

Jacobson, Tor; Roszbach, Kasper

Working Paper

Bank Lending Policy, Credit Scoring and Value at Risk

Sveriges Riksbank Working Paper Series, No. 68

Provided in Cooperation with:

Central Bank of Sweden, Stockholm

Suggested Citation: Jacobson, Tor; Roszbach, Kasper (1998) : Bank Lending Policy, Credit Scoring and Value at Risk, Sveriges Riksbank Working Paper Series, No. 68, Sveriges Riksbank, Stockholm

This Version is available at:

<https://hdl.handle.net/10419/82386>

Standard-Nutzungsbedingungen:

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

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

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

Terms of use:

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

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

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

Bank lending policy, credit scoring and Value at Risk*

Tor Jacobson[†] Kasper Roszbach[‡]

30 July 1998

Abstract

In this paper we apply a bivariate probit model to investigate the implications of bank lending policy. In the first equation we model the bank's decision to grant a loan, in the second the probability of default. We confirm that banks provide loans in a way that is not consistent with default risk minimization. The lending policy must thus either be inefficient or be the result of some other type of optimizing behavior than expected profit maximization. Value at Risk, being a value weighted sum of individual risks, provides a more adequate measure of monetary losses on a portfolio of loans than default risk. We derive a Value at Risk measure for the sample portfolio of loans and show how analyzing this can enable financial institutions to evaluate alternative lending policies on the basis of their implied credit risk and loss rate, and make lending rates consistent with the implied Value at Risk.

JEL Classification: C35, D61, D81, G21, G33

Keywords: Banks, lending policy, credit scoring, Value at Risk, bivariate probit.

*We thank Georgina Bermann, Kenneth Carling and Anders Vredin for their helpful comments and Björn Karlsson and Yngve Karlsson at Upplysningscentralen AB for providing and discussing the data. Roszbach gratefully acknowledges financial support from the Jan Wallanders and Tom Hedelius Foundation for Research in the Social Sciences.

[†]Research Department, Sveriges Riksbank, SE 103 37 Stockholm, Sweden; *Email*: Tor.Jacobson@riksbank.se.

[‡]Department of Economics, Stockholm School of Economics, Box 6501, SE 113 83, Sweden; *Email*: nekr@hhs.se.

1. Introduction

Consumer credit has come to play an increasingly important role as an instrument in the financial planning of households. When current income falls below a household's permanent level and assets are either not available or not accessible for dissaving, credit is a means to maintain consumption at a level that is consistent with permanent income. People expecting a permanent increase in their income but lacking any assets, like students, have a desire to maintain consumption at a higher level than their current income allows. Borrowing can assist them in doing that. Those who accumulate funds in a pension scheme but are unable to get access to them when they experience a temporary drop in current income can also increase their welfare by bridging the temporary fall in income with a loan.

The quantitative importance of consumer credit may be illustrated by the fact that total lending, excluding residential loans, by banks and financial companies to Swedish households amounted to SEK 207 bn., or SEK 22.698 per capita, by the end of 1996. That is the equivalent of 12 percent of Swedish GDP or 22.7 percent of total private consumption. Viewed from the perspective of financial institutions, consumer credit also constitutes a significant part of their activities, making up 25 percent of total lending to the public. If one includes residential loans in total lending, this figure drops to 11 percent. When looking at the risk involved in these loans instead of their volume, their importance is even greater, however. Rules by the Basle Committee on Banking Supervision, that works under the umbrella of the Bank for International Settlements, stipulate an 8 percent capital requirement on consumer credit compared to, for example, 4 percent on residential loans. The above numbers make it clear that lending institutions' decision to grant a loan or not and their choice for a specific loan size can thus greatly affect a great many households' ability to smooth consumption over time, and thereby their welfare.

At a more aggregate level, consumer credit makes up a significant part of financial institutions' assets and the effects of any loan losses on lending capacity will be passed through to other sectors of the economy that rely on borrowing from the financial sector. Consequently, investigating the properties of banks' lending policies is not merely of interest because it enables us to examine how households' ability to smooth consumption is affected; these policies also have indirect implications for welfare, through financial markets. We will restrict ourselves to the first channel, however.

Despite credit's importance, it is common to see households being rationed in financial markets. Stiglitz and Weiss [13] and Williamson [14] contain two

different explanations of this phenomenon. When rationing is the mechanism that allocates resources in credit markets, some applicants will be excluded from credit despite being equally creditworthy as those granted a loan, making the equilibrium that results inefficient.¹ When a lender cannot observe borrowers' probabilities of default, credit scoring models - by enabling a lending institution to rank potential customers according to their default risk - can improve the allocation of resources, from a second best towards the first best equilibrium. Boyes et al. [6] investigate if the provision of credit takes place in an efficient way. For this purpose they estimate a bivariate probit model with two sequential events as the dependent variables: the lender's decision to grant the loan or not, and - conditional on the loan having been provided - the borrower's ability to pay it off or not. The parameters on variables like duration of job tenure, education and credit card ownership are, however, found to carry equal signs in both equations. Variables that increase (decrease) the probability of positive granting decision thus reduce (raise) the likelihood of a default. In addition, *unexplained* tendencies to extend credit, as measured by the regression error, were found to be positively correlated with default frequencies. Both these observations are inconsistent with a policy of default risk minimization.

In this paper, we construct an alternative risk measure for loans and present two problems to which it can be applied. Instead of an unweighted sum, a value weighted sum of all individual default risks is a more suitable measure of the risk on a portfolio of loans for a financial institution to consider when it needs to balance risk and return. Risk, after all, is only of interest because of the expected monetary losses that are associated with it. The focus of this paper will therefore be on Value at Risk rather than default risk.

First, we re-estimate the model of Boyes et al. on a bigger data set that contains both more reliable and more extensive financial and personal information on the loan applicants. This allows us to investigate the robustness of the finding that banks' lending policies are not consistent with default risk minimization. The re-estimated model will become the workhorse for the remainder of the paper. Next, we take Value at Risk as the relevant risk measure and study how marginal changes in a default risk based acceptance rule would shift the size of the bank's loan portfolio, its VaR exposure and average credit losses. Finally, we compare

¹Here, we have in mind the unequal treatment of ex-ante equal people due to an asymmetry in information sets. Several different definitions of credit rationing exist, however. One of them is the so called 'redlining' where less creditworthy people are willing to pay a higher interest rate but do not get a loan.

the risk on the sample portfolio with that on an efficiently provided portfolio of equal size.

The rest of the paper is organized as follows. Section 2 surveys the more recent research on statistical methods for credit scoring. Those who are already familiar with credit scoring or merely interested in the empirical results of the paper can skip this section. The data set and its sources are described in section 3. In section 4, we present the parameter estimates of the econometric workhorse model. In section 5, these estimates are used in the Value at Risk experiments we described earlier. Section 6 summarizes the results and discusses roads for future research.

2. Traditional credit scoring

The starting point of each loan is an application. When lending institutions receive an application for a loan, the process by which it is evaluated and its degree of sophistication can vary greatly. Most continue to use rather naïve, subjective evaluation procedures. This could be a non-formalized analysis of an applicant's personal characteristics or 'scoring with integer numbers' on these characteristics. Some banks, however, have started to use statistical credit scoring models. The objective of most credit scoring models is to minimize the misclassification rate or the expected default rate. To achieve this, various statistical methods are used to separate loan applicants that are expected to pay back their debts from those who are likely to fall into arrears or go bankrupt. Typically, a lending institution will analyse a sample of *granted* loans, their outcomes and the data that has been extracted from the applications and a credit agency's register.

The most commonly used statistical methods have been some form of discriminant analysis (DA) and logistic regression (LR). The DA model assumes that the exogenous variables, x_i , are normally distributed but with different means (and in the case of the 'quadratic discriminant model' even a different variance-covariance matrix) conditional on the group to which the dependent variable belongs. The objective is then to estimate these means (and, in the quadratic model, the covariance matrix) and then predict which of the groups an observation with characteristics x_i is most likely to come from.² DA thus differs from probit and logit

²Depending on the group to which the observation $(y_i, \mathbf{x}_i)$ belongs, the vector of exogenous variables x_i follows either f_0 or f_1 , where f_k are normal densities with mean $\boldsymbol{\mu}_k$ and variance-covariance $\boldsymbol{\Sigma}_k$, $k = 0, 1$. We define $y_i = 0$ and $y_i = 1$ if $\mathbf{x}_i$ is generated by f_0 and f_1 respectively. To be able to classify y_i we will need $Pr(y_i = 0 \mid \mathbf{x}_i)$ and $Pr(y_i = 1 \mid \mathbf{x}_i)$. Define

analysis in that the exogenous variables explicitly determine group membership in the latter two models whereas it is taken as given in the former. In other words, the 'causal' relationship runs from the dependent variable to the explanatory variables in DA, not the other way around. One potential weakness of the DA model is that the underlying assumptions are easy to violate. This occurs, for example, when the x_i are categorical instead of continuous and thus not normally distributed. In addition, discriminant analysis suffers from two other weaknesses that it shares with the other methods reviewed in this section. All of them merely minimize the number of accepted bad loans given an exogenous acceptance rate, without any rule for picking this rate optimally. Beside that, the models can only be estimated on samples of granted loans, which causes a sample selection bias in the parameters estimates.

Logistic regression's advantage over DA is that it does not suffer from the strict distributional assumption for $\mathbf{x}_i$.³ Among practitioners, DA appears to have lost ground to logistic regression. Steenackers and Goovaerts [12] is the latest in a long series of applications of the logistic regression to a sample of personal loans. Their model correctly classifies 62.6 percent of the good loans and 76.6 percent of all bad loans in a holdout sample. In other studies using DA or LR studies, the forecasting accuracy ranges from 54 to 90 percent. Altman, Avery, Eisenbeis and

the marginal probabilities by $p_0 = Pr(y_i = 0)$ and $p_1 = Pr(y_i = 1)$. Then Bayes' rule gives us that $Pr(y_i = 0 | \mathbf{x}_i) = \frac{f_0(\mathbf{x}_i) \cdot p_0}{f_0(\mathbf{x}_i) \cdot p_0 + f_1(\mathbf{x}_i) \cdot p_1}$, which under the normal distribution assumption equals $\Lambda(\beta_0 + \beta_1' \mathbf{x}_i + \mathbf{x}_i' \mathbf{A} \mathbf{x}_i)$, where $\Lambda(z) = \frac{e^z}{1+e^z}$ is the logistic cdf and β_0, β_1 are loglinear functions of $p_0, p_1, \mu_0, \mu_1, \Sigma_0, \Sigma_1$. If $\Sigma_0 = \Sigma_1$ is assumed, as common, then a linear logit model results. The likelihood function needed to estimate the 6 (4 - if one uses priors for p_0 and p_1) parameters is $\ell = \prod_{i=1}^N [f_0(\mathbf{x}_i) \cdot p_0]^{1-y_i} [f_1(\mathbf{x}_i) \cdot p_1]^{y_i}$. With the resulting parameter estimates one can then compute $Pr(y_i = k | \mathbf{x}_i)$, $k = 0, 1$. Altman et al. [1] p.40 show that one can modify this likelihood function to take the perceived cost of the two possible types of misclassification into account.

³The logit model has in common with the linear DA model that $Pr(y_i = 0 | \mathbf{x}_i) = \Lambda(\beta_0 + \beta_1' \mathbf{x}_i)$. Instead of arriving at this point by using Bayes' rule and assuming a normal distribution for the $\mathbf{x}_i$, it drops the assumptions of stochastic $\mathbf{x}_i$ and different group means and variance-covariances and takes the model $y_i^* = \Lambda(\beta_0 + \beta_1' \mathbf{x}_i) + \varepsilon_i$, $y_i = \begin{cases} 0 & \text{if } y_i^* \leq 0 \\ 1 & \text{if } y_i^* > 0 \end{cases}$ as a starting point. As a consequence $f_0(\mathbf{x}_i) = f_1(\mathbf{x}_i)$ and $P(y_i = k | \mathbf{x}_i) = p_k$, causing the likelihood function to take a different form than in the DA model: $\ell = \prod_{i=1}^N [\Lambda(\beta_0 + \beta_1' \mathbf{x}_i)]^{1-y_i} [1 - \Lambda(\beta_0 + \beta_1' \mathbf{x}_i)]^{y_i}$

Sinkev [1] review the literature up to 1980.

In a couple of studies k -nearest-neighborhood, count data or neural network methods have been employed. Henley and Hand [10] apply the k -nearest-neighbor method to credit scoring. A k -NN scoring rule classifies a new case on the basis of a majority vote amongst the k nearest sample elements, as measured by some metric that is defined over the space of explanatory variables. Its strong side is that it is a non-parametric method and thus not liable to any of the specification biases we mentioned before. One of the disadvantages is that the selection of some of the model's parameters involves quite a large degree of arbitrariness.⁴ In addition, the parameter estimates lack a clear interpretation.

Dionne et al. [7] study the costs of defaults by means of a count data model with sample selection effects. Their dependent variable is the number of non-payments under a predetermined repayment scheme. Their sample, however, consists merely of approved loan applications. Because of the sample selection bias in the parameter estimates that this leads to, the methodology in this study is thus of limited interest for the purpose of forecasting the profitability of future applicants.

Arminger et al. [3] compare a classification tree model⁵ and a feedforward

⁴The most popular measure of distance between two points $\mathbf{x}$ and $\mathbf{y}$ is the Euclidean metric $d_E(\mathbf{x}, \mathbf{y}) = \{(\mathbf{x} - \mathbf{y})'(\mathbf{x} - \mathbf{y})\}^{1/2}$. By using a weighting matrix $\mathbf{A}$, we can assign unequal importance to distances in different dimensions in the variable space. For example, $d_P(\mathbf{x}, \mathbf{y}) = \{(\mathbf{x} - \mathbf{y})' \mathbf{A} (\mathbf{x} - \mathbf{y})\}^{1/2}$ allows for larger weights for distances in the direction of variables that are of great importance and smaller weights for variables that affect the probability of belonging to a certain group only marginally. In practice, $\mathbf{A}$ is set equal to $\mathbf{I} + \mathbf{w}\mathbf{w}'$, where $\mathbf{w}$ is the gradient of the iso-probability curves, which is estimated by linear regression from the sample data. $\mathbf{A} = \mathbf{I} + \mathbf{w}\mathbf{w}'$ implies a metric that weighs the Euclidean distance and the distance between the points in the direction orthogonal to the iso-probability curves. One chooses k by trading off a bias against variance, because the former increases with k whereas the latter decreases.

⁵A classification tree splits up a sample into two subsamples, each of which contains only cases from a specific range of values of the dependent variable. The split should have the lowest degree of impurity as measured by the Gini coefficient. Building up an extra branch on the decision tree is done by means of a 2-step routine. First, at an already created node one finds (if the variable is qualitative) the categorization (or the cut-off point, if the variable is metrical) that minimizes the additional impurity for each explanatory variable. The impurity at this node is measured as $p_L \cdot i(t_L) + p_R \cdot i(t_R)$, where p_L is the probability of an individual being classified as a member of the left branch and $i(t_L)$ is the impurity on the left side branch under node t . Then in the second step the actual split is made by choosing the explanatory variable that minimizes the impurity over all exogenous variables at node t . How long the branching continues depends on a misclassification cost function that decreases with the number correct predictions but increases with the tree size. See Arminger et al. [3] p. 297-301.

neural network ⁶ with a logistic regression and find that all three are approximately equally good at predicting loan defaults. A shortcoming that the first two models share, however, is that one cannot quantify the importance of the explanatory variables; in the tree analysis because there are no parameters, and in the neural network because the parameters have no interpretation. Moreover, the neural network was found to be much worse at correctly predicting bad loans than good loans compared to the LR and the tree model.

3. Data

The data set consists of 13,338 applications for a loan at a major Swedish lending institution between September 1994 and August 1995. All loans were granted in stores where potential customers applied for instant credit to finance the purchase of a consumer good. The evaluation of each application took place in the following way. First, the store phoned to the lending institution to get an approval or a rejection. The lending institution then analysed the applicant with the help of a database with personal characteristics and credit variables to which it has on-line access. The database is maintained by Upplysningssentralen AB, the leading Swedish credit bureau which is jointly owned by all Swedish banks and lending institutions. If approval was granted, the store's salesman filled out a loan contract and submitted it to the lending institution. The loan is revolving and administered by the lending institution as any other credit facility. It is provided in the form of a credit card that can only be used in a specific store. Some fixed amount minimum payment by the borrower is required during each month. However,

⁶ An artificial neural network is a non-linear regression model $\mathbf{y} = \nu(\mathbf{x}) + \epsilon$, where the conditional expectation $\nu(\mathbf{x})$ is approximated by a function $\mu(\mathbf{x}; \boldsymbol{\theta})$. The network $\mu(\mathbf{x}; \boldsymbol{\theta})$ is composed of a number of interconnected processing units. These units - called 'perceptrons' when they are functions of linear combinations of their inputs - produce an output that is passed on to another processing unit. Many ANN's are composed of a number of perceptron layers, where the output of each unit of one layer can only be passed on to units of a higher one. Such ANN's are called feedforward networks.

For example, a FfN with one hidden layer (a two-layer-perceptron) would look like $\mu(\mathbf{x}; \boldsymbol{\beta}, \boldsymbol{\gamma}) = \phi(\boldsymbol{\beta}'\boldsymbol{\eta})$, $\boldsymbol{\eta} = (\eta_1, \dots, \eta_K)'$ where $\boldsymbol{\beta}$ is the parameter vector of the output unit, $\eta_k = \psi(\boldsymbol{\gamma}'_k \mathbf{x})$, $k = 1, \dots, K$, and the $\boldsymbol{\gamma}_k$ are the parameter vectors of the K hidden units. The outputs η_k of the hidden units are not observed and can be seen as realizations of latent variables. A common parameterization is to let $\phi(\cdot)$ be a linear model and $\varphi(\cdot)$ the standardized logistic distribution. Finally, predicted values $\hat{y}_i$ are obtained by using the threshold relation

$$\hat{y}_i = \begin{cases} 0 & \text{if } \mu(\mathbf{x}_i; \boldsymbol{\beta}, \boldsymbol{\gamma}) < 0 \\ 1 & \text{if } \mu(\mathbf{x}_i; \boldsymbol{\beta}, \boldsymbol{\gamma}) \geq 0 \end{cases} .$$

Table 1: Definition of variables.

Variable	Definition
<i>AGE</i>	age of applicant
<i>MALE</i>	dummy, takes value 1 if applicant is male
<i>MARRIED</i>	dummy, takes value 1 if applicant is married
<i>DIVORCE</i>	dummy, takes value 1 if applicant is divorced
<i>HOUSE</i>	dummy, takes value 1 if applicant owns a house
<i>BIGCITY</i>	dummy, takes value 1 if applicant lives in one of the three greater metropolitan areas around Göteborg, Malmö and Stockholm.
<i>NRQUEST</i>	number of requests for information on the applicant that the credit agency received during the last 36 months
<i>ENTREPR</i>	dummy, takes value 1 if applicant has taxable income from a registered business
<i>INCOME</i>	annual income from wages as reported to Swedish tax authorities (in 1000 SEK)
<i>DIFINC</i>	change in annual income from wages, relative to preceding year, as reported to Swedish tax authorities (in 1000 SEK)
<i>CAPINC</i>	dummy, takes value 1 if applicant has taxable income from capital
<i>BALINC</i> ⁷	ratio of total collateral free credit facilities actually utilized and <i>INCOME</i> , expressed as percentage.
<i>ZEROLIM</i>	dummy, takes value 1 if applicant has no collateral-free loans outstanding
<i>LIMIT</i>	total amount of collateral free credit facilities already outstanding (in 1000 SEK)
<i>NRLOANS</i>	number of collateral free loans already outstanding
<i>LIMUTIL</i>	percentage of <i>LIMIT</i> that is actually being utilized
<i>LOANSIZE</i>	amount of credit granted (in 1000 SEK)
<i>COAPPLIC</i>	dummy, takes value 1 if applicant has a guarantor

⁷This variable takes value zero when *INCOME* = 0 and is thus actually defined as $DUMMY_{\{income > 0\}} * (BALANCE / INCOME)$.

since the loan is revolving, there is no predetermined maturity of the loan. Earnings on the loan come from three sources: a one-time fee paid by the customer; a payment by the store that is related to total amount of loans granted through it; and interest on the balance outstanding on the card.

For this study, the lending institution provided us with a data file with the personal number of each applicant, the date on which the application was submitted, the size of the loan that was granted, the status of each loan (good or bad) on October 9, 1996, and the date on which bad loans gained this status.

Although one can think of several different definitions of a 'bad' loan, we classify a loan as bad once it is forwarded to a debt-collecting agency. We do not study what factors determine the differences in loss rates, if any, among bad loans. An alternative definition of the set of bad loans could have been 'all customers who have received one, two or three reminders because of delayed payment'. However, unlike 'forwarded to debt-collecting agency', one, two or three reminders were all transient states in the register of the financial institution. Once customers returned to the agreed-upon repayment scheme, the number of reminders was reset to zero. Such a property is rather undesirable if one needs to determine unambiguously which loans have defaulted and which have not.

Upplysningscentralen provided the information that was available on each applicant at the time of application and which the financial institution accessed for *its* evaluation. By exploiting the unique personal number that each resident of Sweden has, the credit bureau was able to merge these two data sets. Before handing over the combined data for analysis, the personal numbers were removed. The database included publicly available, governmentally supplied information, such as sex, citizenship, marital status, postal code, taxable income, taxable wealth, house ownership, and variables reported by Swedish banks like the total number of inquiries made about an individual, the number of unsecured loans and the total amount of unsecured loans. In total there we disposed of some 60 different variables.

A number of the variables in the dataset were not used in the final estimation of the model described in Sections 3 and 4. Among these are the number of months since the most recent change in marital status, citizenship (Swedish, nordic, non-nordic), number of months since immigration, number of houses a person owns (partially), assessed value of all real estate a person has (partial) ownership in, the combined value of all real estate ownership shares, several measures of income, taxable wealth, a large number of entries on the two most recently submitted income-tax return forms, like total taxes due, back tax etc, and a number of

Table 2: Descriptive statistics for all loan applicants ($N = 13338$)

Variable	Rejections ($N = 6899$)				Granted loans ($N = 6439$)			
	mean	stdev	min	max	mean	stdev	min	max
<i>AGE</i>	38.65	12.76	18	84	41.02	12.08	20	83
<i>MALE</i>	.62	.48	0	1	.65	.48	0	1
<i>DIVORCE</i>	.13	.34	0	1	.14	.35	0	1
<i>HOUSE</i>	.34	.47	0	1	.47	.50	0	1
<i>BIGCITY</i>	.41	.49	0	1	.37	.48	0	1
<i>NRQUEST</i>	4.69	2.60	1	10	4.81	2.68	1	19
<i>ENTREPR</i>	.04	.21	0	1	.02	.16	0	1
<i>INCOME</i>	129.93	70.38	0	737.9	189.47	75.70	0	1093.0
<i>DIFINC</i>	5.37	34.06	-438.5	252.6	9.03	34.63	-6226.0	5006.0
<i>CAPINC</i>	.12	.32	0	1	.07	.25	0	1
<i>BALINC</i> ⁸	91.04	894.53	0	41533	31.01	386.15	0	22387
<i>BALINC</i> ⁹	114.01	999.73	1	41533	35.85	431.87	1	22387
<i>ZEROLIM</i>	.15	.36	0	1	<.01	.05	0	1
<i>LIMIT</i>	79.89	93.69	0	1703.0	50.47	51.07	.0	949.2
<i>NRLOANS</i>	2.99	2.42	0	18	3.65	2.04	0	16
<i>LIMUTIL</i>	64.34	38.88	0	278.0	53.22	33.94	0	124.0
<i>COAPPLIC</i>	.16	.36	0	1	.14	.35	0	1

transformations of these variables.

Most of these were disregarded because they lacked correlation with the variables of interest - the loan granting decision and the payment behavior. Examples are back tax and real estate value. Others were disregarded because they displayed extremely high correlation with variables that measured approximately the same thing but had greater explanatory power. The numerous income measures in the dataset and *BALANCE* were eliminated in this way. A number of variables was selected for the statistical analysis because of an actual or supposed covariation with the dependent variables, but omitted from the final model because they did

⁸Only computed for the 6508 rejected and 6372 approved applications with *INCOME* > 0.

⁹Only computed for the 5197 rejected and 5086 approved applications with *BALINC* > 0.

Table 3: Descriptive statistics for granted loans.

Variable	Defaulted loans ($N = 388$)				Good loans ($N = 6051$)			
	mean	stdev	min	max	mean	stdev	min	max
<i>AGE</i>	36.11	11.03	21	75	41.33	12.07	20	83
<i>MALE</i>	.67	.47	0	1	.65	.48	0	1
<i>DIVORCE</i>	.20	.40	0	1	.14	.35	0	1
<i>HOUSE</i>	.28	.45	0	1	.48	.50	0	1
<i>BIGCITY</i>	.41	.49	0	1	.36	.48	0	1
<i>NRQUEST</i>	6.15	2.85	1	14	4.72	2.64	1	19
<i>ENTREPR</i>	.02	.13	0	1	.03	.16	0	1
<i>INCOME</i>	165.36	82.35	0	1093.0	191.01	75.00	0	1031.7
<i>DIFINC</i>	3.52	39.01	-135.0	439.7	9.38	34.30	-622.6	500.6
<i>CAPINC</i>	.04	.20	0	1	.07	.26	0	1
<i>BALINC</i> ¹⁰	39.92	313.51	0	6041	30.44	390.36	0	22387
<i>BALINC</i> ¹¹	46.45	337.81	1	6041	38.33	437.68	1	22387
<i>ZEROLIM</i>	.04	.20	0	1	<.01	.02	0	1
<i>LIMIT</i>	41.44	57.98	0	511.5	51.05	50.54	0	949.21
<i>NRLOANS</i>	2.34	1.64	0	11	3.74	2.04	0	16
<i>LIMUTIL</i>	75.69	33.37	0	124.0	51.78	33.47	0	112.0
<i>LOANSIZE</i>	7.08	3.95	3.0	24.5	7.12	3.83	3.0	30.0
<i>COAPPLIC</i>	.07	.26	0	1	.14	.35	0	1

not gain significance in any of the estimations. Citizenship, immigration related variables and real estate value were among these. Finally, wealth could not be used as an explanatory variable because not a single bad loan concerned a person with positive *taxable* wealth, thereby creating a numerical problem in the gradient of the likelihood function. Wealth up to SEK 900,000 is tax-exempted, making the group of people with *taxable* wealth extremely small in Sweden. Instead we

¹⁰Only computed for the 5988 good and 384 bad loans with *INCOME* > 0.

¹¹Only computed for the 4756 good and 330 bad loans with *BALINC* > 0.

used *taxable* income from capital - which is taxed from the first krona - to create a dummy explanatory variable. Tables 1 and 2 contain definitions and descriptive statistics for the variables that have been selected for the estimation of the final model in Section 4.

Of all applicants, 6,899, or 51.7 percent, were refused credit. The remaining 6,439 obtained a loan ranging from 3,000 to 30,000 Swedish kronor (approximately US\$ 375 - 3750). The lending institution's policy was that no loans exceeding 30,000 kronor were supplied. Although there is an indicated amortization scheme, the loans have no fixed maturity - they are revolving.

On 9 October 1996, the people in the sample were monitored by the lending institution. On that day 388 (6.0 %) of those who obtained a loan had defaulted and been forwarded to a debt collection agency. All other borrowers still fulfilled their minimum repayment obligations at that time. Some descriptive statistics are provided in Tables 2 and 3.

4. Econometric model

In this section we present the econometric model, that will be used as a workhorse in the experiments of Section 5. The model consists of two simultaneous equations, one for the binary decision to provide a loan or not, y_{1i} , and another for the binary outcome, 'default' or 'proper repayment', of each loan, y_{2i} . We let the superscript * indicate an unobserved variable and assume that y_{1i}^* and y_{2i}^* follow

$$\begin{aligned} y_{1i}^* &= \mathbf{x}_{1i} \cdot \boldsymbol{\alpha}_1 + \varepsilon_{1i}, \\ y_{2i}^* &= \mathbf{x}_{2i} \cdot \boldsymbol{\alpha}_2 + \varepsilon_{2i} \quad \text{for } i = 1, 2, \dots, N \end{aligned} \quad (4.1)$$

where the $\mathbf{x}_{ji}$, $j = 1, 2$, are $1 \times k_j$ vectors of explanatory variables.

The disturbances are assumed to be bivariate normally distributed.

$$\begin{pmatrix} \varepsilon_{1i} \\ \varepsilon_{2i} \end{pmatrix} \sim N \begin{pmatrix} 0 & 1 & \rho \\ 0 & \rho & 1 \end{pmatrix}$$

The binary choice variable y_{1i} takes value 1 if the loan was granted and 0 if the application was rejected:

$$y_{1i} = \begin{cases} 0 & \text{if loan not granted} & (y_{1i}^* < 0) \\ 1 & \text{if loan granted} & (y_{1i}^* \geq 0) \end{cases} \quad (4.2)$$

The second binary variable, y_{2i} , takes the value 0 if the loan defaults and 1 if not:

$$y_{2i} = \begin{cases} 0 & \text{if loan defaults} & (y_{2i}^* < 0) \\ 1 & \text{if loan does not default} & (y_{2i}^* \geq 0) \end{cases} \quad (4.3)$$

Due to the fact that one only observes if a loan is good or bad if it was granted, there is not only a censoring rule for (y_{1i}, y_{2i}) but even an *observation* rule. The observation rule is shown in Figure 1.

Figure 1: Observation rule for y_{1i} and y_{2i} . Entries in the 2×2 table show pairs (y_{1i}, y_{2i}) that are observed for all ranges of y_{1i}^* and y_{2i}^* .

	$y_{2i}^* < 0$	$y_{2i}^* \geq 0$	
$y_{1i}^* < 0$	$(0, \cdot)$	$(0, \cdot)$	
$y_{1i}^* \geq 0$	$(1, 0)$	$(1, 1)$	

Because we have three types of observations: no loans, bad loans and good loans, the likelihood function will take the following form:

$$\ell = \prod_{no \text{ loans}} \text{pr}(no \text{ loan}) \cdot \prod_{bad \text{ loans}} \text{pr}(bad \text{ loan}) \prod_{good \text{ loans}} \text{pr}(good \text{ loan}) \quad (4.4)$$

In appendix **A.1.** it is shown that (4.4) implies the following loglikelihood:

$$\begin{aligned} \ln \ell = & \sum_{i=1}^N (1 - y_{1i}) \cdot \ln [1 - \Phi(\mathbf{x}_{1i}\boldsymbol{\alpha}_1)] + \\ & \sum_{i=1}^N y_{1i} \cdot (1 - y_{2i}) \{ \Phi(\mathbf{x}_{1i}\boldsymbol{\alpha}_1) - \Phi_2(\mathbf{x}_{1i}\boldsymbol{\alpha}_1, \mathbf{x}_{2i}\boldsymbol{\alpha}_2; \rho) \} \\ & \sum_{i=1}^N y_{1i} \cdot y_{2i} \ln \Phi_2(\mathbf{x}_{1i}\boldsymbol{\alpha}_1, \mathbf{x}_{2i}\boldsymbol{\alpha}_2; \rho) \end{aligned} \quad (4.5)$$

where $\Phi(\cdot)$ and $\Phi_2(\cdot, \cdot, \rho)$ represent the univariate and bivariate standard normal c.d.f., the latter with correlation coefficient ρ .

The estimated parameters, their standard errors and t-statistics are presented in Table 4. Notice that *LOANSIZE* cannot be used as an explanatory variable in the first equation because no data on this variable is available for rejected applicants. The effect of many variables on the probability of obtaining a loan seems in

Table 4: Bivariate probit MLE of $\hat{\alpha}_1$ and $\hat{\alpha}_2$.

Variable	$P(\text{obtain a loan})$			$P(\text{loan does not default})$		
	$\hat{\alpha}_1$	std. error	t-stat.	$\hat{\alpha}_2$	std. error	t-stat.
<i>CONSTANT</i>	-.2374	.06652	-3.57	2.2900	.1463	15.65
<i>AGE</i>	-.004303	.001166	-3.69	.006892	.002624	2.63
<i>MALE</i>	-.2003	.02823	-7.10	-.02456	.05812	-.43
<i>DIVORCE</i>	-.02588	.03696	-.70	-.2380	.07125	-3.34
<i>HOUSE</i>	.06391	.02759	2.32	-.02019	.05830	.35
<i>BIGCITY</i>	-.2382	.02659	-8.96	-.03724	.05397	-.69
<i>NRQUEST</i>	-.008123	.005153	-1.58	-.1000	.01017	-9.84
<i>ENTREPR</i>	.5223	.06294	8.30	.2065	.1613	1.28
<i>INCOME</i>	.008928	.0001816	49.17	-.002392	.0004968	-4.81
<i>DIFINC</i>	-.002336	.0003445	-6.78	.002233	.0007352	3.04
<i>CAPINC</i>	-.2776	.05066	-5.48	.1189	.1232	.97
<i>BALINC</i>	.00006548	.00001811	3.48	-.00009135	.00005922	-1.54
<i>ZEROLIM</i>	-2.2440	.1049	-21.40	-.6590	.2890	-2.23
<i>LIMIT</i>	-.008381	.0001526	-54.94	.005064	.0004928	10.28
<i>NRLOANS</i>	.08420	.006882	12.23	.2698	.01873	14.41
<i>LIMUTIL</i>	-.007746	.0004370	-17.72	-.01197	.0009266	-12.91
<i>COAPPLIC</i>	.1300	.03395	3.83	.4374	.09715	4.50
<i>LOANSIZE</i>	-	-	-	-.006637	.006794	-.98
ρ	-	-	-	-.9234	.05326	-17.34

Critical values are 1.645, 1.96, and 2.575 for the 10, 5, and 1 percent significance levels.

accordance with the behavior banks commonly display. *INCOME*, *HOUSE*, *ENTREPR*, *NRLOANS* and having a *COAPPLIC*ant confirm their role as important factors that contribute positively while *ZEROLIM*, *LIMIT*, and *LIMUTIL* weigh negatively in the bank's decision. Somewhat surprising are the coefficients on *MALE*, *BIGCITY*, *DIFINC* and *CAPINC*. Men have a significantly smaller chance of being granted a loan as do people living in one of the three metropolitan areas. The same holds for people who have capital income and those who

experienced a rise in income during the last year. The latter effect deserves some further attention, though. Another way to interpret the sign of this parameter would be that people who experience large increases in wage income had quite a low income the preceeding year. Rather than reasoning that a rise in income worsens your chances of getting a loan, one could argue that income uncertainty - embodied in a low income in the year before - does so. If this were the case, then we should expect a similar effect to exist for people experiencing a fall in income. We tested for the presence of such an effect by transforming DIFINC into a variable with absolute values of income changes. We also tried with the standard deviation of income. Neither of these variables gained significance. We therefore interpret the coefficients in *INCOME* and *DIFINC* in the following way. We rewrite $\alpha_y y_t + \alpha_{dy} \Delta y_t$ as $(\alpha_y + \alpha_{dy}) y_t - \alpha_{dy} y_{t-1}$. Current and past income then both have positive and significant coefficients in the equation, with the former carrying the largest weight - as in a calculation of permanent income.

More striking is the fact that only four variables have the equal signs in both equations that one would expect when banks are minimizing default risk. Exploiting the credit facilities one disposes of to a greater extent (a higher *LIMUTIL*) or lacking experience with servicing debt (*ZEROLIM*) reduces an applicant's odds of obtaining a loan from the bank and increases the likelihood of a default. Having more experience in borrowing money and servicing a debt, as reflected by a higher *NRLOANS*, or applying together with a *COAPPLICant* makes it more likely that somebody will receive a loan and also add to the chances that the loan will be paid back.

Four variables have opposite coefficients in the loan granting and default equations. *INCOME*, notwithstanding a large positive weight in the decision to grant a loan by the bank, actually increases a loan's probability of default. Although one should be careful not to rationalize each counter-intuitive finding, we can look for a tentative explanation. Table 3 clearly shows that people who default on their loans have a lower average income than those who do not. This may well lead us to infer - if we disregard the rejected applicants, who have lower incomes than those who were granted a loan - that higher income reduces default risk. Suppose, however, that it is actually the case that other factors than *INCOME* determine a loan's survival. Then the selection of applicants may be taking place on the basis of a negative bivariate correlation between *INCOME* and defaults rather than on grounds of a negative partial correlation, which controls for both the sample selection effect and the correlation with other variables. It may, for example, be the case that people with higher income have other characteristics that are associated

with greater default risk. Similar arguments can be applied to *AGE*, *DIFINC* and *LIMIT*, that have negative weights in the first equation but positive weights in the second. Although one might, for example, expect *LIMIT* to have a negative impact on debt service, one should keep in mind that it is merely the ceiling of the credit facility that a person disposes of. *LIMUTIL* captures the extent to which he or she actually uses it, while *LIMIT* proxies for experience with servicing debt in the same way as *NRLOANS* does.¹² The positive coefficient on *DIFINC* in combination with the negative coefficient on *INCOME* illustrate how popularly assumed relations can lack factual support.

Furthermore, it is worthwhile to take notice of the large number of variables that are significant in only one equation and thus witness of inefficient use of information in the evaluation of applicants. *NRQUEST* is a proxy for people's eagerness to obtain additional credit and as such adds to the probability of a default. In the decision to grant a loan it has no role of importance, however. Being *DIVORCED*, which can bring about a mismatch between financial obligations and income, has an effect similar to that of *NRQUEST*. Five variables, *MALE*, *HOUSE*, *BIGCITY*, *CAPINC* and *BALINC* carry either positive or negative weight in the bank's decision but do not affect a loan's risk of default. Finally, we point out that *LOANSIZE* has no influence whatsoever on default risk. On the margin, an extra credit with a maximum of SEK 30,000 apparently does not affect default probability. Because the average *LIMIT* is between six and seven times the the average *LOANSIZE*, the relevant variable to study in this context is *LIMIT*, the total amount of credit facilities.

The only parameter we have not yet reviewed is the correlation coefficient. The value of -0.9234 implies that non-systematic tendencies to grant loans are almost perfectly correlated with non-systematic increases in default risk.¹³ In other words: the subjective elements - that conflict with the systematic policy described by the first equation in (4.1) - in the bank's lending policy that increase individuals' odds of being granted a loan, are positively related to increases in default

¹²Strong correlation between the variables *BALANCE* and *LIMIT* tended to create numerical problems when trying to use both as explanatory variables. Some test regressions indicated that *LIMIT* and *BALANCE* have opposite effects on *SURVIVAL*, the former a positive and the latter a negative. The coefficient on *LIMIT* in Table 4 is approximately equal to the net effect of *LIMIT* minus *BALANCE*.

¹³Although the value of -0.92324 for ρ is quite close to -1 , and more than twice as large as what Boyes et al. [6] found it to be, this is no symptom of problems with convergence for the algorithm. The correlation coefficient varied between -0.51 and -0.97 depending on the number and the type of variables that we let $\mathbf{x}_{ji}$ consist of.

risk that cannot be unexplained by a systematic relation with the covariates $\mathbf{x}_{2i}$.

If we compare the above results with those in Boyes et al. [6] we can make three observations. Firstly, our results confirm the conclusion in Boyes et al. that banks do not appear to be minimizing default risk. Many of the variables that make the bank approve loan applications are not among those that reduce the probability of default. Secondly, non-systematic tendencies to grant loans are indeed associated with greater default risk. Thirdly, we find that the *size* of a loan does *not* affect default risk. This contradicts the interpretation of Boyes et al. that banks pick out loans with higher default risk because they have higher returns. They suggest that banks actually prefer *bigger*, not riskier, loans for the one reason that they offer higher expected earnings. Because they think of bigger loans as generally also being riskier, maximizing earnings would imply deviating from risk minimization. In this paper we *control* for *LOANSIZE* by including it in the set of explanatory variables and find that it has no significant impact on default risk. Bigger loans are thus not riskier.

As a consequence, the fact that bank behavior is not consistent with risk minimization cannot be ascribed to a disregarded relation between loan size and return. Because all loans in this sample pay the same rate of interest, there remain only two sources of differences in the expected rate of return between loans: survival time - and the amortizations and interest payments that result from it - and the loss rate on bad loans. To get a good forecast of profitability, banks may be evaluating survival and the loss rate simultaneously. In a study of the survival of bank loans Roszbach [11] finds, however, that loans are not provided in a way that is consistent with survival time maximization. As an alternative, banks have been maximizing some other objective than the rate of return on their loan portfolio; for example, the number of customers or lending volume subject to a minimum return constraint, or total profits from a range of financial products. The current organization of information flows in banks and the degree of co-ordination between different departments does not allow for the pursuit of a composite objective such as the return on a range of products. Most important of all, the alternative objectives suggested above are not in agreement with what employees from the lending institution reported to us in a series of interviews on the matter. Rather, the results bear the evidence of a lending institution that has attempted to minimize risk or maximize a simple return function without success.

5. Lending policies and Value at Risk

Risk is of importance to financial institutions only to the extent that it involves expected monetary losses. Estimating individual default risks is merely of limited help, because their linkage with credit losses is unclear. A better way to measure risk is to weigh individual default risks by value, as one does, for example, in the calculation of Value at Risk. Studying Value at Risk not only enables the financial institution to get a measure of the credit risk present in currently administered loans. It also allows an evaluation of the impact of different lending policies on (a specific measure of) risk exposure and creates a better basis for an explicit decision on the implied loss rate. For these reasons, we will shift our attention in the remainder of the paper from the estimation of default risk to the construction of a Value at Risk (VaR) measure. First, we derive a VaR measure using a Monte-Carlo simulation of the bivariate probit model of Section 4. After that, we show how it can be applied in a typical problem that a lending institution may be confronted with when supplying loans.

We define Value at Risk as "the loss that is expected to be exceeded with a probability of only x % during the total holding period of the loan portfolio", where the relevant risk measure x needs to be chosen in advance. If one sets x equal to 5%, for example, a Value at Risk of SEK 10 mn. means that total credit losses on the loan portfolio will be greater SEK 10 mn. with a probability of 5%. Observe that our VaR concept differs from more conventional types in two respects. Firstly, the computed losses concern credit risk rather than market risk. Secondly, Value at Risk is calculated for a specific time horizon. VaR measures over other time horizons require re-estimation of the bivariate probit model.

One of the purposes of this section is to illustrate how using Value at Risk instead of default risk can be auxiliary in optimizing bank lending policy. We will therefore carry out two experiments, in which the bivariate probit model from section 4 will serve as a workhorse. First, we analyse how the Value at Risk is affected by marginal changes in the bank's acceptance rule. Second, we construct a hypothetical portfolio of loans that would be granted if the bank had a default risk based decision rule instead of its current policy. Comparing the distribution of credit losses on this hypothetical portfolio with those on the actual portfolio may supply us with a crude estimate of the efficiency losses that the bank's lending gives rise to.

In the first experiment we study how the bank can affect its Value at Risk exposure by making its acceptance criterion more or less restrictive. Here, we

abandon the bank's current lending policy, as described by the first equation in (4.1) and (4.2). In section 4 we showed that this policy is not consistent with risk minimization. Instead, we construct a default risk based acceptance rule of the form:

$$\left. \begin{array}{l} \text{loan not granted} \\ \text{loan granted} \end{array} \right\} \text{ if } \left\{ \begin{array}{l} pr(y_{2i} = 0) \geq \delta' \\ pr(y_{2i} = 0) < \delta' \end{array} \right. \quad (5.1)$$

By means of a Monte-Carlo simulation similar to the one described above, we can derive the probability distribution of bank credit losses associated with the acceptance/rejection rule (5.1) for any value of the threshold parameter δ' . The Monte-Carlo simulation consists of the following 5 steps:

1. Pick a value for δ' .
2. Draw one observation $\mathbf{x}_{2i}\widetilde{\boldsymbol{\alpha}}_2$ from $N(\mathbf{x}_{2i}\widehat{\boldsymbol{\alpha}}_2, \sigma_{\mathbf{x}_{2i}\widehat{\boldsymbol{\alpha}}_2})$ for $i = 1, 2, \dots, 13338$, where $\sigma_{\mathbf{x}_{2i}\widehat{\boldsymbol{\alpha}}_2}^2 = \mathbf{x}_{2i} \cdot \boldsymbol{\Sigma}_{\widehat{\boldsymbol{\alpha}}_2} \cdot \mathbf{x}_{2i}'$. Number them $i = 1, 2, \dots, 13338$;
3. To determine which applicants will be granted a loan, calculate the expected default probabilities $E[pr(y_{2i} = 0)]$ as

$$E[\widetilde{p}_i] = 1 - \Phi\left(\mathbf{x}_{2i}\widetilde{\boldsymbol{\alpha}}_{2i} / \left(1 + \sigma_{\mathbf{x}_{2i}\widetilde{\boldsymbol{\alpha}}_{2i}}^2\right)^{1/2}\right)$$

and then apply (5.1). Number the approved applications $i = 1, 2, \dots, N_A$;

4. For the N_A approved applications, compute the total credit losses λ on this portfolio as

$$\lambda = \sum_{i=1}^{N_A} E[\widetilde{p}_i] \cdot q_i$$

where q_i is the size of the loan individual i applied for. Because q_i is not available for the rejected applicants, we impute $\bar{q} = \frac{1}{N_{A, true}} \sum_{i=1}^{N_{A, true}} q_i$ in steps 1-4. Here, $N_{A, true}$ is the number of accepted applicants in the original sample.

5. Repeat steps 1-4 M times and compute the approximate probability distribution over losses from the M values one obtains for λ . M should be chosen such that the distribution is invariable for $M' \geq M$.

For our purpose we have picked a series of values δ' in the interval $[\.01, \.20]$. The results from these simulations are displayed in Table 5. The second column of Table 5 shows how expected loan losses increase as the bank relaxes its lending policy. The most restrictive policy, $\delta' = \.01$, results in lending between 36.6 mn. and 37.2 mn. kronor, whereas the most generous policy, $\delta' = \.20$, leads to approximately two and a half times as much lending. As the lending volume grows, losses increase at an accelerating rate. For the most risk averse decision rule loan losses range from .3 to .6 mn. kronor, compared to SEK 3 mn. - up to 9 times as much - on the riskiest loan portfolio. As expected the loss rate, loan losses divided by total lending, rises from .36 percent to 3.35 percent as the acceptance criterion δ is successively relaxed from .01 to .20. Total losses and the loss rate both monotonically increase with δ , at an ever decreasing rate however.

Applying a Value at Risk analysis before selecting a lending policy thus allows the lending institution to decide explicitly on either its aggregate credit risk exposure or its loss rate. Alternatively, it could choose to pick a desirable loss rate conditional on the Value at Risk not exceeding some maximum allowable amount of money. Doing so has several advantages. First, compared to current practice, the risk involved in lending becomes more transparent. Instead of registering loans that have already become non-performing, the financial institution will be able to create provisions for expected losses. This offers gains from both a private and a social perspective. From a private perspective because provisions for loan losses on banks' balance sheets will be forward-looking and only lag *unexpected* events. This should facilitate a correct valuation of the firm. At an aggregate level, there would be less risk for bankruptcy of financial institutions and therefore less risk for financial disturbances to the economy. See for example Bernanke and Gertler [5]. Secondly, unless the bank sets interest rates individually, this methodology also enables a bank to pick a risk-premium on top of the risk free rate of interest that is consistent with average credit risk over the maturity in question. If the loss rate is 2.5 percent, for example, and the average duration of a loan is 3 years, then the bank could charge a risk-premium of approximately .8 percent per annum.

In the second experiment our aim is to produce an estimate of the monetary losses that the inefficiency in the current lending policy gives rise to. Table 5 has given us an impression of how lending volume, loan losses and the loss rate covary, and can help a bank choose one specific efficient lending policy from a larger set. However, before switching to a new policy, a financial institution will first want to quantify the potential gains from doing so. For this purpose, we construct the 'efficient' portfolio of loans that would be granted if the bank used a default risk

Table 5: 95 percent confidence intervals for total loan losses, total lending (both in 1000 SEK) and the loss rate (total credit losses/total lending), all for given rejection threshold δ' .

δ'	loan losses	total lending	loss rate
.01	131 – 137	36,583 – 37,226	.36 – .37
.02	339 – 551	51,124 – 51,763	.66 – .68
.03	548 – 564	59,705 – 60,318	.92 – .94
.04	751 – 773	65,641 – 66,248	1.14 – 1.17
.05	948 – 947	70,083 – 70,683	1.35 – 1.38
.06	1,143 – 1,174	73,678 – 74,276	1.55 – 1.58
.07	1,334 – 1,370	76,672 – 77,225	1.74 – 1.77
.08	1,516 – 1,555	79,138 – 79,703	1.91 – 1.95
.09	1,687 – 1,730	81,181 – 81,730	2.08 – 2.12
.10	1,849 – 1,895	82,916 – 83,449	2.23 – 2.27
.11	2,003 – 2,053	84,411 – 84,923	2.37 – 2.42
.12	2,149 – 2,201	85,699 – 86,190	2.51 – 2.55
.13	2,285 – 2,339	86,812 – 87,279	2.63 – 2.68
.14	2,414 – 2,470	87,785 – 88,230	2.75 – 2.80
.15	2,285 – 2,339	88,641 – 89,066	2.85 – 2.91
.16	2,648 – 2,709	89,389 – 89,793	2.96 – 3.02
.17	2,752 – 2,814	90,034 – 90,416	3.06 – 3.11
.18	2,847 – 2,909	90,586 – 90,945	3.14 – 3.20
.19	2,932 – 2,995	91,057 – 91,398	3.22 – 3.28
.20	3,009 – 3,073	91,467 – 91,789	3.29 – 3.35

based decision rule, instead of its current policy, but preferred a lending volume (approximately) equal to that of the actual portfolio. Executing steps 1-3 in the above Monte-Carlo experiment and picking δ' such that the simulated lending volume equals actual lending gives us the desired portfolio.

By inspecting Table 5 one can already infer that the implied value of δ' will lie between .01 and .02. We find that δ' equals .012. We then repeat steps 2-5 of the Monte-Carlo experiment for both the actual and the 'efficient' portfolio - but do not apply (5.1) in step 3 since we already know which individuals make up our

Table 6: Value at Risk at different risk levels computed for the sample portfolio and an efficiently provided portfolio of equal size (amounts $\times$ thousand SEK).

Portfolio	Risk level		
	1%	5%	10%
<i>Sample</i>	1,513	1,506	1,503
<i>Efficient</i>	263	262	261

sample. From the credit loss distributions that we obtain along these lines, we extract three different Value at Risk measures for each portfolio. These are displayed in Table 6. Credit losses on the two portfolios clearly differ greatly. At the 10 percent risk level, the value at risk amounts to SEK 1,503 thousand for the actual portfolio compared to 261 thousand for the efficient portfolio. At the 1 percent risk level these amounts are 1,513 and 263 respectively. By shifting to a default risk based decision rule and abandoning its current lending policy, the bank can reduce its expected credit losses significantly. Continuing providing loans in the same way as has been done leads to a VaR exposure that is six times higher than with a policy consistent with default risk minimization. Switching to one of the 'efficient' lending policies displayed in Table 5 thus involves large potential benefits for the financial institution.

6. Discussion

In this paper we have applied the bivariate probit model from Boyes et al. [6] to investigate the implications of bank lending policy. With a larger and more extensive data set we confirm earlier evidence that banks provide loans in a way that is not consistent with default risk minimization. It had been suggested that banks prefer *bigger* loans because they offer higher expected earnings. Since bigger loans are generally thought to be riskier, maximizing expected earnings would then imply deviating from risk minimization. However, with the data on the size of all loans that we have at our disposal, size has been shown not to affect the default risk associated with a loan. Banks, even if they are risk averse, are thus not faced with a trade-off between risk and return. The inconsistency in banking behavior can thus not be ascribed to some relation between loan size and return, that earlier models had not accounted for. The banking behavior must thus be either a symptom of an inefficient lending policy or the result of

some other type of optimizing behavior. Banks may, for example, be forecasting survival time, loss rates or both. Another alternative is that they are maximizing another objective than the rate of return on their loan portfolio, e.g. the number of customers, lending volume subject to a minimum return constraint, or total profits from a range of financial products. Current banking technology does not yet allow for the pursuit of a composite objective such as the return on a range of products, however. In addition, the above suggestions are not in agreement with the practices reported to us by the lending institution that provided our data. Rather, the results bear the evidence of a lending institution that has attempted to minimize risk or maximize a simple return function without success.

Value at Risk, being derived from a value weighted sum of all individual risks, provides a more adequate measure of monetary losses on a portfolio of loans than default risk. By means of Monte-Carlo simulation with the bivariate probit model, we have obtained a Value at Risk measure for the sample portfolio of loans. We have also shown how calculating Value at Risk can enable financial institutions to evaluate alternative lending policies on the basis of their implied credit risks and loss rates. An analysis of the VaR involved in lending policies offers both private and social gains. Provisions for loan losses on banks' balance sheets will become more forward-looking. This should facilitate a correct evaluation of the firm. At an aggregate level, the risk of bankruptcy for financial institutions and the likelihood of financial disturbances to the economy would be reduced. Banks would also be able to choose a risk-premium on top of the risk free rate of interest that is *consistent with* average credit risk over the maturity in question.

References

- [1] Altman, E.I., R.B. Avery, R.A. Eisenbeis and J.F. Sinkey, (1981), *Application of classification techniques in business, banking and finance*, JAI Press, Greenwich, CT.
- [2] Amemiya, Y., (1985), *Advanced Econometrics*, Harvard University Press, Cambridge MA.
- [3] Arminger, G., D. Enache and T. Bonne, (1997), Analyzing credit risk data: a comparison of logistic discrimination, classification tree analysis and feed-forward networks, *Computational Statistics* 12, 293-310.
- [4] Bermann, G., (1993), *Estimation and inference in bivariate and multivariate ordinal probit models*, dissertation, Department of statistics, Uppsala University.
- [5] Bernanke, B., and M. Gertler, (1995), Inside the Black Box: The Credit Channel of Monetary Policy Transmission, *Journal of Economic Perspectives* 9(4), 27-48.
- [6] Boyes, W.J., D.L. Hoffman and S.A. Low, (1989), An Econometric Analysis of the Bank Credit Scoring Problem, *Journal of Econometrics* 40, 3-14.
- [7] Dionne, G., M. Artis and M. Guillen, (1996), Count data models for a credit scoring system, *Journal of Empirical Finance* 3, 303-325.
- [8] Gale, D., and M. Hellwig, (1985), Incentive-compatible debt contracts: The one-period problem, *Review of Economic Studies* LII, 647-663.
- [9] Greene, W.E., (1993), *Econometric Analysis*, 2nd edition, Macmillan, New York.
- [10] Henley, W.E., and D.J. Hand, (1996), A k-nearest-neighbor classifier for assessing consumer credit risk, *The Statistician* 45 (1), 77-95.
- [11] Roszbach, K.F., (1998), *Bank lending policy, credit scoring and the survival of loans*, Manuscript, Stockholm School of Economics.
- [12] Steenacker, A., and M.J. Goovaerts, (1989), A credit scoring model for personal loans, *Insurance: Mathematics and Economics* 8, 31-34.

- [13] Stiglitz, J.E., and A. Weiss, (1981), Credit rationing in markets with imperfect information, *American Economic Review* 71, 393-410.
- [14] Williamson, S., (1987), Costly monitoring, loan contracts and equilibrium credit rationing, *Quarterly Journal of Economics* 102 (1), 135-145.

A. Likelihood function

Combining (4.2) – (4.3) and table 1, the likelihood function in equation (4.4) becomes

$$\ell = \prod_{i=1}^N \text{pr}(y_{1i}^* < 0)^{(1-y_{1i})} \cdot \prod_{i=1}^N \text{pr}(y_{1i}^* \geq 0, y_{2i}^* \leq 0)^{y_{1i} \cdot y_{2i}} \times \prod_{i=1}^N \text{pr}(y_{1i}^* \geq 0, y_{2i}^* \geq 0)^{y_{1i} \cdot y_{2i}} \quad (\text{A.1})$$

Substituting for (4.1), (A.1) implies the following loglikelihood function:

$$\begin{aligned} \ln \ell = & \sum_{i=1}^N (1 - y_{1i}) \cdot \ln [\text{pr}(\varepsilon_{1i} < -\mathbf{x}_{1i}\boldsymbol{\alpha}_1)] + \\ & \sum_{i=1}^N y_{1i} \cdot (1 - y_{2i}) \ln [\text{pr}(\varepsilon_{1i} \geq -\mathbf{x}_{1i}\boldsymbol{\alpha}_1 \cap \varepsilon_{2i} \leq -\mathbf{x}_{2i}\boldsymbol{\alpha}_2)] + \\ & \sum_{i=1}^N y_{1i} \cdot y_{2i} \ln [\text{pr}(\varepsilon_{1i} \geq -\mathbf{x}_{1i}\boldsymbol{\alpha}_1 \cap \varepsilon_{2i} \geq -\mathbf{x}_{2i}\boldsymbol{\alpha}_2)] \end{aligned} \quad (\text{A.2})$$

Because of the symmetry property of the bivariate normal distribution, the last line in (A.2) can be rewritten as:

$$\begin{aligned} \text{pr}(\varepsilon_{1i} \geq -\mathbf{x}_{1i}\boldsymbol{\alpha}_1 \cap \varepsilon_{2i} \geq -\mathbf{x}_{2i}\boldsymbol{\alpha}_2) & \Leftrightarrow^{14} \\ \Phi_2(\mathbf{x}_{1i}\boldsymbol{\alpha}_1, \mathbf{x}_{2i}\boldsymbol{\alpha}_2; \rho) \end{aligned} \quad (\text{A.3})$$

Since

$$\text{pr}(y_{1i}^* \geq 0, y_{2i}^* \leq 0) = 1 - \text{pr}(y_{1i}^* < 0) - \text{pr}(y_{1i}^* \geq 0, y_{2i}^* \geq 0)$$

$\forall i$, the loglikelihood function can be written as

$$\begin{aligned} \ln \ell = & \sum_{i=1}^N (1 - y_{1i}) \cdot \ln [1 - \Phi(\mathbf{x}_{1i}\boldsymbol{\alpha}_1)] + \\ & \sum_{i=1}^N y_{1i} \cdot (1 - y_{2i}) \{ \Phi(\mathbf{x}_{1i}\boldsymbol{\alpha}_1) - \Phi_2(\mathbf{x}_{1i}\boldsymbol{\alpha}_1, \mathbf{x}_{2i}\boldsymbol{\alpha}_2; \rho) \} \\ & \sum_{i=1}^N y_{1i} \cdot y_{2i} \ln \Phi_2(\mathbf{x}_{1i}\boldsymbol{\alpha}_1, \mathbf{x}_{2i}\boldsymbol{\alpha}_2; \rho) \end{aligned} \quad (\text{A.4})$$

¹⁴See Greene (1993) p.661 for a summary of results on the bivariate normal cdf.