

doi:10.11918/j.issn.0367-6234.2015.05.007

个人信用评分模型的发展及优化算法分析

姜明辉¹, 许佩¹, 任潇¹, 车凯²

(1.哈尔滨工业大学 管理学院, 150001 哈尔滨; 2.哈尔滨工业大学 计算机科学与技术学院, 150001 哈尔滨)

摘要: 为有效识别与计量个人信用风险, 规避金融危机对商业银行的不利影响, 保持我国信贷和金融市场的正常运转, 对个人信用评分的主要模型及其发展进行了归纳, 阐明个人信用评分研究中仍存在信用样本的有效性、完整性问题、信用指标体系的合理性问题以及模型的选择及适用性问题。鉴于此, 基于相关性分析实现异常样本的预警, 基于蒙特卡洛算法对样本进行补足; 结合统计学模型及人工智能模型, 采用步长遍历算法对指标体系进行优化及显著性加权; 以精确度、稳健性、第1误判率、第2误判率及差异性作为选择指标, 实现评分模型的选择与输出。分析表明: 通过上述优化算法, 将解决个人信用评分中存在的问题, 提高商业银行的风险控制能力。

关键词: 风险控制; 个人信用; 信用评分模型; 优化算法

中图分类号: F830.589 **文献标志码:** A **文章编号:** 0367-6234(2015)05-0040-06

Research on algorithms development and optimization for personal credit scoring

JIANG Minghui¹, XU Pei¹, REN Xiao¹, CHE Kai²

(1.School of Management, Harbin Institute of Technology, 150001 Harbin, China;

2.School of Computer Science and Technology, Harbin Institute of Technology, 150001 Harbin, China)

Abstract: In order to effectively recognize and measure individual credit risk and avoid the negative impact which financial crisis has on the commercial banks, and thus keep the normal operation of the credit and financial markets in China, this paper summarizes the main model of personal credit scoring and its development, and also clarifies the validity and integrity of the credit samples, the rationality of the index system and the selection, applicability of the model in the study of personal credit rating. In view of this, based on correlation analysis, early warning of abnormal samples can be achieved, and the samples are complemented through the Monte Carlo Algorithm. Combined with the statistical model and artificial intelligence model, the index system is optimized and given respective significant weight through traversal algorithm. With accuracy, robustness, first misjudgment rate, second misjudgment rate and difference as the selection index, the selection and output of the scoring model can be realized. Through the above optimization algorithm, the problem in personal credit scoring is solved, and the risk control ability of commercial banks is improved as well.

Keywords: risk management; personal credit; model for credit scoring; optimization algorithm

个人信贷作为银行的主要资产业务之一,其风险水平的控制关系到商业银行对于经济资本的整体要求。因此,能否对个人信用风险进行有效的识别与计量,成为商业银行能否合理控制风险的关键因素。随着我国个人信贷规模和涉及领域日

益扩大,自90年代后期开始,个人信用评分方法开始引起国内的关注。近年来随着我国经济的高速发展,个人住房抵押贷款逐年增加,房地产市场出现过热现象,个人信用贷款的风险也随之增加。因此,在后金融危机时代,研究我国个人信用评分,对有效识别信用风险、规避金融危机的不利影响以及保持我国信贷和金融市场的正常运转,甚至对维持国民经济的持续稳定增长都具有重大的理论和现实意义。

文献[1]指出,金融机构的传统做法是由专

收稿日期: 2014-06-13.

基金项目: 国家自然科学基金(70871030); 黑龙江省自然科学基金(G200914).

作者简介: 姜明辉(1967—),男,教授,博士生导师.

通信作者: 姜明辉, mhjiang2000@sina.com.

家基于自身经验对个人信用进行判断, 并由此形成了“5C”评价法。由于过度依赖于专家的经验, 存在着训练成本高, 主观性和随意性强等问题。正是为了解决这些问题, 个人信用评分模型应运而生, 其实质是基于客观的数学模型, 根据已掌握的客户的相关信息对客户将来可能的信用情况进行判断。模型通过对客户特定个人信息的输入, 将客户分为“好客户”(能够按时还本付息)和“坏客户”(会出现违约情况)两类。

随着国内外对信用评分研究的不断深入, 个人信用评分模型也经历了由统计学方法到非参数方法、运筹学方法再到人工智能方法的演变, 同时由单一模型到组合模型逐渐演进。但是, 已有的个人信用评分模型在我国的实际应用中仍存在诸如样本的有效性、完整性、指标体系的适用性、模型选择的可解释性等问题。鉴于此, 本文提出采用通过对已有样本的规则提取实现异常数据的预警, 结合样本有效性及完整性的改进, 选取解释能力强的单一模型对信用评分的指标体系进行显著性排序, 选取适用的指标显著性权重并综合考虑整体的准确率、两类误判率和差异性, 对现有模型进行优化。

1 个人信用评分模型的发展

1.1 统计学方法

判别分析 (discriminant analysis, DA) 源于对 3 种鸢尾属植物的分类实验并被文献 [2] 首次用来区分信用客户的好坏。判别分析的主要思想是基于某些分类方法来使同类之间距离最小, 异类之间距离最大, 通过建立一个或多个判别方程, 来判断某一变量的类别归属。文献 [3] 认为当变量服从多元椭圆面分布 (多元正态分布是其特例) 时, 线性判别无疑是最优的选择 (忽略样本抽样偏差)。此后, 随着著名的 FICO (fair isaac corporation) 信用评分系统的建立, 作为 FICO 系统的核心方法——判别分析在个人信用评分领域得到了广泛运用。近年来, 文献 [4] 将最新的判别分析方法——SNDA、STDA、SDA、Sparse DA、FDA、MDA 分别应用于个人信用评分, 以总精确度及错分率为判别指标, 指出 SNDA、STDA 和 SDA 在个人信用评分领域表现良好。

回归分析法 (regression analysis, RA) 是起源于遗传学研究的经典统计学方法之一。回归分析法是在大量已知数据的基础上, 来探究一种变量 (自变量) 对另外一种变量 (因变量) 的影响, 并建立描述二者间相关关系的回归方程, 根据已知的自变量的值对因变量的值进行预测。在回归分析

法中, 应用较为广泛的有 Logistic 回归分析、Probit 回归分析及多元线性回归。与判别分析相比, 回归分析的鲁棒性较低, 但回归分析对数据分布的要求相对宽松, 而且能够提供客户的违约概率, 因此获得了大多数学者和银行业的青睐。目前为止, Logistic 回归已经成为最成功且最常用的统计方法之一。文献 [5] 得出 Logistic 回归在分类效果上要优于判别分析的结论。

1.2 非参数方法

最近邻法 (nearest neighbors, NNs) 是首先被用于分类问题的标准非参数方法, 最早由纽约银行应用于信用评分领域。最近邻法中最常用的是 KNN 模型, KNN 模型能够很好的解决概率密度函数的分类和估计问题, 在个人信用评分研究中取得了较好的效果。KNN 模型的基本原理是通过计算寻找与待判样本点距离最近的 k 个信用样本, 再根据 k 个样本的表现, 以投票的方式确定待判样本的信用情况。文献 [6-7] 均指出由于最近邻法不用提前学习和训练模型, 允许动态的更改客户信息, 从而能很好的解决人口漂移问题。有关 KNN 模型较近的应用研究主要关注了“维数祸根” (curse of dimensionality) 问题, 指出最近邻法在应用于高维数据时, 即使样本量很大, 散落在高维空间内的样本点仍十分稀疏, 难以找到相邻的样本。针对该问题, 文献 [8] 提出可以通过非线性的数据投影法来降低数据维度; 文献 [9] 提出可以对最近邻法进行改进, 使用基于排序的最近邻法来解决这一问题。

决策树法 (decision tree, DT) 是近年来被引入信用评分领域的主要非参数方法之一。决策树法以违约的可能上同质性更强为划分标准, 将信用申请者划分为两个子类, 每个子类再次划分为同质性更强的子类, 整个递归过程直到子类达到预设的终止条件为止。决策树算法支持多个参数, 会对所生成的挖掘模型的性能和准确性产生影响。文献 [10] 首次将决策树用于个人信用评估方法中。考虑到样本属性中包括了数值型数据及非数值型数据, 文献 [11] 将 Boosting 算法技术嵌入决策树中, 该尝试取得了更好的判别效果。

数据包络分析法 (data envelopment analysis, DEA) 是在相对效率评价基础上发展的系统分析方法。它是以相对效率概念为基础, 根据多指标投入和多指标产出对相同类型的单位进行相对有效性或效益评价的一种新方法。将数据包络分析法应用于个人信用评估, 可将客户的特征向量视为投入指标, 客户的信用情况视为产出指标进行分

类.数据包络分析法的优点在于能够有效的避免主观因素,减少误差,且建立模型前无须对数据进行量纲一的处理,与个人信用指标的特征一致.文献[12]将 DEA 模型应用于私人融资计划中借款人的信用评分,指出 DEA 有着能够自动生成相对权重等优点.

1.3 运筹学方法

在个人信用评分中应用的运筹学方法主要是线性规划法(linear programming, LP).早在 1965 年,线性规划法即被应用于分类问题.但是直到 1981 年,文献[13]阐明线性规划在判别及分类上的应用及实现之后,该方法才引起了更多学者的关注.线性规划模型通过找到一组权重值,在给定的临界值的条件下,使得所有好客户的得分都在该临界值之上,而所有坏客户的得分都在这个临界值之下从而实现个人信用样本的分类.在线性规划方法应用于个人信用评分的基础上,学者们更关注于与统计学方法在应用效果上的差异,文献[14]通过研究指出统计学方法要优于线性规划的方法.

1.4 人工智能方法

专家系统(expert system),作为人工智能方法应用于个人信用评分最成功的尝试,其核心思想为通过一个包含某特定领域知识的数据库和对信息进行递推的规则,分析新情况并给出专家级的解决方案.文献[15]介绍了 CLUES 专家系统的构建,该系统可以决策是否批准住房抵押贷款申请,后被美全国金融公司采用.

神经网络(artificial neural networks, ANNs)作为最具有代表性的人工智能方法之一,其原理是通过对变量进行线性组合和非线性变化,然后循环修正,进而模拟人类大脑的决策过程,利用神经元相互触发,建立一种学习机制.文献[16]在信用风险评测中引入神经网络的方法.2000 年,Moody's 公司公布了一套上市公司的信用风险评估模型,这套模型的主要方法为神经网络.至此,研究者和实践者开始广泛关注神经网络这一方法,文献[17]将传统的参数和非参数方法和 5 种不同的神经网络算法(包括混合专家系统、失真适应响应和多层感知器等)进行了比较分析,其结果是神经网络的稳定性较好.

支持向量机(support vector machine, SVM)的核心思想是通过某种事先选择的非线性映射将输入向量映射到一个高维特征空间,在此空间中根据区域中的样本计算该区域的决策曲面,由此确定该区域中未知样本的类别.SVM 的出现解决了以往学习方法中存在的小样本、非线性、过学习、

高维数、局部极小等实际问题,在个人信用评分中,支持向量机方法评分精度较高,预测能力强,且受变量限制少,具有很强的泛化能力,因此支持向量机不仅在手写数字识别、文本分类、语音辨识等问题上得到了广泛应用,在个人信用评分领域也成为了研究的热点.文献[18]指出 SVM 算法能够更好的捕捉变量间的非线性关系,并在 SVM 的基础上提出了混合支持向量机算法,通过实证验证了混合支持向量机算法有着更高的精确度,并有效降低了第 2 误判率.

1.5 组合评分模型

正是考虑到上述的单一信用评分方法各有优势,由此引发了学者们对组合方法的尝试.文献[19]总结不同领域的大量相关研究,得出组合模型能够取得更高预测精度的结论,成为组合预测研究的一个里程碑.受此影响,同年《预测杂志》出版了一期组合预测的专刊,进一步激发了学者们对组合方法的热情.组合方法主要分为线性组合和非线性组合.其中权重的确定是问题的关键,权重的确定可分为固定权重和可变权重.到目前为止,比较常用的方法有简单平均法、胜出法、最优法和回归法.

近年来国内在个人信用评分组合方法的研究上也取得了不少成果.文献[20]提出基于贝叶斯算法的投票式组合模型的思想,选择 Logistic 回归、聚类分析和神经网络进行组合,既发挥了这些具有代表性的单一模型优势,同时减少了由于权重确定产生的误差.文献[21]指出现有信用评级中存在的问题,基于粗糙集算法对两个混合模型 FA-RS 和 MEPA-RS 模型进行了深入研究.

在实际应用中,个人信用评分模型选择的决定因素往往来自于多个方面,如线性统计学模型常被有一定历史的评分机构所应用,因为已有的技术比较根深蒂固,而且这些机构也倾向于使用那些已经被使用并通过实践检验的方法.Logistic 回归多被新建机构采用,那些为了防止严重的假设条件违背,或需要违约概率估计的借贷者(尤其是签订新巴塞尔协议的银行)也对其更加青睐.

2 个人信用评分模型应用中的问题

2.1 信用样本有效性及完整性问题

样本有效性是国外成熟的评分模型在我国信用数据中进行应用面对的首要问题.由于我国消费信贷业务发展较晚,信用体系尚未完善,现有的信用数据相当有限,且由于信用信息的提交和纰漏仍不规范,灰色收入等的存在,造成信用样本数据的权威性和有效性面临挑战.对于商业银

行而言,无法对每一位贷款的个体进行数据真实性考察,如何及时的发现信贷业务中存在的“异常数据”,摒弃冗杂的干扰数据,是目前个人信用评分领域需要研究的问题之一。

在信用样本的完整性上,已有的个人信用评分模型都面对着一个不可忽视的数据问题——样本偏差(biased sample).样本偏差来自于非随机性的样本获取过程,表现为样本和总体分布的非一致性,其本质是一种样本选择问题(sample selection).在个人信用评分上,样本偏差表现为拒绝推论(reject inference).拒绝推论就是指在个人信用评分的过程中,银行的评估模型是建立在已接受的信用样本之上,而缺少那些申请被拒绝的样本(拒绝样本)的相应数据.这就导致了银行的信用评分模型所用数据不是随机样本,不能代表整个申请者的“入门总体”(through-the-door population),从而导致评估的偏差.个人信用评分模型的准确性与模型采用的训练数据有着密切的关系,拒绝推论问题的存在也降低了评分模型的价值与精度。

常用的解决样本偏差的方法有外推法(extrapolation)、赋权法(enlargement)和重新赋权法(Re-weighting).外推法是利用已接受的样本建立初始信用评分模型,并用于被拒绝样本的判别,最后利用所有样本建立最终的评分模型.重新赋权法通常与增补法共同使用,通过对已接受的样本重新赋予权重来代表被拒绝的样本.但是,文献[22]认为以上方法都是针对随机性的样本缺失,在解决非随机性的拒绝推论问题时,效果并不理想。

2.2 信用指标体系合理性问题

信用评分指标体系的确定是个人信用评分的第一步,对整个信用评分的精确性及信用风险的有效识别至关重要.目前商业银行在个人信用评分中应用的指标有限且彼此不同.国内学者又偏向于对模型的优化与改进,对指标体系的研究较少,导致我国尚未建立有效、权威的指标体系.而我国的文化习惯和道德标准与国外相差较大,国内不同地区间经济发展水平、人口结构和生活方式,各民族间文化及道德标准也有着较大差异,这就导致同一指标在不同的实际应用中显著性有着较大的变化,因此针对不同的数据样本,对指标体系中的特征变量及变量的权重有所调整,充分适应实际业务需求十分必要.目前优化信用指标体系的方法主要是属性约简法,文献[23]通过 SVM 等方法对指标进行筛选,保留比较重要的指标,构建新的指标体系.但是属性约简的方法在个人信用评分中的应用效果并不理想,因为个人信用评

分指标体系中的指标数量较少,约简后所得的指标体系其有效性和代表性仍有待考证。

2.3 模型选择及适用性问题

目前,无论是学术研究还是商业银行的实践都致力于提高个人信用评分模型的精确性、稳定性及解释性,以便有效地进行风险识别并降低信用风险.但已有的模型各具优缺点.如判别分析法对数据有着较为苛刻的要求,要求信用样本数据服从正态分布,且要求自变量与因变量间存在线性相关关系,但它通过不同的变量组合来探求最小化的特定分离程度,具有良好的解释性;最近邻法不用提前学习和训练模型,从而允许动态的更改客户信息,在解决人口漂移问题上具有优势,如何选择距离公式和确定 k 个相近样本投票权重却是应用中的难点,且对于高维数据,其在样本空间中分布较为稀疏,绝大多数点附近根本没有样本点,导致方法很难使用;决策树法的优点在于能够充分的利用先验信息,受异常数据点影响较小,具有较高的分类精度,缺点则在于对特征属性的权重缺乏判断;传统的神经网络模型具有较高的预测精度但无法处理非数值型数据,而且对初始中心的选取及异常值十分敏感,训练中易于出现过度拟合.同时神经网络“黑箱”化特征决定了其不具解释性.综上所述,统计学模型可以提供假设检验,具有一定的解释性,但与人工智能方法相比,其精确度不高,对数据的要求比较严格;而人工智能方法则正好相反,精确度较高但解释性差.在实际应用的过程中,商业银行的信贷政策也在不断调整,如何根据商业银行的政策及业务需要进行模型选择是目前个人信用评分所面对的一个难题.针对该问题,文献[24]指出,在个人信用评分中应将对模型的研究与对信用评分实际应用的研究进行有效结合;文献[25]认为可以引入商业银行个人信用评分的错分代价(misclassification cost)作为模型选择的标准。

3 优化算法设计

针对上述个人信用评分研究中的问题,本文从样本有效性及完整性、指标体系的合理性及模型的适用性 3 个方面对个人信用评分模型进行优化。

针对信用样本有效性及完整性问题,本文提出通过对已有的样本进行相关性分析,提取样本各特征变量间的相关关系,作为预警规则,对新加入的样本进行识别,实现对异常数据的预警,并通过蒙特卡洛模型生成模拟样本,根据规则进行样本筛选,选取其中的“坏客户”样本进行样本补足。

针对指标体系合理性问题,由于统计学模型理

论基础丰富,解释能力强,稳健性良好,采用统计学模型能够输出个人信用评分指标的显著性,更有效的剖析影响个人信用的相关因素,因此,本文选取了 Fisher 判别分析、Logistic 回归、Probit 回归、多元线性回归 4 种常用的统计学模型,结合投票器的方法对影响个人信用的特征向量进行显著性排序;又由于人工智能方法的判别精度较高,能够有效的识别不良数据,因此,在显著性权重的计算上,采用步长遍历算法,以 BP 神经网络和支持向量机两种精度较高的个人信用评分模型的平均精度为判别标准,输出显著性权重,对个人信用评分指标体系进行显著性加权,提高指标体系的合理性和科学性。

针对模型的选择及适用性问题,本文设计模型选择器,选择器中包括目前个人信用评分中最具有代表性的 5 个模型:Logistic 回归、分类树、Bayes 网络、BP 神经网络和支持向量机,输出每个模型的精确度、稳健性、第 1 误判率、第 2 误判率及差异性作为模型选择的指标,根据实际应用的具体需求,输出适用的单一模型、同类别加强组合模型及差异性组合模型,具体算法设计如图 1 所示。

输入:个人信用评分样本数据;
商业银行政策业务需求选择。
输出:待判样本中“异常点”的输出;
适用的个人信用评分指标体系显著性排序及权重;
适用的个人信用评分模型。
1.“异常点”预警及样本补足。
a)对已有信用样本进行相关性分析及统计分析,提取指标间的相关关系及对应数学表达式;
b)待判样本异常数据预警:对新样本中各指标数据进行识别,若样本中某一指标数据与提取的规则不符,输出;
c)基于蒙特卡洛模型随机生成信用样本,按提取规则进行筛选剔除,将保留的“坏客户”样本加入样本集。
2.信用指标的显著性加权。
a)建立 Fisher 判别分析、Logistic 回归、Probit 回归、多元线性回归 4 种统计回归模型,输出指标显著性排序;
b)采用等权投票的方法决定最终指标显著性顺序 I;
c)设指标显著性权重向量为 $W_k = (w_1, w_2, \dots, w_n)$, 其中 $\sum_{i=1}^n w_i = 100$, 在权重可赋值范围内进行遍历,得到 W_k 所有可能的取值情况;
d)对信用样本中的特征值数据进行显著性加权处理,得到修正指标体系为 $I_k = I * W_k$;
e)采用加权处理后的信用样本,训练 BP 神经网络及支持向量机模型,以模型的平均精度为权重选择依据,选择权重值,输出指标体系显著性顺序及权重。
3.模型的选择与输出。
a)采用补充后的信用样本集及显著性加权后的指标体系建立 Logistic 回归、分类树、Bayes 网络、BP 神经网络、支持向量机五种个人信用评分模型;
b)计算模型的精确性、稳健性、第 1 误判率、第 2 误判率、差异度 5 个指标;
c)选择 3 组模型:精确性最大的模型及稳健性最优的模型;第 1 误判率最小及第 2 误判率最小的模型;差异度最大的两个模型,分别对 3 组模型进行线性组合,得到 3 个组合模型;
e)结合商业银行的政策及业务需求,考虑错分代价的存在,将组合模型与选择的单一模型进行比较,输出最符合商业银行需求的模型。

图 1 优化算法技术路线

4 结 论

1)对个人信用评分模型的发展进行了梳理总结,阐明了个人信用评分模型由统计学方法到非参数方法、运筹学方法再到人工智能方法的演变,同时由单一模型到组合模型的演进过程,指出了各种个人信用评分模型在实际应用中的优势及局限性。

2)结合个人信用评分模型的发展及最新动态,指出个人信用评分研究中仍存在样本有效性及完整性差、指标体系合理性有待提高、模型适用性不明确、难以选择等问题。

3)针对样本的有效性 & 完整性问题,本文以提升样本有效性及完整性、指标体系合理性及模型适用性为目标,通过规则提取及模拟样本的加入实现对我国个人信贷业务中存在的“异常数据”预警,在丰富样本集的同时使样本结构更接近于实际情况,优化样本结构;针对信用指标的合理性问题,本文选取解释性好的统计学模型,结合投票器和步长遍历算法对信用评分指标体系进行显著性加权,避免指标减少的同时充分体现重要的样本属性在评分中的作用;针对模型的选择与适用性问题,通过模型选择器的设计,分别设定不同的标准进行模型的组合,比较单一模型与组合模型,旨在为商业银行基于信贷政策目标选择最适用模型。

参考文献

- [1] THOMAS L C. A survey of credit and behavioural scoring: forecasting financial risk of lending to consumers [J]. International Journal of Forecasting, 2000, 16(2): 149-172.
- [2] DURAND D. Appendix B: Application of the Method of Discriminant Functions to the Good-and Bad-Loan Samples[M]. Cambridge, MA: NBER(Risk Elements in Consumer Instalment Financing, Technical Edition), 1941:125-142.
- [3] HAND D J, HENLEY W E. Statistical classification methods in consumer credit scoring: a review[J]. Journal of the Royal Statistical Society: Series A, 1997, 160(3): 523-541.
- [4] CHEN H, CHEN Y. A comparative study of discrimination methods for credit scoring[C]//Proceedings of the 2010 40th International Conference on Computers and Industrial Engineering (CIE). Piscataway, NJ: IEEE, 2010: 1-5.
- [5] SRINIVASAN V, KIM Y H. Credit granting: a comparative analysis of classification procedures [J]. Journal of Finance, 1987, 42(3): 665-683.

- [6] 姜明辉,王雅林,赵欣,等. k-近邻判别分析法在个人信用评估中的应用[J]. 数量经济技术经济研究, 2004, (2): 143-147.
- [7] HAR-PELED S, INDYK P, MOTWANI R. Approximate nearest neighbor: towards removing the curse of dimensionality[J]. Theory of Computing, 2012, 8(1): 321-350.
- [8] VERLEYSEN M, FRANÇOIS D. The Curse of Dimensionality in Data Mining and Time Series Prediction [M]. Berlin Heidelberg Springer: Computational Intelligence and Bioinspired Systems, 2005.
- [9] HOULE M E, KRIEGL H P, KRÖGER P, et al. Can shared-neighbor distances defeat the curse of dimensionality? [C]//Proceedings of the 22nd International Conference, SSDBM. Berlin Heidelberg: Springer, 2010: 482-500.
- [10] PORTER B W, BAREISS R, HOLTE R C. Concept learning and heuristic classification in weak-theory domains[J]. Artificial Intelligence, 1990, 45(1): 229-263.
- [11] 庞素琳, 巩吉璋. C5.0 分类算法及在银行个人信用评级中的应用[J]. 系统工程理论与实践, 2009, 29(12): 94-104.
- [12] CHENG E W L, CHIANG Y H, TANG B S. Alternative approach to credit scoring by DEA: evaluating borrowers with respect to PFI projects [J]. Building and Environment, 2007, 42(4): 1752-1760.
- [13] FREED N, GLOVER F. Applications and Implementation [J]. Decision Sciences, 1981, 12(1): 68-74.
- [14] NATH R, JACKSON W M, JONES T W. A comparison of the classical and the linear programming approaches to the classification problem in discriminant analysis[J]. Journal of statistical computation and simulation, 1992, 41(1/2): 73-93.
- [15] TALEBZADEH H, MANDUTIANU S, WINNER C F. Countrywide loan-underwriting expert system [J]. AI magazine, 1995, 16(1): 51-64.
- [16] WOLPERT D H. Stacked generalization [J]. Neural networks, 1992, 5(2): 241-259.
- [17] ZHANG R Q, HUANG Z S. Statistical inference on parametric part for partially linear single-index model [J]. Science in China Series A: Mathematics, 2009, 52(10): 2227-2242.
- [18] HUANG C L, CHEN M C, WANG C J. Credit scoring with a data mining approach based on support vector machines[J]. Expert Systems with Applications, 2007, 33(4): 847-856.
- [19] CLEMEN R T. Combining forecasts: A review and annotated bibliography [J]. International journal of forecasting, 1989, 5(4): 559-583.
- [20] 王雪. 投票式组合预测模型在个人信用评估中的应用研究[D]. 哈尔滨: 哈尔滨工业大学, 2011.
- [21] CHEN Y S, CHENG C H. Hybrid models based on rough set classifiers for setting credit rating decision rules in the global banking industry [J]. Knowledge-Based Systems, 2013, 39: 224-239.
- [22] GARCIA S, HAROU P, MONTAGNE C, et al. Models for sample selection bias in contingent valuation: Application to forest biodiversity [J]. Journal of Forest Economics, 2009, 15: 59-78.
- [23] BELLOTTI T, CROOK J. Support vector machines for credit scoring and discovery of significant features[J]. Expert Systems with Applications, 2009, 36(2): 3302-3308.
- [24] MARTIN N. Assessing scorecard performance: A literature review and classification[J]. Expert Systems with Applications, 2013, 40(16): 6340-6350.
- [25] LEE T S, CHEN I F. A two-stage hybrid credit scoring model using artificial neural networks and multivariate adaptive regression splines [J]. Expert Systems with Applications, 2005, 28(4): 743-752.

(编辑 张红)