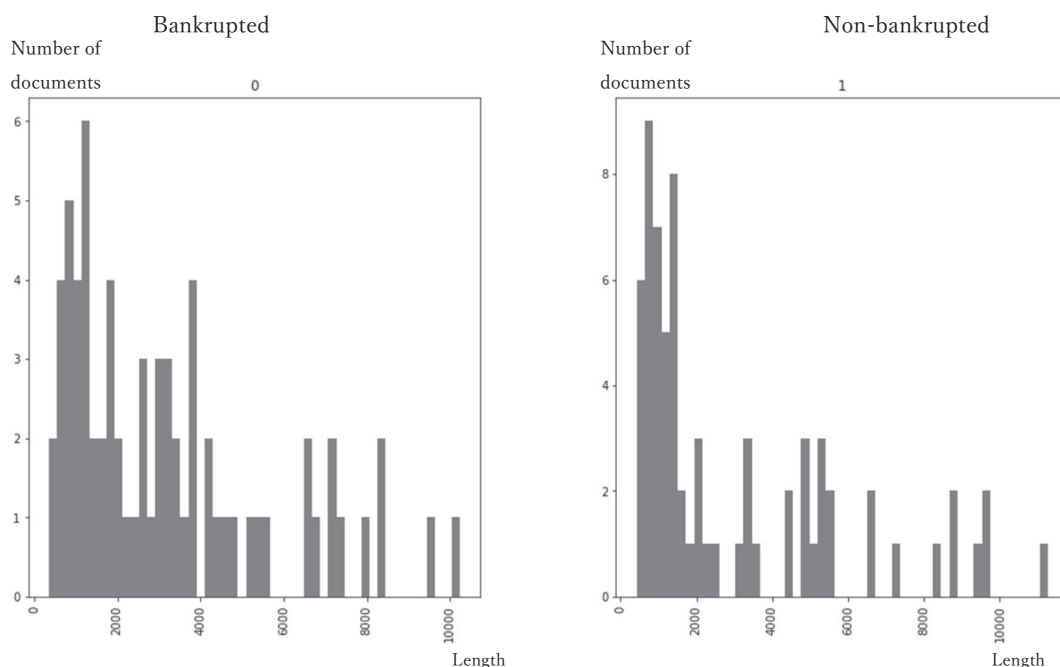
Naïve Bayes is also a widely used learning algorithm, it assumes that all predictors are independent. Naïve Bayes classifiers are computationally efficient and easy to implement. It assigns a probability that a given word or text is positive or negative. Naïve Bayes classifiers typically need lots of training examples to perform well.

Previous studies have shown that the performance of a model is not independent of the used data. It relies heavily on the representation of the data (Goodfellow et al., 2016). A model can perform well on a certain dataset but can have a poor performance on a different dataset. Therefore, in the present study, a benchmark of the performance of the three above-mentioned machine learning techniques is carried out.

3. Text processing and analysis

3.1. Data

The data used in this study is the business risks section from securities reports, imported from EOL Database. The total sample of 138 companies from fourteen industries includes 69 bankrupted and 69 non-bankrupted companies equivalent in terms of asset size, industry, and time frame of collected data.



Source: prepared by the author

Figure 1: Length of documents in both classes

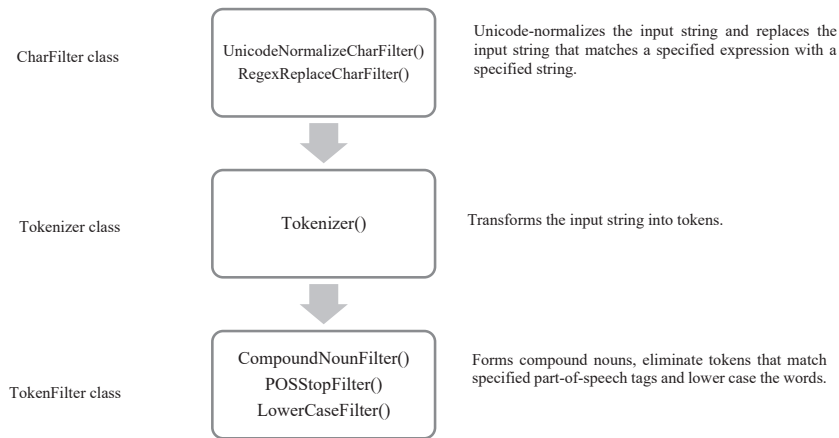
The period ranges from fiscal year 2004 to fiscal year 2017.

3.2. Model construction

Before feeding the data to the model, it is necessary to preprocess the texts first by the means of NLP techniques. The texts are written in Japanese language. Unlike English or German, the words are not separated with a space. Besides, in Japanese, the characters are not limited to alphabet, numbers, and symbols; in addition to those, there are three kinds of characters in Japanese language: kanji, hiragana and katakana. Therefore, the morphological analysis is more complex. In the present study, we used the pure python package Janome version 0.4.1 (Uchida, 2020)¹⁾, which is a Japanese morphological analysis engine (also called tokenizer) including the built-in dictionary and the language model. It uses mecab-ipadic-2.7.0-20070801 as the built-in dictionary.

The first processing step is to clean the texts using the Analyzer framework of Janome (Uchida, 2020). To do so, it is necessary to isolate each word or compound word. Each word or compound word constitutes a token. This implies removing characters that do not bear any useful information, such as punctuations, symbols, or other useless characters (numberings, etc.), and stop words. Stop words are languages that are not useful for the analysis. It can be general words from the language itself (such as “the”, “to”, “I”, etc.), or domain-based, which are vocabularies related to accounting, management, or

1) Janome is an analyzer framework for Japanese language. It is written only in Python, and since we used Python as the programming language in this study, it has the advantage of being easily installed, compared to other analyzers.



Source: prepared by the author

Figure 2: Analyzer framework of Janome

finance (such as “corporation”, “financial statements”, etc.).

The stop words used in the present study was from a programming library called SlothLib, developed by Ohshima et al. (2007) added by domain-based words. Non distinctive words, which are words common to both bankrupted companies' reports and non-bankrupted companies' reports were ignored, hence, removed from the corpus.

The next step is to transform the words into numerical values that mathematical models can understand. The values assigned to each word will be its TF-IDF score. TF-IDF evaluates the originality of a word by analyzing how relevant it is to a collection of documents. TF-IDF does not consider the words order or the relationship between words, it is generally used as a lexical feature.

It is calculated by multiplying term frequency and inverse document frequency:

$$w_{i,j} = tf_{i,j} \times \log \frac{N}{df_i}$$

$tf_{i,j}$ = Number of occurrences of i in j

df_i = Number of documents containing i

N = Total number of documents

In the present study, the model will be trained and evaluated using three (3) classifiers: neural networks, Support Vector Machines and Naïve Bayes.

The classification was performed through a fully connected neural network, using the Keras library. The parameters of the model were chosen through a trial-and-error process²⁾. Hence, the neural network has one (1) hidden layer only.

2) The choice of the parameters and hyperparameters of the model was done based on the best performance. The model showed the most sensitivity and variation in performance when changing the optimization function.

To minimize the error rate of the model in the prediction, optimization function is used. The optimizer chosen is 'rmsprop', and the loss function corresponds to 'binary_crossentropy'. Loss function is used to estimate the error of the model during the training. Using the error backpropagation (Rumelhart et al., 1988), the weights in particular layer are updated in such manner, that the error rate decreases in following evaluation. We used Binary Cross-Entropy (BCE), which can be defined as:

$$\text{BCE} = -(y \log \hat{y} + (1 - y) \log (1 - \hat{y}))$$

Where y is the target value (0 for bankrupted and 1 for non-bankrupted), and $\hat{y}$ is the predicted probability of the class to be 1.

Neural network models have many parameters, and overfitting can easily occur. Overfitting can be alleviated to some extent by regularization. Regularization is any modification we make to a learning algorithm that is intended to reduce its generalization error but not its training error (Goodfellow et al., 2016). A recently proposed alternative regularization method is "dropout" (Srivastava et al., 2014). The dropout method is designed to prevent the network from learning to rely on specific weights. It works by randomly dropping a pre-determined percentage of the neurons in the network (or in a specific layer) in each training example. Ignoring, or "dropping-out" of specific neurons can prevent their over-adaptation, which could lead to over-fitting. The dropout technique is effective in NLP applications of neural networks. It is necessary to set the parameter defining the probability of selection of several neurons to drop out from the network. In this study, we applied a dropout of 0.5 rate in the first layer. Along the dropout technique, we used another common regularization technique called early stopping. This means stopping the training before overfitting occurs but not too early to so that the network has learnt something.

Another way to check if the network generalizes well is to test it. Therefore, a set of the data is hold out for testing (pairs of (*input*, *label*)) but not used in the training. However, it is also necessary to check how well the network has trained, so it is necessary to reserve a third set of data for cross-validation. This third set will be called validation set. Hence, the data was split into training ($n = 92$), validation ($n = 20$) and test ($n = 25$) sets, with respect to industry proportions. If the number of samples in the industry is enough to be split following a 65/10/25 ratio, then those samples were distributed randomly amongst the three sets. The validation set is also useful to determine when to stop the training before it overfits. As we train the network, the validation set helps determine how well the model is generalizing, and at some point, the validation error will start increasing indicating that the network has started to learn the noise and inaccuracies in the data and the function. And that is when early stopping is used (Marsland, 2009).

The output activation function is the sigmoid function (Leshno et al., 1993). It is mostly used because of its non-linearity and simplicity of the computation. The function is defined as:

$$f(x) = \frac{1}{(1+e^{-x})}$$

The network was compared with SVM and Naïve Bayes. The data was split into training ($n = 112$) and test ($n = 25$) sets for both. For the SVM, the kernel function used is the linear kernel. The Naïve

Bayes algorithm was the Multinomial Naïve Bayes.

4. Results and discussion

The point of machine learning is that the algorithm must perform well on new previously unseen inputs, not just on the data used for training. This ability is called generalization. The new data corresponds to the testing set, which is why the evaluation of the models is based on the results using this testing set.

There are several metrics to measure the performance of a machine learning model. The choice of the metrics depends on the domain and the data. In the previous literature, some researchers have used AUC or F1-score, which are useful especially if the dataset is imbalanced. In the present study, since the dataset is balanced (same sample size for bankrupted and non-bankrupted classes) and since in bankruptcy prediction it is more important to spot the risk of bankruptcy than not spot it, the performance of each model was evaluated using accuracy rate (overall percentage of correct classifications), type 1 error rate (percentage of bankrupted companies mistakenly classified by the model as non-bankrupted) and type 2 error rate (non-bankrupted companies mistakenly classified by the model as bankrupted). The results are summarized in Table 2.

Previous studies have shown that there is no universal model for all data, the performance of the model is closely related to each data.

In the present case, the neural network had the best performance. A possible cause is related to the quality of the data. The classes (bankrupted class and non-bankrupted class) are balanced as well. As Batista et al. (2004) argued, learning from imbalanced datasets might be difficult. In the present study, the classes are evenly distributed not only in terms of sample size (68 and 69 samples in each class) but also in terms of the data itself since the samples in the non-bankrupted class are the equivalent of the samples in bankrupted class in terms of industry proportion, data period and asset size of the company. Therefore, it is possible that the network could clearly distinguish and learn the characteristics of each class. Moreover, if the model is trained and validated with samples from service industry only, it will not perform well on a testing set with only data from manufacturing industry, since they do not share the same inherent features. Thus, although the split of the data into three sets is usually done entirely randomly, here, the proportions per industry were controlled so that each set includes samples from each industry. In practice, using different methods depending on the trend of the data may be more suitable though without using such controlled sampling. However, according to the Principles Regarding the Disclosure of Narrative Information published by the Financial Services Agency, to enhance appropriate corporate disclosure practices in Japan, narrative information from annual securities reports is required

Table 2: Summary of results

N = 12 + 13	Confusion matrix	Accuracy	Type 1 error	Type 2 error
Neural networks	$\begin{bmatrix} 10 & 2 \\ 0 & 13 \end{bmatrix}$	92.00%	16.67%	0%
SVM	$\begin{bmatrix} 8 & 4 \\ 0 & 13 \end{bmatrix}$	84.00%	33.33%	0%
Naïve Bayes	$\begin{bmatrix} 7 & 5 \\ 0 & 13 \end{bmatrix}$	80.00%	41.67%	0%

Source: prepared by the author

to be expanded for fiscal year ending on March 31, 2020. Therefore, disclosure practices will be more collective, and quality textual data are expected to become more accessible, making neural networks an exponentially adequate method for bankruptcy prediction using textual data.

The choice of the features might also constitute another reason to why neural networks performed the best. In text classification, semantic quality and statistical quality are generally the main concerns (Sulo et al., 2003; Zhang et al., 2011). Semantic quality refers to the ability of the indexing method to analyze the relationship between words. Statistical quality refers to the ability to classify the term in the domain it belongs to. TF-IDF measures the rarity of each word in a document; and as the objective is to retrieve the terms that contain bankruptcy-related sentiment, TF-IDF helps determine lexical features for bankrupted and non-bankrupted classes.

5. Conclusion

Every year, companies disclose large amounts of textual data, which requires time and effort to process for the human being, hence, the need to involve intelligent models to perform such task efficiently. The present study introduced the use of machine learning and NLP techniques using unstructured qualitative data (business risk information).

The results demonstrated that there is a positive-negative polarity in the business risk section in securities reports since we can effectively classify bankrupted companies from non-bankrupted companies based on the related texts. This suggests that reports regarding business risks carry predictive information, and that the sentiment in the text is linked to the imminent risk of bankruptcy.

We also showed that machine learning techniques can effectively analyze the sentiment in the text. We compared three common techniques in text analysis, and found out that neural networks performed the best, followed by SVM, and Naïve Bayes. However, the implications of these results are closely related to the data from Japanese listed companies used in the present study but are not fully representative of all Japanese companies.

Our findings also showed that TF-IDF, although not as sophisticated as features extracted through word embeddings or LTSM (Long-Term Short Memory), represent an efficient basic metric to extract the most descriptive terms of bankruptcy. This study also compared different classifiers, showing that neural networks are better classifier in this context: business risk reports as data and TF-IDF scores as features. However, neural networks remain computationally expensive and not as transparent and straightforward to explain as SVM or Naïve Bayes.

This study presents other limitations, which prompt future investigations. First, the size of the sample is limited, making it less likely to be representative of all Japanese listed companies. Moreover, models such as neural networks and Naïve Bayes usually perform better the more examples it learns. Second, quantitative, or qualitative information only is not sufficient when evaluating the financial situation of company in practice. Therefore, future work should combine both in the same model, and our study constitutes a useful prior step in building such comprehensive model.

References

- Ahmadi, Z., Martens, P., Koch, C., Gottron, T., & Kramer, S. [2019]. Towards bankruptcy prediction: Deep sentiment mining to detect financial distress from business management reports. *Proceedings - 2018 IEEE 5th International Conference on Data Science and Advanced Analytics, DSAA 2018*, 293–302. <https://doi.org/10.1109/DSAA.2018.00040>
- Altman, E. I. [1968]. Financial Ratios, Discriminant Analysis, and the Prediction of Corporate Bankruptcy. *Source: The Journal of Finance*, 23 (4), 589–609.
- Altman, E. I., Haldeman, R. G., & Narayanan, P. [1977]. ZETA Tin* ANALYSIS A new model to identify bankruptcy risk of corporations. In *Journal of Banking and Finance* (Vol. 1).
- Amani, F. A., & Fadlalla, A. M. [2017]. Data mining applications in accounting: A review of the literature and organizing framework. *International Journal of Accounting Information Systems*, 24, 32–58. <https://doi.org/10.1016/j.accinf.2016.12.004>
- Batista, G. E. A. P. A., Prati, R. C., & Monard, M. C. [2004]. A study of the behavior of several methods for balancing machine learning training data. *ACM SIGKDD Explorations Newsletter*, 6 (1), 20–29. <https://doi.org/10.1145/1007730.1007735>
- Calderon, T. G., & Cheh, J. J. [2002]. A roadmap for future neural networks research in auditing and risk assessment. *International Journal of Accounting Information Systems*, 3 (4), 203–236. [https://doi.org/10.1016/S1467-0895\(02\)00068-4](https://doi.org/10.1016/S1467-0895(02)00068-4)
- Cerchiello, P., Nicola, G., Ronnqvist, S., & Sarlin, P. [2017]. *Deep learning bank distress from news and numerical financial data*.
- Cerchiello, P., Nicola, G., Rönqvist, S., & Sarlin, P. [2018]. Deep Learning for Assessing Banks' Distress from News and Numerical Financial Data. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3292485>
- Chung, W. [2014]. BizPro: Extracting and categorizing business intelligence factors from textual news articles. *International Journal of Information Management*, 34 (2), 272–284. <https://doi.org/10.1016/j.ijinfomgt.2014.01.001>
- Cortes, C., Vapnik, V., & Saitta, L. [1995]. Support-Vector Networks Editor. In *Machine Learning* (Vol. 20). Kluwer Academic Publishers.
- Goldberg, Y. [2016]. A primer on neural network models for natural language processing. *Journal of Artificial Intelligence Research*, 57, 345–420.
- Goodfellow, I., Bengio, Y., & Courville, A. [2016]. Deep learning. *MIT press*. <http://www.deeplearningbook.org>
- Kim, H., Howland, P., & Park, H. [2005]. Dimension Reduction in Text Classification with Support Vector Machines. *Journal of Machine Learning Research*, 6, 37–53. <https://doi.org/10.1021/bi702018v>
- Leshno, M., Lin, V. Y., Pinkus, A., & Schocken, S. [1993]. Multilayer feedforward networks with a nonpolynomial activation function can approximate any function. *Neural Networks*, 6 (6), 861–867. [https://doi.org/10.1016/S0893-6080\(05\)80131-5](https://doi.org/10.1016/S0893-6080(05)80131-5)
- Mai, F., Tian, S., Lee, C., & Ma, L. [2019]. Deep learning models for bankruptcy prediction using textual disclosures. *European Journal of Operational Research*, 274 (2), 743–758. <https://doi.org/10.1016/j.ejor.2018.10.024>
- Marsland, S. [2009]. *Machine Learning – An Algorithmic Perspective*.
- Matin, R., Hansen, C., Hansen, C., & Mølgaard, P. [2019]. Predicting distresses using deep learning of text segments in annual reports. *Expert Systems with Applications*, 132, 199–208. <https://doi.org/10.1016/j.eswa.2019.04.071>
- Muñoz-Izquierdo, N., Laitinen, E. K., Camacho-Miñano, M. del M., & Pascual-Ezama, D. [2020]. Does audit report information improve financial distress prediction over Altman's traditional Z-Score model? *Journal of International Financial Management and Accounting*, 31 (1), 65–97. <https://doi.org/10.1111/jifm.12110>
- Ohlson, J. A. [1980]. Financial Ratios and the Probabilistic Prediction of Bankruptcy. *Journal of Accounting Research*, 18 (1), 109. <https://doi.org/10.2307/2490395>
- Ohshima, H., Nakamura, S., & Tanaka, K. [2007]. *SlothLib: A Programming Library for Researches on the Web*.
- Qu, Y., Quan, P., Lei, M., & Shi, Y. [2019]. Review of bankruptcy prediction using machine learning and deep learning

- techniques. *Procedia Computer Science*, 162 (Itqm 2019), 895–899. <https://doi.org/10.1016/j.procs.2019.12.065>
- Ragini, J. R., Anand, P. M. R., & Bhaskar, V. [2018]. Big data analytics for disaster response and recovery through sentiment analysis. *International Journal of Information Management*, 42, 13–24. <https://doi.org/10.1016/j.ijinfomgt.2018.05.004>
- Rasolomanana, O. [2021]. Corporate Bankruptcy Prediction Using Neural Networks: Case of Japanese companies. *Nenpō Zaimu Kanri Kenkyū (Japan Financial Management Association)*.
- Rumelhart, D. E., Hinton, G. E., & Williams, R. J. [1988]. Learning Representations by Back-Propagating Errors. *MIT Press: Cambridge, MA, USA*, 696–699.
- Sarlin, P. [2013]. *On policymakers' loss functions and the evaluation of early warning systems* (Issue 15).
- Shin, K. S., Lee, T. S., & Kim, H. J. [2005]. An application of support vector machines in bankruptcy prediction model. *Expert Systems with Applications*, 28(1), 127–135. <https://doi.org/10.1016/j.eswa.2004.08.009>
- Shirata, C. Y. [2003]. *Predictors of Bankruptcy after Bubble Economy in Japan: What can you learn from Japan case?* <https://www.researchgate.net/publication/281090087>
- Shirata, C. Y. [2019]. *Bankruptcy prediction model and Corporate rating using AI technology*. Zeimukeiri Kyokai CO., LTD.
- Srivastava, N., Hinton, G., Krizhevsky, A., & Salakhutdinov, R. [2014]. Dropout: A Simple Way to Prevent Neural Networks from Overfitting. In *Journal of Machine Learning Research* (Vol. 15).
- Sulo, Maria, J., & Hidalgo, G. [2003]. *Text Representation for Automatic Text Categorization*.
- Tam, K., & Kiang, M. [1992]. Managerial applications of neural networks: the case of bank failure prediction. *Management Science*, 38(7), 926–947.
- Uchida, T. [2020]. *Janome version 0.4.1 documentation*. <https://mocobeta.github.io/janome/en/>
- Zhang, W., Yoshida, T., & Tang, X. [2011]. A comparative study of TF*IDF, LSI and multi-words for text classification. *Expert Systems with Applications*, 38(3), 2758–2765. <https://doi.org/10.1016/j.eswa.2010.08.066>