

Krichene, Aida

Article

Using a naive Bayesian classifier methodology for loan risk assessment: Evidence from a Tunisian commercial bank

Journal of Economics, Finance and Administrative Science

Provided in Cooperation with:

Universidad ESAN, Lima

Suggested Citation: Krichene, Aida (2017) : Using a naive Bayesian classifier methodology for loan risk assessment: Evidence from a Tunisian commercial bank, Journal of Economics, Finance and Administrative Science, ISSN 2218-0648, Emerald Publishing Limited, Bingley, Vol. 22, Iss. 42, pp. 3-24,
<https://doi.org/10.1108/JEFAS-02-2017-0039>

This Version is available at:

<https://hdl.handle.net/10419/179782>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by/4.0/>

Using a naive Bayesian classifier methodology for loan risk assessment

Naive
Bayesian
classifier
methodology

Evidence from a Tunisian commercial bank

Aida Krichene

Department of Accounting, IHEC Carthage, Tunis, Tunisia

3

Received 7 February 2016
Revised 3 January 2017
Accepted 2 February 2017

Abstract

Purpose – Loan default risk or credit risk evaluation is important to financial institutions which provide loans to businesses and individuals. Loans carry the risk of being defaulted. To understand the risk levels of credit users (corporations and individuals), credit providers (bankers) normally collect vast amounts of information on borrowers. Statistical predictive analytic techniques can be used to analyse or to determine the risk levels involved in loans. This paper aims to address the question of default prediction of short-term loans for a Tunisian commercial bank.

Design/methodology/approach – The authors have used a database of 924 files of credits granted to industrial Tunisian companies by a commercial bank in the years 2003, 2004, 2005 and 2006. The naive Bayesian classifier algorithm was used, and the results show that the good classification rate is of the order of 63.85 per cent. The default probability is explained by the variables measuring working capital, leverage, solvency, profitability and cash flow indicators.

Findings – The results of the validation test show that the good classification rate is of the order of 58.66 per cent; nevertheless, the error types I and II remain relatively high at 42.42 and 40.47 per cent, respectively. A receiver operating characteristic curve is plotted to evaluate the performance of the model. The result shows that the area under the curve criterion is of the order of 69 per cent.

Originality/value – The paper highlights the fact that the Tunisian central bank obliged all commercial banks to conduct a survey study to collect qualitative data for better credit notation of the borrowers.

Keywords ROC curve, Risk assessment, Default risk, Banking sector, Bayesian classifier algorithm

Paper type Research paper

Resumen

Propósito – El riesgo de incumplimiento de préstamos o la evaluación del riesgo de crédito es importante para las instituciones financieras que otorgan préstamos a empresas e individuos. Existe el riesgo de que el pago de préstamos no se cumpla. Para entender los niveles de riesgo de los usuarios de crédito (corporaciones e individuos), los proveedores de crédito (banqueros) normalmente recogen gran cantidad de información sobre los prestatarios. Las técnicas analíticas predictivas estadísticas pueden utilizarse para analizar o determinar los niveles de riesgo involucrados en los préstamos. En este artículo abordamos la cuestión de la predicción por defecto de los préstamos a corto plazo para un banco comercial tunecino.

Diseño/metodología/enfoque – Utilizamos una base de datos de 924 archivos de créditos concedidos a empresas industriales tunecinas por un banco comercial en 2003, 2004, 2005 y 2006. El algoritmo bayesiano de



© Aida Krichene. Published in Journal of Economics, Finance and Administrative Science. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) licence. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and non-commercial purposes), subject to full attribution to the original publication and authors. The full terms of this licence may be seen at <http://creativecommons.org/licenses/by/4.0/legalcode>

Journal of Economics, Finance and
Administrative Science
Vol. 22 No. 42, 2017
pp. 3-24
Emerald Publishing Limited
2077-1886
DOI 10.1108/JEFAS-02-2017-0039

clasificadores se llevó a cabo y los resultados muestran que la tasa de clasificación buena es del orden del 63.85%. La probabilidad de incumplimiento se explica por las variables que miden el capital de trabajo, el apalancamiento, la solvencia, la rentabilidad y los indicadores de flujo de efectivo.

Hallazgos – Los resultados de la prueba de validación muestran que la buena tasa de clasificación es del orden de 58.66% ; sin embargo, los errores tipo I y II permanecen relativamente altos, siendo de 42.42% y 40.47%, respectivamente. Se traza una curva ROC para evaluar el rendimiento del modelo. El resultado muestra que el criterio de área bajo curva (AUC, por sus siglas en inglés) es del orden del 69%.

Originalidad/valor – El documento destaca el hecho de que el Banco Central tunecino obligó a todas las entidades del sector llevar a cabo un estudio de encuesta para recopilar datos cualitativos para un mejor registro de crédito de los prestatarios.

Palabras clave – Curva ROC, Evaluación de riesgos, Riesgo de incumplimiento, Sector bancario, Algoritmo clasificador bayesiano.

Tipo de artículo – Artículo de investigación

1. Introduction

Bank credit risk assessment is widely used at banks around the world. As credit risk evaluation is very crucial, a variety of techniques are used for risk level calculation. In addition, credit risk is one of the main functions of the banking community.

According to the Basel Committee on Banking Supervision, credit risk is most simply defined as the potential that a bank borrower or counterparty will fail to meet its obligations in accordance with the agreed terms (Okan, 2007).

Banks classify clients according to their profiles. While classifying, the financial background of the customers and subjective factors related to them are evaluated. Financial ratios play an important role for risk level calculation, as per Berk *et al.* (2011). These ratios are objective and indicate the financial statement of the business. Balance sheet, income statement and cash flows are some financial statements used for collecting information to calculate objective financial ratios. There are many other subjective factors too; these depend on bank decision strategy and its mission, according to Berk *et al.* (2011).

In February 2006, the Basel Committee on Banking Supervision issued a consultative document for comment. This document was intended to provide banks and supervisors with guidance on sound credit risk assessment and valuation policies and practices for loans regardless of the accounting framework applied. Principle 3 of the document states that “A bank’s policies should appropriately address validation of any internal credit risk assessment models”[1].

The implementation of this principle turns out to be a daily decision based on a binary classification problem, distinguishing good payers from bad payers. Certainly, assessing the insolvency plays an important role, as a good estimate (related to a borrower) can help to decide whether to grant the requested loans or not. The Basel Committee proposes a choice between two broad methodologies for calculating their capital requirements for credit risk, either external mapping approach or internal rating system.

Although the external mapping approach is difficult to apply because of the unavailability of external rating grades, the internal rating system is easy to implement, as numerous methods have been proposed in the literature to develop credit-risk evaluation models[2]. In fact, credit scoring methods are used to evaluate both objective and subjective factors. These techniques spread all around the world in the 1950s (Abramowicz *et al.*, 2003). Through these methods, information collection from a customer is formalized. Besides, the scoring system forms a basis for loan approval. These models include traditional statistical methods [e.g. logistic regression (Steenackers and Goovaerts, 1989)], nonparametric statistical models [e.g. k-nearest neighbour (Henley and Hand, 1997)] and classification trees

(Davis *et al.*, 1992) and neural network (NN) models (Desai *et al.*, 1996; Matoussi and Krichène, 2010). The NN models have served as versatile tools for data analysis in a variety of complex environments. In finance, they have been successfully applied to bankruptcy and loan-default prediction and credit evaluation (West, 2000; Wu and Wang, 2000; Atiya, 2001; Pang *et al.*, 2002; Odom and Sharda, 1990). Recent contributions have proposed using Bayesian classification rules with naive Bayes classifiers. The results of these studies demonstrated their frequent ability to do better than most existing techniques. In this context, Sarkar and Sriram (2001) and Sun and Shenoy (2007) have successfully applied the Bayesian model to bankruptcy prediction. Moreover, Stibor (2010) noted that the naive Bayesian classifier is one kind of Bayesian classifier which is actually known as a simple and effective probability classification method (Jie and Bo, 2011) and works based on applying the theorem of Bayes with strong (naive) independence assumptions. The naive Bayes classifier is particularly appropriate when the dimensionality of the inputs is high. Despite its simplicity, naive Bayes can often outperform more sophisticated classification methods (Hill and Lewicki, 2007). That is why we choose to use them in this study.

Our research question is how banks in Tunisia can develop fairly accurate quantitative prediction models that can serve as very early warning signals for counterparty defaults. Previous works look at business failure prediction from the mid-term and long-term prospects (failure vs non-failure). In our paper, we look at the short-term prospect (payment vs non-payment of the short-term credit at maturity). We also consider the case of a bank that wants to use a prediction model to assess its credit risk (Matoussi *et al.*, 1999; Abid and Zouari, 2000; El-Shazly, 2002; Davutyan and Özar, 2006).

Specifically, we use a naive Bayes classifier model to help the credit-risk manager in explaining why a particular applicant is classified as either bad or good. The NN parameters will be set using an optimization procedure analogous to the gradient – in classical topology – and a feed-forward NN with *ad hoc* connections.

The remainder of this paper is organized as follows: In Section 2, we provide the conceptual framework and empirical modelling supporting our research question and research design, respectively. In Section 3, we describe the data and methodology. In Section 4, we present our results and their interpretations. Finally, Section 5 concludes the paper and presents some limits.

2. Credit risk assessment of banks: conceptual framework and empirical modelling

2.1 Conceptual framework of credit risk problem: agency theory

2.1.1 The problem. One of the most fundamental applications of the agency theory to the lender–borrower problem is the derivation of the optimal form of the lending contract. In a debt market, the borrower usually has better information about the project to be financed and its potential returns and risk. The lender, however, does not have sufficient and reliable information concerning the project to be financed. This lack of information in quantity and quality creates problems before and after the transaction takes place. The presence of asymmetric information normally leads to moral hazard and adverse selection problems. This situation illustrates a classical principal–agent problem.

The principal–agent models of the agency theory may be divided into three classes according to the nature of information asymmetry. First, there are moral hazard models, where the agent receives some private information after signing the contract. Moral hazard refers to a situation in which the asymmetric information problem is created after the transaction occurs. As the borrower has relevant information about the project that the

lender does not have, the lender runs the risk that the borrower will engage in activities that are undesirable from the lender's point of view because they make it less likely that the loan will be paid back. These models are qualified as models with ex-post asymmetric information.

Second, we find adverse selection models, where the agent has private information already before signing the contract. Adverse selection refers to a situation in which the borrower has relevant information that the lender lacks (or vice versa) about the quality of the project before the transaction occurs. This happens when the potential borrowers who are the most likely to produce an undesirable (adverse) outcome (bad credit risks) are the ones who are most active to get a loan and are thus most likely to be selected. In the simplest case, lenders' price cannot discriminate between good and bad borrowers, because the riskiness of projects is unobservable. These models are known as models with ex-ante asymmetric information.

Finally, there are signalling models, in which the informed agent may reveal his private information through a signal which he sends to the principal.

2.1.2 The solution. This problem is traditionally considered in the framework of costly state verification, introduced by [Townsend \(1979\)](#). The essence of the model is that the agent, who has no endowment, borrows money from the principal to run a one-shot investment project. The agent is faced with a moral hazard problem. Should he announce the true value or should he lower the outcome of the project? This situation describes ex-post moral hazard. We can also face a situation of ex-ante moral hazard, where the unobservable effort by the agent during project realization may influence the result of the project. [Townsend \(1979\)](#) showed that the optimal contract which solves this problem is the so-called standard (or simple) debt contract. This standard debt contract is characterized by its face value, which should be repaid by the agent when the project is finished. Another theoretical justification for simple debt contract was considered by [Diamond \(1984\)](#), where the costly state verification was replaced by a costly punishment. [Hellwig \(2000, 2001\)](#) showed that the two models are equivalent only under the risk neutrality assumption. However, when we consider the introduction of risk aversion, the costly state verification model still works, but the costly punishment model does not survive.

To overcome the asymmetric information problem and its consequences on credit risk assessment in the real world, banks use either collateral or bankruptcy prediction modelling or both. The next subsection will deal with this aspect.

2.2 Credit risk assessment and bankruptcy prediction: related studies (works)

After the high number of profile bank failures in Asia, research activity on credit risk took a step further. As a result, the regulators recognize the need and urge banks to utilize cutting-edge technology to assess the credit risk in their portfolios. Measuring the credit risk accurately also allows banks to engineer future lending transactions, so as to achieve targeted return/risk characteristics. The assessment of credit risk requires the development of fairly accurate quantitative prediction models that can serve as very early warning signals for counterparty defaults ([Atiya, 2001](#))[3].

Many researchers proposed two main approaches to deal with credit scoring in the literature. In the first approach, i.e. the structural or market-based models, the default probability derivation is based on modelling the underlying dynamics of interest rates and firm characteristics. This approach is based on the asset value model originally proposed by [Merton \(1974\)](#), where the default process is endogenous, and relates to the capital structure of the firm. Default occurs when the value of the firm's assets falls below some critical level. In the second approach, i.e. the empirical or accounting-based models, instead of modelling

the relationship of default with the characteristics of a firm, this relationship is learned from the data. Raymond (2007), Thomas *et al.* (2002) and Galindo and Tamayo (2000) synthesized some methods used in this context. In this regard, we can cite the work of Beaver (1963) and Altman (1968). Bankruptcy prediction has been studied actively by academics and practitioners. Many models have been proposed and tested empirically. Altman's popular Z-score (Altman, 1968) is an example, which was based on linear discriminant analysis and was used to predict the probability of default of firms. Ohlson's O-score (Ohlson, 1980) is based on generalized linear models. Generalized linear models or multiple logistic regression models have been used either to identify the best determinants of bankruptcy or the predictive accuracy rate of their occurrence. NN models were adapted and used in bankruptcy prediction. Their high power of prediction makes them a popular alternative with the ability to incorporate a very large number of features in an adaptive nonlinear model (Kay and Titterton, 2000).

A lot of research studies have focused on the nonparametric methods class [e.g. k-nearest neighbour (Henley and Hand, 1996), decision trees (Quinlan, 1992) and NNs (McCulloch and Pitts, 1943)], which also have been largely applied in the field of credit scoring. There are also some other approaches that combine several techniques to create a classification model, such as the Support Vector Machine (Lee and Chen, 2005; Lee *et al.*, 2002). Antonietta and Paolo (2003) developed a Bayesian regression model to predict the credit risk of companies classified in different sectors. Maltritz and Molchanov (2008) proposed a Bayesian model to find the variables which are most likely to determine country default risk in emerging markets. The collection provided by Bocker (2010) also includes several studies on Bayesian credit risk modelling. For example, Jacobs and Kiefer (2010) present a step-by-step guide to Bayesian analysis in the default setting, including details on elicitation of expert information. According to Miguéis *et al.* (2012), 'despite the intense study of credit scoring, there is no consensus on the most appropriate classification technique to use'. Baesens *et al.* (2003) revealed that some conflicts can occur when comparing the findings of different studies. However, Thomas *et al.* (2002) also suggested that most methods applied in credit scoring have similar levels of performance. In fact, for banks and financial institutions, the reasons that may motivate the preference for certain methods are their interpretability and the transparency (Martens *et al.*, 2009). According to Miguéis *et al.* (2012), "two aspects of methods for credit scoring are very important: that is the predictive performance, as well as the insights or interpretations that are revealed by the model".

2.3 Empirical research design

2.3.1 Simple naive Bayes classifier algorithm. Banks are in a very competitive environment; therefore, the service quality during credit risk assessment is very important. When a customer demands credit from a bank, the bank should evaluate the credit demand as soon as possible (Berk *et al.*, 2011) to gain competitive advantage. Additionally, for each credit demand, the same process is repeated and constitutes a cost for the bank. Due to the importance of credit risk analysis, most techniques and models are developed by financial institutions to help them decide whether to grant or not to grant credit (Cinko, 2006).

In this section, we briefly review the implementations of the binary classification using the naive Bayes algorithm. In fact, the naive Bayes algorithm is a classification algorithm based on Bayes rule, which assumes the attributes $X_1 \dots X_n$ are all conditionally independent of one another, given Y . The value of this assumption is that it dramatically simplifies the representation of $P(X/Y)$, and the problem of estimating it from the training data, according to Mitchell (2010).

A Bayesian network (BN) represents a joint probability distribution over a set of continuous inputs (attributes) X_i . In this case, designing a naive Bayes classifier is based on the use of [equation \(1\)](#):

$$P(Y = y_k | X_1 \dots X_n) = \frac{P(Y = y_k) \prod_i P(X_i | Y = y_k)}{\sum_j P(Y = y_j) \prod_i P(X_i | Y = y_j)} \quad (1)$$

Given a new occurrence $X^{\text{new}} = (X_1 \dots X_n)$, [equation \(1\)](#) shows how to estimate the probability that Y will take on any given value, given the observed input values of X^{new} and given the distributions $P(Y)$ and $P(X_i/Y)$ estimated from the training data. If we are interested only in the most probable value of Y , then we have the naive Bayes classification rule as shown in [equation \(2\)](#):

$$Y \leftarrow \arg \max_{y_k} P(Y = y_k) \prod_i P(X_i | Y = y_k) \quad (2)$$

However, according to [Mitchell \(2010\)](#), “when the X_i are continuous we must choose some other way to represent the distributions $P(X_i/Y)$ ”. One common approach is to assume that for each possible discrete value y_k of Y , the distribution of each continuous X_i is Gaussian, and is defined by a mean and standard deviation specific to X_i and y_k . To train such a naive Bayes classifier, we must, therefore, estimate the mean and standard deviation of each of these Gaussians:

$$\mu_{ik} = E[X_i | Y = y_k] \quad (3)$$

$$\sigma^2_{ik} = E[(X_i - \mu_{ik})^2 | Y = y_k] \quad (4)$$

For each input X_i and each possible value y_k of Y , note that there are $2nK$ of these parameters, all of which must be estimated independently. Certainly, we have to estimate the priors on Y as well:

$$\pi_k = P(Y = y_k) \quad (5)$$

The above model summarizes a naive Bayes classifier, which assumes that the data X are generated by a mixture of class-conditional (i.e. dependent on the value of the class variable Y) Gaussians. Furthermore, the naive Bayes assumption introduces the additional constraint that the input values X_i are independent of one another within each of these mixture components.

2.3.2 ROC curve as a classifier performance. A receiver operating characteristic (ROC) curve is generally a useful performance-graphing method. In other words, an ROC graph is a method for visualizing, organizing and selecting classifiers based on their performance ([Fawcett, 2006](#)). [Spackman \(1989\)](#) was the earliest adopters of ROC graphs in machine learning. He demonstrated the value of ROC curves in evaluating and comparing algorithms ([Fawcett, 2006](#)). In fact, the use of ROC graphs in the machine-learning community has increased in recent years, as simple classification accuracy is often a poor metric for

measuring performance (Provost and Fawcett, 1997; Provost *et al.*, 1998). Besides, they have properties that make them especially useful for domains with skewed class distribution and unequal classification error costs (Fawcett, 2006; Figure 1).

2.3.3 The criterion of the area under a curve ROC. An ROC curve is a two-dimensional representation of classifier performance. According to Fawcett (2006), “to compare classifiers we may want to reduce ROC performance to a single scalar value representing expected performance”. To do so, many researchers, such as Bradley (1997) and Hanley and McNeil (1982), recommend the use of a common method – which is to calculate the area under the ROC curve, abbreviated AUC. The AUC is defined as a portion of the area of the unit square; its value will always be between 0 and 1.0. However, because random guessing produces the diagonal line between (0, 0) and (1, 1), which has an area of 0.5, no realistic classifier should have an AUC less than 0.5 (Fawcett, 2006).

3. Methodology

The need for models that predict defaults correctly is very high because in commercial banks, credit risk measurement is crucial to discriminate reliable clients from the non-reliable ones. Among the quantitative methods for solving credit risk evaluation problems, the simple Bayesian classifier was applied for estimating the posterior probabilities of default. In fact, Antonakis and Sfakianakis (2009) showed that the posterior probability of an event is the probability of an event after collecting some empirical data. Rosner (2006) demonstrated that the posterior probability is obtained by integrating information from the prior probability with additional data related to the event in question. According to Mileris (2010):

[...] often analysis begins with initial or prior probability estimates for specific events of interest. Then from sources such as a sample we obtain additional information about the events. Given this new information the prior probability values can be updated by calculating revised probabilities, referred to as posterior probabilities.

Anderson *et al.* (2007) demonstrated that the Bayesian theorem provides a means for making these probability calculations.

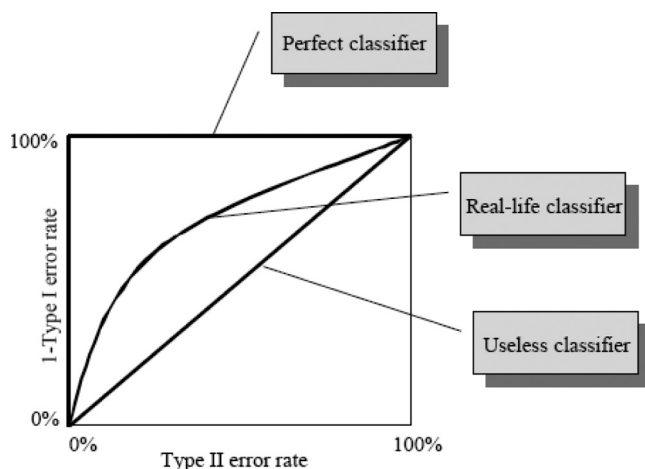


Figure 1.
An example of ROC
curve

Source: Own elaboration and data from Yang (2002), p. 18

In our experiment, we use a sample of bank credit files and divide them into two sub-samples. The first sub-sample is composed of 924 files of short-term loans granted to Tunisian companies in the years 2003, 2004, 2005 and 2006.

3.1 *Sample and data*

Let's recall that our objective is to use the naive Bayes classifier methodology for default prediction of a bank's commercial loans. However, to solve a problem using the Bayesian algorithm, we need to gather data for training purposes. The training data set includes a number of cases, each containing values for a range of input and output variables. The first decision we need to make is which variables to use. The second one concerns the subjects whose behaviour we want to predict. For our case, the variables are indicators of default risk and the subjects are borrowers. The data collected for our investigation came from a large private commercial bank (BIAT). We chose a private bank to avoid the potential inefficiency of the public banking sector, whose decisions are sometimes dictated by government choices. We also chose to work with short-term commercial loans because they represent the largest part of loans and are subject to renewal every year[4].

We have used a database of 924 files of credits granted to industrial Tunisian companies by a commercial bank in the years 2003, 2004, 2005 and 2006. This period was chosen because it corresponds to a central bank instruction, in which it asks banks to provide credit risk classes for their borrowers. In the case of the BIAT, by the end of every quarter, it classifies these files into five clusters, each one corresponding to a risk class. Files without delay of payment correspond to the healthy firms. The four remaining classes correspond to four riskier classes of firms with three months, six months, nine months and one year (or more) delay of payment, respectively. We group these four classes in one class: risky companies.

3.2 *Variables measurement*

3.2.1 *Dependent variable.* Our dependant variable is the probability of default. We use a dummy variable, Y, which equals 0 if the firm is classified as healthy and 1 otherwise. Hence:

- Y = 0 if no delay of payment
- Y = 1 if more there is more than a three-month delay

3.2.2 *Independent variables.* Default risk prediction relies, in general, on a good appraisal of the couple risk-return of a company. Financial ratios drawn from financial statements (balance sheet, income and cash flow statement) are usually used. Financial ratio analysis groups the ratios into categories which tell us about different facets of a company's finances and operations (liquidity, activity or operational, leverage and profitability).

In our experiment, we retain 24 financial and nonfinancial indicators, 22 of them are financial ratios and 2 are not. The financial indicators are inspired from Altman's popular Z-score and recommended textbooks in the field of financial statement analysis and valuation (Berstein and Wild, 1998; Revsine *et al.*, 1999; and Palepu *et al.*, 2000). The financial indicators measure liquidity (working capital, operating activity and cash flow), leverage, long-term solvency and profitability. The nonfinancial variables used in this research are firm size and collateral (Table I).

4. **Empirical results**

4.1 *Descriptive analysis*

To get a better idea about our data before running the naive Bayes classifier models, we will perform a test of mean differences between the two risk classes defined above (Table II). The summary statistics and the mean differences can be seen as an analysis similar to that of

Risk facet	Code	Variable definition	Variable measure
Liquidity indicators	R1	Long-term financing of working capital	Shareholders' equity + Non current liabilities) – Non current assets
	R2	Working capital requirement	Working capital
	R3	Account receivable liquidity	(Shareholders' equity + Non current liabilities) – Non current assets
	R4	Current ratio	Provision for doubtful accounts
			Gross account receivables
			Current assets
			Current liabilities
	R5	Quick ratio	Current assets – inventories
			Current liabilities
	R6	Cash flow ratio	Operating cash flow
			Current liabilities
	R7	Inventory turnover	Sales
	R8	Debt cash flow coverage ratio	Inventories
			Cash flow
Leverage and solvency indicators			Net Income + depreciation
	R9	Liabilities to equity ratio	Total debits = Total debits
			Total liabilities
	R10	Net debt to equity ratio	Shareholders' equity
			Short term debt + long term debt – cash and marketable securities
	R11	Debt to capital ratio	Shareholders' equity
			Short term debt + long term debt
	R12	Long-term debt to assets	Shareholders' equity
			Short term debt + long term debt + shareholders equity
	R13	Long-term debt to tangible assets	Long term debt
			Total assets
	R14	Interest coverage ratio	Long term debt
			Total tangible assets
			Operating income before taxes and interest
			Interest expense

(continued)

Table I.
Variable definition
and measure

Table I.

Risk facet	Code	Variable definition	Variable measure
Profitability indicators	R15	Net profit margin	Net income
	R16	Gross profit margin	$\frac{\text{Total operating revenue}}{\text{Earnings before interest and taxes}}$
	R17	Return on invested capital	$\frac{\text{Total operating revenue}}{\text{Net income}}$
			$\frac{\text{Total income}}{\text{Total assets}}$
Ratios used by the bank	R18	Return on equity (ROE)	Net income
	R19	Fixed asset to debt ratio	$\frac{\text{Stockholders equity}}{\text{Net fixed assets}}$
	R20	Short-term debt to sales ratio	$\frac{\text{Total debt}}{\text{Total sales}}$
	R21	Financial expenses to revenue ratio	$\frac{\text{Financial expenses}}{\text{Total revenues}}$
Other variables	R22	Fixed asset turnover	$\frac{\text{Sales}}{\text{Fixed assets}}$
	V01	collateral	LOG(GUARANTEE)
	V02	Firm size	LOG(TOTAL ASSETS)
Source: Own elaboration			

Ratios	Code	Mean	SD
R2:	0	<i>16.8191</i>	58.85241
	1	<i>8.5990</i>	15.69326
R3:	0	0.0471	0.13891
	1	0.0568	0.14135
R4:	0	<i>2.9623</i>	7.05638
	1	<i>3.2328</i>	8.05572
R6:	0	<i>2.0391</i>	22.41881
	1	<i>-0.6450</i>	38.61664
R7:	0	0.0439	0.10179
	1	0.0757	0.14164
R8:	0	<i>1.4900</i>	2.00318
	1	<i>1.0742</i>	0.91636
R10:	0	0.0452	0.33929
	1	0.0347	0.16519
R11:	0	<i>0.0569</i>	0.15137
	1	<i>0.0166</i>	0.10049
R12:	0	<i>0.4993</i>	2.33091
	1	<i>0.2348</i>	1.14838
R13:	0	0.7708	0.97464
	1	0.7151	0.58013
R14:	0	<i>0.2274</i>	1.06959
	1	<i>0.0588</i>	0.74966
R15:	0	<i>5.0372</i>	55.16137
	1	<i>13.4255</i>	16.99936
R18:	0	<i>0.6227</i>	2.93441
	1	<i>7.4529</i>	54.44136
R19:	0	1.8072	2.80796
	1	1.7822	3.89486
R20:	0	1.1982	2.38237
	1	1.1215	3.14473
R21:	0	<i>0.0634</i>	0.30944
	1	<i>0.2492</i>	2.91846
R22:	0	<i>0.0115</i>	0.43979
	1	<i>-0.0284</i>	0.25000

Notes: 0 corresponds to healthy group; 1 corresponds to risky group; italic data significant repartition of companies of the sample: number of healthy and risky companies

Source: Own elaboration

Table II.
Group means

Beaver (1963). Table II presents the descriptive statistics of our data. When we run the mean difference analysis between the two risk classes (healthy and risky groups), this analysis can give us a flavour of our data, as such an analysis allows us to verify if there is a difference between the two classes in terms of financial ratios. Table II recalculates some summary statistics for the two risk classes.

Table II shows significant mean differences between the two groups for some ratios (R₂, R₄, R₆, R₈, R₁₁, R₁₂, R₁₄, R₁₅, R₁₈, R₂₁ and R₂₂) and no significant differences for others (R₁, R₃, R₅, R₇, R₉, R₁₀, R₁₁, R₁₃, R₁₆, R₁₇, R₁₉ and R₂₀). Globally, they tell us that the liquidity risk does not differentiate the two groups. The leverage and solvency ratios do better in differentiating the two groups. For other indicators (coverage and profitability), the results are mitigated. For example, while return on equity (R₁₈) shows a significant difference, gross profit margin (R₁₆) and return on invested capital (R₁₇) do not.

When we look at the relevance of the mean differences, we realize that, globally, the good indicators are superior in the healthy group, while the bad indicators are higher in the risky group. For example, the mean of cash flow (R_6), working capital requirement (R_2) and leverage and solvency (R_{11} , R_{12} , R_{14} and R_8) ratios is bigger in the healthy group. Current ratio (R_4) and profitability ratios (R_{18} and R_{15}) have a higher mean in the risky group.

Let us see now if Bayes models do a better job in predicting default risk.

4.2 Results of Bayes models

Bayesian classifiers can predict class membership probabilities, such as the probability that a given sample belongs to a particular class. Bayesian classifier is based on Bayes' theorem. Naive Bayesian classifiers assume that the effect of an attribute value on a given class is independent of the values of the other attributes. This assumption is called class-conditional independence.

4.2.1 *Multi-collinearity issues.* Initially, we used a database composed of 32 independent variables. We eliminated eight ratios which posed the problem of multi-collinearity. These ratios are detected by the software. The database used in this work is composed of 24 variables which do not pose the multi-collinearity problem.

4.2.2 *Classification results.* Panels 1 and 2 of [Tables III](#) and [IV](#) show the results for the first and second Bayesian models.

We can see from these results (Panels 1 and 2) that the global good classification rate is getting better when we introduce indicators relating to cash flow, firm size and guarantee. In fact, the best good classification rates are of the order of 59.63 and 63.85 per cent, respectively, for the two models (non-cash flow and full information). A lot of research has examined the criterion of type I and II errors.

According to [Yang \(2002\)](#), type I error rate is also called a rate or credit risk; it is the rate of bad customers being categorized as good. When this happens, the misclassified bad

Table III.
Results for non-cash
flow, and full
information models:
non-cash flow Bayes
model

Companies' type and classification	Healthy	Risky
Healthy companies	331	127
Risky companies	246	220
Total good and bad classification (%)		
Good classification	331 + 220/924 = 59.63%	
Bad classification	127 + 246/924 = 40.37%	
Source: Appendix Figure A1 ; Own elaboration		

Table IV.
Results for non-cash
flow, and full
information models:
full information
Bayes model

Companies' type and classification	Healthy	Risky
Healthy companies	<i>307</i>	<i>151</i>
Risky companies	<i>183</i>	<i>283</i>
Total good and bad classification (%)		
Good classification	283 + 307/924 = 63.85%	
Bad classification	1 83 + 151/924 = 36.15%	
Note: Italic data significant repartition of companies of the sample; number of healthy and risky companies		
Source: Appendix Figure A2 ; Own elaboration		

customers will become default. Therefore, if a credit institution has a high rate, which means the credit granting policy is too generous, it is exposed to credit risk.

Type II error rate is also called a commercial risk; it is the rate of good clients being classified as bad. When this happens, the misclassified good clients are rejected, and the bank, therefore, supports (endures) an opportunity cost caused by the loss of good customers. Bogess (1967) showed that if a credit institution has a high type II error for a long period, which means it follows a restrictive credit granting policy for a long time (Yang, 2002), it may lose its share in the market. The credit institution is, thus, exposed to commercial risk.

In this study, we find that error type I is very high for the first model of the first panel (non-cash flow Bayes model). In fact, this rate is of the order of 52.79 per cent. Table V shows these results. The introduction of cash flow variables (Panel 2) improved the results. The good classification rate got better. In addition, the model based on the full information indicator reduced error type I to 39.27 per cent (cf. Table VI) but error type II increased to 32.97 per cent.

In this research, we would like to assess credit risk using a selection of financial ratios recommended in debt contracts. The predictions on the selection of financial ratios illustrate the relation between financial ratios and credit risk. This evidence is well known in the practitioner and academic literature (Demerjian, 2007). In fact, textbooks emphasize the role of ratios in evaluating credit quality (Lundholm and Sloan, 2004), while academic studies conclude that financial ratios serve to provide signals about borrower credit risk when used as covenants (Smith and Warner, 1979; and Dichev and Skinner, 2002).

Table VI show the working capital requirement (R2), account receivable liquidity(R3), cash flow ratio(R6), debt cash flow coverage ratio (R8), gross profit margin(R16), fixed asset to debt ratio (R19) and guarantee (V02).

In other words, ratios R2, R3, R6, R8, R16, R19 and V02 capture five aspects of credit risk:

- (1) short-term liquidity (working capital and account receivable liquidity);
- (2) operating performance (coverage, debt to cash flow and cash flow);
- (3) profitability;
- (4) leverage; and
- (5) guarantee.

Table V.Criterion of type I
and II errors:

performance measure

Model	Error	
	Type I	Type II
Non-cash flow model	246/466 = 52.79%	127/458 = 27.73%

Source: Appendix Figure A1; Own elaboration

Table VI.Criterion of type II
error

Model	Error	
	Type I	Type II
Full information model	183/466 = 39.27%	151/458 = 32.97%

Source: Appendix Figure A2; Own elaboration

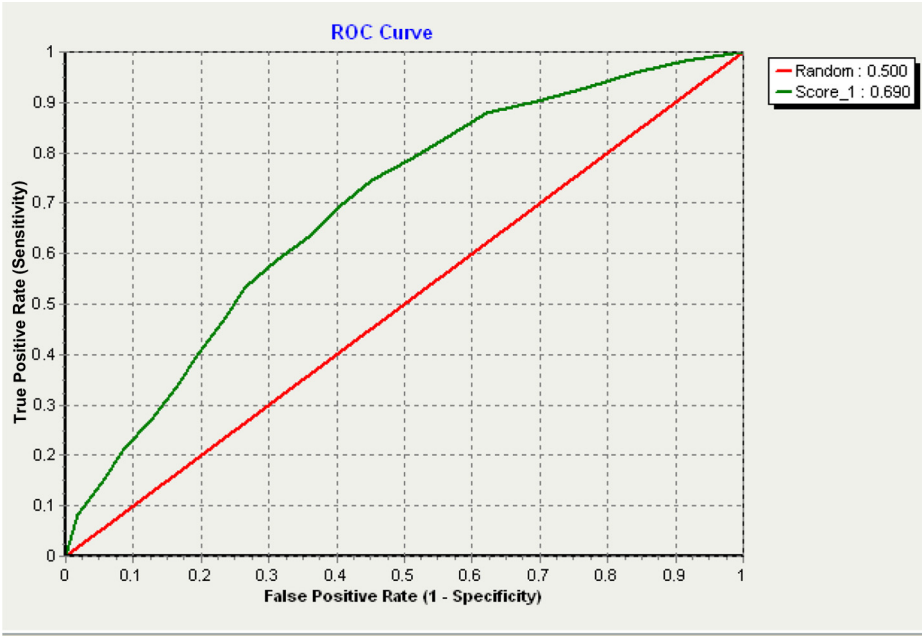


Figure 2.
The ROC curve:
output software
Tanagra

Source: Output Software Tanagra

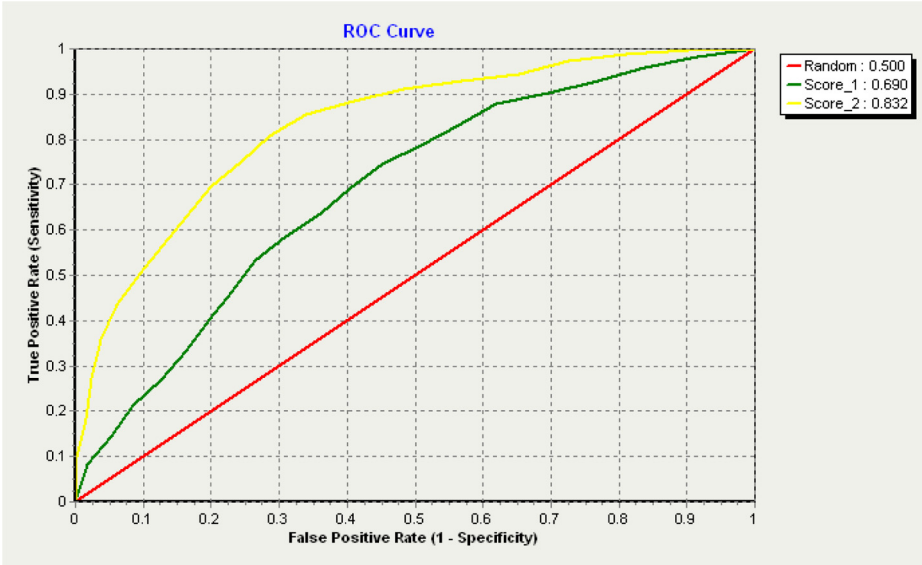


Figure 3.
The curve ROC:
neural network
versus Bayesian
classifier

Source: Output Software Tanagra

Moreover, the evidence also shows that liquidity, cash flow and profitability leverage are the most informative (enlightening) about the failure of corporate borrowers, with a p -value of 0.

The inventory turnover (R7) and interest coverage (R14) ratios, related, respectively, to liquidity and solvency, contribute less than the other ratios cited previously in the prediction of credit risk, with a p -value of the order of 0.014 and 0.038, respectively, for R7 and R14.

Finally, R1, related to long-term financing of working capital, is less informative with a p -value of 0.09.

Regarding the other variables used in this study, R4, R5, R9, R10, R11, R12, R13, R15, R17, R18, R20, R21 and V01, the results are not significant with a very high p -value, varying from 0.82 (R18) to 0.13 (R5 and R20).

4.3 The ROC curve

An ROC curve for the perfect classifier, which orders all bad cases before good cases, is the curve that follows the two axes. It would classify 100 per cent bad cases into class bad and 0 per cent good cases into class bad for some value of the sill. According to [Yang \(2002\)](#):

[. . .] a classifier with a ROC curve which follows the 45° line would be useless. It would classify the same proportion of the bad cases and good cases into the class bad at each value of the threshold; it would not separate the classes at all. Real-life classifiers produce ROC curves which lie between these two extremes.

Classification functions Descriptors	Risky	Healthy	$F(1,922)$	p -value
Intercept	-112.735385	-102.691145	—	—
R1	0.001160	-0.000673	2.866256	0.090793
R2	-0.085074	-0.008751	18.382881	0.000020
R3	4.202971	2.665188	11.653407	0.000669
R4	0.392753	0.438712	1.447920	0.229171
R5	0.484359	0.565610	2.208326	0.137609
R6	10.014929	8.946201	21.092551	0.000005
R7	1.203604	1.831741	5.963674	0.014791
R8	3.972152	2.484139	25.887069	0.000000
R9	0.000016	0.000254	1.319935	0.250902
R10	0.026991	0.061113	0.342810	0.558355
R11	0.000017	0.000255	1.330735	0.248973
R12	0.015294	0.020427	0.324426	0.569099
R13	0.004571	0.010033	1.079883	0.298996
R14	0.048391	0.001563	4.298513	0.038423
R15	-0.072669	0.007240	2.616083	0.106127
R16	-1.460765	1.158819	13.310767	0.000279
R17	0.000558	-0.005868	1.580823	0.208961
R18	0.016025	0.016828	0.046710	0.828938
R19	0.107474	0.057933	8.515640	0.003607
R20	0.077288	0.021156	2.221415	0.136450
R21	0.010592	0.012532	0.931010	0.334854
R22	0.014652	0.008826	9.970419	0.001642
V01	0.508630	0.501425	0.088990	0.765532
v02	21.599527	20.672270	89.931796	0.000000

Source: Output Software Tanagra

Table VII.
Results

To evaluate the performance of the curve, we have to use a measure given by the AUC (Hand, 1997). The curve that has a larger AUC is better than the one that has a smaller AUC (Figure 2).

We can note that the criterion of AUC is of the order of 69 per cent. This score is superior to 50 per cent, but it is not considered a very good score.

By comparing ROC curves, we can learn the difference in classification precision between two or more classifiers. The higher curve will be nearer the perfect classifier and will have more accuracy. In this context, by using the same data and carrying out an NN model with hidden layers, we can compare the classification accuracy between the two models – naive Bayes and the NN model (Figure 3).

We can conclude that the NN model outperforms the naive Bayes model because the criterion of AUC for the NN model is 83.2 per cent (yellow curve) and that of the Bayes model is 69 per cent (green curve).

4.4 Out-of-sample validation

We note that the out-of-sample validation will be done on a second sub-sample, which contains data on 150 files of short-term loans granted to industrial Tunisian companies in 2006. Table VII displays the sub-sample repartition (Table VIII).

Table VIII.
Out of sample
repartition

Companies' type	No. of firms	(%)
Healthy companies	84	56
Risky companies	66	44
Total	150	100
Source: Own elaboration		

Table IX.
Results of the
validation test

Companies' type and classification	Healthy	Risky	Total
Healthy companies	50	34	84
Risky companies	28	38	66
Total	78	72	150
Total good and bad classification (%)			
Good classification	58.66%		
Bad classification	41.34%		
Source: Own elaboration			

Table X.
Criterion of type II
error

Error	
Type I	Type II
28/66 = 42.42%	34/84 = 40.47%
Source: Own elaboration	

Table IX presents the results of the validation test; we can see the corresponding good classification rate of the order of 58.66 per cent. The error types I and II remain relatively high at 42.42 and 40.47 per cent, respectively (Table X).

5. Conclusion

Commercial banks that grant client borrower loans need consistent models that can correctly detect and predict defaults. Moonasar (2007) emphasized that one of the basic tasks which any bank has to deal with, in the current competitive and turbulent business environment, is to reduce its credit risk. Traditionally, we used scoring methods to estimate the creditworthiness of a credit applicant. In fact, the quantitative method known as credit scoring has been developed for the credit assessment problem (Yang, 2002). Credit scoring is basically an application of classification methods, which classify borrowers into different risk groups. The objective of scoring methods is to predict the probability that an applicant or existing borrower will default (Komor'ad, 2002). In credit risk evaluation, credit scoring is a key method that helps financial institutions to make a decision whether or not to grant credit to a customer (Thomas, 2002). According to Moonasar (2007):

[. . .] a common approach of credit scoring is to apply a classification technique on data of previous customers (both good credit customers and delinquent customers) in order to find a relationship between the customers' characteristics and potential failure to service their debt. Institutions use credit scoring techniques (utilizing information from the consumers' past credit history and current economic conditions) to determine which applicants will pay back their liabilities.

An accurate classifier is necessary to discriminate between new potential good and bad credit applicants.

In this paper, we tried to assess the credit risk for a Tunisian bank by modelling the default risk of its commercial loans. We used a database of 924 credit files during the years 2003, 2004, 2005 and 2006. Input variables were classified into two categories: non-cash flow ratios and cash flow ratios.

The main results show that the introduction of cash flow variables improves the prediction quality, and the classification rates passed from 59.63 to 63.85 per cent, respectively, in the non-cash flow and cash flow models. Moreover, collateral played an important role in default risk prediction. In fact, this indicator has an explanatory capacity to predict the credit risk. To evaluate the performance of the model, an ROC curve was plotted. The result shows that the AUC criterion is of the order of 69 per cent. Using the same data, this criterion is improved and passed to 83 per cent when we used the NN methodology.

Our study is, however, incomplete in the sense that we used only quantitative variables and it did not show the importance of qualitative variables based on strategic data in completing the financial analysis and assessing the solidity of a borrower. We note that the Tunisian central bank obliged all commercial banks to conduct a survey study to collect qualitative data for better credit notation of the borrowers.

Notes

1. Sound credit risk assessment and valuation for loans. Consultative document, Bank for International Settlements Press and Communications, Basel (November 2005).
2. Moreover, the subprime crisis, which has shaken down the American and European countries and shown the fragility of the banking sector, throws some doubt on the accuracy and usefulness of agency ratings.

3. To get an idea about the potential impact of the bankruptcy prediction problem, we note that the volume of outstanding debt to corporations in the USA is about \$5tn. An improvement in default prediction accuracy of just a few percentage points can lead to savings of tens of billions of dollars.
4. Loans with maturities of one year or less comprise more than half of all commercial bank loans. For the case of the BIAT, this ratio was around 40 per cent during 2006 and 2007.

References

- Abid, F. and Zouari, A. (2000), "Financial distress prediction using neural networks", available at: <http://ssrn.com/abstract=355980>, doi: 10.2139/ssrn.355980.
- Abramowicz, W., Nowak, M. and Szytkiel, J. (2003), *Bayesian Networks as a Decision Support Tool in Credit Scoring Domain*, Idea Group Publishing.
- Altman, E.I. (1968), "Financial ratios, discriminant analysis and the prediction of corporate bankruptcy", *Journal of Finance*, Vol. 23 No. 4, pp. 589-609.
- Anderson, D.R., Sweeney, D.J., Freeman, J., Williams, T.A. and Shoesmith, E. (2007), *Statistics for Business and Economics*, Thomson Learning EMEA, London.
- Antonakis, A.C. and Sfakianakis, M.E. (2009), "Assessing naive Bayes as a method for screening credit applicants", *Journal of Applied Statistics*, Vol. 36 No. 5, pp. 537-545.
- Antonietta, M. and Paolo, T. (2003), "Bayesian estimate of credit risk via MCMC with delayed rejection", *Economics and quantitative methods*, Department of Economics, University of Insubria.
- Atiya, A.F. (2001), "Bankruptcy prediction for credit risk using neural nets: a survey and new results", *IEEE Transactions on Neural Nets*, Vol. 12 No. 4, pp. 929-935.
- Baesens, B., Van Gestel, T., Viaene, S., Stepanova, M., Suykens, J. and Vanthienen, J. (2003), "Benchmarking state-of-the-art classification algorithms for credit scoring", *Journal of the Operational Research Society*, Vol. 54 No. 6, pp. 627-635.
- Beaver, W. (1963), "Financial ratios as predictors of failure. Empirical research in accounting: selected studies", *Journal of Accounting Research*, Vol. 5, pp. 71-111.
- Berk, B., Hidayet, T. and Utku, C.E. (2011), "Bank credit risk analysis with Bayesian network decision tool", *International Journal of Advanced Engineering Sciences and Technologies*, Vol. 9 No. 2, pp. 273-279.
- Berstein, L.A. and Wild, J.J. (1998), *Financial Statement Analysis: Theory, Application, and Interpretation*, 6th ed., McGraw-Hill.
- Bocker, K. (2010), *Rethinking Risk Measurement and Reporting: Volume II*, Risk Books, London.
- Bogess, W.P. (1967), "Screen-test your credit risks", *Harvard Business Review*, Vol. 45 No. 6, pp. 113-122.
- Bradley, A.P. (1997), "The use of the area under the ROC curve in the evaluation of machine learning algorithms", *Pattern Recognize*, Vol. 30 No. 7, pp. 1145-1159.
- Çinko, M. (2006), "Comparison of credit scoring techniques: İstanbul Ticaret Üniversitesi Sosyal Bilimler", *Dergisi*, Vol. 5 No. 9, pp. 143-153.
- Davis, R.H., Edelman, D.B. and Gamberman, A.J. (1992), "Machine learning algorithms for credit-card applications", *IMA Journal of Management Mathematics*, Vol. 4 No. 1, pp. 43-51.
- Davutyan, N. and Özar, S. (2006), "A credit scoring model for Turkey's micro and small enterprises (MSE's)", *13th Annual ERF Conference*, 16-18 December.
- Demerjian, P.R.W. (2007), *Financial Ratios and Credit Risk: The Selection of Financial Ratio Covenants in Debt Contracts*, Workshop Stephen M. Ross School of Business, University of Michigan, Michigan, MI.
- Desai, V.S., Crook, J.N. and Overstreet, G.A. (1996), "A comparison of neural networks and linear scoring models in the credit union environment", *European Journal of Operational Research*, Vol. 95 No. 1, pp. 24-37.

-
- Diamond, D.W. (1984), "Financial intermediation and delegated monitoring", *Review of Economic Studies*, Vol. 51 No. 3, pp. 393-414.
- Dichev, I. and Skinner, D. (2002), "Large-sample evidence on the debt covenant hypothesis", *Journal of Accounting Research*, Vol. 40 No. 4, pp. 1091-1123.
- El-Shazly, A. (2002), "Financial distress and early warning signals: a non-parametric approach with application to Egypt", *9th Annual ERF Conference, Emirates*, October.
- Fawcett, T. (1997), "Analysis and visualization of classifier performance: comparison under imprecise class and cost distributions", *Proceeding Third International Conference on Knowledge Discovery and Data Mining (KDD-97)*, AAAI Press, Menlo Park, CA, pp. 43-48.
- Fawcett, T. (2006), "An introduction to ROC analysis", *Pattern Recognition Letters*, Vol. 27 No. 8, pp. 861-874.
- Galindo, J. and Tamayo, P. (2000), "Credit risk assessment using statistical and machine learning: basic methodology and risk modeling applications", *Computational Economic*, Vol. 15 Nos 1/2, pp. 107-143.
- Hand, D.J. (1997), "Construction and assessment of classification rules", *Wiley Series in Probability and Statistics*, John Wiley & Sons.
- Hanley, J.A. and McNeil, B.J. (1982), "The meaning and use of the area under a receiver operating characteristic (ROC) curve", *Radiology*, Vol. 143, pp. 29-36.
- Hellwig, M. (2000), "Financial intermediation with risk aversion", *Review of Economic Studies*, Vol. 67 No. 4, pp. 719-742.
- Hellwig, M. (2001), "Risk aversion and incentive compatibility with ex post information asymmetry", *Economic Theory*, Vol. 18 No. 2, pp. 415-438.
- Henley, W.E. and Hand, D.J. (1996), "A k-nearest-neighbour classifier for assessing consumer credit risk", *The Statistician*, Vol. 45 No. 1, p. 77.
- Henley, W.E. and Hand, D.J. (1997), "Statistical classification methods in consumer credit scoring: a review", *Journal of the Royal Statistical Society. Series A (Statistics in Society)*, Vol. 160 No. 3, pp. 523-541.
- Hill, T. and Lewicki, P. (2007), *Statistics: Methods and Applications*, StatSoft, Tulsa, OK.
- Jacobs, M. and Kiefer, N. (2010), "The Bayesian approach to default risk: a guide", in Böcker, K. (Eds), *Rethinking Risk Measurement and Reporting: Volume II*, Risk Books, London, pp. 319-334.
- Jie, L. and Bo, S. (2011), "Naive Bayesian classifier based on genetic simulated annealing algorithm", *Procedia Engineering*, Vol. 23, pp. 504-509.
- Kay, J.W. and Titterton, D.M. (2000), *Statistics and Neural Networks Advances at the Interface*, University of Glasgow.
- Komor'ad, K. (2002), "On credit scoring estimation", Master of Science thesis, Institute for Statistics and Econometrics, Humboldt University, Berlin.
- Lee, T. and Chen, I. (2005), "A two-stage hybrid credit scoring model using artificial neural networks and multivariate adaptive regression splines", *Expert Systems with Applications*, Vol. 28 No. 4, pp. 743-752.
- Lee, T., Chiu, C., Lu, C. and Chen, I. (2002), "Credit scoring using the hybrid neural discriminant technique", *Expert Systems with Applications*, Vol. 23 No. 3, pp. 245-254.
- Lundholm, R. and Sloan, R. (2004), *Equity Valuation and Analysis*, McGraw-Hill/Irwin, New York, NY.
- Maltritz, D. and Molchanov, A. (2008), "Economic determinants of country credit risk: a Bayesian approach", *Proceedings of the 12th New Zealand Finance Colloquium*, Massey University, Palmerston North.
- Martens, D., Van Gestel, T. and Baesens, B. (2009), "Decompositional rule extraction from support vector machines by active learning", *IEEE Transactions on Knowledge and Data Engineering*, Vol. 21 No. 2, pp. 178-191.

- Matoussi, H. and Krichène, A.A. (2010), "Credit risk evaluation of a Tunisian commercial bank: logistic regression versus neural network modelling", *The Journal of Accounting and Management Information Systems*, Vol. 9 No. 1.
- Matoussi, H., Mouelhi, R. and Salah, S. (1999), "La prédiction de faillite des entreprises tunisiennes par la régression logistique", *Revue Tunisienne des Sciences de Gestion*, Vol. 1, pp. 90-106.
- McCulloch, W. and Pitts, W. (1943), "A logical calculus of the ideas immanent in nervous activity", *Bulletin of Mathematical Biophysics*, Vol. 5 No. 4, pp. 115-133.
- Merton, R. (1974), "On the pricing of corporate debt: the risk structure of interest rates", *Journal of Finance*, Vol. 29 No. 2, pp. 449-470.
- Mileris, R. (2010), "Estimation of loan applicants default probability applying discriminant analysis and simple Bayesian classifier", *Economics and Management*, Vol. 15 No. 9, pp. 1078-1084.
- Mitchell, T.M. (2010), "Generative and discriminative classifiers: Naive Bayes and logistic regression", *Machine Learning*, Second edition chapter 3, McGraw Hill.
- Moonasar, V. (2007), "Credit risk analysis using artificial intelligence: evidence from a leading South African banking institution", Research Report: Mbl3.
- Odom, M. and Sharda, R. (1990), "A neural net model for bankruptcy prediction", *Proceeding Intern ate Joint Conference Neural Nets*, San Diego, CA.
- Ohlson, J.A. (1980), "Financial ratios and the probabilistic prediction of bankruptcy", *Journal of Accounting Research*, Vol. 18 No. 1, pp. 109-131.
- Okan, V.S. (2007), "Credit risk assessment for the banking sector of Northern Cyprus", *Banks and Bank Systems*, Vol. 2 No. 1.
- Palepu, K.G., Healy, P.M. and Bernard, V.L. (2000), *Business Analysis and Valuation Using Financial Statements*, 2nd ed., South – Western College Publishing.
- Pang, S.L., Wang, Y.M. and Bai, Y.H. (2002), "Credit scoring model based on neural network", *Proceeding of the First International Conference on Machine Learning and Cybernetics*, Beijing, 4-5 November.
- Provost, F., Fawcett, T. and Kohavi, R. (1998), "The case against accuracy estimation for comparing induction algorithms", in Shavlik, J. (Eds), *Proceeding ICML-98. Morgan Kaufmann, San Francisco, CA*, pp. 445-453, available at: www.purl.org/NET/tfawcett/papers/ICML98-final.ps.gz
- Quinlan, J.R. (1992), *C4.5: Programs for Machine Learning*, Morgan Kaufmann Publishers, California, CA.
- Raymond, A. (2007), *The Credit Scoring Toolkit: Theory and Practice for Retail Credit Risk Management and Decision Automation*, 1st ed., Oxford University Press.
- Revsine, L., Collins, D.W. and Johnson, W.B. (1999), *Financial Statement and Analysis*, Prentice Hall, New Jersey, NJ.
- Rosner, B.A. (2006), *Fundamental of Biostatistics*, Quebecor World, Taunton.
- Sarkar, S. and Sriram, R.S. (2001), "Bayesian models for early warning of bank failures", *Management Science*, Vol. 47 No. 11, pp. 1457-1475.
- Smith, C. and Warner, J. (1979), "On financial contracting", *Journal of Financial Economics*, Vol. 7 No. 2, pp. 117-161.
- Spackman, K.A. (1989), "Signal detection theory: valuable tools for evaluating inductive learning", *Proceeding Sixth Intern ate Workshop on Machine Learning*, Morgan Kaufman, San Mateo, CA, pp. 160-163.
- Steenackers, A. and Goovaerts, M.J. (1989), "A credit scoring model for personal loans", *Insurance: Mathematics and Economics*, Vol. 8 No. 1, pp. 31-34.
- Stibor, T. (2010), "A study of detecting computer viruses in real-infected files in the n-gram representation with machine learning methods", *23rd International Conference on Industrial Engineering and other Applications of Applied Intelligent Systems*, Part I, Cordoba, June 1-4, pp. 509-519.

-
- Sun, L. and Shenoy, P. (2007), "Using Bayesian networks for bankruptcy prediction: some methodological issues", *European Journal of Operational Research*, Vol. 180 No. 2, pp. 738-753.
- Thomas, L.C., Edelman, D.B. and Crook, J.N. (2002), *Credit Scoring and its Applications*. Society for Industrial Mathematics, 1st ed., Philadelphia.
- Thomas, L.C. (2002), "A survey of credit and behavioral scoring: forecasting financial risk of lending to consumers", *International Journal of Forecasting*, Vol. 16 No. 1, pp. 149-172.
- Townsend, R.M. (1979), "Optimal contracts and competitive markets with costly state verification", *Journal of Economic Theory*, Vol. 21 No. 2, pp. 265-293.
- Miguéis, V.L., Benoit, D.F. and Van den Poel, D. (2012), "Enhanced decision support in credit scoring using Bayesian binary quantile regression", Working Paper.
- West, D. (2000), "Neural network credit scoring", *Computer and Operations Research*, Vol. 27 No. 11, pp. 1131-1152.
- Wu, C. and Wang, X.M. (2000), "A neural network approach for analyzing small business lending decisions", *Review of Quantitative Finance and Accounting*, Vol. 15 No. 3, pp. 259-276.
- Yang, L. (2002), "The evaluation of classification models for credit scoring", Working Paper No. 02/2002 Edit, Matthias Schumann University of Göttingen Institute of Computer Science.

Further reading

- Berg, D. (2005), "Bankruptcy prediction by generalized additive models", Statistical Research Report No. 1, University of Oslo.
- Schmidt-Mohr, U. (1997), "Rationing versus collateralization in competitive and monopolistic credit markets with asymmetric information", *European Economic Review*, Vol. 41 No. 7, pp. 1321-1342.
- Wong, B.K. and Selvi, Y. (1995), "A bibliography of neural networks business applications research", *Expert Systems*, Vol. 12 No. 3, pp. 253-262.

Corresponding author

Aida Krichene can be contacted at: aidakrichene@yahoo.fr

Supervised Learning 1 (Naive bayes continuous)	
Parameters	
Parameters	
Lambda for laplacian	0.0
Homoscedasticity assumption	1
Results	

Classifier performances

Error rate			0.4037			
Values prediction			Confusion matrix			
Value	Recall	1-Precision		RISKY	HEALTHY	Sum
RISKY	0.4721	0.3660	RISKY	220	246	466
HEALTHY	0.7227	0.4263	HEALTHY	127	331	458
			Sum	347	577	924

Figure A1.
Non-cash flow Bayes
model

Supervised Learning 1 (Naive bayes continuous)	
Parameters	
Parameters	
Lambda for laplacian	0.0
Homoscedasticity assumption	1
Results	

Classifier performances

Error rate			0.3615			
Values prediction			Confusion matrix			
Value	Recall	1-Precision		RISKY	HEALTHY	Sum
RISKY	0.6073	0.3479	RISKY	283	183	466
HEALTHY	0.6703	0.3735	HEALTHY	151	307	458
			Sum	434	490	924

Figure A2.
Full information
Bayes model