

CORVINUS UNIVERSITY OF BUDAPEST

**APPLICATION OF DYNAMIC FINANCIAL VARIABLES IN  
BANKRUPTCY PREDICTION**

PHD. THESIS

Supervisor: Prof. Dr. Miklós Virág, CSc

Tamás Nyitrai

Budapest, 2015

**Tamás Nyitrai**

Application of dynamic financial variables in bankruptcy prediction

## **Department of Enterprise Finances**

Supervisor: Prof. Dr. Miklós Virág, CSc

**Corvinus University of Budapest**

**Doctoral School of Management and Business**

**Administration**

**Application of dynamic financial variables in bankruptcy  
prediction**

PhD Thesis

Tamás Nyitrai

Budapest, 2015

## TABLE OF CONTENTS

|                                                                               |    |
|-------------------------------------------------------------------------------|----|
| Table of contents .....                                                       | 5  |
| Table of figures .....                                                        | 8  |
| List of tables .....                                                          | 9  |
| 1. Introduction .....                                                         | 10 |
| 1.1. Motivation for selecting this topic .....                                | 10 |
| 1.2 The conceptual framework of bankruptcy prediction .....                   | 13 |
| 1.3. Content of the thesis.....                                               | 18 |
| 2. Literature review .....                                                    | 21 |
| 2.1. Possible ways of reviewing the literature.....                           | 21 |
| 2.2. Chronological review of the literature on bankruptcy prediction .....    | 22 |
| 2.3. „Cross-sectional” review of the literature on bankruptcy prediction..... | 26 |
| 2.3.1. Sample selection.....                                                  | 28 |
| 2.3.1.1. Activity of firms under consideration .....                          | 28 |
| 2.3.1.2. Time interval covered by the sample .....                            | 31 |
| 2.3.1.3. Geographical aspects of sample selection.....                        | 34 |
| 2.3.1.4. Problems of dichotomous classification.....                          | 34 |
| 2.3.1.5. The number of instances in the classes .....                         | 37 |
| 2.3.2. Data preprocessing .....                                               | 39 |
| 2.3.2.1. Calculating problems .....                                           | 39 |

|                                                                                              |    |
|----------------------------------------------------------------------------------------------|----|
| 2.3.2.2. Industry effects.....                                                               | 40 |
| 2.3.2.3. Preprocessing data for modelling.....                                               | 43 |
| 2.3.2.4. The question of static bankruptcy prediction models .....                           | 46 |
| 2.3.3. Feature selection.....                                                                | 48 |
| 2.3.3.1. Difficulties of feature selection.....                                              | 48 |
| 2.3.3.2. Other data sources for bankruptcy prediction .....                                  | 49 |
| 2.3.3.3. Feature selection methods .....                                                     | 52 |
| 2.3.4. Constructing bankruptcy model .....                                                   | 55 |
| 2.3.4.1. Multivariate statistical methods .....                                              | 57 |
| 2.3.4.2. Nonparametric methods based on machine learning .....                               | 60 |
| 2.3.4.3. Hybrid methods and ensembles .....                                                  | 70 |
| 2.3.4.4. Results of methodological comparative studies.....                                  | 73 |
| 2.3.5. Performance evaluation of the model .....                                             | 76 |
| 2.3.5.1. Measuring the classification performance .....                                      | 76 |
| 2.3.5.2. Validation procedures .....                                                         | 80 |
| 2.4. The Hungarian literature of bankruptcy prediction.....                                  | 85 |
| 2.4.1. The beginning of bankruptcy prediction in Hungary .....                               | 85 |
| 2.4.2. Neural networks in the Hungarian bankruptcy prediction .....                          | 86 |
| 2.4.3. Appearance of state-of-the-art methods in the Hungarian bankruptcy<br>prediction..... | 87 |
| 3. Research hypotheses .....                                                                 | 90 |
| 3.1. Open research questions in the literature .....                                         | 90 |

|                                                                                 |     |
|---------------------------------------------------------------------------------|-----|
| 3.1.1. Dynamic financial variables in bankruptcy prediction .....               | 91  |
| 3.1.2. Handling outliers .....                                                  | 93  |
| 3.2. Research hypotheses examined in the thesis .....                           | 95  |
| 4. Data used in the empirical research .....                                    | 98  |
| 4.1. General sampling issues .....                                              | 98  |
| 4.2. The process of data gathering .....                                        | 100 |
| 5. Empirical research.....                                                      | 106 |
| 5.1. The examination of the first research hypothesis.....                      | 106 |
| 5.2. The examination of the second research hypothesis .....                    | 111 |
| 5.3. The examination of the third research hypothesis .....                     | 114 |
| 5.4. The summary of the results obtained in the empirical examinations.....     | 120 |
| 6. Summary .....                                                                | 122 |
| 6.1. The contribution of the thesis to the development of bankruptcy prediction | 123 |
| 6.2. The limitations of the analysis presented in the thesis .....              | 126 |
| 6.3. Possible future research directions.....                                   | 128 |
| Appendix .....                                                                  | 131 |
| References .....                                                                | 133 |
| Own publications related to this topic.....                                     | 152 |

## TABLE OF FIGURES

|                                                                                      |    |
|--------------------------------------------------------------------------------------|----|
| Figure 1. The process of bankruptcy prediction model building .....                  | 27 |
| Figure 2. The effect of changing the distribution of financial ratios .....          | 41 |
| Figure 3. Categorization of classification methods used in bankruptcy prediction ... | 56 |
| Figure 4. A simple example for decision trees.....                                   | 60 |
| Figure 5. The general structure of neural networks .....                             | 63 |
| Figure 6. The basic idea of SVM .....                                                | 65 |
| Figure 7. Classification of SVM – two dimensional case.....                          | 66 |
| Figure 8. Confusion matrix .....                                                     | 76 |
| Figure 9. ROC curve .....                                                            | 78 |
| Figure 10. The time series of Cash flow/Debts ratio .....                            | 93 |
| Figure 11: A time series containing outlier value .....                              | 94 |

## LIST OF TABLES

|                                                                                                               |     |
|---------------------------------------------------------------------------------------------------------------|-----|
| Table 1. Publications comparing classification methods in bankruptcy prediction...                            | 74  |
| Table 2. The name and formula of financial ratios used in the empirical research .                            | 103 |
| Table 3. The average hit rates obtained by applying tenfold cross-validation using CHAID decision trees ..... | 113 |
| Table 4. A time series containing outlier value.....                                                          | 115 |
| Table 5. The time series containing outlier value after the replacement .....                                 | 117 |
| Table 6. The hit rates of models developed in the empirical analysis .....                                    | 119 |
| Table F1. Descriptive statistics of financial ratios used in the empirical analyses ..                        | 131 |

# **1. INTRODUCTION**

In my work, I summarize the results of my research conducted during my PhD studies. I have studied at Management and Business Administration Doctoral School at Corvinus University of Budapest specialized in corporate finance. My research on the field of bankruptcy prediction, which belongs to the most important problems of business classification (Hu-Tseng [2007]), was supervised by Prof. Dr. Miklós Virág. In the introduction of my thesis, you can read about the followings:

- motivation for selecting this topic;
- presentation of terms used in my work;
- summary of the content of the thesis.

## **1.1. Motivation for selecting this topic**

Enterprises are continuously forming and closing down. This is a natural phenomenon of the economy. There are two forms of closing down a firm: voluntarily and involuntary. In the first case, the owners of the business entity voluntarily decide on stopping the activity, pay the debt of the enterprise and then divide the assets left in it. In the legal system of Hungary, the voluntary dissolution makes this possible.

However, there are also some involuntary types of terminating business activity. These happen when a business entity is unable to fulfill its payment obligations when they come due. In this case, the creditors of the firm can initiate liquidation proceeding against the debtor. During this proceeding, debts of the firm are paid

under the control of court. If the assets of the debtor doesn't cover the debts, creditors may lose their claims. This case creates loss to owners as well because they lose the sum of their capital invested in the firm.

It should be emphasized that insolvency and failure of firms may create loss not only for creditors and owners. In such cases, the management and the employees of the firm also have to bear the cost of losing their jobs, furthermore, customers and suppliers of the failed firm also belong to this group because they lose their business partners.

Costs created by the insolvency may be significant for every stakeholder of the failed firm. In the case of large number of bankruptcies, the costs can be measurable at the level of the whole economy. Because of this reason, bankruptcy prediction is one of the most important topics in the finance and accounting literature, according to Kim-Kang [2012]. Similarly, according to Shetty et al. [2012], bankruptcy of firms is one of the major threats for the economy because great number of failures may lead to economic and social problems.

Though to different extent, but every economy had to face the consequences mentioned above during the crisis started in 2008. According to some researchers, the basis of this recession was that creditors didn't properly evaluate the risk of their contracts (Riberio et al. [2012]). Because of the reasons mentioned above, Cao [2012a] thinks that the recession has lifted the importance of predicting financial distress and bankruptcy to a level never seen before.

Chen et al. [2013] have similar opinion. They think that this topic has been investigated for a long time, however, the crisis raised the attention of the research community to the questions still unanswered in this field. In the years of recession, it

became important to identify firms at potential risk of bankruptcy in order to minimize the costs associated with the failure.

Besides the recession, new impetus has been given to the development of bankruptcy prediction by initiating the Basel II Accord because financial institutions have become interested in applying the most accurate credit scoring method in their internal rating systems in order to meet the prescriptions of the Accord (Brown-Mues [2012]). Bankruptcy prediction models play an important role in the internal rating systems because a good model makes it possible to minimize the amount of money to be held in the form of regulatory capital and to increase the profit coming from lending (Sánchez-Lasheras et al. [2012]).

On the basis of the above mentioned facts, bankruptcy prediction can be considered as a “hot topic” in the field of business sciences (Wang-Ma [2012]). This is reflected in the huge number of studies published on this field in the last 50 years. The main reasons for this phenomenon are the followings:

- the development of computer sciences supporting this research area;
- the accessibility of financial data for private firms;
- the great number of open research questions.

The trends described above can be observed in Hungary as well, but the research possibilities haven't been exploited yet. On the ground of this, I have also decided to choose this topic. During my research, I also tried to take advantage of the developments described above.

## **1.2 The conceptual framework of bankruptcy prediction**

From grammatical point of view, the name of this field can be considered as a compound word. Its elements should be defined exactly in order to confine the domain of the work.

In the case of word for word translation, the field deals with the bankruptcy of economic entities. However, this statement is not completely right. For this topic, “financial distress prediction” and “business failure prediction” are also commonly used terms in the international literature.

The problem of the traditional name (bankruptcy prediction) is that it restricts the domain of the investigation to the bankruptcy of business entities (typically joint businesses). The two other names presented above interpret the object of the prediction in a wider sense. The drawback of these can be found in their subjectivity, because “business failure” and “financial distress” are not exact terms, so they should be defined exactly.

The literature of bankruptcy prediction is not uniform with respect to the target of the prediction. Consensus appears only in the fact that the aim of the publications in this field is to predict the occurrence of events which can create loss to the relevant interest groups (typically creditors and owners) of the companies.

In the case of setting up a company, the owners invest and risk their own private wealth in the hope of reaching more profit over their investment than it is possible elsewhere. When the debtor becomes insolvent, the creditors may lose the amount of their credit. This event creates loss to owners as well because the subdivision of the capital invested in the company by the owners is overtaken by the claims of the creditors in the hierarchy to be applied in the case of liquidation.

When a company voluntarily ends its operation, the debt of the firm will be paid in the dissolution proceeding and then the capital left in the company will be divided among the owners. In this case, loss can be created only to the owners if the market value of the capital invested in the company is less than the amount invested by the owners earlier; however, this loss can be “planned” because the owners took the risk of losing their private wealth invested in the company in the hope of gaining higher return than it was obtainable elsewhere at the time of investing. That is why voluntarily termination of firms is neither subject of bankruptcy prediction nor of my thesis.

Starting of liquidation or bankruptcy proceeding against a company can be viewed as an event with financial relevance because both can be linked to the solvency of firms. Firms unable to meet its financial obligations for a long time can ask for moratorium in order to enter into an arrangement with creditors according to Section (1) of Paragraph (1) of Act XLIX of 1991 on bankruptcy and liquidation procedure. So, the aim of bankruptcy procedure is that the company gains some time to meet its financial obligations and then continues its operation. In contrast, the purpose of liquidation procedure is to pay the debts of the company and eliminate the debtor without successor.<sup>1</sup>

The ground of starting both of the above mentioned procedures is the fact that the debtor unable or unwilling to meet its financial obligations as they come due. Because of it, bankruptcy and liquidation procedures constitute relevant risk for creditors with respect to their claims and for the owners as well who may lose their capital invested in the company. In an extreme case, it could happen that the whole

---

<sup>1</sup> Source: Section (3) of Paragraph (1) of Act XLIX of 1991 on bankruptcy and liquidation process.

amount of debtors' claim and capital invested by the owners get lost due to the insolvency of the debtor.

In my thesis, the subjects of the prediction are insolvent firms registered in Hungary against which bankruptcy or liquidation procedure has been initiated. It should be noted that the consequence of such procedures are not obvious: initiation of bankruptcy procedure doesn't necessarily mean the survival of the firms, and similarly, a liquidation procedure can be ended up with surviving and paying debts. Consequently, under the term of "bankruptcy prediction", I understand insolvency prediction in the rest of my work.

This definition is wider than the word for word translation of bankruptcy prediction, but at the same time, this approach means restriction as well since bankruptcy and liquidation procedures are initiated against companies which haven't met their financial obligations for some time, that is, the initiation of these procedures can be considered as a final outcome of the insolvency of firms. Consequently, insolvent firms against which none of the above-mentioned proceedings has been initiated were excluded from the investigation presented in the thesis. In the literature of bankruptcy prediction, this group is considered as firms in financial distress (Gilbert et al. [1990]). Predicting the future for such firms is difficult because of the absence of objective information about their financial status. The initiation of bankruptcy and liquidation procedures can be checked in the official registration, however the term of "financial distress" hasn't been objectively defined in the literature, so the existence of such a circumstance cannot be identified unambiguously.

As a part of the presentation of the conceptual framework of bankruptcy prediction, the term "prediction" should also be defined properly. From methodological

perspective, this field can be considered as a dichotomous classification problem whose aim is to distinguish between two types of firms: failed and non-failed as accurate as possible. In the rest of my work, I also use the terms of “normal” and “bankrupt” for the above-mentioned two kinds of companies.

As it has been noted, occurrence of insolvency can objectively be observed only from the data of Trade Register, so firms unable to meet their financial obligations for some time belong to the healthy group. These enterprises are insolvent for some time, but it is open to doubt whether either of the aforementioned procedures will be initiated against them. It depends on the fact whether the firm is able to solve its financial problems in short run or how long creditors of the firm can or willing to wait for receiving their claims. This kind of firms is especially difficult to classify correctly since they generally have weak financial data, but they are still operating and have real opportunity to do so in the future.

Taking into account the aforementioned problem, bankruptcy prediction models are called as “early warning systems” (EWS) by numerous authors in the literature. Traditional bankruptcy prediction models classify firms into two groups: distressed and non-distressed which can be problematic because of the aforementioned difficulties of predicting the future of firms having financial troubles. EWS models differs from traditional bankruptcy prediction models only in that their classification is not so exact. For example, in the framework of EWS, if a company is classified into bankrupt group by a model, it doesn't mean that the company will certainly go bankrupt in the future. Based on this classification, it can only be concluded that the firm under consideration, with respect to its financial data, is more similar to bankrupt firms than to the operating group (Virág-Hajdu [1996]). This also means that the company exhibits substantial risk of future insolvency which should be taken

into consideration by the creditors. Based on the idea described above, Agarwal-Taffler [2007] think that bankruptcy prediction models can be considered as similar tools to fever thermometers since neither of them precisely diagnoses the illness of the patient or predict the recovery, but gives signs about the existence and the severity of the problem and raises the attention to the importance of intervention.

This approach distances itself from predicting insolvency as an aim of research, however, it's not possible to objectively measure the performance of EWS models which is unacceptable in science. For this reason, the aim of my thesis is also to predict future insolvency of Hungarian enterprises, however I would like to emphasize that I also agree with the early warning approach and its core concept, namely that the predictions of bankruptcy models can be interpreted only as early warning signals of possible future insolvency.

In order to achieve the aforementioned research aims, bankruptcy prediction models are developed which represent statistical relationships between independent variables used for modelling and the fact of insolvency as dependent variable. In the former group, ratio type financial variables extractable from accounting information system of firms are commonly used by researchers and practitioners as well. Following this trend, I also use these financial ratios in the thesis.

The relationship to be explored between financial variables and future bankruptcy was examined earlier by means of classification techniques developed in the field of mathematical statistics, but recently, methods of artificial intelligence are also commonly used for this purpose by researchers. The methodological background and use of these methods are presented in the subsection 2.3.4 in more detail.

Based on the above detailed facts, bankruptcy prediction can be considered as a multidisciplinary topic located on the border of corporate finance and statistics (data mining). Research in this field tries to predict future solvency of firms by using financial ratios (and occasionally other factors) as explanatory variables in a suitable multivariate classification method.

### **1.3. Content of the thesis**

The first steps in scientific research of bankruptcy prediction date back to the mid-60s when financial data of public firms listed on the US Stock Exchange were examined in order to find out whether it is possible to predict future bankruptcy of companies under consideration based on the values of commonly used financial ratios. Predicting bankruptcy of public firms has still been an active area of research because data for these firms are easy to obtain, furthermore, another important issue is the reliability of data for firms traded on the stock exchange since their financial statements have to be audited.

A more recent area of research in this field is the assessment of probability of bankruptcy for micro, small and medium sized companies (SME) where availability of data is a challenging issue because of the private operation. Furthermore, it should also have been raised the attention to the reliability of data because most of private firms in Hungary is not obliged to audit their financial statements. In spite of these difficulties, there's a need from the creditors for evaluating the probability of future bankruptcy of such smaller enterprises as well (Riberio et al. [2012]).

Financing SMEs is also very relevant from macroeconomic point of view because this kind of firms play important role with respect to the most crucial economic

indicators, such as employment, GDP, innovation etc. (Koyuncugil-Ozbulbas [2012]). However, bankruptcy models developed on the basis of data for public firms are not necessarily suitable to predict insolvency of smaller private enterprises.

With respect to the fact that the proportion of public companies is relatively small in our country, the aforementioned aspect is important in the Hungarian literature of bankruptcy prediction. In this country, it is impossible to develop and use bankruptcy prediction models based on data of public companies. The scientific examination of bankruptcies of Hungarian private firms was also restricted by the availability of data until 2009.

From 2009, companies registered in Hungary are obliged to publish their financial statements in electronic form retroactively until 2000. Financial statements are available online for anyone on the website maintained by the Ministry of Public Administration and Justice.<sup>2</sup>

The ground of my research conducted during my PhD studies was the free availability of data which made it possible to scientifically examine the bankruptcy of Hungarian enterprises on the basis of much bigger database than any other used previously in the Hungarian bankruptcy prediction literature. The results of this research work are summarized in the following sections.

During my doctoral studies, I have paid special attention to reviewing the Hungarian and international publications in this field since relevant research questions can be raised only in the light of the existing scientific results. In bankruptcy prediction, presentation of the literature is usually carried out chronologically. In my thesis, I try to do this in such a way that publications will be assigned to the corresponding phase

---

<sup>2</sup> <http://e-beszamolo.kim.gov.hu/kereses-Default.aspx>

of bankruptcy model building process. The underlying principle of such a presentation is that the literature of bankruptcy prediction has become so diverse that each phases of the model building can be considered as a standalone research topic. I hope that this new approach simultaneously gives comprehensive picture about the recent research questions of bankruptcy prediction as well as about their evolution.

Based on the reviewed literature, I'm going to outline the research questions studied in the thesis. Answers given to those in my work, hopefully, can go further than the results of the last 50 years of bankruptcy prediction and raise new future research questions.

The rest of my thesis is organized as follows: in the next section, publications and their most important research results of the international and Hungarian literature on bankruptcy prediction will be presented. Based on those, in section 3, research hypotheses of my work will be constructed, then, in the next section, I give an overview about the applied sampling method and the resulting database used to verify the hypotheses. Section 5 presents the empirical examinations and their results and finally, the last section summarizes the results of the thesis, the conclusions coming from them, furthermore, possible future research directions are also highlighted at the end of my work.

## **2. LITERATURE REVIEW**

This section wishes to briefly review the literature of bankruptcy prediction. Though, with respect to its size, it is the biggest part of the thesis, it should be noted that this review cannot be viewed as a full summary. The number of publications in this field has been increased dramatically in the last ten years, so, due to space limitations, it is not possible to present every research question and every empirical investigation related to each aspects of bankruptcy prediction. Instead of giving a full review, I present only some “representative” empirical examinations in detail related to the research directions discussed in the thesis. Studies dealing with similar questions are referred in the following subsections. In those articles, the Reader can find a lot of interesting empirical research containing similar conclusions with those discussed in my work. In these papers, references to other investigations can also be found.

### **2.1. Possible ways of reviewing the literature**

Besides the citations from the authors in the introduction, practical importance of bankruptcy prediction is also proven by the fact that the number of publications in this field is so high that standalone studies deal with summarizing and reviewing the existing literature. Abdou-Pointon [2011] present the actual state of credit scoring literature, raise the attention to the most challenging issues and outline possible future research directions as well. Similarly, Jones [1987] devoted a study to collect the methods applicable in the field of bankruptcy prediction. Balcean-Ooghe [2006] summarized the possible problems which may be arisen when these techniques are applied.

The literature of bankruptcy prediction is often reviewed chronologically. This approach is adequate and common in scientific research, however, in my opinion, it is less and less suitable to give comprehensive picture for the Reader about the actual state and the main research questions of bankruptcy prediction due to the dynamic increase in the number of publications in this field.

In my opinion, this purpose can be achieved more easily by applying an approach which gathers the research results around the main phases of bankruptcy model building. The ground of this approach is the fact that newer publications deal with research questions corresponding to one (or some) phase of the process of bankruptcy model building, in contrast to the pioneering works of this field where this process was “globally” viewed.

In my work, the reviewed literature is assigned to each phase of bankruptcy model building. However, with respect to the practice of international journals, where the application of the chronological approach is common, the most important milestones of bankruptcy prediction are reviewed in chronological order as well.

## **2.2. Chronological review of the literature on bankruptcy prediction**

Already in the first half of the 20th century, uncertainty related to future solvency of firms had relevant effect on business decisions. The first analysis conducted via scientific tools was done by Beaver [1966] in the mid-60s. He compared the most popular financial ratios calculated from data for public manufacturing firms. 30 variables were investigated for the five years period prior to bankruptcy. Conclusions drawn from this experiment have still been playing important role in the literature of

bankruptcy prediction. So, almost every studies published in this field refers to this pioneering work in their introduction.

The research effort of Beaver should be highlighted because he was the first to have scientifically dealt with the predictability of future insolvency of firms. It should also be emphasized that most of his conclusions has still been valid and raises further fruitful future research directions which can contribute to enhancing the predictive power of bankruptcy prediction models. My research can also be considered as an example for this since the main hypothesis of my work has been formulated on the basis of his investigations.

The method applied by Beaver [1966] is known as univariate discriminant analysis in the literature. The essence of the methodology is that it investigates only one financial variable to distinguish between healthy and bankrupt companies. This approach was heavily criticized in the literature later because its background is not “scientific enough”. The cited author itself also emphasized that there is another serious drawback when a company is classified into the healthy group according to the values of some financial variables but into the bankrupt group when different financial ratios are under consideration.

The possible conflict of different kinds of financial variables can be solved by using multivariate statistical methods. First, multivariate discriminant analysis (MDA) was employed in the work of Altman [1968] which is known worldwide and which can also be found in the famous handbook of corporate finance (Brealey-Myers [2005]) commonly used at universities. Further methodological issues of this approach are discussed in the subsection 2.3.4.

MDA gained popularity in the literature because of its capability of answering the three main research questions of bankruptcy prediction (Virág [2004]):

- which financial ratios have statistically significant discriminating power between healthy and bankrupt groups of companies;
- which weights should be assigned to these variables in models;
- in what way these weights can be determined objectively.

The bankruptcy prediction model developed by Altman [1968] exhibited 95 % prediction accuracy on a sample consisting of 33 healthy and 33 bankrupt manufacturing firms from the United States. Due to this excellent classification performance, discriminant analysis became the primary tool in bankruptcy prediction till the end of the 1970s (see for example the work of Deakin [1972] and Altman et al. [1977]).

The popularity of Altman's [1968] model is proven by the fact that it has still been the subject of comparative studies (see for example the works from Altman [2002], Philosphov-Philosphov [2002], Hillegeist et al. [2004], Neves-Vieira [2006], Parnes [2007], Agarwal-Taffler [2008], Sueyoshi-Goto [2009]). Numerous studies report surprisingly good predictive power related to this model published almost 50 years ago.

From the end of the 70s until now, the development of this research field has been dominated by the modernization of mathematical-statistical methods capable of solving classification problems and computer science providing its background. The next important step was the appearance of logistic regression in the field of bankruptcy prediction in the work of Martin [1977]. After his work, this method became common due its relative flexibility, since it doesn't require to meet so

rigorous assumptions like normal distribution of independent variables and equal variance-covariance matrices among the groups under consideration (Ohlson [1980], Zavgren [1985]).

Next to logistic regression, similar reasons led to the application of probit regression (Zmijewski [1984], Sjøvoll [1999], Bongini et al. [2000], Bernhardsen [2001], Grunert et al. [2005]) and decision trees (Frydman et al. [1985]). The latter don't require to meet any statistical assumption. Thanks to this property, they have been popular and commonly used until now.

Neural networks simulating the operating mechanism of nervous systems were firstly applied by Odom-Sharda [1990] in bankruptcy prediction. Their work can be considered as a milestone in the evolution of this topic not only for this reason. By the appearance of neural networks, in addition to the traditional statistical methods (discriminant analysis, logistic regression) and decision trees, it became possible to apply procedures of artificial intelligence as well. This possibility raised the following research question: Which method can be considered as the best in bankruptcy prediction?

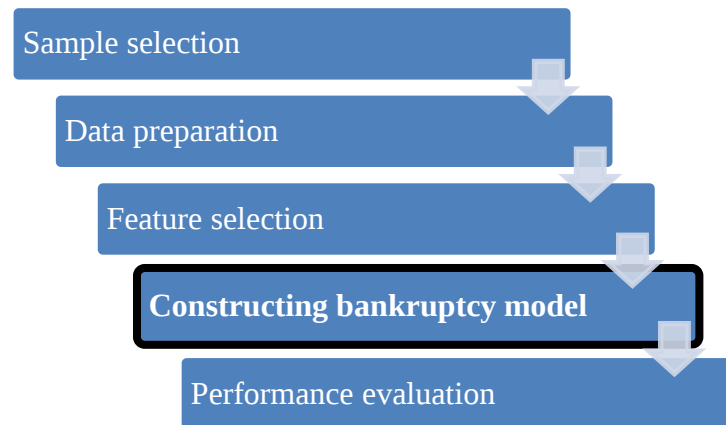
There has been no unambiguous answer to this question in the literature. According to empirical research results, each methods exhibit different forecasting ability when applied on different datasets (Oreski et al. [2012]). Marqués et al. [2012a] came to the conclusion that it is very possible that the "best method" doesn't exist. In spite of this, comparative studies seeking this "best method" has been dominating the literature since the early years of the 90s. With respect to the huge number of studies on this topic (Sánchez-Lasheras et al. [2012]), this kind of comparative analyses can be considered as the mainstream research direction in bankruptcy prediction.

It should be emphasized that, in addition to the mainstream, there are also other alternative research directions in bankruptcy prediction which try to enhance the predictive power of the models but not by applying a newer or better classification algorithm. The number of studies related to this research direction is much lower and they have been published in a scattered manner in the almost 50 years of literature of bankruptcy prediction. So, they cannot be embedded directly in the continuous research trend which is typical of the mainstream dominated by the development of classification methods. Because of this fact, it is worth considering the evolution of bankruptcy prediction from a “cross sectional” point of view in which, besides the mainstream, these alternative research directions playing similarly important role can also appear.

### **2.3. „Cross-sectional” review of the literature on bankruptcy prediction**

I think that this approach is more useful than the chronological one because more comprehensive picture can be given about the literature of bankruptcy prediction by applying this “cross-sectional” point of view. The process of bankruptcy prediction model building is shown in Figure 1. In this section, I present the main stages of this process along with the corresponding research questions and the most relevant empirical results which can be found in the literature related to them.

**Figure 1. The process of bankruptcy prediction model building**



In the figure, the most important steps of building bankruptcy prediction models can be seen in the time order of the execution. Every studies on bankruptcy prediction can be assigned to one<sup>3</sup> of those shown in Figure 1 dependent on the decision of the author(s) about the aim of his/her/their research. In the figure, methodological comparative studies, which can be considered as mainstream research direction in this field, are categorized into the phase called as “Constructing bankruptcy model” highlighted by black line. Parameter optimization of data mining techniques can also be classified in this group.

It should be noted, that literature review categorized according to the phases of model building doesn’t require to give up the application of the chronological order. It is reasonable to chronologically present the studies related to each phases in order to let the development of each phases’ literature to be outlined in addition to the mainstream research direction.

---

<sup>3</sup> More complex studies can also be assigned to more phases simultaneously.

### **2.3.1. Sample selection**

In Figure 1, sample selection is shown as the first phase of bankruptcy model building. Here, two aspects should be taken into consideration: the number of instances and the number of independent variables available to characterize them. To build statistical models, it is essential to have enough observation from bankrupt and operating groups as well. Along with these trivial requirements, there are several other research questions which can be assigned to the tasks of sample selection. The most important aspects of it will be presented in the following subsections.

#### ***2.3.1.1. Activity of firms under consideration***

An important factor may be the activity of the firms in the sample. Financial companies are often excluded from the databases of bankruptcy models with respect to their special characteristics. Specialized bankruptcy prediction models should be used for these firms. These models have specialized literature as well (see for example the early works of Pantalone-Platt [1987a], Tam-Kiang [1992], Sarkar-Sriram [2001], furthermore Kolari et al. [2002]; or more recent studies from Ruzgar et al. [2008], Boyacioglu et al. [2009], and Quek et al. [2009]). As exceptions can be mentioned the research by Jones-Hensher [2004] and Duan et al. [2012] who didn't excluded financial firms from the database used to construct bankruptcy prediction models.

Negative impacts of the financial crisis started in 2008 caused several bankruptcies in the banking sector as well. Because of it, the need for actualizing predictive models able to forecast bankruptcy of financial institutions has increased from the side of regulating authorities. Data for banks became bankrupt during the crisis was

used to construct model by Ioannidis et al. [2011], Serrano-Cinca-Guittérrez-Nieto [2013]. In Japan, Harada et al. [2013] tried to predict bankruptcy of financial institutions. The failure of insurance companies was investigated by Segovia-Vargas [2003] and Salcedo-Sanz et al. [2004].

Classification techniques commonly used in this field can be employed not only for predicting insolvency. For example, using financial ratios, Jagtiani et al. [2003] tried to predict whether the capital adequacy ratio of banks in the following year will reach the level set by the regulator or not. Bellotti et al. [2011] examined the predictability of Fitch ratings for banks. For banks from Tunisia, the ratio of defaulted claims was assessed by Feki et al. [2012] by means of classification methods.

Consensus can be observed in the literature about the fact that specialized models should be used when dealing with the risk of bankruptcy operating in the financial industry. However, in the case of other industries, similar consensus doesn't exist. Some researchers use samples consisting of firms only from a particular industry or industries which can be considered as similar to each other, at the same time, other scholars pay less attention to the activity of firms. This issue was examined by Gaganis et al. [2007]. They raised the question of whether it is necessary to develop specialized bankruptcy prediction models for each industry or it is sufficient to use a "general" model without taking into consideration the firm's activity. Due to their contradictory empirical results, they couldn't draw unambiguous conclusions with respect to this question. Chava-Jarrow [2004] also examined this issue. The activity of firms was taken into consideration via dummy variables in the models, but the empirical results were only slightly better than those of models without these dummy variables. In contrast, Neophytou-Mar Molinero [2005] found significantly better

performance when the companies' activity was taken into consideration. Based on their empirical results, they concluded that different models are needed for companies operating in different industries. Their results were confirmed by Glennon-Nigro [2005] who found that companies in the retail sector exhibit greater risk than others in the service industry.

The ratio of failed firms in different industries at macroeconomic level was investigated in Germany by Rösch [2003] who found substantial differences between the sectors. He explained these deviations (among others) with the fact that different activities have different risk levels. Failure rates of different industries were modeled with macroeconomic variables by Harada-Kageyama [2013]. They found slight differences between the models which indicate the relative importance of firms' activity related to the risk of bankruptcy.

Other researchers (such as for example Manjón-Antolin – Arauzo-Carod [2008], Chen [2012]) think that samples for bankruptcy prediction should be restricted to only one sector. The necessity of this is explained by Huynh et al. [2010] with the fact that if the sample for constructing bankruptcy prediction model consists of companies from different industries, it is possible that one finds bankrupt and operating companies with completely same financial ratios. For such cases, only the activity of the firms can have discriminatory power between the failed and the healthy company showing the same financial profile. According to Beaver et al. [2005], public utilities and financial companies are necessary to exclude from databases used to build bankruptcy prediction models. Specialized models are needed for these firms. According to Demers-Joos [2007], traditional and high-tech companies are necessary to distinguish before developing bankruptcy prediction

model. Another example for this is the research of Bose [2006] who examined only dotcom enterprises.

In spite of the fact that consensus cannot be observed in the literature related to the effectiveness of restricting samples to only one sector, this practice is very often in the international literature of bankruptcy prediction. For example, Foreman [2003] investigated firms only from the telecommunication industry of the US. Pindado-Rodrigues [2004] used data only for shoes manufacturing firms. Other examples: electronic companies listed on Taiwanese stock exchange (Lin et al. [2009], Yeh et al. [2010]), firms from the construction industry (Chen [2012]), IT enterprises from India were also examined (Shetty et al. [2012]) separately.

Deawelheyns-Van Hulle [2006] examined another interesting aspect related to sample selection. Their empirical results showed significant difference between the predictability of bankruptcy of independent firms and subsidiaries. So, this issue also deserves attention in the practice of modelling.

#### ***2.3.1.2. Time interval covered by the sample***

The state of the whole economy also has significant impact on the accuracy of bankruptcy prediction models. Foundation-stones of bankruptcy prediction had been laid down in Anglo-Saxon countries in the 60s and 70s and after this time data for firms listed on the stock exchanges of the USA and Western Europe were used by the researchers of this field. Since these regions weren't affected by relevant recessions during these decades, researchers didn't have the possibility to examine the effect of the state of macroeconomy. The need for this kind of research was emerged at first in the Far East. Sung et al. [1999] observed that the performance of bankruptcy

prediction models developed on data from booming years decreased significantly during the recession affected only the countries of East Asia in the second half of the 90s. Based on their results, it can be concluded that unique models are needed for the years of prosperity as well as for the periods of recession in order to predict future insolvency of firms as accurate as possible.

Another aspect of the recessions was examined by Bongini et al. [2000]. They studied what kind of firms are more vulnerable to economic depression. Downturns also have positive impact which is known as market cleaning effect. This means that inefficient firms exit from the market during the years of recessions. One can think that smaller firms with lower equity have higher probability to go bankrupt during these difficult times. In contrast, the authors cited above found that more companies in relatively good financial position also fell prey to the recession restricted only to countries in East Asia. According to the analysis conducted by these authors, companies operating in conglomerates had the highest probability to survive the downturn hit this region in the late 1990s.

Empirical examinations were conducted on the basis of the recession in East Asia by Nam-Jinn [2000] and Nam et al. [2008] as well. Economic circumstances altered by the impacts of the recession were incorporated into bankruptcy prediction model by the latter cited authors. The predictive power of their model significantly increased due to this modification.

The most recent economic depression has had negative impacts on the stock exchanges of the USA and Western Europe as well. Using data of French enterprises, similar conclusions were drawn by Du Jardin-Séverin [2012] than the authors examining the impact of the crisis restricted only to the countries of East Asia. The

French authors cited above found that models based on data from booming years lose much from their predictive power when applied on data from the recession. However, they examined whether the relationship exists inversely or not. They found that, in the years of prosperity, the performance of models based on data from the recession is inferior to those developed on data from the years of economic growing. These empirical research results suggest that different bankruptcy prediction models are needed according to the actual state of the whole economy.

Conclusions drawn by Pompe-Bilderbeek [2005] are somewhat contradictory to those presented above because they think that inferior results of models developed on data from the booming years and applied in the period of the recession can be explained by the fact that the predictability of bankruptcy is much lower during an economic decline rather than by the problems arising from the fitness of the models.

With respect to the effect of the state of macroeconomy on bankruptcy prediction models, new aspect was highlighted by Gersbach-Lipponer [2003] who conducted simulating examinations in order to explore the influence of the correlation between bankruptcies of firms on the risk of the banking activity. According to the authors, a negative macroeconomic shock can increase not only the firm level probability of bankruptcy. With respect to fact that companies in the economy cannot be considered as completely independent units, correlation between their bankruptcies may be a relevant issue. In the case of recession, this correlation may increase which can cause higher risk in the whole banking sector as well.

#### ***2.3.1.3. Geographical aspects of sample selection***

It was mentioned in the chronological review of the literature that the seminal model of Altman [1968] has been used very often as reference. Although in some cases very promising prediction results are reported by the authors conducting such comparative studies, consensus seems to be existed in the fact that a model developed on data from a particular country cannot be used with the same accuracy on another database containing data for firms from different countries. This hypothesis was confirmed by the research of Ooghe et al. [1999]. Databases of companies from 12 countries were examined by Moro et al. [2011] who found different relationships between the event of bankruptcy and the values of financial ratios in each database under consideration. Similar findings can be read in the article of Ioannidis et al. [2011] in which bank failures were investigated. Sueyoshi-Goto [2009] raised the attention to the country specific aspects of the operation of firms when examined Japanese companies. Based on the above mentioned research results, it can be concluded that this issue cannot be left out when building bankruptcy prediction models.

#### ***2.3.1.4. Problems of dichotomous classification***

The prediction itself is the most critical aspect of the science of bankruptcy prediction, as it was mentioned in the introduction. In my work, I also try to predict future insolvency of firms based on financial data for solvent and insolvent companies being in the available sample. More formally, bankruptcy prediction can be considered as a dichotomous classification problem which means that companies are categorized only into two mutually exclusive groups (bankrupt, healthy).

The main problem of this approach is that it doesn't take into account the fact that these groups are not necessarily mutually exclusive with respect to the values of attributes used to distinguish them. The reason is that previously healthy companies may become insolvent in the future. It is a common experience that insolvency of firms is not a sudden event, in most cases it has early signs (Lin et al. [2011]). It can be assumed that the event of bankruptcy is the final outcome of a longer or shorter period of financial problems. In such a period, it is a real assumption that firms in financial trouble have similar or even weaker financial indicators than the bankrupt companies in the sample. This kind of enterprises can be considered as an interim group between the healthy and bankrupt classes.

Similarly, firms become insolvent because of an unexpected event (for example natural disaster) represent another special group when classifying firms in a dichotomous manner. In this case, it should be taken into consideration that a company may become bankrupt even with relatively good financial indicators which don't reflect the problems caused by the unexpected event because financial statements of Hungarian enterprises are available with more than 6 months lag. However, it should be noted that this "group" cannot be observed so frequently in the real life than the example discussed in the previous paragraph.

These two "interim" groups were categorized by Altman [1968] into the so-called grey zone because such companies cannot be unambiguously classified into one of two main classes (bankrupt, healthy). This problem was addressed by Gilbert et al. [1990] who suggested a three-class approach instead of the dichotomous classification in bankruptcy prediction. Besides bankrupt and healthy companies, they examined another group where companies are still operating but have financial problems as well. Their research results showed that discrimination between

companies in good financial situation and bankrupt companies is much easier than discrimination between bankrupt and operating companies having financial problems. Based on these results, Anandarajan et al. [2001] considered only those companies as operating which could continue their operations after negative events like default on payments of debt or business year with negative cash flow.

The *raison d'être* of multiclass approaches is obvious, however, their practical feasibility are restricted because the term of “financial problems” hasn’t been objectively defined. Though, subjective definitions are available but their application cannot be considered as scientific because of the lack of objectivity. In spite of this, several authors have carried out empirical research in this direction: see for example the works of Altman et al. [1994], Slowinski-Zouponidis [1995], Bioch-Popova [2001] and Jones-Hensher [2007]. The latter applied a four-class approach in bankruptcy prediction.

Different possible solutions have been raised by researchers in order to solve the problem arising from the absence of objective definition. Yang [2001] and Du Jardin [2010] tried to find significantly different groups within the two main classes (bankrupt, operating). They argued that if such groups are identifiable, then they can help in identifying companies in the “interim” stage between healthy and bankrupt classes without objective definition. Sueyoshi-Goto [2009] identified the grey zone of dichotomous classification models, then they developed a separate model to classify the firms belonging to this special group. Companies with at least three years with negative profit were excluded from the analysis by Alfaro et al. [2008] with referring to the assumption that this kind of firms possibly have financial troubles, and consequently have financial ratios which are more similar to bankrupt companies

than to healthy firms, so inclusion of them in the database would have negative impact on the predictive performance of the models.

One can find contradictory finding in the literature in this regard as well. Gruszczynski [2004] found better hit rate by using a dichotomous classification model than in the case when companies having financial difficulties constituted a separate class in the database.

#### ***2.3.1.5. The number of instances in the classes***

The number of healthy and bankrupt companies in the sample can also be classified to the questions related to sample selection. From statistical point of view, it can be reasonable to use representative (proportional) samples but it would have several negative characteristic in the field of bankruptcy prediction:

- besides the number of instances from healthy and bankrupt groups of companies in the sample, there are other relevant aspects with respect to the risk of bankruptcy such as the age, size and activity of the companies under consideration. The sample should be proportional to these aspects as well but it is very difficult to collect a sample which is simultaneously representative of these four issues and, at the same time, large enough so that statistical inferences can be drawn based on it;
- the proportion of insolvent companies is usually much lower in the population than that of solvent firms. Because of it, even in the case of a relatively large representative sample, the number of instances from the bankrupt group could be low. Hence, the sample wouldn't contain enough information to properly distinguish the two classes;

- state-of-the-art data mining methods tend to be specialized to the characteristics of the majority group (Chen et al. [2013]), so the literature suggests the use of samples containing the groups in the same ratio when applying this techniques;
- it is a common experience that financial ratios for failed firms exhibit greater variance than those of healthy ones. This fact also contradicts to using proportional samples in order to reduce the negative impacts coming from the error of sampling.

Because of the reasons listed above, oversampling of failed firm is a common practice in the literature. One way of implementing this is to represent failed firms in the sample in a higher ratio than they are present in the population. Others insist on using representative samples. In this case, failed firms are overrepresented (often by weighting) in order to apply the most recent classification methods of artificial intelligence. Empirical research was carried out in this field by Kennedy et al. [2013] among others. Their results showed that overrepresenting failed firms in the sample didn't increase the performance of classic statistical methods (e.g. logistic regression) but they found significant increase in the discriminating ability for data mining algorithms. Similar conclusions were drawn by Horta-Camanho [2013] as well who examined data of Portuguese construction companies. Another solution was suggested by Sánchez-Lasheras et al. [2012] who balanced their sample by underrepresenting the majority class. Operating companies with similar financial ratios were replaced by a representative case in the work of the lastly cited authors.

### **2.3.2. Data preprocessing**

Ratio type financial ratios calculable from the accounting information system of firms have constituted the explanatory variable set for bankruptcy prediction from the 1960s (Chen [2012]). Although, these variables are often capable of discriminating bankrupt and healthy companies, their drawbacks also have to be emphasized. The most important problem is that accounting data are backward-looking. Those reflect the state of the company at the end of a financial year, hence the predictive capability of models based on them are restricted (Lin et al. [2014]).

Financial ratios have often been used in their “raw” form, that is, in the form as they were calculated from data of income statement and balance sheet without any corrections. Predictive power of financial ratios is further constrained by this approach. Nowadays, consensus can be observed in the literature related to the importance of data preprocessing before modeling. Hence, this issue has become an independent research area of bankruptcy prediction.

#### ***2.3.2.1. Calculating problems***

In spite of its importance, there are only very few publications which deal with the calculating problems like those mentioned by Kristóf [2008]. He raised two important and often occurring problems related to ratio type financial variables:

- zero value in denominator;
- the case when the numerator and the denominator are also negative.

Although, only the first makes financial ratios impossible to calculate, the latter could also have serious consequences. Let's take for example the ratio of return on equity (ROE) for a company whose earning after interest and taxes and its equity

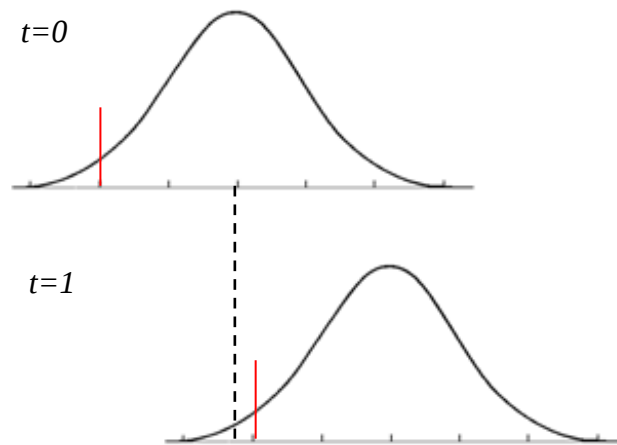
were also negative. In this case, the ratio reflects completely misleading picture about the profitability of the firm. If the problem remains unsolved, reliability of models containing this variable becomes considerably questionable.

#### ***2.3.2.2. Industry effects***

Performance stability of models over time has constituted a serious problem in the literature of bankruptcy prediction for several decades. Researchers often found that models developed on data from some time period show good performance over the time period from which data for model building comes from, but at the same time, they often reported that these models lose very much from their predictive power when applied on data from another (later) time period.

To solve this problem, an effective approach was proposed by Platt-Platt [1990]. They assumed that instability of models over time can be attributed to the fact that the distribution of financial ratios may vary with the lapse of time. Because of it, an arbitrary value of a variable close to the mean value of an industry at the time of model building could be extreme low or high some time later if the distribution of the financial ratio under consideration is changing over this period. Assuming normal distribution, this can be visualized on the figure below.

**Figure 2. The effect of changing the distribution of financial ratios**



Source: Own figure based on Platt-Platt [1990]

As it can be seen in Figure 2, when the mean from a normal distribution increases *ceteris paribus*, then the former mean of  $t=0$  will appear on the edge of the distribution at the time  $t=1$ , that is, this unchanged value can be considered as very low compared to the mean of the industry at  $t=1$ . In the figure, the left edges of the distributions bounded by vertical lines indicate the regions in which financial ratios of failed firms can be typically observed.

Such a change can significantly decrease the predictive power of a model based on data from  $t=0$  when it is applied on data from  $t=1$ . Assuming that in the case of the financial ratio under consideration the mean value is typical among healthy companies, then this model will classify the firm having the financial ratio equals to the mean of  $t=0$  as healthy in spite of the fact that this value at  $t=1$  is already typical of bankrupt ones. According to Platt-Platt [1990], this problem can be solved by employing industry relative ratios which can be calculated by the following formula:

$$\frac{\text{ratio of the firm}}{\text{industry mean}} \times 100$$

This formula compares the value of the financial ratio to the mean of the industry in which the firm under consideration is operating. The advantage of this approach is that changes to the mean over time don't cause problems since the industry relative ratios reflect the deviation of the financial ratios from the current mean of the industry, hence they are invariant to changes of the mean (Platt-Platt [1990]).

In addition to the fact that industry relative ratios facilitate solving the problems arising from time instability of data, it should be mentioned another important characteristic of them as well. Developing bankruptcy prediction models on the basis of samples consisting of companies from different industries has still been a common practice in the most recent literature. Comparability of financial ratios for firms from different industries, however, is questionable because textbooks dealing with financial analysis teach as a principle that the values of financial ratios cannot be considered as absolute criteria; they can objectively be judged only in comparison with some benchmark (Virág et al. [2013]). The average value in an industry can be an excellent mean for this purpose. The application of industry relative ratios can allow compiling databases for bankruptcy prediction models from data of firms operating in different industries.

In spite of their advantageous characteristics, industry relative ratios are rarely applied in the literature. After reading more than three hundreds of scientific articles on this field, I have found only three cases (Dewaelheyns-Van Hulle [2004], Hillegeist et al. [2004], Berg [2007]) where financial ratios were compared to the mean of the industry. Possible reasons for this relative underrepresentativity in the literature:

- industry means of financial ratios are not necessarily available for researchers and practitioners;
- some financial ratio (typically profitability rates) can be measured only on interval scale where the ratio type of industry relative variables may be problematic;
- comparability of deviations from the industry mean across industries may also be open to doubt.

#### ***2.3.2.3. Preprocessing data for modelling***

Statistical properties of financial ratios were already examined by Beaver [1966] as well. After analyzing the values of financial ratios, he concluded that the distribution of financial ratios are usually asymmetric. This issue seriously affects the performance of statistical models commonly used in bankruptcy prediction. With respect to the fact that the occurrence of outliers and deviations from normal distribution are general characteristics rather than rare exceptions of databases used for building bankruptcy prediction models (McLeay-Omar [2000]), hence solving these problems is very important in order to maximize the performance of predictive models.

Mathematically independence of explanatory variables is one of the most frequent assumptions when applying statistical models. With respect to the fact that financial ratios are calculated from the same balance sheet and income statement, independence of financial ratios cannot be a realistic assumption. To solve this problem, principal component analysis is commonly used because this method condenses the information captured in different financial ratios into artificial

variables which are mathematically independent. Among others, this approach was applied by Li-Sun [2011a] and Xiaosi et al. [2011].

One cannot also find consensus in the literature regarding the question of handling the above-mentioned problem of multicollinearity by principal components. According to Wang [2004] and Huang et al. [2012], serious drawback of this approach is that it doesn't make distinction between the two groups (healthy and bankrupt). The cited authors think that the method of principal component analysis is inappropriate in the case when different groups can be identified in the population under consideration. They tried to find methodological solution to this problem. Their proposed algorithms could enhance the performance of the models compared to those using the classic principal components. The effectiveness of principal components in condensing the information of explanatory variables was also questioned by Erdal-Ekinci [2012] who found better predictive performance when using the original financial ratios in their model than in the case of models developed on principal components showing eigen values greater than 1.<sup>4</sup>

Values showing relevant deviation from the other values are called outliers in the literature of statistics. Observations having outlier values may significantly decrease the performance of statistical models. There are several possible solutions to solve this problem:

- the simplest approach is to exclude observations having outlier value from the database for developing model. This practice can be observed in the literature even nowadays (Min et al. [2011], Cao [2012b], Sánchez-Lasheras et al. [2012], Fedorova et al. [2013]);

---

<sup>4</sup> The use of this rule of thumb is common in mathematical statistics. See also for example the work of Chen [2011].

- according to McLeay-Omar [2000] and Balcean-Ooghe [2006]<sup>5</sup>, relevant information for the model can be lost by removing outliers, hence, instead of removing, transformation of variables (taking a logarithm, square root transformation, etc.) is a more preferred solution;
- categorization of variables can also be applied for this purpose (Sun-Shenoy [2007]). This approach makes bins within the range of financial ratios according to some point of view, and then only the ids of the bins are used, as a variable measured on ordinal scale, in modelling. An advantage of the method is the fact that outliers cannot cause problems because they appear in the first or in the last bin. At the same time, one loses the information contained in the exact differences between each pair of observations which could play important role regarding to the performance of models;
- winsorizing is another commonly used approach to avoid the distorting effect of outliers. The essence of this is that extreme values are replaced by the corresponding percentiles of the distribution. Values greater (smaller) than the 99.5<sup>th</sup> (0.5<sup>th</sup>) percentile were set to the 99.5<sup>th</sup> (0.5<sup>th</sup>) percentile by Duan et al. [2012]. Similarly, the first and last percentiles were used by Cubiles-De-La-Vega et al. [2013], the 5<sup>th</sup> and 95<sup>th</sup> percentiles were chosen by Pindado-Rodrigues [2004] for this purpose.
- problems arising from outliers are often solved by transforming the range of the variables into the [0;1] interval<sup>6</sup>. This approach was applied for example by Hu [2009], Li-Sun [2011a], Swiderski et al. [2012], Riberio et al. [2012], Kennedy et al. [2013], Saberi et al. [2013].

---

<sup>5</sup> The hypothesis of the authors was proven later by Yu et al. [2014] as well who found that firms in their training sample can be classified extremely well by examining only the fact whether the value of Equity/Assets is outlier or not.

<sup>6</sup> Transforming the range of variables into the [-1, 1] interval is also a common practice in the literature, see for example the work by Zhong et al. [2014].

However, it is not obvious what value can be considered as an outlier because this term hasn't been objectively defined in the literature. Application of statistical rules of thumb is common in the absence of unambiguous definition. Standardizing financial ratios by their means and standard deviations is often part of data preprocessing tasks (see for example the work of Hájek [2011]). Standardized values outside the range of five or three standard deviations around the mean are often considered as outliers. The work of Van Gestel et al. [2006] can be mentioned for an example for the latter.

#### ***2.3.2.4. The question of static bankruptcy prediction models***

From the early works of bankruptcy prediction until now, input variables for models have been selected from the ratio type financial ratios coming from the accounting information systems of firms. However, performance of models based on them are not stable.<sup>7</sup> This problem was the first which raised the attention to the limitations coming from the static bankruptcy prediction models based on financial ratios. To solve this problem, researchers suggested the regular re-estimation of models (Pantalone-Platt [1987b]), but this approach hasn't been an effective solution.

The importance and actuality of this research question was also indicated by Abdou-Pointon [2011] in their literature review where dynamization of bankruptcy prediction models were considered as an important future research direction. To fill this gap, some researchers searched for methodological solutions. At first, survival models appeared which relaxed the restrictions coming from the static nature of

---

<sup>7</sup> This problem was discussed in Section 2.3.2.2 in more detail.

statistical methods (Hillegeist et al. [2004]). More sophisticated methods were applied by Du Jardin-Séverin [2012] who assumed that, due to changes in the macroeconomy, distribution of financial ratios are not stable over time but the direction and the volume of the changes in financial ratios are more stable. Based on this idea, the cited authors tried to predict future solvency of firms on the basis of the evolution of financial ratios over a six year time horizon and not on the basis of the values of financial ratios observed in the most recent year. Time series of financial ratios were examined by the authors for several hundreds of French firms. Since they found similarities between them, they clustered these time evolutions to reveal the typical changes to financial ratios in the case of bankrupt and operating companies. With respect to the fact that bankruptcy prediction models are based on values of several financial ratios, the cited authors used Kohonen maps to visualize the typical bankruptcy routes (in the words of the cited authors: “trajectories”) which are capable of predicting future insolvency as well. By this dynamic approach, they could achieve significantly higher hit rates than by the static models based on the well-known classification algorithms. A similar examination was conducted later by Chen et al. [2013] as well.

Other researchers also tried to relax the static nature of models by dynamizing financial ratios. In addition to the static values, Berg [2007] applied their changes from one year to another as input variables for his models. He examined the possibility of using “multi-period” bankruptcy prediction models as well. This kind of models were developed on the basis of financial ratios from the year  $t-1$ ,  $t-2$  and  $t-3$ , where  $t$  is used to denote the year of bankruptcy. Such a multi-period model showed higher predictive power than the static models. Duan et al. [2012] also carried out empirical examinations in this direction. They compared the most recent

values of financial ratios to the average of these variables from the past 12 months and found significant differences between bankrupt and operating public companies.

### **2.3.3. Feature selection**

The use of ratio type financial indicators<sup>8</sup> as independent variables in models is a common practice in bankruptcy prediction. The source of them are data of accounting documents to be made mandatorily by the firms, more precisely the balance sheet and the income statements. According to Du Jardin [2010], the number of possible predictors is almost infinite. However, it is not indifferent which variables will be included in the models because, as stated by Beaver [1966] in his pioneering work as well, there is a substantial difference between the predictive capabilities of each financial ratios. According to Nikolic et al. [2013], the aim of feature selection is to filter out redundant and irrelevant variables in order to decrease the complexity and the computing burden of the model, furthermore to increase its classification performance. Due to its importance, this topic has become a standalone research direction in bankruptcy prediction.

#### ***2.3.3.1. Difficulties of feature selection***

The absence of underlying theory which can determine the set of explanatory variables to be used for predicting future insolvency of firms is a serious deficiency of bankruptcy prediction (Nikolic et al. [2013]). The need for theoretical framework

---

<sup>8</sup> Deficiencies of these variables have also been realized in the literature. Value based ratios appeared as an alternative later such as economic value added (see Virág et al. [2013] for more details). These variables were applied in bankruptcy prediction models by Pasaribu [2008] but according to his results these variables cannot be considered as effective predictors in predicting future insolvency of firms.

was also observed by Beaver [1966]. To fill in this gap, he tried to develop a simple theoretical model in which companies were considered as reservoirs of liquid assets. In this concept, bankruptcy is the case when the reservoir drains. The source of water for the reservoir is the income of the company, the expenses are symbolized by the outflows. Although, this model is rather simple and doesn't determine the set of explanatory variables for bankruptcy prediction models, Beaver's concept has been referred a lot in the literature: for example, Laitinen-Laitinen [2000] also used this theory to select input variables of their model.

Using financial ratios applied with success in previous studies is a common approach in the absence of theoretical models in the field of bankruptcy prediction (Du Jardin [2010]). However, different authors have found different significant variables when classifying failed and operating companies in the last 50 years of literature in this field (Lin et al. [2011]).

#### ***2.3.3.2. Other data sources for bankruptcy prediction***

In the case of listed companies, market data are also available as input variables for bankruptcy prediction models. There hasn't been consensus in the literature regarding the question of whether financial or market variables can be considered as better input factors in predicting future insolvency. Accounting ratios, cash flow variables<sup>9</sup>, furthermore market returns and their standard deviations were compared by Mossman et al. [1998] and concluded that accounting variables are effective predictors one year before failure, cash flow variables two and three years before bankruptcy. It was an interesting result that market data weren't proven to be

---

<sup>9</sup> The cited authors considered ratio type financial variables based on data from balance sheet and income statement as accounting variables, furthermore operating and tax paying cash flow calculable from cash flow statement were considered as cash flow variables.

significant at any time horizon examined by the cited authors. Agarwal-Taffler [2008] compared Altman's [1968] model to other models based on market data and found better performance in the case of Altman's [1968] accounting-based model.

In contrary, Atiya [2001] developed more accurate models when market data were also incorporated in addition to the accounting ratios. Complementary relationship was also experienced between the two types of variables by Das et al. [2009] in modelling CDS (credit default swap) prices.

Beaver et al. [2005] addressed the question of whether accounting ratios have lost from their predictive power during the 40 years elapsed after publishing his pioneering work. They found that accounting variables don't capture extra information compared to market data. Hillegeist et al. [2004] experienced better forecasting ability in the case of models developed solely on market variables than those developed only on accounting ratios. Similar conclusion was also drawn by Chawa-Jarrow [2004].

In addition to the financial ratios calculable on the basis of accounting data, the importance of non-financial variables was also highlighted by several authors in the literature. These factors can significantly enhance the predictive power of the models (Barniv et al. [2002], Balcean-Ooghe [2006], Lensberg et al. [2006], Sohn-Kim [2007], Manjón-Antolin – Arauzo-Carod [2008]). Dambolena-Khuory [1980] raised the attention to the fact that standard deviation of some financial ratio may also have significant discriminating power between bankrupt and operating companies.

Macroeconomic variables were proven to be insignificant predictors of failure in the case of banks (Pantalone-Platt [1987b]) as well as in the case of enterprises (Duffie et al. [2007]). In contrary, Carling et al. [2007] found significant variables among

macroeconomic factors when examined a credit portfolio at a Swedish commercial bank. The age of the companies were suggested as potential explanatory variable by Sjøvoll [1999]. Both of the variables mentioned in this paragraph were found to be significant by Bonfim [2009]. His research revealed that in the case of younger firms different factors should be taken into consideration than in the case of older ones.

The incorporation of macroeconomic variables is straightforward in the framework of survival models which constitute a class among the panel econometric models. By applying this method on data from the period of the crisis in East Asia, Nam et al. [2008] found several significant macroeconomic factors. The relationship between macroeconomic variables and the risk of bankruptcy has been proven by empirical investigations conducted on national level data (Rösch [2003], Koopman-Lucas [2005], Harada-Kageyama [2011]), hence their application is useful in the context of survival models.

The use of non-financial variables was also proven to be beneficial beyond the field of bankruptcy prediction. Gaganis [2009] tried to model fraudulent activities regarding financial statements with accounting ratios as explanatory variables. The cited author found significant increase in the performance of his model when non-financial variables were also included. Default events of car credits were examined by Sinha-Zhao [2008]. Besides the standard financial ratios, they used subjective opinions of a financial expert. The inclusion of the latter enhanced the performance of the models based on only objective information.

According to Balcean-Ooghe [2006], financial indicators don't reflect all of the information regarding future insolvency of firms, hence they think that it is important to use qualitative independent variables in bankruptcy prediction models. This idea

attracted a vast amount of research in the literature. Some example: Slowinski-Zouponidis [1995], Anandarajan et al. [2001], Becchetti-Sierra [2003], Grunert et al. [2005], Alfaro et al. [2008], Huynh et al. [2010], Blanco et al. [2013], Horta-Camanho [2013]. The cited authors emphasized the incorporation of the following qualitative factors into bankruptcy prediction models: variables measuring the efficiency of the management, the position of the firm on the market, form of the ownership, geographical location of the activity and the bank financing the company. Beynon-Peel [2001] and Demers-Joos [2007] used the auditing firms of the observations as qualitative predictors in their models.

#### ***2.3.3.3. Feature selection methods***

According to Du Jardin [2010], restricting the possible set of input variables to those applied successfully in prior research is a frequent practice in bankruptcy prediction. This approach can be considered as a re-estimation of the older models on more recent data. Though, the results are often very promising, some researchers, for example Lin [2009] among others, argue that it is worth developing new models using all of the available data. Predictors of previous models are only possible alternatives in this approach.

Feature selection methods are categorized into two main groups by Feki et al. [2012]. Procedures in the first group are called as filters. This group consists of standalone feature selection methods which seek explanatory variables capable of discriminating the two groups under consideration. Simplicity and low computing cost are the advantages of this kind of algorithms which, at the same time, don't guarantee to find

the optimal subset of variables for an arbitrarily chosen classification algorithm and therefore don't warrant the highest predictive power.

The second group is called as wrapper algorithms which seek the optimal subset of variables for the chosen classification method, hence its predictive performance can be maximized, but at the cost of much higher computational time and effort. Comparative analysis were conducted by Li-Sun [2011b], Oreski et al. [2012] and Lin et al. [2014] recently on this new research direction. In spite of the promising results of the lastly cited authors, due to their lower computational cost, filter algorithms are playing dominant role in bankruptcy prediction nowadays.

In order to select significant features, stepwise methods (forward and backward) are commonly applied in the case of traditional statistical models<sup>10</sup> (discriminant analysis, logistic regression). These feature selection algorithms are available as built-in functions in the statistical analysis software where the user can specify the level of significance for entering and removing variables into the model. Decision trees are also popular for similar reasons. The process of developing the tree is also automatic in this case because the branches of the tree are generated according to variables which have the greatest discriminating power between the two groups.

Traditional statistical techniques don't have parameters to be optimized, so it's possible to use the above-mentioned "built-in" feature selection algorithms. However, parameter optimization is necessary for data mining procedures based on machine learning concepts. The performance of these methods depends on the values of the parameters (Ping-Yongheng [2011]) but optimal values of them are impossible to objectively determine. They should be set according to the problem and database

---

<sup>10</sup> Theoretical background of classification methods used in bankruptcy prediction is overviewed in Section 2.3.4.

under consideration. This issue makes difficult to use forward and backward feature selection algorithms in the case of machine learning methods.

At first, the problem was solved by using filter methods. The essence of this is to use significant variables selected by a filter procedure as input variables to data mining algorithms. One of the simplest approach was chosen by Li-Sun [2011a] who used financial ratios which differed significantly according to the results of the t-test for comparing means from two independent samples. The nonparametric counterpart of this test was used by Pindado-Rodrigues [2004] for the same purpose. Significant features of logistic regression were used by Kim-Sohn [2010] and Telmoudi et al. [2011] for the SVM and rough set theory. Yazici [2011] used significant features of discriminant analysis for his model developed with neural networks. Min et al. [2011] used generalized additive model to select input variables for SVM. Classification and regression trees (CART) was used by Brezigar-Masten – Masten [2012] to find the optimal subset of explanatory variables.

Though, the application of the approach discussed in the previous paragraph has been common to date, later the attention of the research community turned to genetic algorithms which try to find the global optimum of any optimization problem by simulating the natural selection process. In the context of bankruptcy prediction, the task is to maximize the predictive performance of the model by using different subsets of possible explanatory variables. The possible subsets of independent factors represent the population of genes and they compete for surviving which depends on their capability of effectively solving the problems determined by the user. Similarly to the evolution in the nature, genes which are unable to effectively solve the problem perish and will be replaced by the descendant of the genes which were able to solve the task with at least satisfactory level. The new generation is generated by

crossover and mutation of their parents' genes. This approach has shown very good results in finding the set of explanatory variables for data mining algorithms (Back et al. [1996], Lensberg et al. [2006], Hájek [2011], Oreski-Oreski [2014]). It should be mentioned the work of Chi-Hsu [2012] who also applied this approach to determine the independent factors for their logit model.

Du Jardin [2010] investigated the most commonly applied classification methods in bankruptcy prediction (neural networks, discriminant analysis and logistic regression) in order to determine the most appropriate feature selection method for them. In the case of parametric methods, stepwise algorithm was revealed to be the most effective. In the case of neural networks, the best discriminating power was found when they were developed on the feature set selected by genetic algorithm.

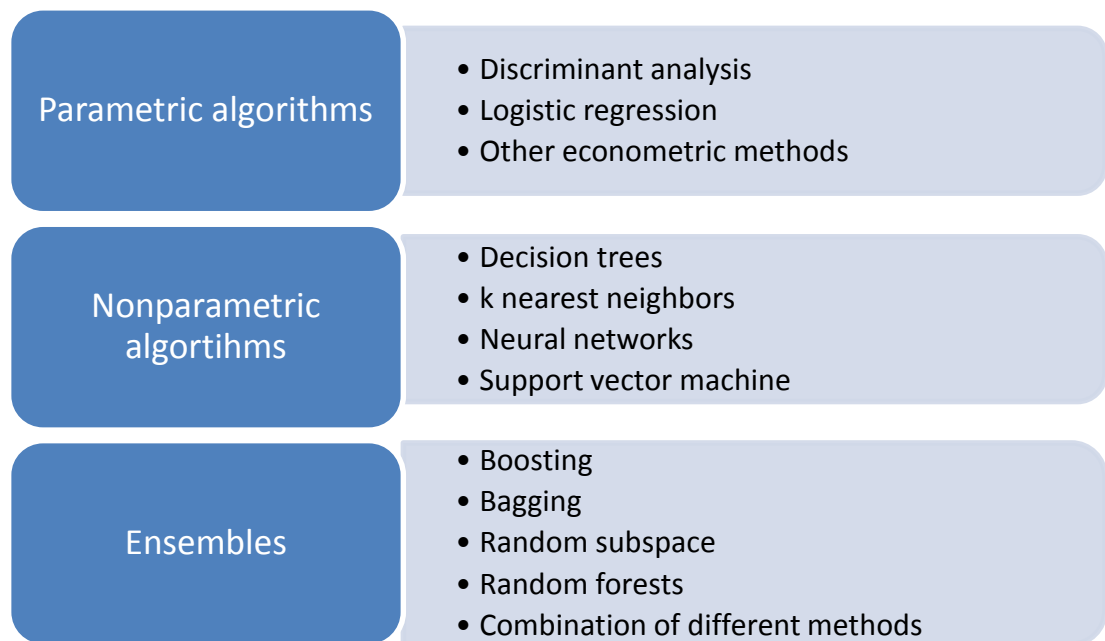
Later, feature selection methods optimized for each data mining algorithms were tried to develop, but these were too complex and didn't become common in the literature. To mention some examples, a feature selection method was developed for SVM by Hardle et al. [2009]; variable selection based on information gain was tried to be integrated into ensemble methods by Wang et al. [2014]. Ensemble methods will be overviewed in more detail in the next subsection.

#### **2.3.4. Constructing bankruptcy model**

After sampling, selecting the possible explanatory variables and preprocessing data for model building, the next phase in bankruptcy modelling is to choose a classification algorithm. The choice is not simple because, according to Du Jardin [2010], there has been developed more than 50 methods in the last 50 years' of

literature in this field for predicting future insolvency of firms. These methods can be categorized into three main groups (Kim-Kang [2012]).

**Figure 3. Categorization of classification methods used in bankruptcy prediction**



The number of available methods raised the question of what technique is worth using in order to maximize the predictive power of models. The importance of this question is indicated by the fact that most of the studies in bankruptcy prediction deals with comparing the performance of different classification algorithms (Sánchez-Lasheras et al. [2012]). In spite of this, there hasn't been consensus in the literature to date regarding a question of what method can be considered as the best in this field. Researchers examining different datasets often find that the performance of each methods varies considerably according to the data under consideration.

Due to space limitations, in this section, I present the theoretical background only for those classification methods which are most frequently used in the literature. I place great emphasis on outlining the most recent data mining methods which are less common in the Hungarian literature of bankruptcy prediction. In addition, other classification methods having smaller weight in the literature will also be overviewed briefly. The Reader can find more information about these techniques in the cited references.

#### **2.3.4.1. Multivariate statistical methods**

Linear discriminant analysis was the first multivariate method that appeared in bankruptcy prediction. The general form of the model is the following:

$$Z = \beta_0 + \sum_{i=1}^k \beta_i x_i$$

where  $x_i$  denotes the independent variables,  $\beta_i$  stands for the estimated coefficients of each predictors<sup>11</sup>. The output of the model is called as discriminant score (Z) which condenses into a single value the information inherited in independent variables. The classification is based on this output value and a cut value controllable by the modeler which divides the possible range of Z into two distinct areas. If the Z score of an observation falls below (above) the cut value, the observation is classified into to the first (second) group. The independent variables of the model have the highest discriminating power (measured by Wilks- $\lambda$ ) among the possible input variables between the two groups under consideration.

---

<sup>11</sup> for more details see Virág et al. [2013]

The advantages of discriminant analysis:

- accessibility in statistical software;
- fast and easy-to-use;
- allow analyzing the relationship between the event of bankruptcy and the predictor variables;
- independent variables can automatically be selected via stepwise algorithms.

Drawbacks of the method:

- its assumptions (multidimensional normality, equal variance-covariance matrices, independence of predictor variables) are generally not met (Du Jardin [2010]);
- its classification performance is usually lower than those of other methods because of the violation of the above mentioned assumptions (Wang-Ma [2011]);
- it is not capable of directly<sup>12</sup> measuring the risk of bankruptcy in a cardinal<sup>13</sup> way.

Due to the drawbacks listed above, attention of researchers turned to logistic regression which was firstly applied in bankruptcy prediction by Martin [1977]. The general form of the model is shown below:

---

<sup>12</sup> The output of discriminant analysis can mathematically be converted into probability of bankruptcy as shown by Deakin [1972].

<sup>13</sup> The risk of bankruptcy can be measured only on ordinal scale by using discriminant analysis because the output of this technique can take any positive or negative value without limits. That is, one cannot see how riskier a company is compared to another one. Only a ranking can be established based on the output of discriminant analysis. The difference between two outputs is difficult to interpret. However, in the case of logistic regression, the probability of bankruptcy is easily and directly obtainable. This measure can tell us how riskier a firm is compared to another one because the range of the probability is between 0 and 1, so these probabilities can directly be compared.

$$\ln\left(\frac{p}{1-p}\right) = \beta_0 + \sum_{i=1}^k \beta_i x_i$$

where  $x_i$  denotes the independent variables,  $\beta_i$  stands for the estimated coefficients of each predictors. The dependent variable on the left side of the equation is often called as “logit”.  $p$  stands for the probability of bankruptcy.

Higher classification performance can be achieved by this method compared to discriminant analysis because logistic regression doesn't require the fulfilment of such rigorous assumptions than the discriminant analysis does. The advantages of logistic regression are the same as those of discriminant analysis, but the probability of bankruptcy can directly be determined by using a logistic model (Li-Sun [2011a]). Due to this and its relatively high classification performance, it has been one of the most popular methods in bankruptcy prediction to date (Nikolic et al. [2013]).

Disadvantage of both statistical models is the assumption of linear relationship between the dependent and independent variables which is usually violated in the real life. The nature of the relationship is unknown, but it is certain that there is nonlinear relationship between the occurrence of bankruptcy and the values of financial ratios (Du Jardin [2010]). This relationship can be expressed by incorporating the transformations and interactions of the variables, but the number of possible combinations is infinite, hence its implementation is restricted in the real life (Blanco et al. [2013]).

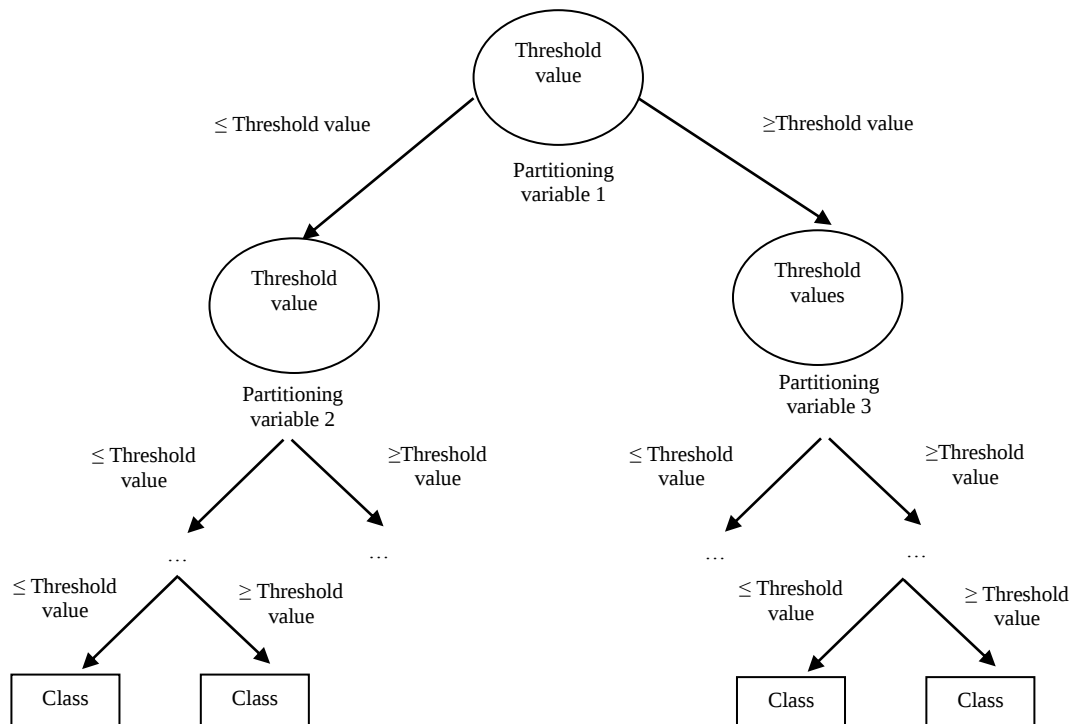
Finally, the application of econometric methods in bankruptcy prediction should also be mentioned. Survival models belonging to panel econometric procedures have been commonly used, see for examples the works of Orbe et al. [2001], Cantner et al. [2006], Bharath-Shumway [2008], Löffler-Maurer [2011], Lyandres-Zhdanov

[2013]. The extended version of logistic regression was used by Heiss-Köke [2004]. Application of Markov chains was tried by Elliott et al. [2014] in this field.

#### 2.3.4.2. Nonparametric methods based on machine learning

Based on the drawbacks of parametric methods discussed in the previous section, the literature of bankruptcy prediction turned to nonparametric algorithms from the mid-1980s. Decision trees appeared at first from this group. Several types of decision trees have been developed and become popular. The difference between the types of decision trees is in the method used to generate the tree. General form of decision trees is shown in the figure below.

**Figure 4. A simple example for decision trees**



Source: Kristóf [2008]

Advantages of univariate and multivariate methods are unified in the algorithms of decision trees. The first partition is made according to the variable along with the most homogenous groups can be established. Then, in the following steps, the procedure seeks those variables and their appropriate values inside the partitions generated previously on which the observations can be further partitioned in order to generate more homogenous groups with respect to dependent variable (bankrupt, operating).

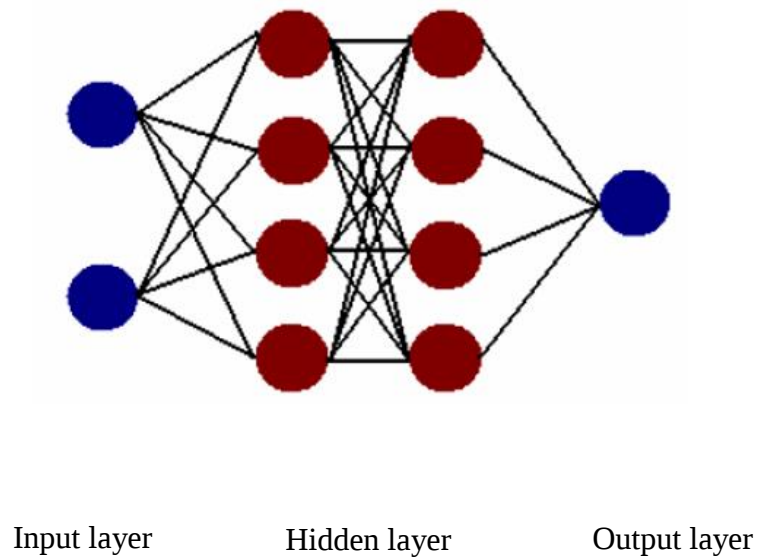
Among decision trees, classification and regression trees (CART) and CHAID method (Koyuncugil-Ozgulbas [2012]) are frequently used. The latter performs the partitioning on the basis of the results of  $\chi^2$ -based independence test between the value of the decision variable and the frequencies in the groups generated in the range of independent variables. Application of C4.5 procedure is also a popular alternative in the literature. This method generates splits where the information gain ratio is maximal (Olmeda-Fernandez [1997], Baesens et al. [2003], Kiang [2003], Brown-Mues [2012], Abellán-Mantas [2014]). Rough set theory (RST) is another method of increasing interest in the literature. The methodological background of this algorithm is presented in more detail in the following works: Ahn et al. [2000], McKee [2003] and McKee-Lensberg [2002].

The main problem of methods based on any statistical theory is the fact that they assume a predetermined functional form that describe the relationship between the probability of bankruptcy and the explanatory variables. However, the relation cannot be defined in advance by using an explicit mathematical formula. In order to overcome this difficulty, the research of bankruptcy prediction turned to methodologies capable of extracting this complex relationship based on available data.

Simple methods such as k nearest neighbors can be applied for this purpose. This method classifies any observation into the class to which k of its nearest neighbors also belong. The distance between the observations is often measured by Euclidian distance. In spite of its relative simplicity, this method exhibits considerably good classification performance as it is shown in the work of Paleologo et al. [2010] and García et al. [2012] among others. This method constitutes the heart of the Case Based Reasoning procedure which classifies the observations into the class to which the most similar cases to the one under consideration also belong (Wang-Ma [2011]). There has been several research reporting excellent predictive performance of case based reasoning in the literature, see for example Vukovic et al. [2012] and Cao [2012a].

Classification methods presented up to this point play only secondary role besides neural networks and support vector machine (SVM) in the recent literature of bankruptcy prediction. Neural networks simulate the learning process of the human brain. Their general structure is visualized in the following figure.

**Figure 5. General structure of neural networks**



Source: Kristóf [2008]

Neural networks consist of three layers:

- the input layer contains the independent variables;
- neurons of the hidden layer process the information passed by the input neurons. This information processing means a mathematical transformation executed on the weighted sum of the input variables;
- the output neuron gives (after similar mathematical transformations) the predicted probability of bankruptcy of the observations under consideration.

In Figure 5, weights are symbolized by lines between the layers. These weights are continuously changing during the learning of the network such that the difference between the output of the network and the observed value of the decision variable will be minimal. The end of the learning is determined by a stopping criterion defined by the user. The learning process may terminate when

- the number of learning cycles reaches its limit defined by the user;
- changes in the values of weights are sufficiently low;
- error on the testing sample starts increasing.

The greatest advantage of neural networks is their capability of approximating any continuous function with any desired precision (Cao [2012a]). The only requirement is that the number of the neurons in the hidden layer should be sufficiently large. For this reason, the networks are enough to have only one hidden layer. The instances of the training sample can perfectly be classified if the user appropriately determines the topology of the network. This property is simultaneously the advantage and the drawback of the algorithm since perfect classification of the training sample is the typical signal of overfitting which means that the idiosyncrasies of the training sample were learnt by the network, so its capability to correctly classify other observations is quite limited. Overfitting is avoidable to some extent. It is a common approach that researchers use a testing sample which is independent of the training sample and stop the learning of the network if its error on the testing sample starts increasing (Blanco et al. [2013]). Accepting the merits of this approach, a word should be said about the critique of McKee-Greenstein [2000] who argued that the neural network trained by this approach may be overfitted for recognizing observations of the selected test sample, so their capability of classifying instances other than those in the training and testing samples available for model building remain questionable.

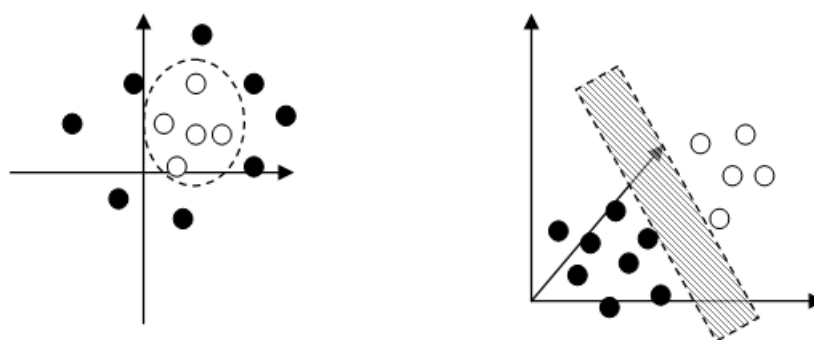
Another frequently occurring problem that the training process of neural networks tends to stop at a local minimum of the error function of the model. To best of my knowledge, there hasn't been any exact solution for avoiding this problem (Blanco et al. [2013]).

In spite of the disadvantages presented above, neural networks have been very popular in the literature up to date. The method lived its “golden age” during the 90s and in the first decade of the new millennium. Its prestige has been decreased in the last 5-10 years because of its problems presented earlier and the emergence of new methods trying to overcome those (Jeong et al. [2012]). According to the lastly cited authors, neural networks can be considered to the methodologies with the highest predictive power henceforward but their performance is underestimated in the current literature because researchers are often not careful enough in determining the optimal parameters.

To overcome the listed difficulties of the method and to enhance its predictive power, researchers of the machine learning community tried to improve the algorithm. These improvements occasionally appeared (Neves-Vieira [2006], Hájek [2011], Pendharkar [2011]) but they haven’t become widespread.

Leading position of neural networks in the field of bankruptcy prediction has been taken over by support vector machine (SVM) in the last decade. Its essence is visualized in the following figure.

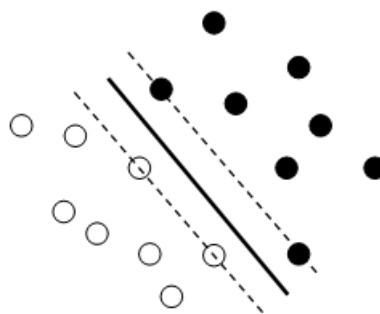
**Figure 6. The basic idea of SVM**



The left side of Figure 6 shows a situation which is common in bankruptcy prediction: the aim is to explore such a nonlinear relationship which is capable of discriminating observations denoted by two kinds of circles in the figure. This is done in SVM by projecting the observations into a higher dimensional feature space via kernel functions where the two groups under consideration are linearly separable. The hyperplane capable of perfectly classifying the observations generated in the higher dimensional feature space is equivalent to the nonlinear function in the original space.

The main difference between neural networks and SVM is that the former is based on the principle of empirical risk minimization, but the latter was developed by following the principle of structural risk minimization (Lin et al. [2011]). The empirical risk minimization means that the models try to minimize the error measured on training sample. This deficiency is trying to be eliminated by the approach which terminates the training process of the classifier when the error on the test sample starts increasing. In contrary, the principle of structural risk minimization aims to minimize to whole error in order to maximize the predictive performance and avoid overfitting. This is visualized in the case of two dimensions in Figure 7.

**Figure 7. Classification of SVM – two dimensional case**



The area free of observations is bounded by dashed lines in the figure. Theoretically, there is infinite number of lines in this area which capable of perfectly discriminating the two groups. Among them, the line farthest from both groups is considered as optimal by SVM since, based on the data of the training sample, observations not covered by the training sample can correctly be classified most likely by this line. This is the essence of the principle of structural risk minimization. Due this property, SVM has such an excellent generalization capability (Lin et al. [2011]) which outperforms the predictive power of neural network in several cases. Another important characteristic that, thanks to the principle of structural risk minimization, SVM effectively avoids overfitting, so it is capable of developing models with good performance even on the basis of small samples (Cao [2012b]), furthermore it is robust to outliers (Lin et al. [2011]). The former advantage was investigated by Shin et al. [2005] and Doumpos et al. [2005] who found that predictive power of SVM remains relatively stable while performance of NN is decreasing as the size of the sample is getting smaller.

In the case seen in Figure 7, the two groups are perfectly separable by a linear function. It is very often in the real life that this cannot be done without errors even in higher dimensional feature space. However, the method can be extended to such cases by introducing a  $C$  parameter into the mathematical form of the model which is not detailed here due to space limitations to penalize misclassification. The value of this parameter cannot theoretically be optimized. It should be set by the user depending on the extent to which error in training sample is tolerated. By increasing the value of  $C$ , the error on training sample is decreasing but at the cost of possible emergence of overfitting which may degrade the predictive performance of the model.

Like any other mathematical model, SVM also has its own limitations. As it was previously mentioned, in order to find the optimal separating hyperplane, observations are projected into a higher dimensional feature space by using kernel functions. Among them, the followings are used most frequently:

- polynomial:

$$k(x_i, x_j) = (x_i \cdot x_j)^d$$

- radial basis (Gauss) function:

$$k(x_i, x_j) = e^{\frac{-\|x_i - x_j\|^2}{2\sigma^2}}$$

- hyperbolic tangent:

$$k(x_i, x_j) = \tanh(\kappa x_i \cdot x_j + c)$$

- ANOVA:

$$k(x_i, x_j) = \sum_{k=1}^n e^{[-\sigma(x_i^k - x_j^k)^2]^d}$$

In the kernels, observations are denoted by  $x_i$  and  $x_j$ , the other parameters, which control the dimensionality and complexity of the projection, should be determined by the user in line with the characteristics of the dataset under consideration. This empirical optimization task is often done by the so-called grid search approach during which the user determine  $n$  different values for parameter  $C$  and  $m$  different values for the parameter of the applied kernel function and then develops a model for each of the elements of this  $m \times n$  matrix of which the parameter setting resulting the highest performance is selected for model building (see for an example the work by Yeh et al. [2010]). Another possibility is to use genetic algorithms to find the optimal parameters for SVM (Hsieh et al. [2012]).

Another problem that running time of the algorithm is increasing in sample size, so significant effort has been done in order to simplify the mathematical background of the method and to enhance its classification ability (Yang [2007], Peng et al. [2008], Yu et al. [2011], Kim et al. [2012]).

Neural networks and SVM are categorized into the methods of artificial intelligence in the literature. Their advantage is their capability of extracting the complex nonlinear relationship between the dependent variable and the explanatory factors without any assumption on the functional form of the relationship. In other word, these techniques try to discover the relation on the basis of the available data (Marqués et al. [2012b]). Thanks to this property, classification performance of both methods is substantially higher than it is in the case of traditional statistical models. However, their common drawback that both method can be considered as a black box which means that the modeler know only the input and the resulting output but he/she doesn't receive answer to the question of how much is the weight of each variables in making the predictions. According to Martens et al. [2010], the problem is so severe that, in spite of their excellent classification accuracy, it is not likely that these methods become popular as decision support systems for lending decisions. Similar conclusion was drawn by Lee-Choi [2013] who think that the absence of interpretability seriously restricts their applicability in making managerial decisions. Probably, this is the main reason for the fact that banks have still been using the simpler traditional statistical methods in their internal rating systems (Riberio et al. [2012]).

The problem arising from the absence of interpretability has been unsolved for SVM, however in the case of NN there have been attempts at filling this research gap. Combination with fuzzy approach was proposed by Akkoc [2012], Trinkle-Baldwin

[2007] tried to calculate variables based on the weights of a developed network in order to interpret it. In spite of the promising results, this research direction doesn't play a dominant role in the literature, so neural networks are still categorized to the black box type methods.

After discussing the most important methods used in bankruptcy prediction, it should be mentioned the work of Serrano-Cinca – Gutiérrez-Nieto [2013] who developed models with the classification algorithms frequently used in this field then classified the methods themselves based on their predictions in order to examine their possible similarities. Building on their results, it can be concluded that methods based on similar theoretical foundations tend to give similar predictions. The cited authors identified the following clusters among the classifiers they examined:

- methods generating decision trees;
- classification algorithms based on linear separation (linear discriminant analysis, SVM);
- logistic regression and neural networks.

#### ***2.3.4.3. Hybrid methods and ensembles***

In the previous sections, theoretical basis, advantages and drawbacks of classification methods commonly used in bankruptcy prediction were presented. With respect to their mathematical nature, all methods suffer from some defects and they have also advantageous properties. This fact raises the possibility of combining different methods in order to enhance the performance of individual classifiers. This approach is an increasing trend in bankruptcy prediction (Cao [2012a]). Research in this field can be categorized into two main groups:

- ensemble methods: a *particular* classification algorithm is applied several times, then the final prediction is made on the basis of the outputs given by the single models developed previously;
- hybrid models: combine the outputs of *different* classification methods in order to make more accurate predictions.

According to Kim-Kang [2012], ensemble methods are machine learning techniques which try to enhance the learning capability of weak classifiers by their repeated application. The two most popular ensemble methods are the Boosting and the Bagging.

The essence of Bagging is to choose  $n$  random subsample with replacement from the available dataset, then a particular classification algorithm will be applied on every subsample, that is, as a result,  $n$  predictions will be available for every observation in the dataset. The final prediction of the Bagging model is generated as the average of the predictions given by the  $n$  models. In the Bagging procedure, every instance has equal probability to be incorporated in each subsample.

In contrary, it is true only for the first subsample in the case of Boosting where observations misclassified (correctly classified) by the model developed on the first subsample have higher (lower) probability to be chosen in the next subsample. The number of subsamples can be determined by the user, so the probability of inclusion in the subsamples is varying according to the performance of the models built on data of the previous steps. The aim of the Boosting procedure is to correct the mistakes of the previous models which is an advantageous property but classification performance of Boosting is often lower than it is in the case of Bagging. The main reason for this is the fact that Boosting is prone to overweight outlier observations

which are very often unidentifiable and this can lead to the distortion of the model and to the degradation of its discriminating power (Kim-Kang [2012]).

The effectiveness of ensemble methods roots in the diversity of models developed on the subsamples. In the case of the two previously discussed ensembles, diversity comes from the random samples drawn by the entire set of observations. It is also possible to construct an ensemble when random subsamples are generated from the available set of independent variables. The final prediction is based on the average of the outputs given by the models containing different subsets of independent variables but based on the full set of observations. There are two popular method in this group: Random Subspace and Random Forests. The methodological background of these procedures is not described here due to space limitations. For more details, please see the cited works in this subsection.

Different kinds of ensemble methods are combined in order to further enhance the performance of the models in the current literature. Marqués et al. [2012a] found that the best combination is the simultaneous application of Bagging and Random Subspace. This combination was also applied by Wang-Ma [2012] in order to enhance the predictive power of SVM. The combination of Boosting and Random Subspace had also been examined earlier by the same authors (Wang-Ma [2011]).

It can be questionable which classification method is worth using in ensembles. The answer was tried to be given by Marqués et al. [2012b]. They applied several basic classification algorithm on several databases. They obtained the best result with the C4.5 procedure, so its application is reasonable in ensembles. This result is in line with the research finding of Kim-Kang [2012] who stated that those classification methods are effective in ensembles whose performance are unstable. The cited

authors analyzed the correlation of predictions obtained on different subsamples and found that the outputs of neural networks and SVM are relatively stable, in contrast to those of decision trees where they found higher standard deviation. They also assumed that the application of decision trees in ensembles is more efficient than the usage of other methods. This assumption was verified by their empirical investigations as well.

The second approach (hybrid methods) combines the outputs of different classification algorithms in order to enhance the performance obtainable via a single classifier. The idea was suggested earlier by Olmeda-Fernandez [1997], but it has become an active research direction only nowadays. Later, it was also suggested by Kiang [2003] that in the case when different methods misclassify different observations, it can be worth combining their results. These suggestions were confirmed by the empirical results of Yim-Mithcell [2005] who used the predictions of discriminant analysis, logistic regression and probit regression as inputs to the neural network which showed higher accuracy than the single classifiers. Recently, Cao [2012b] has conducted similar research. He also found better predictive capability in the case of models which combined the outputs of several single classifiers.

#### ***2.3.4.4. Results of methodological comparative studies***

Nonparametric procedures discussed in the previous sections started replacing statistical methods in the research of bankruptcy prediction from the 90s (Feki et al. [2012]). The mainstream research direction dealing with the question of which classifier achieves the best predictive performance was also emerged during the last

decade of the 20<sup>th</sup> century. In spite of the fact that more and more can be read that it is very likely that this “best classifier” doesn’t exist (Marqués et al. [2012a]), most of the studies have still been dealing with comparing the hit rates of different classifiers (Sánchez-Lasheras et al. [2012]). Some of them are compiled here in a non-exhaustive list below.

Table 1. Publications comparing classification methods in bankruptcy prediction

| Author(s)                   | Applied methods <sup>14</sup> | Best performing method <sup>15</sup> |
|-----------------------------|-------------------------------|--------------------------------------|
| Salchenberger et al. [1992] | NN, LA                        | NN                                   |
| Tam-Kiang [1992]            | NN, DA, LA, KNN               | NN                                   |
| Coats-Fant [1993]           | NN, DA                        | NN                                   |
| Altman et al. [1994]        | NN, DA                        | NN                                   |
| Wilson-Sharda [1994]        | NN, DA                        | NN                                   |
| Alici [1995]                | NN, DA, LA                    | NN                                   |
| Back et al. [1996]          | NN, DA, LA                    | NN                                   |
| Leshno-Spector [1996]       | NN, DA                        | NN                                   |
| Tan [1999]                  | NN, PA                        | NN                                   |
| Zhang et al. [1999]         | NN, LA                        | NN                                   |
| Fan-Palaniswami [2000]      | NN, DA, SVM, LVQ              | SVM                                  |
| McKee-Greenstein [2000]     | NN, LA, ID3                   | ID3                                  |
| Neophytou et al. [2000]     | NN, DA, LA                    | NN                                   |
| Huang et al. [2004]         | NN, LA, SVM                   | SVM                                  |
| Min-Lee [2005]              | NN, DA, LA, SVM               | NN, SVM                              |
| Ding et al. [2008]          | NN, DA, LA, SVM               | SVM                                  |
| Yoon et al. [2008]          | NN, DA, LA, CART, C5          | NN, SVM                              |
| Kim-Sohn [2010]             | NN, LA, SVM                   | SVM                                  |
| Lee-To [2010]               | NN, SVM                       | SVM                                  |
| Moro et al. [2011]          | LA, SVM                       | SVM                                  |
| Xiaosi et al. [2011]        | NN, LA, SVM                   | SVM                                  |
| Bae [2012]                  | NN, DA, LA, SVM, C5, BC       | SVM                                  |
| Erdal-Ekinci [2012]         | NN, SVM                       | SVM                                  |
| Lee-Choi [2013]             | DA, NN                        | NN                                   |

---

<sup>14</sup> Meaning of abbreviations:  
DA: discriminant analysis  
NN: neural networks  
LA: logistic regression  
C5: extended version of C4.5  
BC: Bayes classification  
CART: classification and regression trees  
ID3: former version of C4.5  
LVQ: learning vector quantization  
KNN: k nearest neighbor

<sup>15</sup> There was no significant difference between the methods where two methods are indicated

The list in Table 1 cannot be considered as full because the number of publications related to the mainstream research direction is enormous. However, this short list is capable of reflecting the trends observable in the literature. Until the new millennium, neural networks was compared to traditional statistical methods. The findings consistently indicated the superiority of the former over the latter. SVM appeared in bankruptcy prediction in the early years of 2000s. From this time, the comparison of these two machine learning techniques was the main research direction. SVM exhibited higher accuracy than neural networks in the majority of the cited publications. According to Yu et al. [2014], studies based on the application of machine learning can be considered as the mainstream research direction in bankruptcy prediction. However, it should be emphasized an important issue related to the trends observable in the literature. As Serrano-Cinca – Gutiérrez-Nieto [2013] highlighted, studies showing the fact that state-of-the-art classification algorithms have better discriminating ability than those of older ones are overrepresented in the literature. Furthermore, it also should be noted that classification performance of machine learning algorithms depend to a great extent on their parameters to be set by the user. In light of this, however, it is questionable how much effort was made by the researcher to find the appropriate parameters maximizing the predictive performance of the methods.

Table 1 doesn't contain research results from the last 4-5 years. The most recent trend is the application of ensemble algorithms and hybrid methods. The accuracy of them are usually higher than it is obtainable by using single classifiers. To sum up, these techniques can be considered as the state-of-the-art methods in bankruptcy prediction nowadays.

### 2.3.5. Performance evaluation of the model

Evaluating model performance is a basic requirement in science as well as in practice in order to see whether the performance of the model is higher than it is obtainable by random guessing, furthermore in order to judge the performance of a newer model compared to the extant models. A wide range of performance indicators is available in the literature. In the thesis, only the most important and most frequently used measures are presented in detail due to space limitations.

#### 2.3.5.1. Measuring the classification performance

Since bankruptcy models are generally developed for predictive purposes, the most important measure of their performance is the hit rate which is the ratio of the correctly classified observations to the total number of instances. This measure can be calculated on the basis of values in the confusion matrix. The general form of this matrix is shown below.

**Figure 8. Confusion matrix**

|                |          | Predicted class |         |
|----------------|----------|-----------------|---------|
|                |          | Bankrupt        | Healthy |
| Observed class | Bankrupt | a               | b       |
|                | Healthy  | c               | d       |

The hit ratio of a model can be calculated by using the formula below

$$\frac{a + d}{a + b + c + d}$$

Based on the values in the matrix, further performance metrics are also calculable. Special attention is devoted in the literature to the type I error which can be determined by the following ratio:

$$\frac{b}{a + b}$$

The proportion of bankrupt observations incorrectly classified by the model is measured by this indicator. According to Du Jardin [2010], bankruptcy prediction should focus on minimizing this type of error because classifying a bankrupt company as normal can result in bad lending decision causing serious losses if the company under consideration will go bankrupt in the future. Much lower cost is associated to the Type II error which can be calculated with the formula below.

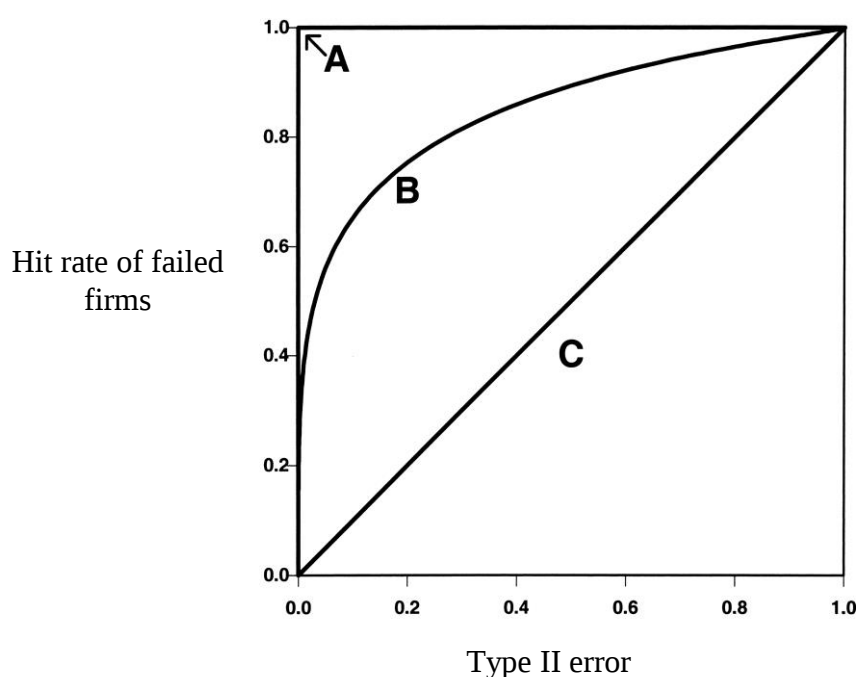
$$\frac{c}{c + d}$$

This measure expresses the proportion of healthy firms classified as bankrupt by the model relative to the total number of healthy observations. This type of error is less costly from investment point of view. In the context of lending decisions, it can lead to not financing a company. In such a case, the lender loses the amount of possible profit which is much lower than the cost of Type I error where the lender could lose even the full amount of the provided credit. In the literature, only estimations can be found for the ratio of the Type I and Type II errors. In the absence of exact information, most of the researchers doesn't pay attention to the costs of misclassification but they are aware of the inequality of the costs associated to these two types of errors and try to develop models minimizing the more costly Type I error.

Some authors should also be mentioned because they emphasized the fact that the two types of errors are extremely asymmetric in bankruptcy prediction, hence they aimed to minimize the expected cost of misclassification rather than maximize the hit rate of the model. In this respect, among others the works by Nanda-Pendharkar [2001] and Chen et al. [2011] can be mentioned.

Classification performance of bankruptcy models is often visualized and measured by ROC<sup>16</sup> curve, which is shown on the figure below.

**Figure 9. ROC curve**



ROC curve denoted by B expresses the relationship between Type II error and the hit rate of bankrupt firms. Models achieving the highest hit rate among bankrupt companies at any given level of Type II error are preferred when comparing different

<sup>16</sup> Receiver Operating Characteristic

models. In this case, curve B diverges from the diagonal line (C) which represents the random guessing. The greater the divergence, the higher the classification performance. To measure this, area under the ROC curve, which is the area between curve B and the diagonal line in this example, is a commonly used performance indicator (see for example the work by Blanco et al. [2013]). The area under the ROC curve is maximal for curve A which is the case of the perfect model.

In addition to these measures, Brier score is also used which is calculable only in the case when the probability of bankruptcy is quantifiable by the classification algorithm. Brier score can be obtained with the following formula (Lin et al. [2011]):

$$BS = \frac{\sum_{i=1}^n (p_i - o_i)^2}{n}$$

Where  $p_i$  denotes the probability of bankruptcy estimated by the model,  $o_i$  stands for the observed class of the company (1 in the case of bankrupt, 0 in the case of operating firm),  $n$  denotes the number of observations.

The presented performance measures are the most popular metrics in the literature when assessing and comparing different models. However, there are also other tools for evaluating the performance of the models which are also often used in the literature but I don't have the possibility to describe them in more detail in this work due to space limitations. The Gini coefficient and the CAP<sup>17</sup> curve are also worth mentioning because these metrics are commonly used in the practice of default modelling as well as in the literature.

---

<sup>17</sup> Cumulative Accuracy Profiles

### **2.3.5.2. Validation procedures**

Performance measures discussed above can be applied on the training sample used to develop the model and on another set of observations which weren't included in the training sample. This dataset is called as test sample in the literature. Since, bankruptcy models aim to predict future insolvency, their predictive power can be assessed by the error on the test samples. It should be emphasized that the error rate on a particular test sample is only an estimation since this error rate is dependent on which observations are present in the given test sample. The assessment of the predictive power of the models is called as validation in the literature, so the test sample is often called as validation sample as well.

Ex post (ex ante) validation is the case when the performance measures are calculated on the training (test) sample. In the most rigorous form of ex ante validation, the predictive performance of the model developed on data from a given period of time is evaluated on a test set from a later period. For example, if a model is developed on a training sample which covers the period 2010-2012, then a test sample should consist of observations for example from 2013 (or from a later period) in the case of this validation scheme which can also be named as intertemporal validation.

In the first decades of research in bankruptcy prediction, researchers often faced with the problem that the performance of the models is much lower on data coming from a later period of time (Platt-Platt [1990]). Instability<sup>18</sup> of financial ratios over time was marked as the most likely explanation for this phenomenon. The problem has still been existed as it was highlighted by Berg [2007], however, this question is often ignored in the mainstream research of bankruptcy prediction.

---

<sup>18</sup> See the subsection 2.3.2.1. for more details

It should be noted that there are two main modelling approach in bankruptcy prediction. The first stands closer to the practical model building process used for rating debtors. In this context, bankruptcy prediction can be considered as a data fitting task (Darayseh et al. [2003]) which aims to develop as accurate predictive models as possible. The second modelling approach is more common in the academic world since researchers not necessarily want to develop a ready-to-use model for practical application. Instead, the purpose is often to propose and demonstrate a new or different concept and its usage in bankruptcy prediction. In such cases, the authors wish to prove the fact that the predictive power of models developed by applying the concept proposed by them is not attributable to a particular training and testing sets chosen by the researchers, but the proposed concept is efficient irrespective of the way that the training and testing samples were generated. In the case of this kind of empirical investigations, the dataset is divided into training and testing samples several times. This multiple validation scheme is conducted in the following three main forms in the literature.

#### 1. Leave-one-out (Jackknife) validation

In line with its name, the procedure uses  $n-1$  instances as training sample and the remaining observation as test “sample” in the case of a dataset consisting of  $n$  companies. The method is repeated  $n$ -times so that each instance will be used once as a testing case. The predictive performance of the model developed with the examined concept is estimated by using the predictions made on each observations when they were used as testing cases.

If the number of the available observations is low and it is not possible to generate test sample(s) then this method is to be preferred. It barely decreases the number of instances available for the training sample, so it is possible to generate more accurate models when the sample size is very limited. This is another advantageous property of the method (Wang [2004]) especially for classification algorithms which performance is more sensitive to the size of the training sample (Leshno-Spector [1996]). The drawback of this validation scheme lies in its high computational cost which may cause problems in the case of large datasets and more sophisticated classification algorithms. These are the reasons why this procedure is rarely used in practice.

## 2. Cross-validation

The most frequent technique used to estimate the predictive power of bankruptcy models is the cross-validation scheme in which the entire database is divided into  $n$  equal parts.  $n-1$  parts are used as training sample and the remaining for testing purposes. In this scheme, every part is used once as the testing sample. At the end of the process, after every part has been used as testing sample, the average of the accuracies obtained on each test samples is used as the estimation for the predictive power of the model. There has been a vast amount of research using this approach form 2000 (see for example the works by Ahn et al. [2000] and Harris [2013]).

The estimation obtained by using cross-validation may be sensitive to the percentiles used to divide the database into  $n$  equal parts. In order to mitigate the possible bias coming from arbitrarily chosen percentiles, cross-validation procedure can be applied multiple times ( $k$  times). In this case, percentiles dividing the database into  $n$  equal

parts are randomly chosen  $k$  times, then the whole cross-validation process is repeated  $k$  times. The predictive performance of the model is evaluated on the basis of the average classification accuracy obtained on the  $k \times n$  testing samples. This method was used by Alfaro et al. [2008] for example.

Finally, it should be noted that there isn't any guideline in the literature for the values of  $k$  and  $n$ . These should be determined by the researcher who has to take into account that the size of the testing sample should be large enough so that statistical conclusions on the predictive performance of the model can be drawn. In the reviewed literature,  $n$  varied usually between 5 and 50.

### 3. Multiple random division

In addition to the methods discussed above, partition of database into training and testing samples in a user defined ratio is a very common approach in the literature. An example for this is the study by Min et al. [2011]. They divided the dataset into training and testing set in 80:20 ratio 100 times by using different cut point for the division. Predictive power of the model was evaluated on the basis of the average accuracy ratio obtained on the 100 testing samples. Similar to the cross-validation, there isn't objective guideline neither for the ratio of the division nor for the number of divisions. Hu-Tseng [2007] and Hu [2009] examined the following partitioning ratios with different classifiers: 80:20, 70:30, 60:40, 50:50. They highlighted that the difference between the estimated predictive powers of the same classifier may reach the level of 5 percentage points depending on the ratio applied for division. Optimal ratio, however, wasn't found and proposed by the authors because their results were contradictory. So, only one guideline remains for practice and scientific research: the

size of the testing set should be large enough so that conclusion can be drawn based on them, furthermore, there must be enough instances in the training sample as well in order to have enough information for building adequate models.

Though, cross-validation and multiple division are frequently used, they are sometimes subject to criticism. Both methods estimate the predictive performance of a model on the basis of the results obtained on test samples which were taken into consideration to some extent during the model building. However, this can raise the issue that these models may be overfitted to the given testing samples. To avoid this problem, Chen et al. [2011] divided their database into three parts: training, testing and validation with the ratio of 60:20:20. The 60 % of the dataset is used for developing a model. The testing sample is used to optimize the parameters of the classifier in order to maximize its performance which is tested on the remaining 20 % of the observations called as validation set. The advantage of this approach is that observations used for estimating the predictive power of the model are not present in any form in the model building, so the probability of sample specific results could be minimal.

Finally, it should have been noted that there hasn't been consensus in the literature regarding the choice between validation techniques presented above. The most popular procedure is the cross-validation followed by the multiple random division. Application of leave-one-out method is very rare due to its extremely high computational cost. It is an interesting issue that researchers usually don't justify why they choose one or another validation method. This practice can raise the question of whether there is a difference between the two most popular validation methods. Do one or another give more optimistic estimate for the predictive power of the model or not? I empirically investigated this question. The details of this research

cannot be described here due to space limitations. If the Reader is interested in this topic, I recommend one of my works (Nyitrai [2014]).

## **2.4. The Hungarian literature of bankruptcy prediction**

The possibility of scientific research of bankruptcy prediction in Hungary was facilitated by the appearance of bankruptcy law in 1991 (Virág et al. [2013]). The first Hungarian bankruptcy prediction model was published by Virág [1993]. The literature on bankruptcy prediction has become increasingly relevant in the Hungarian journals related to business sciences in the last two decades. In the following subsections, I briefly present the main findings coming from the Hungarian publications in this field. Since the Hungarian literature is narrower than the international, furthermore it focuses primarily on the mainstream research direction, it is not possible to review the Hungarian literature by following the “cross-sectional” approach presented earlier in the thesis. Instead, I briefly describe the most important findings of Hungarian studies published in scientific journals in chronological order.

### **2.4.1. The beginning of bankruptcy prediction in Hungary**

Bankruptcy prediction was pioneered by Virág-Hajdu [1996] in Hungary. They examined 17 financial ratios for 156 Hungarian manufacturing firms over the period of 1990-1991. They applied linear discriminant analysis and logistic regression as classification tools and achieved 77.9 % and 81.8 % accuracy respectively.

Later, the same authors developed a set of bankruptcy models which consisted of one model for the whole Hungarian economy, 10 models for the main economic sections

and 30 specialized models for the most important sectors of our country. This set of models was based on data for more than 10000 Hungarian enterprises and was capable of achieving especially high hit rates. Another important finding of the cited authors that they emphasized that predictions of bankruptcy models can be considered as early warning signals rather than as predictions literally. They argued that the predicted probabilities express only the extent to which each observation is similar to bankrupt or operating companies (see Virág et al. [2013] for more details).

#### **2.4.2. Neural networks in Hungarian bankruptcy prediction**

The Hungarian literature of bankruptcy prediction has also been affected by the development of classification methods. In the middle of the first decade of the new millennium, studies comparing different classification methods, which constitute the mainstream research direction in the international literature, appeared in Hungary as well. The results obtained in our country are in line with those of the international tendencies.

Machine learning methods based on artificial intelligence have started replacing for statistical methods from the 1990s in the international literature. This phenomenon dates back to the early years of the new millennium in Hungary. A comparative study based on data of the first Hungarian bankruptcy prediction model was conducted by Virág-Kristóf [2005] who examined the question of whether neural networks are real alternatives to traditional statistical methods, similarly to the trends observable in the international literature.

Later, the same authors conducted another empirical investigation on the basis of the same database. They compared the performance of four widely used classification<sup>19</sup> algorithms in which industry adjusted financial ratios proposed by Platt-Platt [1990] were used as input variables. They concluded that traditional statistical methods are outperformed by the methods of artificial intelligence in the case of Hungarian enterprises as well (Virág-Kristóf [2006]).

#### **2.4.3. Appearance of state-of-the-art methods in the Hungarian bankruptcy prediction**

Combination of different methods constitutes a new and developing trend in bankruptcy prediction. This approach appeared in the Hungarian literature as well in the work of Virág-Kristóf [2009] at the end of the first decade of the new millennium. In their analysis conducted on a small sample, firms as coordinates in a high dimensional feature space were projected into a lower dimensional space by multidimensional scaling, then they applied logistic regression to classify the observations in this reduced space. The cited authors found outstanding classification performance even in the case of this simple method.

The Hungarian literature of bankruptcy prediction has reached the level of the international literature with respect to examined research questions as well as the applied methods. In addition to the comparative studies related to the mainstream research direction, studies emphasizing the importance of data preprocessing tasks also appeared thanks to the work by Kristóf-Virág [2012]. They examined the predictive ability of discretized variables via CHAID method and the factors

---

<sup>19</sup> discriminant analysis, logistic regression, decision trees, neural networks

produced by principal component analysis as input variables to the commonly used classification algorithms. Based on the results, the cited authors concluded that effects of data preprocessing tasks under consideration depend on the applied classification algorithm, that is, for example, using principal components instead of the original financial ratios had different impact on the performance of neural networks than in the case of decision trees. They also found that the application of discretization via CHAID and principal components (or both of them) never decreased the performance of the classification algorithms under consideration. Data preprocessing had positive impact on the discriminating performance but to different extent depending on the chosen classification method.

SVM belonging to the state-of-the-art methods also appeared in the Hungarian literature of bankruptcy prediction in the work of Virág-Nyitrai [2013]. They were the first in Hungary who applied this method to data for Hungarian firms. The cited authors used the database of the first Hungarian bankruptcy prediction model since numerous empirical analysis had been conducted by different classification algorithms on this dataset. So, the choice of this database was reasonable because it allowed direct comparison of the performance of SVM with those of traditional statistical methods and neural networks. The findings of the cited authors are in line with the trends observable in the international literature because the SVM showed higher hit rates on the training and the testing samples under consideration than the traditional methods of mathematical statistics and neural networks on this database.

Similarly, the database of the first Hungarian bankruptcy prediction model was used to compare the performance of rough set theory (RST) with other classifiers by Virág-Nyitrai [2014]. As it was stated in the sections presenting the methodological issues of bankruptcy prediction, algorithms of machine learning are typically black

boxes which property significantly restricts their applicability in the real life in spite of their relatively high classification power. Interpretability is a very important issue regarding bankruptcy prediction models, hence the application of decision trees were recommended by Martens et al. [2010] in this field. Similar question was examined by Virág-Nyitrai [2014] because the RST method also generates easily interpretable “if-then” rules. Using the dataset of the first Hungarian bankruptcy prediction model, the cited authors addressed the question of whether the modeler has to give up the interpretability of the model to achieve higher classification performance. They found that the RST method was capable of providing the same classification performance than the SVM did in their earlier work, that is, they didn’t find a trade-off between the interpretability and the predictive power of bankruptcy prediction models.

### **3. RESEARCH HYPOTHESES**

In the first part of the thesis, I described the main terms of the chosen research topic, then in the second part, I overviewed the evolution of the literature of bankruptcy prediction by applying a new “cross-sectional” approach. In this part, I present the research questions wished to examine which may be relevant to the future development of this field based on the tendencies observable in the reviewed literature.

#### **3.1. Open research questions in the literature**

As it could be seen in the literature review, the development of this field is dominated by the progress of methodology which is indicated by the number of publications dealing with the comparison of different classification methods. I think that this research direction can be considered as a mainstream topic in bankruptcy prediction. Based on the results published in the international literature, I think that the possibility of developing model with higher classification performance by using a more recent and more sophisticated methodology is limited, the incremental increase in predictive performance to the extant and widely used methods is usually marginal.

At the same time, it should also be emphasized that only a marginal increase in predictive performance could generate substantial benefit (Wang-Ma [2012]) to those who apply these models. Because of it, it should be taken all the possible opportunities to enhance the performance of the bankruptcy prediction models. The main research hypothesis of my work is that the predictive power of models can be increased not only from the methodological side. This can also be possible by

exploiting the information content of financial variables to a greater extent. An alternative to do this is to take into account the information captured in the time series of static financial ratios.

### **3.1.1. Dynamic financial variables in bankruptcy prediction**

The process of bankruptcy model building depicted in Figure 1 was the guide line for the literature review of the thesis. Comparative studies analyzing different classification algorithms have given very much attention in the scientific literature. This can be considered as the mainstream research line in this field. Compared to this, other stages of bankruptcy prediction model building are understudied.

It is especially true for the data preprocessing which has just started getting greater emphasis in the last decade. In the literature review, I cited numerous authors who think that the static nature of bankruptcy prediction models is a serious problem, so filling this gap can be considered as a fruitful future research direction. A simple approach was proposed by Berg [2007] who used the year-to-year changes of financial ratios in addition to their static values as input variables for his models. He tried to relax the static nature of models by using financial ratios not only from the most recent year, but he used financial data from the first, second and third year prior to bankruptcy as explanatory variables.

A more sophisticated methodology was applied by Du Jardin-Séverin [2012] to dynamize bankruptcy prediction models. They visualized the time series of financial ratios by using self-organizing maps and then classified them in order to identify the typical evolution of financial ratios in the case of bankrupt and operating companies.

In my opinion, similar dynamization can be conducted by generating such variables which can express the value of a financial ratio from the most recent year in the mirror of those from the previous period. I propose the following formula which can be a possible tool for this purpose:

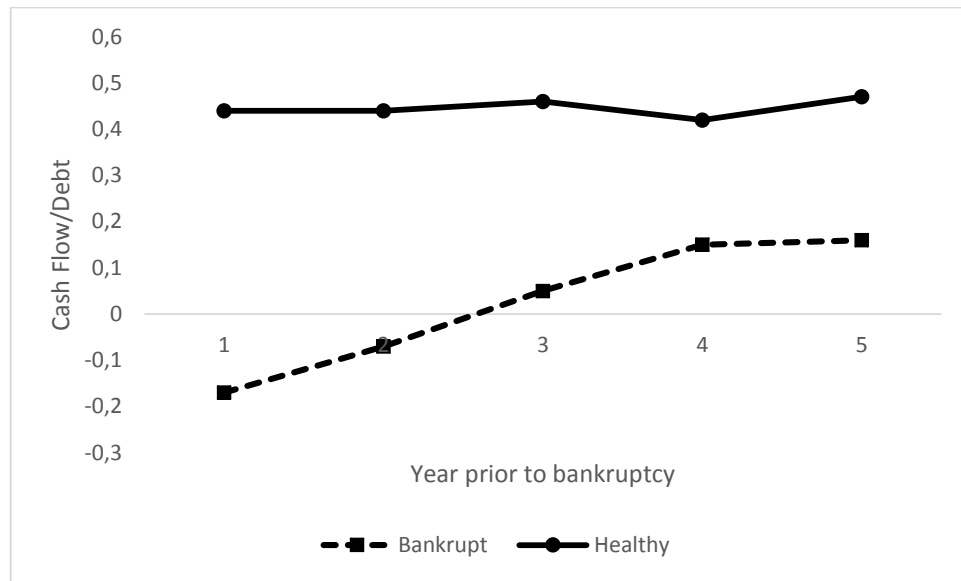
$$\frac{X_{i,t-1} - X_{i,\min_{[t-2;t-n]}}}{X_{i,\max_{[t-2;t-n]}} - X_{i,\min_{[t-2;t-n]}}}$$

where

- $X$       the value of a financial ratio,
- $i$       id for a firm,
- $t$       the year for which prediction to be made,
- $n$       the number of years for which the firm was observed.

The variable defined above is called as dynamic financial variable (for short dynamic variable) in the followings. I assume that these variables are capable of effectively discriminating between bankrupt and non-bankrupt companies. The base of this hypothesis is the seminal work of Beaver [1966] who visualized the time evolution of Cash flow/Debt ratio over a five year period before bankruptcy for a bankrupt and for a healthy company in a figure similar to the following.

**Figure 10. The time series of Cash flow/Debts ratio**



Source: own construction based on Beaver [1966]

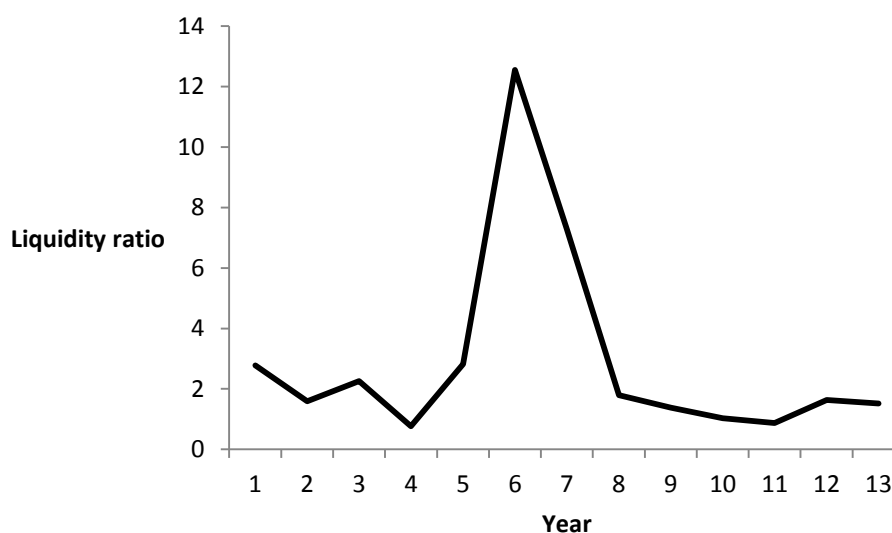
In the figure, time series of Cash flow/Debt ratio for a bankrupt (dashed line) and a non-bankrupt (continuous line) firm were visualized by the cited author. As one can see in the figure, the ratio for the healthy firm is relatively stable and high, in contrast, it is continuously decreasing for the failed company as it is coming closer to bankruptcy. Dynamic financial variables proposed in the thesis are capable of expressing this time evolution for each financial ratio.

### **3.1.2. Handling outliers**

To calculate the dynamic variables, it is necessary to have a database consisting of continuous time series of financial ratios for each observation so that their trends can be analyzed. It should be emphasized that time series of financial ratios for Hungarian companies is not likely to be so tendentious as it could be seen in Figure

10. Examining longer time series, it is likely to see such years which breaks the trend formed by the data of another years, that is, outlier years may also be observable in the time series. An example drawn from my research database is shown in Figure 11 where the time evolution of the liquidity ratio for a healthy Hungarian firm can be seen.

**Figure 11: A time series containing outlier value**



As it can be seen in the figure, the liquidity ratio of the firm under consideration is relatively stable over the examined period. The only exception is the sixth year when the variable had an extreme high (outlier) value. Conclusions based on this time series would be biased if one takes the value of the outlier year into consideration.

To avoid this, I think it is worth replacing for the outlier values in each financial ratio time series by the value which is the closest to the outlier but it cannot be considered as outlier. This replacement is not worth conducting by the data from the most recent

year because failing companies often exhibit extreme high/low values one year prior to bankruptcy in the case of some financial ratios. Their correction, though, would make the database easy to handle statistically but at the cost of losing information which could be relevant in identifying bankrupt companies.

Before replacing, it should be defined which value can be considered as outlier. In the absence of objective definition, I think it is worth using statistical rules of thumb. It is a common approach to consider a value as outlier if it is outside the range of  $\pm$  three standard deviations around the mean of the given variable. However, the application of this “rule” is grounded only when the number of observations is relatively high (more hundreds or thousands). Given the length of the time series for each companies, I think it is more appropriate to use a “more rigorous outlier definition” in order to more effectively filter out outlier observations. Hence, I examined the effects of the case when an observation is considered as outlier if its value is outside the range of  $\pm$  2 standard deviations around the mean of the time series of the given variable for each firms. The method proposed here for handling outliers will be presented in detail on the basis of a numerical example in Section 5.3.

### **3.2. Research hypotheses examined in the thesis**

There is a lot of open research question in the field of bankruptcy prediction of which numerous example was presented in the literature review. Given the empirical nature of this science, unambiguous answers cannot be expected for any questions. However, the more studies come to similar conclusions, the more likely is that those are the true answers to the given question. Due to space limitations, I had to restrict

my thesis to the most relevant topics of which I could examine only a few research questions during my PhD studies.

Building on the reviewed literature, I examined the following hypotheses in my work.

1. In the case of my research sample, there are such dynamic financial variables which have statistically significant discriminating power between the bankrupt and healthy companies.
2. Using decision trees developed with CHAID procedure as classification algorithm for bankruptcy prediction models, the predictive power of models containing static financial ratios from the most recent year in combination with the dynamic variables is significantly higher than those of models in which only static financial ratios from the most recent year were used as input variables.
3. In the case of the available sample, predictive power of models containing dynamic financial variables and developed with CHAID algorithm can be increased if outlier values are replaced by such values which are the closest to them in the same time series but not outliers.

During the examination of the third hypothesis, the whole time series of each financial ratios were standardized by using the mean and the standard deviation of the time series in the period  $[t-2; t-n]$ . Observations in each time series outside the range  $(-2, 2)$  were considered as outliers. The calculations are presented in detail on the basis of a numerical example in Section 5.3.

Examination of the hypotheses was carried out by using the CHAID procedure generating decision tree. One of the reasons for the choice is that this algorithm is a

built-in method in SPSS which is a commonly used software in empirical researches. Another reason is that interpretable models are preferred in the literature with respect to the applicability of research results in the real life (Martens et al. [2010]). Decision trees were recommended by the cited authors because these algorithms don't impose any statistical restriction on the examined dataset, furthermore classification is done via easily interpretable "if-then" rules, and finally their classification performance is relatively high. Based on the above mentioned reasons, I have chosen the CHAID procedure which also generates decision tree.

## **4. DATA USED IN THE EMPIRICAL RESEARCH**

As it was mentioned in the introduction, the scientific research of bankruptcy prediction was restricted by data availability in Hungary until the end of the first decade of the new millennium. This obstacle has been eliminated since 2009 because, from this year, enterprises registered in Hungary has been obliged to publish their financial statements on a website determined by the law. Financial statements are freely available on this site where it is possible to search based on the name or the register number of firms.

Taking this opportunity, I wished to examine my hypotheses based on a database consisting of as much observation as possible for the most recent data. I describe the method of sampling and the resulting dataset in this section.

### **4.1. General sampling issues**

In order to examine the addressed research questions, I collected my own sample. The first question was related to the size of the dataset. I maximized the number of observations at 1000 instances that can be considered as a small sample compared to the total number of firms operating in Hungary. The choice of this relatively small sample size can be justified by the manual way of data gathering which was a very time-consuming task especially from publicly available databases. Furthermore, it should be emphasized that data were collected for all the available years until 2001, not only for the most recent year.

Next to the sample size, the second important question was related to the proportion of bankrupt and non-bankrupt firms in the dataset. I represented the two groups evenly in my database. The main reason for this decision was the fact that application of a representative sample would have resulted in a database containing very few instances from the bankrupt group. This issue would have raised the problem of whether such a dataset provides enough information for the classification algorithm to develop bankruptcy prediction model which is capable of identifying bankrupt companies with an acceptable accuracy. Since, in bankruptcy prediction, the proper classification of failed firms is more important than the healthy companies (Du Jardin [2010]), I chose the ratio of 50-50 % for bankrupt and operating companies during data gathering in order to avoid the above mentioned problem.

The third question regarding to sampling is the operationalization of the terms used in the research hypotheses. In the context of bankruptcy prediction, it should be defined which firm can be considered as bankrupt and which as healthy. It is of special importance since, as it was discussed in the introduction describing the terms used in the thesis, the term of bankruptcy is usually not literally used in this field.

Observations under liquidation or bankruptcy proceeding at the time of sampling were considered as bankrupt companies. Firms which were not under liquidation, bankruptcy or voluntary dissolution proceedings at the time of sampling were considered as healthy observations. These facts were observed on the basis of the Hungarian Trade Register which is available online at [www.e-cegjegyzek.hu](http://www.e-cegjegyzek.hu) provided by the Ministry for Public Administration and Justice.

## 4.2. Process of data gathering

The Company Gazette was the starting information source of data gathering. Some of its issues were randomly chosen. The observations of the sample were selected from the companies publishing announcement in these issues by following the aspects below.

1. The most essential aspect of sampling was that financial statements (balance sheet and income statement) of firms must be available at least three consecutive years starting from the most recent year in case of operating and starting from the year before bankruptcy in the case of failed firms. The main reason behind it that the primary purpose of the research is to examine the discriminating power of variables which reflect the value of the most recent financial ratios in the mirror of those from previous years. To do this, it is necessary to have data over at least a three year period.
2. Firms whose financial ratios didn't exhibit standard deviation over time were excluded from the analysis because dynamic financial variables cannot be calculated in such cases.
3. Companies which hadn't been realizing sales<sup>20</sup> for at least two consecutive years were also excluded from the experiment. The reason for it that these firms are likely not doing real business activity, so their inclusion would have had distorting impact on the performance of the models.

Financial ratios for firms included in the sample were calculated on the basis of data in financial statements published in accordance with the requirements of the

---

<sup>20</sup> Companies with at most one year without sales were allowed to be included in the sample since I experienced several times during the data gathering that bankrupt companies often have zero sales in the year before their bankruptcy. So, excluding such cases would also exclude an important signal for impending failure which can help in identifying and predicting bankruptcy.

Hungarian accounting act. The financial statements were accessed at <http://e-beszamolo.im.gov.hu/kereses-Default.aspx> which is provided by the Ministry mentioned earlier.

It was a problematic issue that in the case of the discussed public databases it is not possible to search based on the activity and/or the size of the companies. Because of it, I didn't have the opportunity to take these aspects into account during the sampling. However, it is not unique to my research because studies dealing with the bankruptcy of small and medium sized companies are usually not restricted to one or more particular industry.

As it was previously mentioned, the characteristics of industries may have considerable impact on the values of financial ratios which can decrease the predictive power of models based on a dataset consisting of observations from different industries. Since I didn't have the opportunity to take into account the industries of the firms during the sampling, so models developed in the thesis are not suitable to using in practice for predictive purposes. However, this is not necessarily a problem because the aim of the thesis is to demonstrate the usage of dynamic financial variables in bankruptcy prediction and not to develop models for real life application. The sample available for this purpose is quite heterogeneous but this fact could also be advantageous from the viewpoint of the aim of the thesis: if the incorporation of dynamic variables into the models based on such a heterogeneous database will be proven to be useful, then it is likely that dynamic variables can also improve the performance of those models which are based on a more homogenous sample.

Since it was not possible to search in the publicly available databases used for sampling on the basis of the activity and the size of the companies, the sample was impossible to be restricted with respect to these factors a priori. However, having finished the sampling, these factors can also be incorporated into the models as independent variables. Furthermore, after sampling it could have been possible to construct models for companies from different industries (for firms from different groups based on their size) but the number of observations in each industries (groups based on the size of companies) would have been too small for building meaningful bankruptcy prediction models. For these reasons, I used the whole database regardless of the size and the activity of each companies.

The situation was similar in the case of the age of the firms. The publicly available databases served as the source of my data don't contain the age of the companies. Instead, I used the number of years for which the companies published their financial statements on the website regulated legally for this purpose. In my opinion, this variable can be a good proxy for the age of the firms.

Independent variables for bankruptcy models, in line with the traditions of the literature for decades, were chosen from the financial ratios calculable on the basis of data in the balance sheets and income statements published by the firms. The primary source applied in choosing the explanatory variables was the set of financial ratios frequently used in prior studies in Hungary (Virág et al. [2013]), furthermore I relied on my own considerations as well. The name and the formula of financial ratios applied in the experiments are shown in Table 2. Closing values in balance sheets and income statements were used to calculate financial ratios listed below.

Table 2. The name and formula of financial ratios used in the empirical research

| Variable name                   | Formula                                               |
|---------------------------------|-------------------------------------------------------|
| Liquidity ratio I               | Current assets/Current liabilities                    |
| Liquidity ratio II              | (Current assets – Inventories)/Current liabilities    |
| Ratio of cash                   | Cash/Current assets                                   |
| Cash flow/Liabilities           | (After-tax profit + Depreciation)/Liabilities         |
| Cash flow/Current liabilities   | (After-tax profit + Depreciation)/Current liabilities |
| Equity financing ratio          | (Fixed assets + Inventories)/Equity                   |
| Assets turnover                 | Net sales/Assets                                      |
| Inventory turnover              | Net sales/Inventories                                 |
| Receivable turnover (time)      | Receivables/Net sales                                 |
| Indebtedness                    | Liabilities/Assets                                    |
| Equity ratio                    | Equity/Assets                                         |
| Creditworthiness                | Liabilities/Equity                                    |
| Return on sales                 | After-tax profit/Net sales                            |
| Return on assets                | After-tax profit/Assets                               |
| Receivables/Current liabilities | Receivables/Current liabilities                       |
| Ratio of working capital        | (Current assets – Current liabilities)/Assets         |
| Size                            | Natural logarithm of Assets                           |
| Years                           | Number of observed years                              |

Return on equity (ROE) is a commonly used independent variable in bankruptcy prediction models. However, this ratio often raises the problem of double negative division (Kristóf [2008]). This is the case when the numerator and the denominator of the ratio are also negative. There hasn't been preferred solution to this problem in the literature, so this ratio was excluded from the analysis.

Another problematic issue related to financial ratios is the case when the denominator is zero. It is a common solution to this problem that these cases are considered as missing values and then replaced by the average or some extreme percentiles of the other observations. In my opinion, this approach doesn't necessarily bring consistent data into the models. I wish to illustrate this with the following example.

Let us consider an enterprise which always meets its financial obligations when they come due or just before the end of the financial year, hence it doesn't have current liabilities at the balance sheet date which makes impossible the calculation of the

liquidity ratio<sup>21</sup>. Assume that the firm in the example also has a substantial amount of current assets which allows financing a liquidity shock occurring in the future. If the liquidity ratio of such a firm was replaced by the average of the other examples in the sample, then this company would be considered as an “average” firm with respect to its liquidity, but it is not true based on its data. Another solution is the replacement by some extreme percentile. In this case, however, the replacement value may be sample specific.

To avoid the above mentioned problems, for the cases where the denominator would be zero, the zero value was replaced by one<sup>22</sup>. So the liquidity ratio of the firm in the previous example takes a very high value, indicating very high liquidity for the company. Such observations exhibiting extreme high or low values (outliers) are often excluded from the databases for building bankruptcy prediction models because they distort the performance of statistical models. However, in agreement with the finding that deviation from normal distribution and presence of outliers are basic characteristics of bankruptcy prediction rather than rare exceptions (McLeay-Omar [2000]), it is reasonable to use such classification methods that are capable of handling outliers.

In the case of my research, keeping outliers within the model was especially important due to the sample size, since their exclusion would drastically have reduced the number of observations available for model building. Furthermore, in my opinion, models developed without outlier observations wouldn't have enough capability of identifying outlier observations when applied on a dataset containing outliers.

---

<sup>21</sup> Liquidity ratio = Current assets/Current Liabilities

<sup>22</sup> Since data in Hungarian financial statements are expressed in 1000 HUF, so the replacement by one means 1000 HUF.

The intention of keeping outlier observations within the model also had influence on choosing the classification method used to examine the hypotheses. The advantage of decision trees that their classification performance is not affected by the presence of outliers in the database (Twala [2010]).

The database resulted in the sampling process described above consists of 1000 Hungarian enterprises for which financial data are available over the 2001-2012 period (this means altogether 7592 firm-year observations). The sample is divided evenly between bankrupt and non-bankrupt firms. These observations are described by 17 financial ratios presented in Table 2 and by the number of years for which the data were available in the public databases used as the source of data. For each of the 17 financial ratios, dynamic variables defined in Section 3 were also calculated. The purpose of these variables is to express the data of the most recent year compared to those from the previous period. In the case of healthy (bankrupt) companies, financial ratios for the most recent year (the year preceding bankruptcy) are called as static financial ratios in the rest of the thesis. The resulting database consists of 35 variables. Their descriptive statistics are given in Table F1. in the Appendix.

## **5. EMPIRICAL RESEARCH**

After reviewing the literature presented earlier, I had identified numerous possible research directions of which I chose the dynamization of bankruptcy prediction models as the main topic of my thesis. I proposed a formula in Section 3 for this purpose. The three research hypotheses presented there were examined based on the sample described in Section 4. In this part, I summarize the results of my empirical research following the examined hypotheses.

### **5.1. The examination of the first research hypothesis**

Among the research hypotheses presented in Section 3, the first was as follows:

„In the case of my research sample, there are such dynamic financial variables which have statistically significant discriminating power between the bankrupt and healthy companies.”

Examination of this hypothesis was carried out by using a simple statistical method: the T-test for comparing means from two independent samples. It should be emphasized that result of this test is only valid if the following requirements of the test are met:

- populations to be compared are independent;
- the distribution of financial ratios within the populations to be compared is normal;
- the standard deviation of financial ratios for failed and non-failed firms are approximately equal.

Due to the sampling characteristics, the first assumption is valid in this case, but the two other assumptions are certainly not since it is well-documented in the literature that financial ratios are usually not normally distributed and their variances are not equal for bankrupt and non-bankrupt firms. In spite of these, using the t-test is a common practice in bankruptcy prediction which is proven by the study of Liang et al. [2015] who recommend this method for selecting independent variables in the case of data mining algorithms. They didn't examine (or didn't report) whether the restrictive assumptions of the test are met in the case of their empirical examination or not. For this reason, I followed the typical practice of the mainstream research direction on bankruptcy prediction and I neither examined the validity of these assumptions in the case of my research database.

Descriptive statistics for the examined financial ratios and for the dynamic variables are included in Table F1. in the Appendix where dynamic variables calculated from the time series of static financial ratios are denoted by the prefix "D\_". Significant differences at the significance level of 5 % between the means of bankrupt and non-bankrupt companies are highlighted by \*.

According to the results, of the examined variables 13 revealed to be significantly different between bankrupt and healthy companies. If the number of observed years is not taken into account, then it can be seen that out of the 12 significant features 7 are static and 5 are dynamic. That is, we can see a considerable increase in the number of variables which have statistically significant discriminating power between the two groups.

Based on the fact that dynamic variables were calculated for each of 17 financial ratios, that is, a dynamic variable belongs to every financial ratio, three groups can be distinguished based on the results of the t-test:

- only the static financial ratios from the most recent year were significant (Cash ratio, Ratio of working capital);
- financial ratios where the static and dynamic variables were also significant (Receivables turnover (time), Indebtedness, Equity ratio, ROA);
- only the dynamic variables were significant (Cash Flow/Liabilities, Inventories turnover).

The results suggest that the possible set of variables which are capable of discriminating bankrupt and healthy companies in Hungary has increased. The results also showed that dynamic variables are complementary to the static financial ratios rather than replace for them because there were financial ratios where only the static values were significant, at the same time, there were other cases where only the dynamic variable showed significant difference between the bankrupt and non-bankrupt companies.

It is interesting to note that liquidity ratios and the size of the company weren't proven to be significantly different. It is also conspicuous that the number of significant features (13) is relatively small compared to the total number of variables (35), furthermore one can find extreme high values among the descriptive statistics of the independent variables which can be considered as unrealistic by the Reader. Two main reasons must be mentioned for this: on one hand, the sample is quite heterogeneous with respect to the size and the activity of the firms in it. On the other hand, data preprocessing could also lead to such extreme values because in the cases

where the denominator of the ratio would have been zero, I replaced the zero value by one which resulted in very large values. This can be observed particularly in the case of turnover ratios where the denominator is the amount of the Net sales which was often very low or zero in the case of bankrupt companies and for the firms having financial problems. The presence of such extreme high or low values can seem to be distorting, but, in my opinion, this approach may help in exploiting the information content of financial variables to a greater extent and hence makes it possible to develop bankruptcy prediction models with higher predictive power.

Based on the mean values of significant variables in Table F1. in the Appendix, conclusions can be drawn on the financial state of the Hungarian companies. General observation that the magnitudes of the averages in the two groups (bankrupt, healthy) are in line with the preliminary expectations, namely, the mean of the Cash ratio is higher in the case of non-bankrupt companies than for the bankrupt firms. Similar tendency can be observed for the other variables as well.

In his pioneering work, Beaver [1966] compared the predictive ability of the financial ratios most frequently used in financial analysis. He found that CF/Debt is the “best” variable with this respect. Though the static type of this variable is not significant, but the dynamic variable calculated from the time series of CF/Liabilities, which is a very similar ratio to that used by Beaver [1966], proved to be significantly different between the bankrupt and operating Hungarian companies. As it can be seen in Table F1., the average value of the dynamic CF/Liabilities for the non-bankrupt group is between the minimum and the maximum values of the previous period while this average for the failed firms is much lower than the minimum value of the previous period. That is, examining the change of this ratio

from the previous period to the most recent year can help in identifying the bankrupt companies.

Similar remarks can be done on the turnover ratios, it is especially true for the Inventories and Receivables turnover. Based on the results, very high values can be observed for the latter ratio in the case of bankrupt companies which indicates a considerable increase from the previous period to the most recent year. Similar trend can also be seen for Indebtedness. That is, for Hungarian firms, sharp increase can be observed in the case of Receivables turnover and Indebtedness in the year just before bankruptcy relative to the previous period.

Another interesting result that the mean of ROA and the ratio of working capital were negative in both groups (bankrupt, non-bankrupt). This result can be attributable to the period of sampling because the last year of the time series for the firms in the sample came from the years 2009-2012 which is the deepest period of the most recent economic recession which seriously affected Hungary as well. In such hard times, it is not contradictory to the previous expectations that the average of the profitability ratios is negative even for healthy firms as well as for bankrupt ones. In such circumstances, discriminating ability of the dynamic variables is manifested in the fact that dynamic variables take more negative (lower) values for bankrupt than for the non-bankrupt firms indicating that the decrease from the previous period to the most recent year, for example in the case of profitability, is much greater for the bankrupt companies than for the healthy ones.

Lastly, some words should be devoted to the only one non-financial variable, the number of observed years which also showed significant difference between the two groups. Based on the results, it can be concluded that healthy firms in the sample

released on average more financial statements than the failed firms which suggests that younger firms have higher probability to go bankrupt. This finding is consistent with the conclusions which can be read in the international literature but it should be emphasized that this variable could be more affected by the characteristics of the sampling method presented earlier.

Based on the results obtained by using the available sample, it can be concluded that there are such variables among the dynamic variables proposed in this work which exhibited statistically significant differences between bankrupt and healthy firms in Hungary, hence the first hypothesis of the thesis has been accepted.

## **5.2. The examination of the second research hypothesis**

The second hypothesis wished to examine was formulated in Section 3 as shown below:

“Using decision trees developed with CHAID procedure as classification algorithm for bankruptcy prediction models, the predictive power of models containing static financial ratios from the most recent year in combination with the dynamic variables is significantly higher than those of models in which only static financial ratios from the most recent year were used as input variables.”

It should be emphasized that the aim of the thesis is to propose dynamic financial variables for bankruptcy prediction and to give empirical evidence on their usefulness. So, developing models which are ready-to-use in practice was beyond the scope of my research. Because of it, I didn't deal with practical questions often occurring in practice such as predicting the exact probability of default for each company or the calibration of these probabilities. The heterogeneity of the available

sample also restricted the possibility and the usefulness of some kind of examinations. However, this heterogeneity can help in proving the *raison d'être* of dynamic financial variables in bankruptcy prediction.

The empirical experiments were carried out by using SPSS Statistics 20. I ran the CHAID procedure by employing its default settings with only one exception: with respect to the size of the sample, the minimal number of instances for producing a new parent node was set to 10, the minimal number of observations for generating a new child node was set to 5.

To avoid the distorting impact coming from the division of the sample into training and testing sets on the reliability of the conclusions, I used tenfold cross-validation which means that the whole sample consisting of 1000 instances was divided into ten equal parts. Models were developed by using 90 % of the entire dataset and its predictive power was assessed on the remaining 10 % of observations. The procedure was repeated ten times such that every decile (part) will be used once as testing set. Consequently, ten decision trees were developed. The predictive performance measured by the hit rate<sup>23</sup> on the corresponding training and testing sets was recorded for each decision trees, then these results were averaged. The rationale behind using cross-validation is that the predictive power estimated on an arbitrarily chosen test set may be biased. In order to avoid or mitigate the possible problems coming from dividing the dataset into training and testing samples, I used the average of the hit rates obtained on the ten testing sets as an estimation for the predictive performance. The results are shown in Table 3.

---

<sup>23</sup> The hit rate of the models means the ratio of correctly classified instances to the total number of observations.

Table 3. The average hit rates obtained by applying tenfold cross-validation using CHAID decision trees

| Sample   | Group        | Input variables         |                   |                              |
|----------|--------------|-------------------------|-------------------|------------------------------|
|          |              | Static financial ratios | Dynamic variables | Static and dynamic variables |
| Training | Bankrupt     | 85.7%                   | 83.2%             | 84.6%                        |
|          | Non-bankrupt | 80.1%                   | 77.8%             | <b>83.8%</b>                 |
|          | Total        | 82.9%                   | 80.5%             | <b>84.2%</b>                 |
| Testing  | Bankrupt     | 76.6%                   | 75.0%             | <b>76.8%</b>                 |
|          | Non-bankrupt | 71.0%                   | 71.6%             | <b>74.6%</b>                 |
|          | Total        | 73.8%                   | 73.3%             | <b>75.7%</b>                 |

Based on the results obtained after tenfold cross-validation, it can be observed that when the dynamic variables were applied alone as independent variables for decision trees, the predictive performance didn't increase compared to the performance obtained by using only static variables. Moreover, in most of the cases, weaker predictive power was observed when only dynamic variables were used compared to the models based solely on static financial ratios. But, if the two sets of variables (static and dynamic) were simultaneously used as independent variables for decision trees, the predictive power was higher than it was when only static variables were used in the models. To illustrate this, the hit rates which are higher than those obtained on the basis of only one set of variables are highlighted in bold. In total, the average extent of the increment was 1.3 percentage point in the case of the training samples and 1.9 percentage point for the testing sets. The obtained results verify the assumption made in the previous section, namely that synergic relationship can be found between the static and dynamic variables in bankruptcy prediction models.

Based on the presented results, the second research hypothesis was also accepted.

### **5.3. The examination of the third research hypothesis**

The last hypothesis examined in the thesis was formulated as follows:

“In the case of the available sample, predictive power of models containing dynamic financial variables and developed with CHAID algorithm can be increased if outlier values are replaced by such values which are the closest to them in the same time series but not outliers.”

With regard to this hypothesis, it should be emphasized the fact discussed in Section 3 that objective definition is not available in the literature for which values can be considered as outliers. In the absence of such definition, statistical rules of thumb are frequently used in statistical analyses. It is a common practice to consider a value as outlier if its standardized value is outside the range of three or five standard deviations around its mean.

In my opinion, this approach is adequate only in the case when the number of observations is sufficiently large. However, the time series of financial ratios analyzed in this research is relatively short, the number of items in each time series is varied between 3 and 12. In the case of such low number of observations, the values mentioned above for defining outliers (3 and 5 standard deviations) could be too “slight” to effectively filter out the distorting effect of outliers on the dynamic variables. So, I had decided to take a more rigorous “rule” for defining outliers. Having standardized the time series of each financial ratio for every firm in the database, standardized values outside the range of two standard deviations around the mean of the given time series were considered as outliers. The potential distorting impact of outliers on the dynamic variables, the way and the result of the replacement proposed earlier in the thesis are shown in the following tables.

Table 4. A time series containing outlier value

| Year | ROA     | Standardized ROA |
|------|---------|------------------|
| 2011 | -4.9112 | -2.3599          |
| 2010 | -1.8360 | -0.5296          |
| 2009 | 0.1153  | 0.6317           |
| 2008 | 0.1671  | 0.6626           |
| 2007 | 0.0403  | 0.5871           |
| 2006 | -0.0924 | 0.5081           |
| 2005 | -0.1939 | 0.4477           |
| 2004 | 0.0284  | 0.5800           |
| 2003 | -0.4214 | 0.3123           |
| 2002 | -5.5220 | -2.7234          |
| 2001 | -1.7470 | -0.4766          |

The table shows the time series of ROA for a company which went bankrupt in 2012. It can be seen that in the year just before bankruptcy the financial ratio was much lower than in the previous 8 years. In order to keep the information content of this extreme low value within the model, data for the most recent year (2011) wasn't taken into account when calculating the mean and the standard deviation of the time series. The data preprocessing was carried out according to the following steps:

- the mean and the standard deviation were calculated based on data from the period 2001-2010;
- the whole time series (2001-2011) was standardized by using the mean and standard deviation calculated in the previous step; the results can be seen in the last column of Table 4.

It can be observed that the ROA for this firm was continuously decreasing from 2008. The extent of the decline was rapid as the firm approached its bankruptcy. At the same time, it must be seen that the financial ratio under consideration had been lower in 2002 than it was in 2011. If one would calculate the dynamic variable

presented in Section 3 for this time series, the following calculation should be carried out:

$$\frac{X_{i,t-1} - X_{i,\min_{[t-2;t-n]}}}{X_{i,\max_{[t-2;t-n]}} - X_{i,\min_{[t-2;t-n]}}} = \frac{-4.9112 - (-5.522)}{0.1671 - (-5.522)} = 0.107$$

This result can be interpreted as the ROA value of the firm was relatively low compared to the values of the previous ten years. In other words, the ratio for the most recent year was only 10 % higher than the minimum value of the previous ten years if one considers the whole range of the period 2001-2010 as a basis for comparison. With respect to the fact, that the company had “survived” such a business year (2002) when the ROA ratio was lower than that of the most recent year, the value of the dynamic variable didn’t help in identifying the potential failure because this company was liquidated in the next year.

It may be assumed that the low value of 2002 may be attributed to the starting difficulties of the company or to another particular event which didn’t result in the liquidation of the firm, so taking it into account when calculating the dynamic variable may bias the information content of the dynamic variable. Furthermore, it can also be observed that the standardized value of this year (2002) is outside the range of two standard deviations. Hence, it is reasonable to replace this value with another value which is the closest to it but not outlier (lies within the range of two standard deviations around the mean of the time series).

It should be emphasized that the purpose of the concept presented here is to exploit as much information as possible from financial ratios. It is especially true for financial data of the most recent year. According to my experience, bankrupt companies often exhibit extreme low/high values in the year prior to bankruptcy,

hence I wished to keep these, often outlier, values in their original form within the models. For this reason, similarly to the process of the standardization presented earlier, I didn't use the values of the most recent years in the replacement. In other words, the data of the most recent year is completely left out from the data preprocessing in order to keep them in their original form and to preserve their original information content. In line with this, the value of the year 2002 was replaced by the value from 2010. Having finished the replacement, the standardization was carried out again by using the new values of the time series. The resulting time series is shown in Table 5.

Table 5. The time series containing outlier value after the replacement

| Year | ROA     | Standardized ROA |
|------|---------|------------------|
| 2011 | -4.9112 | -5.2845          |
| 2010 | -1.8360 | -1.5345          |
| 2009 | 0.1153  | 0.8448           |
| 2008 | 0.1671  | 0.9080           |
| 2007 | 0.0403  | 0.7534           |
| 2006 | -0.0924 | 0.5916           |
| 2005 | -0.1939 | 0.4678           |
| 2004 | 0.0284  | 0.7389           |
| 2003 | -0.4214 | 0.1905           |
| 2002 | -1.8360 | -1.5345          |
| 2001 | -1.7470 | -1.4260          |

After the replacement, every standardized value in the time series lies between -2 and 2, so further handling of outliers is not necessary. If one wishes to calculate the value of the dynamic variable proposed in the thesis, the following calculation has to be carried out:

$$\frac{-4.9112 - (-1.836)}{0.1671 - (1.836)} = -1.535$$

This result shows that the value of ROA realized in the last observed year is much lower than any other value from the previous period. Unlike the result obtained without the replacement, this dynamic value provides more information for predicting the insolvency of this firm occurred in the following year.

The presented example also constitutes an empirical example for the usefulness of applying a more rigorous definition for identifying outliers than the statistical rules of thumb. In the case of the company under consideration, the standardized value in the year which was later replaced had been -2.72, that is, this value wouldn't have been replaced if I had followed the traditional "three standard deviations" rule in identifying outliers. Consequently, by applying a "slighter" definition for outliers, the predictive power of the dynamic variable would also have been lower, as it is was illustrated earlier.

The third hypothesis of my work deals with the question of whether the example presented above can be considered as a unique case; that is, is it possible to enhance the predictive power of dynamic variables in bankruptcy models if the replacement discussed above is carried out for every financial ratio of every firm in the database?

In order to examine this hypothesis, the replacement procedure had been executed where it was necessary in the database and the resulting dynamic variables were used as input variables in decision trees developed by applying CHAID procedure within the same tenfold cross-validation framework than that was used during the examination of the previous hypothesis. The settings used in CHAID were also the same as in the second hypothesis. Having finished the ten model fitting processes,

the obtained model performances (hit rate) were averaged. The results are shown in Table 6.

Table 6. The hit rates of models developed in the empirical analysis

| Dataset           | Sample   | Group        | Input variables         |                   |                              |
|-------------------|----------|--------------|-------------------------|-------------------|------------------------------|
|                   |          |              | Static financial ratios | Dynamic variables | Static and dynamic variables |
| Original          | Training | Bankrupt     | 85.7%                   | 83.2%             | 84.6%                        |
|                   |          | Non-bankrupt | 80.1%                   | 77.8%             | 83.8%                        |
|                   |          | Total        | 82.9%                   | 80.5%             | 84.2%                        |
|                   | Testing  | Bankrupt     | 76.6%                   | 75.0%             | 76.8%                        |
|                   |          | Non-bankrupt | 71.0%                   | 71.6%             | 74.6%                        |
|                   |          | Total        | 73.8%                   | 73.3%             | 75.7%                        |
| After replacement | Training | Bankrupt     | 85.7%                   | 80.0%             | 86.9%                        |
|                   |          | Non-bankrupt | 80.1%                   | 81.8%             | 83.1%                        |
|                   |          | Total        | 82.9%                   | 80.9%             | 85.0%                        |
|                   | Testing  | Bankrupt     | 76.6%                   | 71.6%             | 80.2%                        |
|                   |          | Non-bankrupt | 71.0%                   | 72.4%             | 75.8%                        |
|                   |          | Total        | 73.8%                   | 72.0%             | 78.0%                        |

In the first column of Table 6, “Original” refers to the case when outlier values of the time series weren’t handled. The results documented in these rows are the same as in Table 3. The rows called “After replacement” show the performance of models developed with dynamic variables for which the outlier values were replaced by those which were the closest to them but not outliers.

In this case, one can find similar patterns as earlier in Table 3. Namely, dynamic variables are not recommended to use alone even when the outlier values had been handled (replaced) because the average performance of models developed on only dynamic variables were usually weaker than that of models based on solely static financial ratios. Incremental increase can be observable only in the case when the

two types of variables (static and dynamic) were used together in the models. The increase in the hit rates after replacement was higher than in the case without handling outliers: the average increment was 2.1 percentage points for the training sample and 4.2 percentage points for the testing sample. These increments were 1.3 percentage point and 1.9 percentage point respectively when outlier values weren't replaced.

Obtaining the results presented above, I also accepted the third hypothesis of the thesis because the predictive power of models containing dynamic variables after replacing outliers were higher than that of models containing dynamic variables without handling outliers.

#### **5.4. The summary of the results obtained in the empirical examinations**

Based on the presented results, it can be concluded that irrespective of handling outliers in the time series of the financial ratios, predictive power of bankruptcy models developed by using CHAID procedure can considerably be increased when the values expressing the financial status of the firm in the most recent year relative to the previous period (dynamic variables) are also incorporated into the models in addition to the static values of financial ratios. The empirical investigations revealed that it is reasonable to replace for the outlier values in the time series of the financial ratios by the values which are the closest to them but not outliers.

It is an interesting result that irrespective of handling outliers, it is true that the hit rate of bankrupt firms was higher than that of healthy companies. This suggests that, during the period of the recent recession in Hungary, it was easier to identify bankrupt firms than the healthy ones. A possible explanation for this can be found in

the table of the Appendix where it can be seen that firms, irrespective of their financial health, had negative profitability and working capital ratios on average during the examined period. Furthermore, based on the average values of dynamic variables, it can be seen that financial performance in the most recent year was usually weaker than it had been in the previous period even for the healthy companies. In such circumstances, it is less surprising that the models could identify the failed firms with greater confidence than the healthy ones since the operating companies may also have financial problems in such times of need like the first years of the recent crisis were.

## 6. SUMMARY

My thesis was started with presenting the motivation of selecting this topic, then I clarified the most important terms used in the dissertation. Following this, the Hungarian and the international literature of bankruptcy prediction were reviewed. I pointed out the most relevant research directions and open research questions which can be outlined based on the literature of which the topic of the dynamization of bankruptcy prediction models was chosen. This task was tried to be executed by applying a formula which compares the financial ratio of the most recent year to those of the previous period, allowing to take into account the dynamics of financial ratios in the framework of the widely used classification techniques without applying more sophisticated methodologies for this purpose.

The predictive power of the dynamic variables was assessed on the basis of a sample consisting of 1000 Hungarian enterprises. The data was manually collected by the author. Significant effort was made to mitigate the distorting impact of outlier values in the time series of financial ratios. Based on the results of the empirical investigations, my main research hypothesis was confirmed, namely, the predictive power of bankruptcy models can be enhanced not only by using more sophisticated classification algorithms. This is possible by exploiting the information content of the financial ratios to a greater extent. In the case of decision trees generated by using the CHAID algorithm, I experienced considerable increase in the predictive performance of the models when the dynamic variables proposed in the thesis were also incorporated into the models as input variables in addition to the static financial ratios commonly used in the literature of bankruptcy prediction.

The aims of the closing section of my work:

- to locate the thesis within the literature of bankruptcy prediction and to present the contribution of my research to the development of this science;
- to discuss the limitations of the empirical research presented in the dissertation;
- to outline possible future research directions opening on the basis of the results of my work.

### **6.1. The contribution of the thesis to the development of bankruptcy prediction**

The literature review was the most extensive part of the thesis. In spite of the vast amount of literature discussed in my work, it should be emphasized that the presented review cannot be considered as complete. The reason for it that the number of articles published in scientific journals has been increasing dramatically in the last five years. During reviewing this large literature, I aspired for giving as comprehensive picture as possible about the development, the current state and the most important research directions.

Similarly to other disciplines, it is adequate to chronologically review the literature of bankruptcy prediction as well. Almost all articles published on this topic follows this guiding principle. The common practice is to start with recalling the milestones of bankruptcy prediction, then studies related more closely to the examined topic are presented in more detail. I also applied this guiding principle but I organized the most important researches around the main steps of the bankruptcy model building process. This approach is based on the fact that current studies of this field usually formulate and examine research questions only about one (or some) steps of this

model building process. Because of this practice, it is not an overstatement to say that the main steps have their own (sub)literatures, however these literatures cannot be considered as independent of each other. In my work, I presented the main research questions, the most important empirical investigations and the corresponding results related to each steps. I named this approach as the cross-sectional review of the literature. With respect to the fact that I haven't found similar approach for presenting the literature yet, this kind of review can be mentioned as a contribution to the development of this research area.

Based on the studies which have been published so far, I identified several fruitful directions for future research. Among them, the static nature of bankruptcy model is one of the questions which is well-known in the literature but researchers have started dealing with this issue only in the last couple of years. To solve the problems arising from the static nature of models, mainly complex methodological solutions cited earlier in the thesis have been suggested in the literature, but despite the promising results, probably because of their complexity, these approaches haven't become common in bankruptcy prediction. Based on this, I tried to find an effective and simple way for incorporating the time trend of financial ratios into bankruptcy models. To perform this task, I proposed the use of dynamic variables which were formally defined in Section 3.

Based on the results of the empirical analyses, I concluded that the usefulness of dynamic variables is proven in the case of my research sample – containing data for 1000 Hungarian enterprises – described in Section 4 because the predictive power of CHAID decision trees considerably increased when the dynamic variables expressing the time trends of financial ratios were also used as independent variables in addition to the static financial data. During the empirical investigations, I pointed out that it is

reasonable to replace the outlier values in the time series of the financial ratios with those values which are the closest to them but not outliers when applying the proposed dynamic variables in bankruptcy models. The obtained results confirmed my main hypothesis, namely that the predictive performance of bankruptcy models can be enhanced not only by using the state-of-the-art classification algorithms as it is suggested by the mainstream research of bankruptcy prediction. There is a real possibility to improve the models in alternative ways for which I also found an example. Consequently, I locate my work within these alternative research directions. With respect to the fact that I haven't found similar attempt at dynamizing bankruptcy models, I consider the application of the proposed dynamic variables as an aspect of my thesis which can mostly contribute to the development of bankruptcy prediction.

In addition to their capability of improving the performance of bankruptcy prediction models, the dynamic variables may fill in some research gap related to the theoretical background of bankruptcy prediction as well. It was mentioned that bankruptcy prediction doesn't have such an exact theoretical framework as it is usual behind the economic models. The most cited theoretical construction was published by Beaver [1966] who interpreted the companies as reservoirs of liquid assets. In this context, bankruptcy is the case when the reservoir dries out. It could be an important aspect that the companies cannot be considered as completely the same reservoirs. However, the firms can be imagined as lakes of liquid assets which are different in size and in their geographical locations, so their desiccation depends on several factors. What is more important, due to these differences, the lakes cannot be compared directly. In this framework, the water line of the lake (as the financial position of the firm) can be judged more objectively if the lastly measured water line

of the lake is compared to the previous water lines which had been measured on this given lake in the past.

So, the application of the proposed dynamic variables may bring improvement in the field of bankruptcy prediction from theoretical point of view as well. However, it should be added that this concept doesn't constitute a general theoretical framework, it can be only a supplement to the model created by Beaver [1966]. It cannot determine which variables should be used in bankruptcy prediction models. Constructing such a model which is able to do this remains one of the biggest challenges for future research.

## **6.2. The limitations of the analysis presented in the thesis**

Though, the results of the empirical research are promising, the limitations of the presented analysis should also be highlighted.

The impact of the sampling manually conducted for the empirical analysis may be questionable on the presented results. As it have been mentioned earlier, restricting the sample on the basis of the size, age or activity of the firms was not possible during the data gathering process because search on the ground of these parameters was not possible in the publicly available databases serving as sources of data for the thesis. Consequently, the models developed in this work are based on a considerable heterogeneous dataset, so it is desirable to conduct similar analyses on more homogenous databases as well.

Though, the research dataset may be considered as a large one (it consists of data for more than 7500 firm-year observations), it can be said that the empirical analysis is based on a relatively small sample because the 1000 firms in the database can be

viewed as a small number compared to the total number of companies operating in Hungary. A similar investigation on a larger sample can be a good verification of the conclusions which were drawn on the basis of the results obtained by using the available dataset. It would be interesting to conduct a similar study on data for public companies from the USA or from the Far East which are frequently used in the literature of bankruptcy prediction because it is questionable whether the predictive power of the proposed dynamic variable can be considered as unique to the Hungarian firms or these type of variables are worth using for data from another countries as well.

The empirical research has another limitation coming from the sampling strategy applied in the presented data gathering process. As it has been mentioned earlier, only those firms could be included in the sample for which financial data were available for at least three consecutive years. Consequently, the concept of dynamic variables are not applicable to enterprises which are younger than 3 years. So, future solvency of this kind of companies can be predicted only on the basis of static financial ratios and other qualitative information.

Some words also have to be said about the applied classification method. In my work, only decision trees based on CHAID algorithm were applied. The choice was reasoned by the fact that this algorithm generates models fast and easily with relatively high classification performance. Furthermore, another important aspect, which played very crucial role in my decision on choosing this method, of the decision trees that their performance is not affected by the presence of outliers which, in my opinion, inherit essential information for predicting the event of bankruptcy because failing firms often have extreme high/low financial ratios as they approaching bankruptcy. However, this issue raises the following question: to what

extent could be the results attributable to the application of the CHAID classification method? To answer this question, the dynamic financial variables should also be applied in other classification procedures as input variables.

Finally, it has to be mentioned that my work is primarily an empirical attempt at demonstrating the usefulness of dynamic variables, so developing models which are ready-to-use in real life to rate companies was beyond the scope of my thesis. Consequently, numerous aspects of practical credit scoring were left out from the analysis such as calculating and calibrating the probability of default for each observations or the stability tests of the developed models. Another avenue for future research is to investigate whether the concept of dynamic variables presented in the thesis is able to support these practical issues as well.

### **6.3. Possible future research directions**

It is common in science that a research raises more question than it wished to answer. I also had to experience this fact in the case of my research.

As the title of the thesis indicates, the aim of my work was to examine the applicability of dynamic variables in bankruptcy prediction. The presented results suggest that this dynamic concept has its own *raison d'être* in assessing future solvency of firms. However, the calculation of the proposed variable is based mainly on intuition rather than sound scientific foundations. Consequently, the presented analysis have several aspects which need to be further explored in the future. These investigations can contribute to enhancing the predictive ability of the dynamic variables in bankruptcy prediction models. Some of these possible future research directions are highlighted in this section.

1. The dynamic variables compare the values of the most recent financial ratios to those from the entire previous period. However, it is questionable whether it is worth taking into account the complete time series of financial ratios in the calculation. It is possible that using a particular period, for example the time series of the last 5 years, as a basis for comparison would result in more predictive dynamic variables. The application of shorter time series can be useful from the side of the high computational burden related to data preprocessing because the shorter the time series, the lower the probability of occurring outlier values which are worth being handled in the way as it was proposed in the thesis earlier.

2. In my work, the distorting impact of outliers on the dynamic variables was presented on the basis of a real life example. The results of the empirical analyses showed that it is worth handling such extreme values in order to enhance the predictive power of the dynamic variables. However, the identification of outliers was carried out by using my own intuition in the absence of objective definition. It can constitute a distinct research topic how to define outliers in order to maximize the predictive performance of bankruptcy prediction models containing the presented dynamic variables.

3. The concept of dynamic variables is based on the principle of relativity which means that the values of financial ratios are not absolute criteria, these values can objectively be judged only in the mirror of a suitable benchmark (Virág et al. [2013]). Among the possible benchmarks, the mean value of the industry is often mentioned. The concept of the dynamic variables could open new research directions in this field as well because, knowing the mean value of the industry, it is possible to analyze the position of the firms compared to the tendencies observable in their industries. The idea of dynamic variables can also be applied to the industry relative

ratios (Platt-Platt [1990]) discussed in Section 2.3.2.2. The resulting dynamic industry relative variables could further enhance the predictive ability of bankruptcy prediction models.

## APPENDIX

Table F1. Descriptive statistics of financial ratios used in the empirical analyses

| Variable name                   | Group    | Mean     | Standard deviation |
|---------------------------------|----------|----------|--------------------|
| Liquidity ratio I               | Healthy  | 39.82    | 381.86             |
|                                 | Bankrupt | 16.38    | 224.52             |
| D_Liquidity ratio I             | Healthy  | 1.38     | 9.11               |
|                                 | Bankrupt | 22.70    | 348.44             |
| Liquidity ratio II              | Healthy  | 38.96    | 381.77             |
|                                 | Bankrupt | 16.09    | 224.47             |
| D_Liquidity ratio II            | Healthy  | 1.40     | 9.17               |
|                                 | Bankrupt | 19.21    | 312.11             |
| Ratio of cash*                  | Healthy  | 0.32     | 0.33               |
|                                 | Bankrupt | 0.15     | 0.27               |
| D_Ratio of cash                 | Healthy  | 1.00     | 8.13               |
|                                 | Bankrupt | 2.63     | 25.54              |
| Cash flow/Liabilities           | Healthy  | 4.31     | 139.08             |
|                                 | Bankrupt | -0.92    | 11.74              |
| D_Cash flow/Liabilities*        | Healthy  | 0.83     | 14.19              |
|                                 | Bankrupt | -2.28    | 25.69              |
| Cash flow/Current Liabilities   | Healthy  | 4.46     | 139.14             |
|                                 | Bankrupt | -1.28    | 12.70              |
| D_Cash flow/Current Liabilities | Healthy  | 0.73     | 14.25              |
|                                 | Bankrupt | -0.69    | 8.54               |
| Equity financing ratio          | Healthy  | 2.65     | 24.46              |
|                                 | Bankrupt | 105.98   | 2359.90            |
| D_Equity financing ratio        | Healthy  | 58.55    | 1259.06            |
|                                 | Bankrupt | -2.92    | 76.36              |
| Assets turnover                 | Healthy  | 4.76     | 28.33              |
|                                 | Bankrupt | 14.49    | 171.90             |
| D_Assets turnover               | Healthy  | 0.90     | 4.99               |
|                                 | Bankrupt | 9.13     | 109.87             |
| Inventory turnover              | Healthy  | 65334.50 | 688444.08          |
|                                 | Bankrupt | 18316.01 | 72091.13           |
| D_Inventory turnover *          | Healthy  | 54.23    | 685.75             |
|                                 | Bankrupt | 693.64   | 5675.54            |
| Receivables turnover*           | Healthy  | 112.44   | 1184.25            |
|                                 | Bankrupt | 2023.27  | 17177.09           |
| D_Receivables turnover (time)*  | Healthy  | 94.90    | 1361.02            |
|                                 | Bankrupt | 2827.33  | 20538.81           |

|                                   |          |           |           |
|-----------------------------------|----------|-----------|-----------|
| Indebtedness*                     | Healthy  | 17.64     | 244.48    |
|                                   | Bankrupt | 153.17    | 1206.91   |
| D_Indebtedness *                  | Healthy  | 1.36      | 11.79     |
|                                   | Bankrupt | 263.86    | 2800.30   |
| Equity ratio*                     | Healthy  | -16.83    | 244.53    |
|                                   | Bankrupt | -164.61   | 1230.10   |
| D_Equity ratio *                  | Healthy  | -0.35     | 11.77     |
|                                   | Bankrupt | -298.95   | 3120.05   |
| Creditworthiness                  | Healthy  | 7.51      | 84.73     |
|                                   | Bankrupt | 212.76    | 4758.64   |
| D_Creditworthiness                | Healthy  | 13.59     | 267.85    |
|                                   | Bankrupt | 0.02      | 14.98     |
| Return on sales                   | Healthy  | -323.96   | 7036.36   |
|                                   | Bankrupt | -1549.67  | 12425.06  |
| D_Return on sales                 | Healthy  | -48.92    | 660.10    |
|                                   | Bankrupt | -20988.26 | 358038.57 |
| Return on assets*                 | Healthy  | -3.17     | 63.44     |
|                                   | Bankrupt | -92.61    | 923.00    |
| D_Return on assets *              | Healthy  | -0.31     | 8.16      |
|                                   | Bankrupt | -159.18   | 1353.89   |
| Receivables/Current liabilities   | Healthy  | 21.35     | 249.91    |
|                                   | Bankrupt | 7.47      | 134.47    |
| D_Receivables/Current liabilities | Healthy  | 1.45      | 9.79      |
|                                   | Bankrupt | 1.54      | 28.50     |
| Ratio of working capital*         | Healthy  | -14.76    | 241.86    |
|                                   | Bankrupt | -134.64   | 1130.71   |
| D_Ratio of working capital        | Healthy  | -0.10     | 10.92     |
|                                   | Bankrupt | -208.37   | 2596.76   |
| Size                              | Healthy  | 9.65      | 2.59      |
|                                   | Bankrupt | 9.61      | 2.86      |
| D_Size                            | Healthy  | 0.47      | 1.03      |
|                                   | Bankrupt | -5.23     | 91.05     |
| Years*                            | Healthy  | 8.39      | 2.89      |
|                                   | Bankrupt | 6.80      | 2.82      |

## REFERENCES

- Abdou, H. – Pointon, J. [2011]: Credit scoring, statistical techniques and evaluation criteria: a review of the literature. *Intelligent Systems in Accounting, Finance and Management*. Vol. 18., p. 59-88., DOI: <http://dx.doi.org/10.1002/isaf.325>.
- Abellán, J. – Mantas, C. J. [2014]: Improving experimental studies about ensembles of classifiers for bankruptcy prediction and credit scoring. *Expert Systems with Applications*. Vol. 41., p. 3825-3830, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.12.003>.
- Agarwal, V. – Taffler, R. [2008]: Comparing the performance of market-based and accounting-based bankruptcy prediction models. *Journal of Banking & Finance*. Vol. 32., p. 1541-1551, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2007.07.014>.
- Agarwal, V. – Taffler, R. J. [2007]: Twenty-five years of the Taffler z-score model: does it really have predictive ability? *Accounting and Business Research*. Vol. 37., No. 4., p. 285-300, DOI: <http://dx.doi.org/10.1080/00014788.2007.9663313>.
- Ahn, B. S. – Cho, S. S. – Kim, C. Y. [2000]: The integrated methodology of rough set theory and artificial neural networks for business failure prediction. *Expert Systems with Applications*. Vol. 18., No. 2., p. 65-74, DOI: [http://dx.doi.org/10.1016/S0957-4174\(99\)00053-6](http://dx.doi.org/10.1016/S0957-4174(99)00053-6).
- Akkoc, S. [2012]: An empirical comparison of conventional techniques, neural networks and the three stage hybrid adaptive neuro fuzzy inference system (ANFIS) model for credit scoring analysis: The case of Turkish credit card data. *European Journal of Operational Research*. Vol. 222., p. 168-178, DOI: <http://dx.doi.org/10.1016/j.ejor.2012.04.009>.
- Alfaro, E. – García, N. – Gámez, M. – Elizondo, D. [2008]: Bankruptcy forecasting: An empirical comparison of Adaboost and neural networks. *Decision Support Systems*. Vol. 45., p. 110-122, DOI: <http://dx.doi.org/10.1016/j.dss.2007.12.002>.
- Alici, Y. [1995]: Neural networks in corporate failure prediction: The UK experience. In: Refenes, A. N. – Abu-Mostafa, Y. – Moody, J. – Weigend, A. (eds.): *Processing Third International Conference of Neural Networks in the Capital Markets*, London, October 1995, p. 393-406.
- Altman, E. I. – Marco, G. – Varetto, F. [1994]: Corporate distress diagnosis: Comparisons using linear discriminant analysis and neural networks (the Italian experience). *Journal of Banking and Finance*. Vol. 18., p. 505-529, DOI: [http://dx.doi.org/10.1016/0378-4266\(94\)90007-8](http://dx.doi.org/10.1016/0378-4266(94)90007-8).
- Altman, E. I. [1968]: Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *The Journal of Finance*. Vol. 23., No. 4., p. 589-609, DOI: <http://dx.doi.org/0.1111/j.1540-6261.1968.tb00843.x>.

Altman, E. I. [1977]: ZETA analysis. A new model to identify bankruptcy risk of corporations. *Journal of Banking & Finance*. Vol. 1., p. 29-54, DOI: [http://dx.doi.org/10.1016/0378-4266\(77\)90017-6](http://dx.doi.org/10.1016/0378-4266(77)90017-6).

Altman, E. I. [2002]: *Corporate distress prediction models in a turbulent economic and BASEL II environment*. Salomon Center for the Study of Financial Institutions <New York, NY> / Credit & Debt Markets Research Program.

Anandarajan, M. – Lee, P. – Anandarajan, A. [2001]: Bankruptcy prediction of financially stressed firms: An examination of the predictive accuracy of artificial neural networks. *Intelligent Systems in Accounting, Finance and Management*. Vol. 10., p. 69-81, DOI: <http://dx.doi.org/10.1002/isaf.199>.

Atiya, A. F. [2001]: Bankruptcy Prediction for Credit Risk Using Neural Networks: A Survey and New Results, *IEEE Transactions on Neural Networks*, Vol. 12, No. 4, p. 929-935.

Back, B. – Laitinen, T. – Sere, K. – van Wezel, M. [1996]: *Choosing bankruptcy predictors using discriminant analysis, logit analysis, and genetic algorithms*. Technical Report No. 40. Turku Centre for Computer Science, Turku.

Bae, J. K. [2012]: Predicting financial distress of the South Korean manufacturing industries. *Expert Systems with Applications*. Vol. 39., p. 9159-9165, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.058>.

Baesens, B. – Setiono, R. – Mues, C. – Vanthienen, J. [2003]: Using neural network rule extraction and decision tables for credit-risk evaluation. *Management Science*. Vol. 49., No. 3., p. 312-329, DOI: <http://dx.doi.org/10.1287/mnsc.49.3.312.12739>.

Balcean, S. – Ooghe, H. [2006]: 35 years of studies on business failure: an overview of the classic statistical methodologies and their related problems. *The British Accounting Review*. Vol. 38., p. 63-93., DOI: <http://dx.doi.org/10.1016/j.bar.2005.09.001>.

Barniv, R. – Agarwal, A. – Leach, R. [2002]: Predicting bankruptcy resolution. *Journal of Business Finance & Accounting*. Vol. 29., No. 3-4., p. 497-520, DOI: <http://dx.doi.org/10.1111/1468-5957.00440>.

Beaver, W. H. – McNichols, M. F. – Rhie, J. W. [2005]: Have financial statements become less informative? Evidence from the ability of financial ratios to predict bankruptcy. *Review of Accounting Studies*. Vol. 10., p. 93-122, DOI: <http://dx.doi.org/10.1007/s11142-004-6341-9>.

Beaver, W. H. [1966]: Financial ratios as predictors of failure. *Journal of Accounting Research*. Vol. 4., Supplement, p. 71-111, DOI: <http://dx.doi.org/10.2307/2490171>.

Becchetti, L. – Sierra, J. [2003]: Bankruptcy risk and productive efficiency in manufacturing firms. *Journal of Banking & Finance*. Vol. 27., p. 2099-2120, DOI: [http://dx.doi.org/10.1016/S0378-4266\(02\)00319-9](http://dx.doi.org/10.1016/S0378-4266(02)00319-9).

Bellotti, T. – Matousek, R. – Stewart, C. [2011]: A note comparing support vector machines and ordered choice models' predictions of international banks' ratings. *Decision Support Systems*. Vol. 51., p. 682-687, DOI: <http://dx.doi.org/10.1016/j.dss.2011.03.008>.

Berg, D. [2007]: Bankruptcy prediction by generalized additive models. *Applied Stochastic Models in Business and Industry*. Vol. 23., p. 129-143, DOI: <http://dx.doi.org/10.1002/asmb.658>.

Bernhardsen, E. [2001]: *A model of bankruptcy prediction*. Working Paper, Financial Analysis and Structure Department, Research Department, Norges Bank, Oslo. <http://www.norges-bank.no/Upload/import/publikasjoner/arbeidsnotater/pdf/arb-2001-10.pdf>

Beynon, M. J. – Peel, M. J. [2001]: Variable precision rough set theory and data discretisation: an application to corporate failure prediction. *Omega*. Vol. 29., p. 561-576, DOI: [http://dx.doi.org/10.1016/S0305-0483\(01\)00045-7](http://dx.doi.org/10.1016/S0305-0483(01)00045-7).

Bharath, S. T. – Shumway, T. [2008]: Forecasting default with the Merton distance to default model. *The Review of Financial Studies*. Vol. 21., No. 3., p. 1339-1369, DOI: <http://dx.doi.org/10.1093/rfs/hhn044>.

Bioch, J. C. – Popova, V. [2001]: *Bankruptcy prediction with rough sets*. No. ERS-2001-11-LIS). Erasmus Research Institute of Management (ERIM).

Blanco, A. – Pino-Mejías, R. – Lara, J. – Rayo, S. [2013]: Credit scoring models for the microfinance industry using neural networks: Evidence from Peru. *Expert Systems with Applications*. Vol. 40., p. 356-364, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.07.051>.

Bonfim, D. [2009]: Credit risk drivers: Evaluating the contribution of firm level information and of macroeconomic dynamics. *Journal of Banking & Finance*. Vol. 33., p. 281-299, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2008.08.006>.

Bongini, P. – Ferri, G. – Hahm, H. [2000]: Corporate bankruptcy in Korea: Only the strong survive? *The Financial Review*. Vol. 35., No. 4., p. 31-50, DOI: <http://dx.doi.org/10.1111/j.1540-6288.2000.tb01428.x>.

Bose, I. [2006]: Deciding the financial health of dot-coms using rough sets. *Information & Management*. Vol. 43., p. 835-846, DOI: <http://dx.doi.org/10.1016/j.im.2006.08.001>.

Boyacioglu, M. A. – Kara, Y. – Baykan, Ö. K. [2009]: Predicting bank failures using neural networks, support vector machines and multivariate statistical methods: A comparative analysis in the sample of savings deposit insurance fund (SFID) transferred banks in Turkey. *Expert Systems with Applications*. Vol. 36., p. 3355-3366, DOI: <http://dx.doi.org/10.1016/j.eswa.2008.01.003>.

Brealey, R. A. – Myers, S. C. [2005]: *Modern Vállalati Pénzügyek*. Panem, Budapest.

Brezigar-Masten, A. – Masten, I. [2012]: CART-based selection of bankruptcy predictors for the logit model. *Expert Systems with Applications*. Vol. 39., p. 10153-10159, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.125>.

Brown, I. – Mues, C. [2012]: An experimental comparison of classification algorithms for imbalanced credit scoring data sets. *Expert Systems with Applications*. Vol. 39., p. 3446-3453, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.09.033>.

Cantner, U. – Dressler, K. – Krüger, J. J. [2006]: Firm survival in the German automobile industry. *Empirica*. Vol. 33., p. 49-60, DOI: <http://dx.doi.org/10.1007/s10663-006-9006-z>.

Cao, Y. [2012a]: MCELCCh-FDP: Financial distress prediction with classifier ensembles based on firm life cycle and Choquet integral. *Expert Systems with Applications*. Vol. 39., p. 7041-7049, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.01.043>.

Cao, Y. [2012b]: Aggregating multiple classification results using Choquet integral for financial distress early warning. *Expert Systems with Applications*. Vol. 39., p. 1830-1836, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.08.067>.

Carling, K. – Jacobson, T. – Lindé, J. – Roszbach, K. [2007]: Corporate credit risk modeling and the macroeconomy. *Journal of Banking & Finance*. Vol. 31., p. 845-868, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2006.06.012>.

Chava, S. – Jarrow, R. A. [2004]: Bankruptcy prediction with industry effects. *Review of Finance*. Vol. 8., No. 4., p. 537-569, DOI: <http://dx.doi.org/10.1093/rof/8.4.537>.

Chen, J. H. [2012]: Developing SFNN models to predict financial distress of construction companies. *Expert Systems with Applications*. Vol. 39., p. 823-827, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.07.080>.

Chen, M. Y. [2011]: Predicting corporate financial distress based on integration of decision tree classification and logistic regression. *Expert Systems with Applications*. Vol. 38., p. 11261-11272, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.02.173>.

Chen, N. – Riberio, B. – Vieira, A. – Chen, A. [2013]: Clustering and visualization of bankruptcy trajectory using self-organizing map. *Expert Systems with Applications*. Vol. 40., p. 385-393, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.07.047>.

Chen, N. – Riberio, B. – Vieira, A. S. – Duarte, J. – Neves, J. C. [2011]: A genetic algorithm-based approach to cost-sensitive bankruptcy prediction. *Expert Systems with Applications*. Vol. 38., p. 12939-12945, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.04.090>.

Chi, B. W. – Hsu, C. C. [2012]: A hybrid approach to integrate genetic algorithm into dual scoring model in enhancing the performance of credit scoring model. *Expert Systems with Applications*. Vol. 39., p. 2650-2661, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.08.120>.

Coats, P. – Fant, L. [1993]: Recognizing financial distress patterns using a neural network tool. *Financial Management*. Vol. 22., No. 3., p. 142-155.

Cubiles-De-La-Vega, M. D. – Blanco-Oliver, A. – Pino-Mejías, R. – Lara-Rubio, J. [2013]: Improving the management of microfinance institutions by using credit scoring models based on statistical learning techniques. *Expert Systems with Applications*. Vol. 40., No. 17., p. 6910-6917, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.06.031>.

Dambolena, I. G. – Khoury, S. J. [1980]: Ratio stability and corporate failure. *The Journal of Finance*. Vol. 35., No. 4., p. 1017-1026, DOI: <http://dx.doi.org/10.1111/j.1540-6261.1980.tb03517.x>.

Darayseh, M. – Waples, E. – Tsoukalas, D. [2003]: Corporate failure for manufacturing industries using firms specifics and economic environment with logit analysis. *Managerial Finance*. Vol. 28., No. 8., p. 23-36, DOI: <http://dx.doi.org/10.1108/03074350310768409>.

Das, S. R. – Hanouna, P. – Sarin, A. [2009]: Accounting-based versus market-based cross-sectional models of CDS spreads. *Journal of Banking & Finance*. Vol. 33., p. 719-730, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2008.11.003>.

Deakin, E. B. [1972]: A discriminant analysis of predictors of failure. *Journal of Accounting Research*. Vol. 10., No. 1., p. 167-179, DOI: <http://dx.doi.org/10.2307/2490225>.

Demers, E. – Joos, P. [2007]: IPO failure risk. *Journal of Accounting Research*. Vol. 45., No. 2., p. 333-371, DOI: <http://dx.doi.org/10.1111/j.1475-679X.2007.00236.x>.

Dewaelheyns, N. – Van Hulle, C. [2006]: Corporate failure prediction modeling: Distorted by business groups' internal capital markets? *Journal of Business Finance & Accounting*. Vol. 33., Vol. 5-6., p. 909-931, DOI: <http://dx.doi.org/10.1111/j.1468-5957.2006.00009.x>.

Ding, Y. – Song, X. – Zen, Y. [2008]: Forecasting financial condition of Chinese listed companies based on support vector machines. *Expert Systems with Applications*. Vol. 34., p. 3081-3089, DOI: <http://dx.doi.org/10.1016/j.eswa.2007.06.037>.

Doumplos, M. – Gaganis, C. – Pasiouras, F. [2005]: Explaining qualifications in audit reports using a support vector machine methodology. *Intelligent Systems in Accounting, Finance and Management*. Vol. 13., p. 197-215, DOI: <http://dx.doi.org/10.1002/isaf.268>.

Du Jardin, P. – Séverin, E. [2012]: Forecasting financial failure using a Kohonen map: A comparative study to improve model stability over time. *European Journal of Operational Research*. Vol. 221., p. 378-396, DOI: <http://dx.doi.org/10.1016/j.ejor.2012.04.006>.

Du Jardin, P. [2010]: Predicting bankruptcy using neural networks and other classification methods: The influence of variable selection techniques on model accuracy. *Neurocomputing*. Vol. 73., p. 2047-2060, DOI: <http://dx.doi.org/10.1016/j.neucom.2009.11.034>.

Duan, J. C. – Sun, J. – Wang, T. [2012]: Multiperiod corporate default prediction – A forward intensity approach. *Journal of Econometrics*. Vol. 170., p. 191-209, DOI: <http://dx.doi.org/10.1016/j.jeconom.2012.05.002>.

Duffie, D. – Saita, L. – Wang, K. [2007]: Multi-period corporate default prediction with stochastic covariates. *Journal of Financial Economics*. Vol. 83., p. 635-665, DOI: <http://dx.doi.org/10.1016/j.jfineco.2005.10.011>.

Elliott, R. J. – Siu, T. K. – Fung, E. S. [2014]: A double HMM approach to Altman Z-scores and credit ratings. *Expert Systems with Applications*. Vol. 41., p. 1553-1560, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.08.052>.

Erdal, H. I. – Ekinici, A. [2012]: A comparison of various artificial intelligence methods in the prediction of bank failures. *Computational Economics*. Vol. 42., No. 2., p. 199-215, DOI: <http://dx.doi.org/10.1007/s10614-012-9332-0>.

Fan, A. – Palaniswami, M. [2000]: A new approach to corporate loan default prediction from financial statements. In: *Proceedings of Computational Finance/Forecasting Financial Markets Conference CF/FFM-2000*, London.

Fedorova, E. – Gilenko, E. – Dovzhenko, S. [2013]: Bankruptcy prediction for Russian companies: Application of combined classifiers. *Expert Systems with Applications*. Vol. 40., p. 7285-7293, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.07.032>.

Feki, A. – Ishak, A. B. – Feki, S. [2012]: Feature selection using Bayesian and multiclass support vector machines approaches: Application to bank risk prediction. *Expert Systems with Applications*. Vol. 39., p. 3087-3099, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.08.172>.

Foreman, R. D. [2003]: A logistic analysis of bankruptcy within the US local telecommunications industry. *Journal of Economics and Business*. Vol. 55., p. 135-166, DOI: [http://dx.doi.org/10.1016/S0148-6195\(02\)00133-9](http://dx.doi.org/10.1016/S0148-6195(02)00133-9).

Frydman, H. – Altman, E. I. – Kao, D. L. [1985]: Introducing recursive partitioning for financial classification: The case of financial distress. *The Journal of Finance*. Vol. 40., No. 1., p. 269-291, DOI: <http://dx.doi.org/10.1111/j.1540-6261.1985.tb04949.x>.

Gaganis, C. – Pasiouras, F. – Spathis, C. – Zopounidis, C. [2007]: A comparison of nearest neighbours, discriminant analysis and logit models for auditing decisions. *Intelligent Systems in Accounting, Finance and Management*. Vol. 15., p. 23-40, DOI: <http://dx.doi.org/10.1002/isaf.283>.

Gaganis, C. [2009]: Classification techniques for the identification of falsified financial statements: A comparative analysis. *Intelligent Systems in Accounting, Finance and Management*. Vol. 16., p. 207-229, DOI: <http://dx.doi.org/10.1002/isaf.303>.

García, V. – Marqués, A. I. – Sánchez, J. S. [2012]: On the use of data filtering techniques for credit risk prediction with instance-based models. *Expert Systems with Applications*. Vol. 39., p. 13267-13276, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.05.075>.

Gersbach, H. – Lipponer, A. [2003]: Firm defaults and the correlation effect. *European Financial Management*. Vol. 9., No. 3., p. 361-377, DOI: <http://dx.doi.org/10.1111/1468-036X.00225>.

Gilbert, L. R. – Menon, K. – Schwarz, K. B. [1990]: Predicting bankruptcy for firms in financial distress. *Journal of Business Finance & Accounting*. Vol. 17., No. 1., p. 161-171, DOI: <http://dx.doi.org/10.1111/j.1468-5957.1990.tb00555.x>.

Glennon, D. – Nigro, P. [2005]: Measuring the default risk of small business loans: A survival analysis approach. *Journal of Money, Credit and Banking*. Vol. 37., No. 5., p. 923-947.

Grunert, J. – Norden, L. – Weber, M. [2005]: The role of none-financial factors in internal credit ratings. *Journal of Banking & Finance*. Vol. 29., p. 509-531, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2004.05.017>.

Gruszczynski, M. [2004]: Financial distress of companies in Poland. *International Advances in Economic Research*. Vol. 10., No. 4., p. 249-256, DOI: <http://dx.doi.org/10.1007/BF02295137>.

Hájek, P. [2011]: Municipal credit rating modeling by neural networks. *Decision Support Systems*. Vol. 51., p. 108-118, DOI: <http://dx.doi.org/10.1016/j.dss.2010.11.033>.

Harada, K. – Ito, T. – Takahashi, S. [2013]: Is the distance to default a good measure in predicting bank failures? A case study of Japanese major banks. *Japan and the World Economy*. Vol. 27., p. 70-82, DOI: <http://dx.doi.org/10.1016/j.japwor.2013.03.007>.

Harada, N. – Kageyama, N. [2011]: Bankruptcy dynamics in Japan. *Japan and the World Economy*. Vol. 23., p. 119-128, DOI: <http://dx.doi.org/10.1016/j.japwor.2011.01.002>.

Hardle, W. – Lee, Y. J. – Schafer, D. – Yeh, Y. R. [2009]: Variable selection and oversampling in the use of smooth support vector machines for predicting the default risk of companies. *Journal of Forecasting*. Vol. 28., p. 512-534, DOI: <http://dx.doi.org/10.1002/for.1109>.

Harris, T. [2013]: Quantitative credit risk assessment using support vector machines: Broad versus narrow default definitions. *Expert Systems with Applications*. Vol. 40., p. 4404-4413, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.01.044>.

Heiss, F. – Köke, J. [2004]: Dynamics in ownership and firm survival: Evidence from corporate Germany. *European Financial Management*. Vol. 10., No. 1., p. 167-195. DOI: <http://dx.doi.org/10.1111/j.1468-036X.2004.00244.x>.

Hillegeist, S. A. – Keating, E. K. – Cram, D. P. – Lundstedt, K. G. [2004]: Assessing the probability of bankruptcy. *Review of Accounting Studies*. Vol. 9., p. 5-34., DOI: <http://dx.doi.org/10.2139/ssrn.307479>.

Horta, L. M. – Camanho, A. S. [2013]: Company failure prediction in the construction industry. *Expert Systems with Applications*. Vol. 40., p. 6253-6257, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.05.045>.

Hsieh, T. J. – Hsiao, H. F. – Yeh, W. C. [2012]: Mining financial distress trend data using penalty guided support vector machines based on hybrid of particle swarm optimization and artificial bee colony algorithm. *Neurocomputing*. Vol. 82., p. 196-206, DOI: <http://dx.doi.org/10.1016/j.neucom.2011.11.020>.

Hu, Y. C. [2009]: Bankruptcy prediction using ELECTRE-based single-layer perceptron. *Neurocomputing*. Vol. 72., p. 3150-3157, DOI: <http://dx.doi.org/10.1016/j.neucom.2009.03.002>.

Hu, Y-C. – Tseng, F-M. [2007]: Functional-link net with fuzzy integral for bankruptcy prediction. *Neurocomputing*, Vol. 70, p. 2959-2968, DOI: <http://dx.doi.org/10.1016/j.neucom.2006.10.111>.

Huang, S. C. – Tang, Y. C. – Lee, C. W. – Chang, M. J. [2012]: Kernel local Fisher discriminant analysis based manifold-regularized SVM model for financial distress prediction. *Expert Systems with Applications*. Vol. 39., p. 3855-3861, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.09.095>.

Huang, Z. – Chen, H. – Hsu, C. H. – Chen, W. H. – Wu, S. [2004]: Credit rating analysis with support vector machines and neural networks: a market comparative study. *Decision Support Systems*. Vol. 37., p. 543-558, DOI: [http://dx.doi.org/10.1016/S0167-9236\(03\)00086-1](http://dx.doi.org/10.1016/S0167-9236(03)00086-1).

Huynh, K. P. – Petrunia, R. J. – Voia, M. [2010]: The impact of initial financial state on firm duration across entry cohorts. *The Journal of Industrial Economics*. Vol. 58., No. 3., p. 661-689, DOI: <http://dx.doi.org/10.1111/j.1467-6451.2010.00429.x>.

Ioannidis, C. – Pasiouras, F. – Zopounidis, C. [2010]: Assessing bank soundness with classification techniques. *Omega*. Vol. 38., p. 345-357, DOI: <http://dx.doi.org/10.1016/j.omega.2009.10.009>.

Jagtiani, J. – Kolari, J. – Lemieux, C. – Shin, H. [2003]: Early warning models for bank supervision: Simpler could be better. *Federal Reserve Bank of Chicago Economic Perspectives*. Vol. 27., p. 49-60.

Jeong, C. – Min, J. H. – Kim, M. S. [2012]: A tuning method for the architecture of neural network models incorporating GAM and GA as applied to bankruptcy prediction. *Expert Systemes with Applications*. Vol. 39., p. 3650-3658, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.09.056>.

Jones, F. L. [1987]: Current techniques in bankruptcy prediction. *Journal of Accounting Literature*. Vol. 6., p. 131-164.

Jones, S. – Hensher, D. A. [2004]: Predicting firm financial distress: A mixed logit model. *The Accounting Review*. Vol. 79., No. 4., p. 1011-1038, DOI: <http://dx.doi.org/10.2308/accr.2004.79.4.1011>.

Jones, S. – Hensher, D. A. [2007]: Modelling corporate failure: A multinomial nested logit analysis for unordered outcomes. *The British Accounting Review*. Vol. 39., p. 89-107, DOI: <http://dx.doi.org/10.1016/j.bar.2006.12.003>.

Kennedy, K. – Mac Namee, B. – Delany, S. J. [2013]: Using semi-supervised classifiers for credit scoring. *Journal of the Operational Research Society*. Vol. 64., p. 513-529, DOI: <http://dx.doi.org/10.1057/jors.2011.30>.

Kiang, M. Y. [2003]: A comparative assessment of classification methods. *Decision Support Systems*. Vol. 35., p. 441-454, DOI: [http://dx.doi.org/10.1016/S0167-9236\(02\)00110-0](http://dx.doi.org/10.1016/S0167-9236(02)00110-0).

Kim, G. – Wu, C. H. – Lim, S. – Kim, J. [2012]: Modified matrix splitting method for the support vector machine and its application to the credit classification of companies in Korea. *Expert Systems with Applications*. Vol. 39., p. 8824-8834, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.007>.

Kim, H. S. – Sohn, S. Y. [2010]: Support vector machines for default prediction of SMEs based on technology credit. *European Journal of Operational Research*. Vol. 201., p. 838-846, DOI: <http://dx.doi.org/10.1016/j.ejor.2009.03.036>.

Kim, M. J. – Kang, D. K. [2012]: Classifiers selection in ensembles using genetic algorithms for bankruptcy prediction. *Expert Systems with Applications*. Vol. 39., p. 9308-9314, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.072>.

Kolari, J. – Glennon, D. – Shin, H. – Caputo, M. [2002]: Predicting large US commercial bank failures. *Journal of Economics and Business*. Vol. 54., No. 4., p. 361-387, DOI: [http://dx.doi.org/10.1016/S0148-6195\(02\)00089-9](http://dx.doi.org/10.1016/S0148-6195(02)00089-9).

Koopman, S. J. – Lucas, A. [2005]: Business and default cycles for credit risk. *Journal of Applied Econometrics*. Vol. 20., p. 311-323, DOI: <http://dx.doi.org/10.1002/jae.833>.

Koyuncugil, A. S. – Ozgulbas, N. [2012]: Financial early warning system model and data mining application for risk detection. *Expert Systems with Applications*. Vol. 39., p. 6238-6253, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.12.021>.

Kristóf, T. [2008]: *Gazdasági szervezetek fennmaradásának és fizetőképességének előrejelzése*. Ph. D. thesis, Corvinus University of Budapest, Budapest.

Kristóf, T. – Virág, M. [2012]: Data reduction and univariate splitting – Do they together provide better corporate bankruptcy prediction? *Acta Oeconomica*. Vol. 62., No. 2., p. 205-227, DOI: <http://dx.doi.org/10.1556/AOecon.62.2012.2.4>.

Laitinen, E. K. – Laitinen, T. [2000]: Bankruptcy prediction. Application of the Taylor's expansion in logistic regression. *International Review of Financial Analysis*. Vol. 9., No. 4., p. 327-349, DOI: [http://dx.doi.org/10.1016/S1057-5219\(00\)00039-9](http://dx.doi.org/10.1016/S1057-5219(00)00039-9).

Lee, M. C. – To, C. [2010]: Comparison of support vector machine and back propagation neural network in evaluating the enterprise financial distress. *International Journal of Artificial Intelligence & Applications*. Vol. 1., No. 3., p. 31-43, DOI: <http://dx.doi.org/10.5121/ijaia.2010.1303>.

Lee, S. – Choi, W. S. [2013]: A multi-industry bankruptcy prediction model using back-propagation neural network and multivariate discriminant analysis. *Expert Systems with Applications*. Vol. 40., p. 2941-2946, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.12.009>.

Lensberg, T. – Eilifsen, A. – McKee, T. E. [2006]: Bankruptcy theory development and classification via genetic programming. *European Journal of Operational Research*. Vol. 169., p. 677-697, DOI: <http://dx.doi.org/10.1016/j.ejor.2004.06.013>.

Leshno, M. – Spector, Y. [1996]: Neural network prediction analysis: The bankruptcy case, *Neurocomputing*, Vol. 10, No. 2., p. 125-147, DOI: [http://dx.doi.org/10.1016/0925-2312\(94\)00060-3](http://dx.doi.org/10.1016/0925-2312(94)00060-3).

Li, H. – Sun, J. [2011a]: Principal component case-based reasoning ensemble for business failure prediction. *Information & Management*. Vol. 48., p. 220-227, DOI: <http://dx.doi.org/10.1016/j.im.2011.05.001>.

Li, H. – Sun, J. [2011b]: Predicting business failure using support vector machines with straightforward wrapper: A re-sampling study. *Expert Systems with Applications*. Vol. 38., p. 12747-12756, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.04.064>.

Liang, D. – Tsai, C. F. – Wu, H. T. [2015]: The effect of feature selection on financial distress prediction. *Knowledge-Based Systems*. Vol. 73., p. 289-297, DOI: <http://dx.doi.org/10.1016/j.knosys.2014.10.010>.

- Lin, F. – Liang, D. – Chen, E. [2011]: Financial ratio selection for business crisis prediction. *Expert Systems with Applications*. Vol. 38., p. 15094-15102, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.05.035>.
- Lin, F. – Liang, D. – Yeh, C. C. – Huang, J. C. [2014]: Novel feature selection methods to financial distress prediction. *Expert Systems with Applications*. Vol. 41., p. 2472-2483, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.09.047>.
- Lin, R. H. – Wang, Y. T. – Wu, C. H. – Chuang, C. L. [2009]: Developing a business failure prediction model via RST, GRA and CBR. *Expert Systems with Applications*. Vol. 36., p. 1593-1600, DOI: <http://dx.doi.org/10.1016/j.eswa.2007.11.068>.
- Lin, T. H. [2009]: A cross model study of corporate financial distress prediction in Taiwan: Multiple discriminant analysis, logit, probit and neural networks models. *Neurocomputing*. Vol. 72., p. 3507-3516, DOI: <http://dx.doi.org/10.1016/j.neucom.2009.02.018>.
- Löffler, G. – Maurer, A. [2011]: Incorporating the dynamics of leverage into default prediction. *Journal of Banking & Finance*. Vol. 35., p. 3351-3361, DOI: <http://dx.doi.org/10.1016/j.jbankfin.2011.05.015>.
- Lyandres, E. – Zhdanov, A. [2013]: Investment opportunities and bankruptcy prediction. *Journal of Financial Markets*. Vol. 16., p. 439-476, DOI: <http://dx.doi.org/10.1016/j.finmar.2012.10.003>.
- Manjón-Antolín, M. C. – Arauzo-Carod, J. M. [2008]: Firm survival: methods and evidence. *Empirica*. Vol. 35., p. 1-24, DOI: <http://dx.doi.org/10.1007/s10663-007-9048-x>.
- Marqués, A. I. – García, V. – Sánchez, J. S. [2012a]: Two-level classifier ensembles for credit risk assessment. *Expert Systems with Applications*. Vol. 39., p. 10916-10922, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.03.033>.
- Marqués, A. I. – García, V. – Sánchez, J. S. [2012b]: Exploring the behaviour of base classifiers in credit scoring ensembles. *Expert Systems with Applications*. Vol. 39., p. 10244-10250, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.092>.
- Martens, D. – Van Gestel, T. – De Backer, M. – Haesen, R. – Vanthienen, J. – Baesens, B. [2010]: Credit rating prediction using ant colony optimization. *Journal of Operational Research Society*. Vol. 61., No. 4., p. 561-573, DOI: <http://dx.doi.org/10.1057/jors.2008.164>.
- Martin, D. [1977]: Early warning of bank failure: A logit regression approach. *Journal of Banking & Finance*. Vol. 1., No. 3., p. 249-276, DOI: [http://dx.doi.org/10.1016/0378-4266\(77\)90022-X](http://dx.doi.org/10.1016/0378-4266(77)90022-X).
- McKee, T. E. – Greenstein, M. [2000]: Predicting bankruptcy using recursive partitioning and a realistically proportioned data set. *Journal of Forecasting*. Vol. 19., No. 3., p. 219-230, DOI: [http://dx.doi.org/10.1002/\(SICI\)1099-131X\(200004\)19:3](http://dx.doi.org/10.1002/(SICI)1099-131X(200004)19:3).

McKee, T. E. – Lensberg, T. [2002]: Genetic programming and rough sets: A hybrid approach to bankruptcy classification. *European Journal of Operational Research*. Vol. 138., No. 2., p. 436-451, DOI: [http://dx.doi.org/10.1016/S0377-2217\(01\)00130-8](http://dx.doi.org/10.1016/S0377-2217(01)00130-8).

McKee, T. E. [2003]: Rough sets bankruptcy prediction models versus auditor signalling rates. *Journal of Forecasting*. Vol. 22., p. 569-586, DOI: <http://dx.doi.org/10.1002/for.875>.

McLeay, S. – Omar, A. [2000]: The sensitivity of prediction models to the non-normality of bounded and unbounded financial ratios. *British Accounting Review*. Vol. 32., p. 213-230, DOI: <http://dx.doi.org/10.1006/bare.1999.0120>.

Min, J. H. – Jeong, C. – Kim, M. S. [2011]: Tuning the architecture of support vector machine: The case of bankruptcy prediction. *International Journal of Management Science*. Vol. 17., No. 1., p. 19-43.

Min, J. H. – Lee, Y. C. [2005]: Bankruptcy prediction using support vector machine with optimal choice of kernel function parameters. *Expert Systems with Applications*. Vol. 28., p. 603-614, DOI: <http://dx.doi.org/10.1016/j.eswa.2004.12.008>.

Moro, R. – Hardle, W. – Aliakbari, S. – Hoffmann, L. [2011]: *Forecasting corporate distress in the Asian and Pacific region*. Working Paper, No. 11-08, Economics and Finance Working Paper Series, 2011 May, Internet: <http://www.brunel.ac.uk/about/acad/ss/depts/economics>

Mossman, C. E. – Bell, G. G. – Swartz, L. M. – Turtle, H. [1998]: An empirical comparison of bankruptcy models. *The Financial Review*. Vol. 33., No. 2., p. 35-54, DOI: <http://dx.doi.org/10.1111/j.1540-6288.1998.tb01367.x>.

Nam, C. W. – Kim, T. S. – Park, N. J. – Lee, H. K. [2008]: Bankruptcy prediction using a discrete-time duration model incorporating temporal and macroeconomic dependencies. *Journal of Forecasting*. Vol. 27., p. 493-506, DOI: <http://dx.doi.org/10.1002/for.985>.

Nam, J. H. – Jinn, T. [2000]: Bankruptcy prediction: Evidence from Korean listed companies during the IMF crisis. *Journal of International Financial Management and Accounting*. Vol. 11., No. 3., p. 178-197, DOI: <http://dx.doi.org/10.1111/1467-646X.00061>.

Nanda, S. – Pendharkar, P. [2001]: Linear models for minimizing misclassification costs in bankruptcy prediction. *International Journal of Intelligent Systems in Accounting, Finance and Management*. Vol. 10., p. 155-168, DOI: <http://dx.doi.org/10.1002/isaf.203>.

Neophytou, E. – Charitou, A. – Charalambous, C. [2000]: *Predicting Corporate Failure: Empirical Evidence for the UK*. Department of Accounting and Management Science, University of Southampton, Southampton

Neophytou, E. – Mar Molinero, C. [2005]: Financial ratios, size, industry and interest rate issues in company failure: An extended multidimensional scaling analysis. *Kent Business School Working Papers 10*. University of Kent, Kent.

Neves, J. C. – Vieira, A. [2006]: Improving bankruptcy prediction with hidden layer learning vector quantization. *European Accounting Review*. Vol. 15., No. 2., p. 253-271. DOI: <http://dx.doi.org/10.1080/09638180600555016>.

Nikolic, N. – Zarkic-Joksimovic, N. – Stojanovski, D. – Joksimovic, I. [2013]: The application of brute force logistic regression to corporate credit scoring models: Evidence from Serbian financial statements. *Expert Systems with Applications*. Vol. 40., p. 5932-5944, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.05.022>.

Nyitrai, T. [2014]: Validációs eljárások a csődelőrejelző modellek teljesítményének megítélésében. *Statisztikai Szemle*, Vol. 92., No. 4., p. 357-377.

Odom, M. D. – Sharda, R. [1990]: A neural network model for bankruptcy prediction. In: *Proceeding of the International Joint Conference on Neural Networks, San Diego, 17-21 June 1990, Vol. II.*, IEEE Neural Networks Council, Ann Arbor, p. 163-171.

Ohlson, J. A. [1980]: Financial ratios and the probabilistic prediction of bankruptcy. *Journal of Accounting Research*. Vol. 18., No. 1., p. 109-131, DOI: <http://dx.doi.org/10.2307/2490395>.

Olmeda, I. – Fernández, E. [1997]: Hybrid classifiers for financial multicriteria decision making: The case of bankruptcy prediction. *Computational Economics*. Vol. 10., No. 4., p. 317-335, DOI: <http://dx.doi.org/10.1023/A:1008668718837>.

Ooghe, H. – Claus, H. – Sierens, N. – Camerlynck, J. [1999]: *International comparison of failure prediction models from different countries: an empirical analysis*. Department of Corporate Finance, University of Ghent, Ghent.

Orbe, J. – Ferreira, E. – Núñez-Antón, V. [2001]: Modelling the duration of firms in Chapter 11 bankruptcy using a flexible model. *Economic Letters*. Vol. 71., No. 1., p. 35-42, DOI: [http://dx.doi.org/10.1016/S0165-1765\(01\)00358-5](http://dx.doi.org/10.1016/S0165-1765(01)00358-5).

Oreski, S. – Oreski, D. – Oreski, G. [2012]: Hybrid system with genetic algorithm and artificial neural networks and its application to retail credit risk assessment. *Expert Systems with Applications*. Vol. 39., p. 12605-12617, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.05.023>.

Oreski, S. – Oreski, G. [2014]: Genetic algorithm-based heuristic for feature selection in credit risk assessment. *Expert Systems with Applications*. Vol. 41., p. 2052-2064, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.09.004>.

Paleologo, G. – Elisseeff, A. – Antonini, G. [2010]: Subagging for credit scoring models. *European Journal of Operational Research*. Vol. 201., p. 490-499, DOI: <http://dx.doi.org/10.1016/j.ejor.2009.03.008>.

Pantalone, C. C. – Platt, M. B. [1987a]: Predicting failure of savings & loan associations. *Real Estate Economics*. Vol. 15., No. 2., p. 46-64, DOI: <http://dx.doi.org/10.1111/1540-6229.00418>.

Pantalone, C. C. – Platt, M. B. [1987b]: Predicting commercial bank failure since deregulation. *New England Economic Review*. 1987 July/August, p. 37-47.

Parnes, D. [2007]: Applying credit score models to multiple states of nature. *The Journal of Fixed Income*. Vol. 17., No. 3., p. 55-71, DOI: <http://dx.doi.org/10.3905/jfi.2007.700304>.

Pasaribu, R. B. F. [2008]: Financial distress prediction in Indonesia Stock Exchange. *Journal of Economics, Business and Accounting*. Vol. 11., No. 2., p. 153-172.

Pendharkar, P. C. [2011]: Probabilistic approaches for credit screening and bankruptcy prediction. *Intelligent Systems in Accounting, Finance and Management*. Vol. 18., p. 177-193, DOI: <http://dx.doi.org/10.1002/isaf.331>.

Peng, Y. – Kou, G. – Shi, Y. – Chen, Z. [2008]: A multi-criteria convex quadratic programming model for credit data analysis. *Decision Support Systems*. Vol. 44., p. 1016-1030, DOI: <http://dx.doi.org/10.1016/j.dss.2007.12.001>.

Philosophov, L. V. – Philosophov, V. L. [2002]: Corporate bankruptcy diagnosis: An attempt at a combined prediction of the bankruptcy event and time interval of its occurrence. *International Review of Financial Analysis*. Vol. 11., p. 375-406, DOI: [http://dx.doi.org/10.1016/S1057-5219\(02\)00081-9](http://dx.doi.org/10.1016/S1057-5219(02)00081-9).

Pindado, J. – Rodrigues, L. F. [2004]: Parsimonious models of financial insolvency in small companies. *Small Business Economics*. Vol. 22., No. 1., p. 51-66, DOI: <http://dx.doi.org/10.1023/B:SBEJ.0000011572.14143.be>.

Ping, Y. – Yongheng, L. [2011]: Neighborhood rough set and SVM based hybrid credit scoring classifier. *Expert Systems with Applications*. Vol. 38., p. 11300-11304, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.02.179>.

Platt, H. D. – Platt, M. B. [1990]: Development of a class of stable predictive variables: The case of bankruptcy prediction. *Journal of Business Finance & Accounting*. Vol. 17., No. 1., p. 31-51, DOI: <http://dx.doi.org/10.1111/j.1468-5957.1990.tb00548.x>.

Pompe, P. P. M. – Bilderbeek, J. [2005]: Bankruptcy prediction: the influence of the year prior to failure selected for model building and the effects in a period of economic decline. *Intelligent Systems in Accounting, Finance and Management*. Vol. 13., p. 95-112, DOI: <http://dx.doi.org/10.1002/isaf.259>.

Quek, C. – Zhou, R. W. – Lee, C. H. [2009]: A novel fuzzy neural approach to data reconstruction and failure prediction. *Intelligent Systems in Accounting, Finance and Management*. Vol. 16., p. 165-187, DOI: <http://dx.doi.org/10.1002/isaf.299>.

Riberio, B. – Silva, C. – Chen, N. – Vieira, A. – Das Neves, J. C. [2012]: Enhanced default risk models with SVM+. *Expert Systems with Applications*. Vol. 39., p. 10140-10152, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.02.142>.

Rösch, D. [2003]: Correlations and business cycles of credit risk: Evidence from bankruptcies in Germany. *Financial Markets and Portfolio Management*. Vol. 17., No. 3., DOI: <http://dx.doi.org/10.1007/s11408-003-0303-2>.

Ruzgar, N. S. – Unsal, F. – Ruzgar, B. [2008]: *Predicting bankruptcies using rough set approach: The case of Turkish banks*. American Conference on Applied Mathematics, Harvard, Massachusetts, USA, March 24-26, 2008. Internet: <http://www.naun.org/main/NAUN/ijmmas/mmmas-60.pdf>

Saberi, M. – Mirtalaie, M. S. – Hussain, F. K. – Azadeh, A. – Hussain, O. K. – Ashjari, B. [2013]: A granular computing-based approach to credit scoring modeling. *Neurocomputing*. Vol. 122., p. 100-115, DOI: <http://dx.doi.org/10.1016/j.neucom.2013.05.020>.

Salcedo-Sanz, S. – Deprado-Cumplido, M. – Segovia-Vargas, M. J. – Pérez-Cruz, F. – Bousoño-Calzón, C. [2004]: Feature selection methods involving support vector machines for prediction of insolvency in non-life insurance companies. *Intelligent Systems in Accounting, Finance and Management*. Vol. 12., p. 261-281, DOI: <http://dx.doi.org/10.1002/isaf.255>.

Salchenberger, L. – Cinar, E. – Lash, N. [1992]: Neural networks: A new tool for predicting thrift failures. *Decision Sciences*. Vol. 23., No. 4., p. 899-916, DOI: <http://dx.doi.org/10.1111/j.1540-5915.1992.tb00425.x>.

Sánchez-Lasheras, F. – De Andrés, J. – Lorca, P. – De Cos Juez, F. J. [2012]: A hybrid device for the solution of sampling bias problems in the forecasting of firms' bankruptcy. *Expert Systems with Applications*. Vol. 39., p. 7512-7523, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.01.135>.

Sarkar, S. – Sriram, R. S. [2001]: Bayesian models for early warning of bank failures. *Management Science*. Vol. 47., No. 11., p. 1457-1475, DOI: <http://dx.doi.org/10.1287/mnsc.47.11.1457.10253>.

Segovia-Vargas, M. J. – Gil-Fana, J. A. – Heras-Martínez, A. – Vilar-Zanón, J. L. – Sanchis-Arellano, A. [2003]: *Using rough sets to predict insolvency of Spanish non-life insurance companies*. Documentos de trabajo de la Facultad de Ciencias Económicas y Empresariales 03-02, Universidad Complutense de Madrid, Facultad de Ciencias Económicas y Empresariales, No. 03-02.

Serrano-Cinca, C. – Guitérez-Nieto, B. [2013]: Partial least square discriminant analysis for bankruptcy prediction. *Decision Support Systems*. Vol. 54., p. 1245-1255, DOI: <http://dx.doi.org/10.1016/j.dss.2012.11.015>.

Shetty, U. – Pakkala, T. P. M. – Mallikarjunappa, T. [2012]: A modified formulation of DEA to assess bankruptcy: An application to IT/ITES companies in India. *Expert Systems with Applications*. Vol. 39., p. 1988-1997. DOI: <http://dx.doi.org/10.1016/j.eswa.2011.08.043>.

Shin, K. S. – Lee, T. S. – Kim, H. J. [2005]: An application of support vector machines in bankruptcy prediction model. *Expert Systems with Applications*. Vol. 28., p. 127-135, DOI: <http://dx.doi.org/10.1016/j.eswa.2004.08.009>.

Sinha, A. P. – Zhao, H. [2008]: Incorporating domain knowledge into data mining classifiers: An application in indirect lending. *Decision Support Systems*. Vol. 46., p. 287-299, DOI: <http://dx.doi.org/10.1016/j.dss.2008.06.013>.

Sjovoll, E. [1999]: *Assessment of credit risk in the Norwegian business sector*. Norges Bank. Internet: <http://www.norges-bank.no/Upload/import/publikasjoner/arbeidsnotater/pdf/arb-1999-09.pdf>

Slowinski, R. – Zopounidis, C. [1995]: Application of the rough set approach to evaluation of bankruptcy risk. *Intelligent Systems in Accounting, Finance and Management*. Vol. 4., No. 1., p. 27-41, DOI: <http://dx.doi.org/10.1002/j.1099-1174.1995.tb00078.x>.

Sohn, S. Y. – Kim, H. S. [2007]: Random effects logistic regression model for default prediction of technology credit guarantee fund. *European Journal of Operational Research*. Vol. 183., p. 472-478, DOI: <http://dx.doi.org/10.1016/j.ejor.2006.10.006>.

Sueyoshi, T. – Goto, M. [2009]: Methodological comparison between DEA (data envelopment analysis) and DEA-DA (discriminant analysis) from the perspective of bankruptcy assessment. *European Journal of Operational Research*. Vol. 199., p. 561-575, DOI: <http://dx.doi.org/10.1016/j.ejor.2008.11.030>.

Sun, L. – Shenoy, P. P. [2007]: Using Bayesian networks for bankruptcy prediction: Some methodological issues. *European Journal of Operational Research*. Vol. 180., p. 738-753, DOI: <http://dx.doi.org/10.1016/j.ejor.2006.04.019>.

Sung, T. K. – Chang, N. – Lee, G. [1999]: Dynamics of modeling in data mining: Interpretive approach to bankruptcy prediction. *Journal of Management Information Systems*. Vol. 16., No. 1., p. 63-85.

Swiderski, B. – Kurek, J. – Osowski, S. [2012]: Multistage classification by using logistic regression and neural networks for assessment of financial condition of company. *Decision Support Systems*. Vol. 52., p. 539-547, DOI: <http://dx.doi.org/10.1016/j.dss.2011.10.018>.

Tam, K. Y. – Kiang, M. Y. [1992]: Managerial applications of neural networks: The Case of Bank Failure Predictions. *Management Science*. Vol. 38., No. 7., p. 926-947, DOI: <http://dx.doi.org/10.1287/mnsc.38.7.926>.

Tan, C. N. W. [1999]: *An Artificial Neural Networks Primer with Financial Applications Examples in Financial Distress Predictions and Foreign Exchange Hybrid Trading System*. School of Information Technology, Bond University, Australia

Telmoudi, F. – El Ghourabi, M. – Limam, M. [2011]: RST-GCBR-clustering-based RGA-SVM model for corporate failure prediction. *Intelligent Systems in Accounting, Finance & Management*. Vol. 18., p. 105-120, DOI: <http://dx.doi.org/10.1002/isaf.323>.

Trinkle, B. S. – Baldwin, A. A. [2007]: Interpretable credit model development via artificial neural networks. *Intelligent Systems in Accounting, Finance and Management*. Vol. 15., p. 123-147, DOI: <http://dx.doi.org/10.1002/isaf.289>.

Twala, B. [2010]: Multiple classifier application to credit risk assessment. *Expert Systems with Applications*. Vol. 37., p. 3326-3336, DOI: <http://dx.doi.org/10.1016/j.eswa.2009.10.018>.

Van Gestel, T. – Baesens, B. – Van Dijcke, P. – Garcia, J. – Suykens, J. A. K. – Vanthienen, J. [2006]: A process model to develop an internal rating system: Sovereign credit ratings. *Decision Support Systems*. Vol. 42., p. 1131-1151, DOI: <http://dx.doi.org/10.1016/j.dss.2005.10.001>.

Virág M. – Nyitrai T. [2014]: Is there a trade-off between the predictive power and the interpretability of bankruptcy models? The case of the first Hungarian bankruptcy prediction model, *Acta Oeconomica*, Vol. 64, No. 4., p. 419-440, DOI: <http://dx.doi.org/10.1556/AOecon.64.2014.4.2>.

Virág, M. – Hajdu, O. [1996]: Pénzügyi mutatószámokon alapuló csődmodell-számítások, *Bankszemle*, Vol. 15., No. 5., p. 42-53.

Virág, M. – Kristóf, T. – Fiáth, A. – Varsányi, J. [2013]: *Pénzügyi elemzés, csődelőrejelzés, válságkezelés*. Kossuth Kiadó, Budapest.

Virág, M. – Kristóf, T. [2005]: Az első hazai csődmodell újraszámítása neurális hálók segítségével, *Közgazdasági Szemle*, Vol. 52., No. 2., p. 144-162.

Virág, M. – Kristóf, T. [2006]: Iparági rátákon alapuló csődelőrejelzés sokváltozós statisztikai módszerekkel, *Vezetéstudomány*, Vol. 37., No. 1., p. 25-35.

Virág, M. – Kristóf, T. [2009]: Többdimenziós skálázás a csődelőrejelzésben. *Vezetéstudomány*. Vol. 40., No. 1., p. 50-58.

Virág, M. [1993]: *Pénzügyi viszonyyszámokon alapuló teljesítmény-megítélés és csődelőrejelzés*. Candidate thesis, Budapest.

Virág, M. [2004]: A csődmodellek jellegzetességei és története. *Vezetéstudomány*. Vol. 35., No. 10., p. 24-32.

- Virág, M. – Nyitrai, T. [2013]: Application of support vector machines on the basis of the first Hungarian bankruptcy model. *Society and Economy*. Vol. 35., No. 2., p. 227-248, DOI: <http://dx.doi.org/10.1556/SocEc.35.2013.2.6>.
- Vukovic, S. – Delibasic, B. – Uzelac, A. – Suknovic, M. [2012]: A case-based reasoning model that uses preference theory functions for credit scoring. *Expert Systems with Applications*. Vol. 39., p. 8389-8395, DOI: <http://dx.doi.org/10.1016/j.eswa.2012.01.181>.
- Wang, G. – Ma, J. – Yang, S. [2014]: An improved boosting based on feature selection for corporate bankruptcy prediction. *Expert Systems with Applications*. Vol. 41., p. 2353-2361, DOI: <http://dx.doi.org/10.1016/j.eswa.2013.09.033>.
- Wang, G. – Ma, J. [2011]: Study of corporate credit risk prediction based on integrating boosting and random subspace. *Expert Systems with Applications*. Vol. 38., p. 13871-13878, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.04.191>.
- Wang, G. – Ma, J. [2012]: A hybrid ensemble approach for enterprise credit risk assessment based on Support Vector Machine. *Expert Systems with Applications*. Vol. 39., p. 5325-5331, DOI: <http://dx.doi.org/10.1016/j.eswa.2011.11.003>.
- Wang, Z. [2004]: Financial ratio selection for default-rating modeling: A model-free approach and its empirical performance. *Journal of Applied Finance*. Spring/Summer, p. 20-35.
- Wilson, R. L. – Sharda, R. [1994]: Bankruptcy prediction using neural networks. *Decision Support Systems*. Vol. 11., No. 5., p. 545-557, DOI: [http://dx.doi.org/10.1016/0167-9236\(94\)90024-8](http://dx.doi.org/10.1016/0167-9236(94)90024-8).
- Xiaosi, X. – Ying, C. – Haitao, Z. [2011]: The comparison of enterprise bankruptcy forecasting method. *Journal of Applied Statistics*. Vol. 38., No. 2., p. 301-308, DOI: <http://dx.doi.org/10.1080/02664760903406470>.
- Yang, R. Z. [2001]: *A new method for company failure prediction using probabilistic neural network*. Department of Computer Science, Exeter University.
- Yang, Y. [2007]: Adaptive credit scoring with kernel learning methods. *European Journal of Operational Research*. Vol. 183., p.1521-1536, DOI: <http://dx.doi.org/10.1016/j.ejor.2006.10.066>.
- Yazici, M. [2011]: Combination of discriminant analysis and artificial neural networks in the analysis of credit card customers. *European Journal of Finance and Banking Research*. Vol. 4., No. 4., p. 1-10.
- Yeh, C. C. – Chi, D. J. – Hsu, M. F. [2010]: A hybrid approach of DEA, rough set and support vector machines for business failure prediction. *Expert Systems with Applications*. Vol. 37., p. 1535-1541, DOI: <http://dx.doi.org/10.1016/j.eswa.2009.06.088>.

Yim, J. – Mitchell, H. [2005]: A Comparison of Corporate Distress Prediction Models in Brazil: Hybrid Neural Networks, Logit Models and Discriminant Analysis, *Nova Economia Belo Horizonte*, Vol. 15, No. 1., p. 73-93.

Yoon, J. – Kwon, Y. S. – Lee, C. H. [2008]: Bankruptcy prediction for small businesses using credit card sales information: Comparison of classification performance. *Proceedings of the 9th Asia Pacific Industrial Engineering & Management Systems Conference*. Nusa Dua, Bali Indonesia, December 3-5, 2008

Yu, L. – Yao, X. – Wang, S. – Lai, K. K. [2011]: Credit risk evaluation using weighted least squares SVM classifier with design of experiment for parameter selection. *Expert Systems with Applications*. Vol. 38., p. 15392-15399, DOI: <http://dx.doi.org/10.1016/j.dss.2007.12.001>.

Yu, Q. – Miche, Y. – Séverin, E. – Lendasse, A. [2014]: Bankruptcy prediction using extreme learning machine and financial expertise. *Neurocomputing*. Vol. 128., p. 296-302, DOI: <http://dx.doi.org/10.1016/j.neucom.2013.01.063>.

Zavgren, C. V. [1985]: Assessing the vulnerability to failure of American industrial firms: A logistic analysis. *Journal of Business Finance & Accounting*. Vol. 12., No. 1., p. 19-45, DOI: <http://dx.doi.org/10.1111/j.1468-5957.1985.tb00077.x>.

Zhang, G. – Hu, M. – Patuwo, B. [1999]: Artificial neural networks in bankruptcy prediction: General framework and cross-validation analysis. *European Journal of Operational Research*. Vol. 116., No. 1., p. 16-32, DOI: [http://dx.doi.org/10.1016/S0377 2217\(98\)00051-4](http://dx.doi.org/10.1016/S0377 2217(98)00051-4).

Zhong, M. – Miao, C. – Shen, Z. – Feng, Y. [2014]: Comparing the learning effectiveness of BP, ELM, I-ELM and SVM for corporate credit ratings. *Neurocomputing*. Vol. 128., p. 285-295, DOI: <http://dx.doi.org/10.1016/j.dss.2005.10.001>.

Zmijewski, M. E. [1984]: Methodological issues related to the estimation of financial distress prediction models. *Journal of Accounting Research*. Vol. 22. Supplement, p. 59-82. DOI: <http://dx.doi.org/10.2307/2490859>.

## OWN PUBLICATIONS RELATED TO THIS TOPIC

Nyitrai, T. [2014]: Növelhető-e a csőd-előrejelző modellek előrejelző képessége az új klasszifikációs módszerek nélkül? *Közgazdasági Szemle*, Vol. 61., No. 5., p. 566-585.

Nyitrai, T. [2014]: Validációs eljárások a csődelőrejelző modellek teljesítményének megítélésében. *Statisztikai Szemle*, Vol. 92., No. 4., p. 357-377.

Nyitrai, T. [2015]: Hazai vállalkozások csődjének előrejelzése a csődeseményt megelőző egy, két, illetve három évvel korábbi pénzügyi beszámolók adatai alapján. *Vezetéstudomány*, Vol. 46., No. 5., p. 55-65.

Virág M. – Nyitrai T. [2014]: Is there a trade-off between the predictive power and the interpretability of bankruptcy models? The case of the first Hungarian bankruptcy prediction model, *Acta Oeconomica*, Vol. 64, No. 4., p. 419-440, DOI: <http://dx.doi.org/10.1556/AOecon.64.2014.4.2>.

Virág, M. – Nyitrai, T. [2014]: Metamódszerek alkalmazása a csődelőrejelzésben. *Hitelintézeti Szemle*, Vol. 13, No. 4., p. 180-195.

Virág, M. – Nyitrai, T. [2015]: Csődelőrejelző modellek dinamizálása – a szakértői tudás megjelenítése a csődelőrejelzésben. *Vezetés és szervezet társadalmi kontextusban. Tanulmányok Dobák Miklós 60. születésnapja tiszteletére*. Editors: Gyula Bakacsi – Károly Balaton. Akadémiai Kiadó, Budapest, ISBN 978 963 05 9634 3, p. 284-304.

Virág, M. – Nyitrai, T. [2013]: Application of support vector machines on the basis of the first Hungarian bankruptcy model. *Society and Economy*. Vol. 35., No. 2., p. 227-248, DOI: <http://dx.doi.org/10.1556/SocEc.35.2013.2.6>.