

# 基于贝叶斯优化和融合模型的个人信用评价研究

贺新楠

河北工业大学理学院, 天津

收稿日期: 2025年5月9日; 录用日期: 2025年6月2日; 发布日期: 2025年6月10日

## 摘要

金融科技的迅速发展推动个人信贷市场规模扩大, 信用风险评估成为金融机构风险管控的核心环节。针对这一问题, 本研究提出了一种基于贝叶斯优化和Blending融合模型的信用评估方法。通过数据预处理和特征选择构建高质量训练数据基础, 超参数优化使用贝叶斯优化方法构建模型。为了评估模型性能, 选取Logistic回归、随机森林、LightGBM和XGBoost模型进行了对比实验, 其中Blending融合模型在所有模型中表现最佳。实验结果表明, 本文提出的 Blending融合模型在个人信用风险评估中具有明显优势, 能够有效整合不同模型的特点, 提供更精准的信用风险预测。此外, 贝叶斯优化在调参效率与模型性能上均优于网格搜索与随机搜索, 其全局寻优特性尤其适用于高维数据场景。

## 关键词

集成模型, 超参数优化, 融合模型, 信用风险

# Research on Personal Credit Evaluation Based on Bayesian Optimization and Fusion Models

Xinnan He

College of Science, Hebei University of Technology, Tianjin

Received: May 9<sup>th</sup>, 2025; accepted: Jun. 2<sup>nd</sup>, 2025; published: Jun. 10<sup>th</sup>, 2025

## Abstract

The rapid development of financial technology drives the expansion of personal credit market scale,

and credit risk assessment becomes the core link of risk control of financial institutions. To address this problem, this study proposes a credit assessment method based on Bayesian optimization and Blending fusion model. The high-quality training data base is constructed through data preprocessing and feature selection, and the hyperparameter optimization uses Bayesian optimization to construct the model. In order to evaluate the model performance, Logistic regression, random forest, LightGBM and XGBoost models are selected for comparative experiments, in which the Blending fusion model performs the best among all models. The experimental results show that the Blending fusion model proposed in this paper has obvious advantages in personal credit risk assessment, and can effectively integrate the features of different models to provide more accurate credit risk prediction. In addition, Bayesian optimization outperforms grid search and random search in terms of tuning efficiency and model performance, and its global optimization characteristics are especially suitable for high-dimensional data scenarios.

## Keywords

Integrated Models, Hyperparameter Optimization, Fusion Models, Credit Risk

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

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



Open Access

## 1. 绪论

伴随我国经济体制改革的深化与市场经济的成熟，金融领域逐步向个人开放多元化的借贷服务，业务种类与覆盖范围持续拓展。进入新世纪后，在普惠金融理念推广与数字技术赋能的共同作用下，叠加社会消费观念转型的驱动效应，我国信用消费市场规模呈现指数级增长态势。住房购置、汽车分期、教育融资及信用卡服务等消费信贷产品加速渗透，逐步成为居民金融活动的重要组成部分。鉴于该领域庞大的市场潜力与可观收益预期，消费信贷已然成为金融机构竞逐的核心业务板块[1]。值得注意的是，个人信贷业务凭借客群基数大、单笔额度低、风险分散性强等特征，在利率市场化进程中展现出强劲增长动能。相较于波动较大的企业信贷，个人业务的客户黏性更高且存款稳定性更强，尤其在宏观经济下行周期中更具利润平滑功能。对金融机构而言，如何构建科学高效的风险管理机制是信贷业务可持续发展的关键命题[2]。

在传统信用风险管理实践中，人工经验评估法长期占据主导地位。随着数据挖掘技术的发展，机器学习模型为信用风险评估带来了新方法。其中，逻辑回归和集成模型在该领域应用广泛。Chen 等[3]开发了 XGBoost 梯度增强框架，通过正则化项约束与并行计算优化，显著降低算法资源消耗。庞素琳[4]等创新性融合 C5.0 决策树与 Boosting 算法，有效提升商业银行个人信用评估的判别准确率。针对样本不平衡问题，杨海江等[5]改进 AdaBoost 算法的权重更新机制，实证表明该策略可降低误判损失达 18.6%。萧超武等[6]则构建随机森林集成分类器，其预测效能较 SVM、KNN 等单一模型提升 23% 以上。李学峰[7]基于违约数据实证发现，XGBoost 模型的分类性能超越 Logistic 回归与随机森林达 12.3 个百分点。

综上所述，个人信用风险评价领域的研究方法经历了从传统统计方法到机器学习集成模型，再到融合模型的演变过程。逻辑回归作为传统的统计方法，具有简单易懂、可解释性强等优点，但在处理复杂数据关系时存在局限；集成模型通过组合多个基础模型，有效提高了预测准确性和泛化能力，但在模型复杂度和可解释性方面仍有待改进。

## 2. 研究方法

### 2.1. Blending 融合模型

Blending 融合是在 Stacking 融合的基础上改进过后的算法。stacking 可以令融合本身向着损失函数最小化的方向进行，同时 stacking 使用自带的内部交叉验证来生成数据，可以深度使用训练数据，让模型整体的效果更好。但也存在一些问题，stacking 融合需要巨大的计算量，需要的时间和算力成本较高，以及 stacking 融合在数据和算法上都过于复杂，因此融合模型过拟合的可能性太高。针对 stacking 存在的这两个问题，学者提出了 Blending 方法。Blending 的核心思路其实与 Stacking 完全一致：使用两层算法串联，具有多个基学习器，有且只有一个元学习器，且基学习器负责拟合数据与真实标签之间的关系、并输出预测结果、组成新的特征矩阵，然后让元学习器在新的特征矩阵上学习并预测。

### 2.2. 贝叶斯优化

贝叶斯优化是一种基于概率模型的序列优化框架，其核心机制在于通过整合历史评估信息指导后续参数搜索策略。该算法通过建立目标函数的概率代理模型，将先验知识与迭代过程中获得的观测数据相结合，动态更新对目标函数形态的认知，从而有效缩小参数探索范围并提升搜索效率[8]。

## 3. 实验设置

### 3.1. 数据来源

本文采用 Lending Club 平台 2019 年发布的 P2P 借贷数据集进行实证分析。该数据集作为美国知名网络借贷平台的核心业务数据，在个人信用风险评估及违约预测研究领域具有重要的应用价值[9]。数据集完整记录了平台多年运营中的借贷交易信息。

### 3.2. 数据预处理

本文研究的是用户是否存在违约风险，关注的变量是贷款状态，具体情况如表 1 所示。

在接下来的分析中，将完全结清的用户定义为非违约客户用 0 表示，将坏账、逾期 31~120 天、处于宽限期和逾期 16~30 天用户定义为违约客户用 1 表示。使用上述方法把因变量转化成为了只含 0, 1 标记值的变量，使得问题转化为一个判断用户是否有违约风险的二分类问题。最终数据集中定义为 0 的非违约客户和定义为 1 的违约客户比例为 61840:24503，存在一定的样本不平衡。

**Table 1.** Loan\_status variable information

**表 1.** Loan\_status 贷款状态变量信息表

| Loan_status        | 贷款状态        | 样本数     |
|--------------------|-------------|---------|
| Current            | 还款中         | 431,763 |
| Fully Paid         | 完全结清        | 61,840  |
| Charged Off        | 坏账          | 13,487  |
| Late (31~120 days) | 逾期 31~120 天 | 6228    |
| In Grace Period    | 处于宽限期       | 3695    |
| Late (16~30 days)  | 逾期 16~30 天  | 1015    |
| Default            | 违约          | 78      |

由于数据集存在样本不平衡，而划分数据集需要特别注意以确保训练集和测试集都能合理地反映数据的分布，本文采用分层抽样来划分数据集。分层抽样是一种常用的划分方法，它确保每个类别在训练集和测试集中按比例表示。这有助于保持原始数据集中各类别的分布。

将数据集的 80% 作为训练集，20% 作为测试集。最终得到样本量为 69,074 的训练集和样本量为 17,269 的测试集。

为了构建符合贷前风控场景的信用评估模型并避免数据泄露，对数据集进行无效值处理：首先剔除贷后变量，因其包含贷款发放后的还款行为信息，引入未来变量将导致模型过拟合；其次删除与信用评价无关的冗余变量，其缺乏风险解释力；第三移除 Lending Club 内部评价指标，防止评级结果提前泄露目标变量信息；最后过滤唯一值变量，因其无预测区分度。通过上述处理共删除 25 个非相关或干扰性特征，数据集中剩余了 125 个变量。以上处理有效提升模型泛化能力与业务解释性，为后续构建信用评价模型奠定数据基础。

对于缺失值，考虑删除或使用适当的填充方法进行处理。本文以 0.8 为阈值，删除缺失率大于 0.8 的变量，使用随机森林模型填充缺失值小于 0.8 的变量。随机森林适用于连续型和分类型变量，能处理非线性关系，通过集成多棵树降低过拟合风险，对非线性关系和复杂交互作用建模能力强。

在特征工程处理中，分类型变量可根据其数学特性分为两类：顺序型变量和名义型变量。本研究采用不同的编码策略：对于顺序型特征，通过数值映射保留其序数信息；对于名义型特征，则采用独热编码方法处理。

### 3.3. 特征选择

本研究选用随机森林算法进行特征筛选。作为一种集成学习框架，随机森林通过构建多个决策树模型，能够有效评估各个特征对预测结果的贡献度，在分类和回归任务中均有广泛应用。

基于随机森林模型对数据集进行特征选择，根据特征重要性选取前 30 个特征。在实证分析中，目标变量通常受到多维特征的交互影响，然而特征间的统计关联性(如多重共线性)可能导致模型参数估计偏差并降低统计推断的可信度。因此，在通过随机森林模型筛选特征变量后，计算相关系数大于 0.8 的变量。在此基础上，根据特征重要性删除其中一个变量，优先保留信息量更高的特征。经过处理，最终得到 23 个特征变量。

### 3.4. 模型训练

采用 Blending 融合算法构建贷前违约预测模型，Blending 作为一种简单高效的融合技术，其核心思想是利用一个线性模型来组合多个基模型在验证集上的预测结果，进而生成最终预测。

Blending 融合模型的实现步骤如下：

#### (1) 数据集划分

将原始数据集分为训练集和验证集。通常，训练集用于训练基模型，验证集用于生成元特征。

#### (2) 训练基模型

选择多个不同的基模型，本文采用 Logistic 回归模型、随机森林模型、LightGBM 模型和 XGBoost 模型。使用训练集对每个基模型进行训练，其中随机森林、LightGBM 模型和 XGBoost 模型的超参数优化采用贝叶斯优化方法。

#### (3) 生成元特征

利用训练好的基模型对验证集进行预测，得到各模型的预测结果。将这些结果作为新特征，构成元特征矩阵，作为后续线性模型的输入。

#### (4) 训练线性模型

选择线性回归或逻辑回归等线性模型，本文选用逻辑回归模型。基于元特征矩阵和验证集标签进行训练。线性模型学习如何最优组合基模型预测，以降低预测误差。

#### (5) 模型预测

对于新的测试数据，先用各基模型预测，得到预测结果后输入到训练好的线性模型中，线性模型综合判断并输出最终预测。

### 3.5. 模型结果

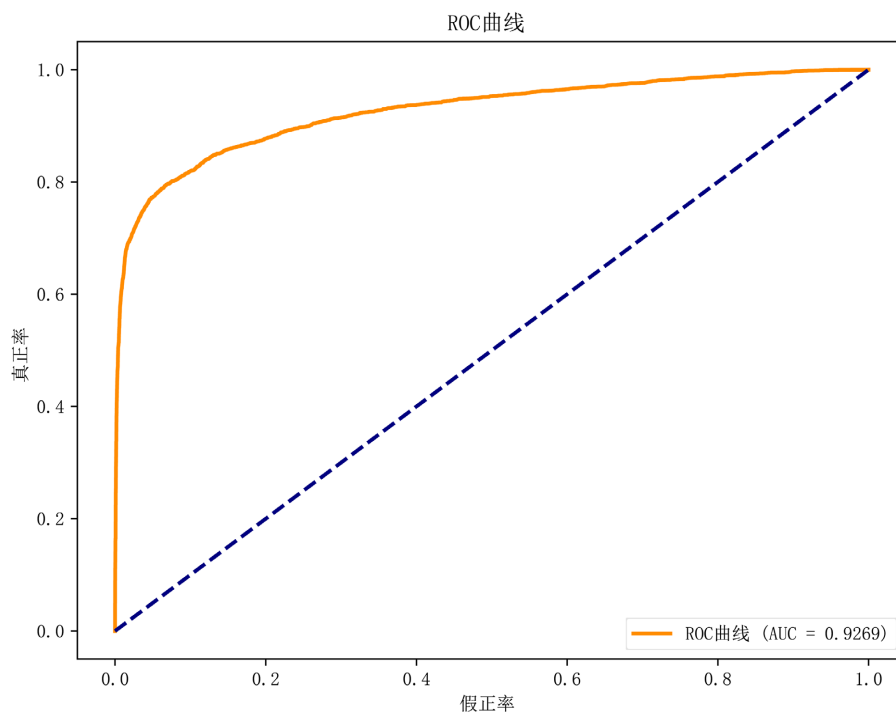
构建 Blending 模型后，可以得到其在测试集上的各项评估指标如表 2，由表 2 可知 Blending 模型负样本分类性能显著优于正样本，其准确率与召回率分别达到 0.92 和 0.94，这表明模型在识别低风险客户方面表现出色。正样本召回率(0.78)与准确率(0.84)降低，相比负样本，模型在正样本上的表现稍逊一筹，但仍处于可接受范围。

**Table 2.** Evaluation performance of Blending model on test set

**表 2.** Blending 模型在测试集上的评估效果

| 样本划分与评估指标 | 准确率  | 召回率  | 样本数目   |
|-----------|------|------|--------|
| 负样本       | 0.92 | 0.94 | 12,368 |
| 正样本       | 0.85 | 0.78 | 4901   |

绘制 Blending 融合模型在测试集上的 ROC 曲线和 KS 曲线分别如图 1 和图 2 所示，其中 AUC 为 0.9269，表示模型具有较高的分类能力，KS 统计量为 0.7272，说明模型具有很强的区分正负样本的能力。



**Figure 1.** Blending model ROC curve

**图 1.** Blending 模型 ROC 曲线

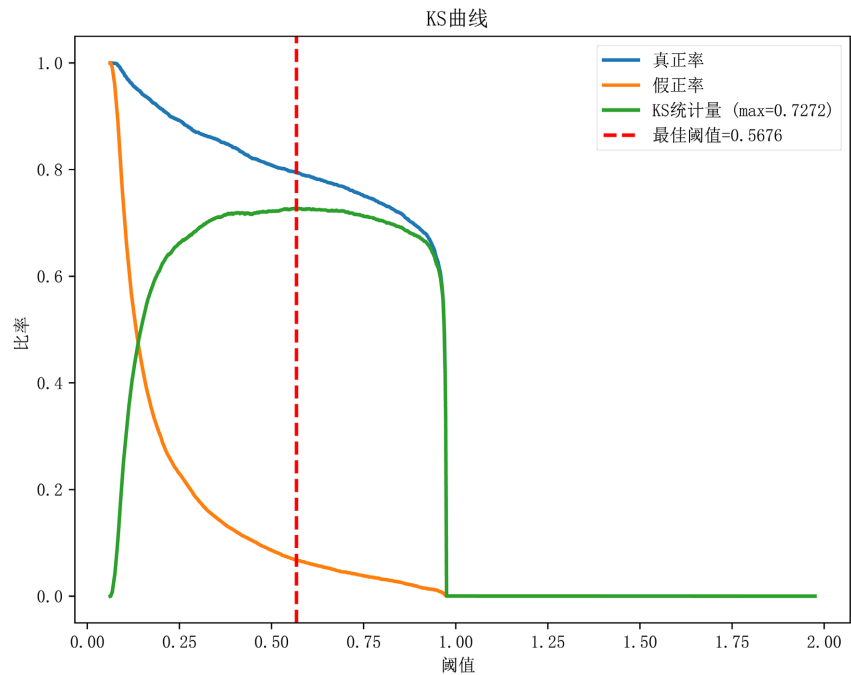


Figure 2. Blending model KS curve  
图 2. Blending 模型 KS 曲线

3.6. 模型分析

3.6.1. 基模型贡献度分析

贡献度分析需要根据每个基模型在元模型中的权重，基模型单独的性能，比如它们的 AUC、F<sub>1</sub> 分数等指标，表现好的模型可能贡献更大。

本文中基模型贡献度分析的实现步骤如下：

- (1) 查看元模型(逻辑回归)的系数，系数绝对值大小反映贡献度。
- (2) 分析基模型单独的性能指标，性能好的模型可能贡献更大。
- (3) 使用排列重要性方法评估每个基模型对元模型性能的影响。

Table 3. Comprehensive contribution evaluation  
表 3. 综合贡献度评估

| 模型        | AUC    | KS     | 权重     | 排列重要性  | 标准化权重  | 标准化排列重要性 |
|-----------|--------|--------|--------|--------|--------|----------|
| 逻辑回归      | 0.9186 | 0.7215 | 1.2700 | 0.0103 | 0.3949 | 0.2308   |
| 随机森林      | 0.9240 | 0.7262 | 2.3507 | 0.0301 | 0.9950 | 0.8859   |
| Light GBM | 0.9266 | 0.7266 | 2.3597 | 0.0336 | 1.0000 | 1.0000   |
| XGBoost   | 0.9259 | 0.7254 | 0.5589 | 0.0034 | 0.0000 | 0.0000   |

分析表 3 可知：

(1) 基模型性能指标

在这几个模型中，LightGBM 的 AUC 值为 0.9266 最高，说明其区分正负样本的能力相对最优；逻辑回归的 AUC 值 0.9186 最低。LightGBM 的 KS 值 0.7266 最高，显示其在区分正负样本分布上表现最佳。

## (2) 权重相关指标

权重反映在 **Blending** 模型中各基模型的相对重要程度。**LightGBM** 和随机森林的权重较高, 分别为 2.3597 和 2.3507, 说明在综合模型中它们的作用相对突出; **XGBoost** 的权重仅为 0.5589, 在模型集成中所占比重较小。标准化权重是对权重进行标准化处理后的值。**LightGBM** 标准化权重为 1.0000, 表明在标准化体系下它的权重相对最高, **XGBoost** 标准化权重为 0.0000, 说明其在标准化衡量下权重最低。

## (3) 特征重要性指标

排列重要性是通过打乱特征值并观察模型性能下降程度来衡量特征重要性。**LightGBM** 的排列重要性为 0.0336 最高, 意味着其特征对模型的重要程度相对较高; **XGBoost** 的排列重要性仅为 0.0034 最低。标准化排列重要性是排列重要性经过标准化后的结果。**LightGBM** 的标准化排列重要性为 1.0000, 在标准化体系下特征重要性最高, **XGBoost** 的为 0.0000 最低。

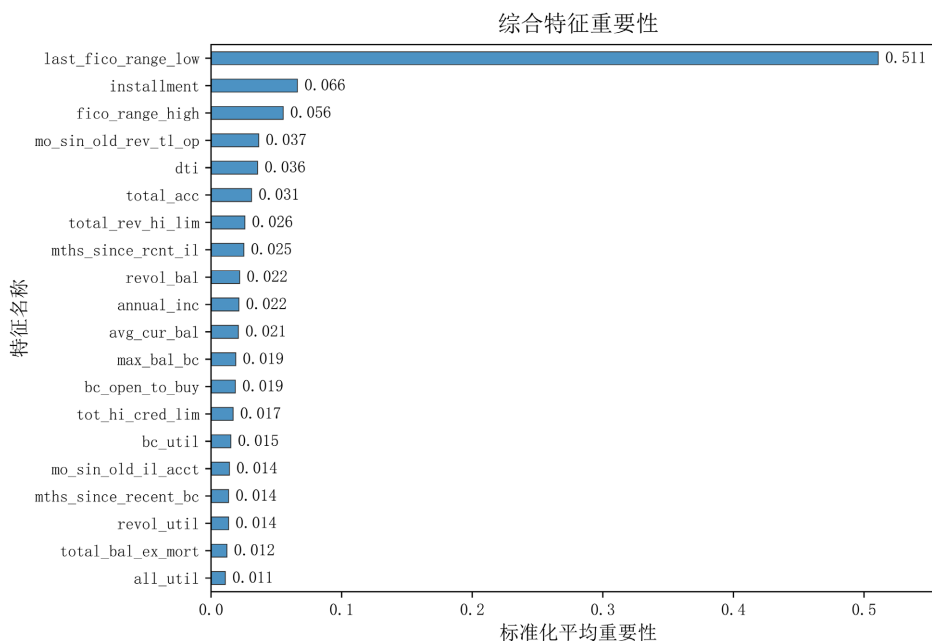
## (4) 综合分析

从各项指标综合来看, **LightGBM** 在模型性能(AUC、KS)、权重以及特征重要性方面均表现出色, 是对 **blendind** 模型贡献较大的基模型; 随机森林整体表现也较好, 尤其在权重方面突出; 逻辑回归各项指标表现相对较为均衡但都不突出; **XGBoost** 在多个关键指标上表现欠佳, 对 **Blending** 模型的贡献相对较小。

### 3.6.2. 特征重要性分析

分析特征影响需要分别看各个基模型的特征重要性, 或者结合元模型的结构来分析。由于元模型是逻辑回归, 只能看到各个基模型的预测结果作为输入特征, 而不是原始特征。所以需要从基模型入手, 分析各个基模型的特征重要性, 再进行综合分析。

对于每个基模型, 在训练完成后提取它们的特征重要性, 综合各个基模型的特征重要性来得到一个总体的认识。比如, 对每个特征, 计算它在不同基模型中的重要性排名或分数, 然后取平均或其他统计量。由于基模型性能差异较小(AUC 相差小于 0.05), 同时避免因单个模型异常影响整体判断, 直接取平均来计算综合特征重要性, 可视化结果如图 3。



**Figure 3.** The importance of comprehensive features in Blending models

**图 3.** Blending 模型综合特征重要性

由图 3 可知 last\_fico\_range\_low (借款人最近一次信用评分所属范围的下限)重要性最高, 在标准化平均重要性上远高于其他特征, 对模型结果影响显著, 可能在信用评估等场景起关键作用。Installment (若贷款发放, 借款人每月需偿还的金额)和 fico\_range\_high (贷款发放时借款人信用评分所属范围的上限)重要性次之, 在模型中也有一定影响力。其余特征重要性相对较低, 且彼此间差距较小, 对模型结果影响程度相对有限。

3.7. 模型对比分析

为了评估模型性能, 选取 Logistic 回归、随机森林、LightGBM 和 XGBoost 模型进行了对比实验。实验采用准确率、AUC 和 KS 值三项指标综合评价模型性能。表 4 为各个模型评估结果对比。

Table 4. Comparison of results from various models  
表 4. 各模型结果对比

| 模型                   | ACC    | AUC    | KS 统计量 |
|----------------------|--------|--------|--------|
| Logistic 回归(网格搜索)    | 0.8813 | 0.9185 | 0.7190 |
| Logistic 回归(贝叶斯优化)   | 0.8817 | 0.9186 | 0.7197 |
| 随机森林(随机搜索)           | 0.8815 | 0.9205 | 0.7189 |
| 随机森林(贝叶斯优化)          | 0.8842 | 0.9206 | 0.7226 |
| LightGBM (贝叶斯优化)     | 0.8840 | 0.9227 | 0.7220 |
| XGBoost (贝叶斯优化)      | 0.8835 | 0.9235 | 0.7216 |
| Blending 融合模型(网格搜索)  | 0.8890 | 0.9245 | 0.7208 |
| Blending 融合模型(贝叶斯优化) | 0.8969 | 0.9269 | 0.7272 |

从表 4 可以看出, 在个人信用风险评估中, 不同模型展现出不同的性能特点:

(1) 单一模型

Logistic 回归模型无论是通过网格搜索还是贝叶斯优化进行超参数调整, 其 ACC 分别为 0.8813 和 0.8817, AUC 为 0.9185 和 0.9186, KS 统计量分别为 0.7190 和 0.7197, 表现相对稳定, 但提升幅度较小, 说明在处理信用风险这类复杂问题时, 传统 Logistic 回归模型有一定局限性, 贝叶斯优化对其性能提升效果不明显。

随机森林模型采用随机搜索时 ACC 为 0.8815, AUC 为 0.9205, KS 统计量为 0.7189; 经贝叶斯优化后, ACC 提升至 0.8842, AUC 提升为 0.9206, KS 统计量升至 0.7226。这表明贝叶斯优化对随机森林在信用风险评估中的分类准确率有一定提升, 但在区分正负样本能力上略有波动, 整体性能有所改善。

LightGBM 和 XGBoost 模型中, LightGBM 的 ACC 为 0.8840, AUC 为 0.9227, KS 统计量为 0.7220; XGBoost 的 ACC 为 0.8835, AUC 为 0.9235, KS 统计量为 0.7216。两者在信用风险评估任务中表现相近, XGBoost 在 AUC 上略胜一筹, 说明其在排序能力上稍强, 但两者在准确率和区分度上基本持平, 均优于前两种模型。

(2) 融合模型

Blending 融合模型的 ACC 最高, 为 0.8969, AUC 为 0.9269, KS 统计量为 0.7272。Blending 融合模型在所有模型中表现最佳, 说明这种融合方式能更有效地整合不同模型的特点, 在信用风险评估中提供更可靠的预测结果。

同时，贝叶斯优化超参数与传统网格搜索和随机搜索相比，显著缩短了调参时间，并且实现了更优的模型性能。

综合来看，融合模型在个人信用风险评估中具有明显优势，尤其是 **Blending** 融合模型，能够为金融机构提供更精准的信用风险预测，帮助其更好地识别高风险客户，优化信贷决策。在实际应用中，建议优先考虑采用融合模型进行信用风险评估。同时，针对不同数据特点和业务需求，可进一步探索和优化融合策略，以持续提升模型性能。

## 4. 总结和展望

### 4.1. 总结

在普惠金融快速发展的背景下，个人信用风险评估作为金融风险管理的核心环节，正面临传统评估方法难以应对海量数据与复杂风险特征的挑战。本研究以 **Lending Club** 公开数据集为基础，系统性地探索了数据处理与建模技术在信用风险评估中的应用，旨在提升评估的准确性与可靠性，为金融科技领域的风险管理提供理论支持与实践指导。

传统逻辑回归模型因其线性假设和对特征独立性的依赖，在处理复杂数据和非线性关系时表现出明显的局限性。实验结果显示，逻辑回归模型在 **ACC**、**AUC** 和 **KS** 统计量等关键指标上表现相对稳定，但提升幅度有限。即使通过贝叶斯优化进行超参数调整，其性能提升效果仍不显著。这一结果表明，传统逻辑回归模型在应对信用风险评估这类复杂问题时，难以捕捉数据中的非线性特征和高阶交互关系。

相比之下，集成模型(如随机森林、**LightGBM** 和 **XGBoost**)在 **AUC** 和 **KS** 值等关键指标上显著优于传统逻辑回归模型。这表明集成模型通过结合多个弱学习器的优势，能够更有效地捕捉数据中的复杂特征和非线性关系，从而显著提升信用风险评估的准确性与可靠性。

融合模型通过整合多种单一模型的优势，进一步提升了信用风险评估的性能。实验结果表明，**Blending** 融合模型在所有模型中表现最佳，其 **ACC** 达到 0.8969，**AUC** 为 0.9269，**KS** 统计量为 0.7272。这一结果表明，**Blending** 融合模型能够更有效地整合不同模型的特点，尤其是在处理类别不平衡问题和复杂特征关系时，展现出更强的泛化能力和预测准确性。

这一发现对金融机构的实际应用具有重要意义。通过采用融合模型，金融机构可以更精准地识别高风险客户，优化信贷决策，从而降低违约风险并提升整体风险管理水平。同时，研究建议在实际应用中优先考虑融合模型，并根据具体数据特点和业务需求进一步优化融合策略，以持续提升模型性能。

将贝叶斯优化方法引入信用风险评估领域，与传统网格搜索和随机搜索相比，显著缩短了调参时间，同时实现了更优的模型性能。例如，随机森林模型在贝叶斯优化后，**ACC** 从 0.8815 提升至 0.8842，**KS** 统计量从 0.7189 提升至 0.7226，表明贝叶斯优化能够更高效地探索超参数空间，找到更优的参数组合。这为信用风险评估中的模型调优提供了高效的新思路，尤其适用于复杂数据环境下的模型优化。

综上所述，本研究通过系统性对比多种模型性能、引入贝叶斯优化方法以及构建融合模型，为个人信用风险评估提供了科学的方法论支持。研究结果验证了集成学习与超参数优化在信用风险评估中的有效性，为金融科技领域的风险管理实践提供了理论依据和实践指导。

### 4.2. 研究不足与展望

本研究提出了一种基于贝叶斯优化和 **Blending** 融合模型的信用评估方法。基于 **Lending Club** 平台发布的公开数据集，研究不仅开展了违约预测模型的构建与评估工作，还深入分析了超参数优化方法对模型预测性能的影响机制。尽管本研究通过模型对比与优化，验证了集成学习与超参数优化在信用风险评估中的有效性，但仍存在一些局限性。首先，研究数据来源于 **Lending Club** 数据集，它在跨国应用上存

在诸多限制。其局限性体现在多方面：样本选择偏差使其主要覆盖美国特定人群，难以代表低收入或信用记录薄弱群体；高度依赖美国特有的信用评估体系，如 FICO 评分；还受美国经济周期、利率政策及监管框架影响。信用体系、经济环境、文化行为以及法律监管等方面的差异，都使得模型在其他国家可能失效。为解决模型跨国适用问题，需采取一系列调整策略。在特征工程上，要适配本地指标，引入宏观经济变量；利用数据增强与迁移学习，通过预训练和调整模型或者生成合成数据来应对；采用领域适应技术减少分布差异，用工具提升模型解释性；进行本地化验证与迭代优化模型。综上，Lending Club 数据集虽对美国 P2P 借贷研究价值高，但跨国适用性受限。实现模型跨国迁移需先诊断差异，再适应性改造，最后动态优化，兼顾全球经验与本地洞察，不能将其视为通用标准。其次，模型的可解释性不足，尤其是在融合模型中，复杂的元学习机制增加了模型解释的难度。未来研究可从以下几个方向展开深入探索：

#### (1) 多源数据融合

整合来自不同金融机构和非传统数据源(如社交媒体、交易记录)的信息，以丰富特征集并提升模型的泛化能力。

#### (2) 深度学习技术应用

探索深度学习模型在信用风险评估中的应用潜力，尤其是在处理高维数据和复杂特征关系时的优势。

#### (3) 动态信用评估系统

开发能够实时更新和适应市场变化的动态信用评估系统，以应对快速变化的金融环境。

#### (4) 模型可解释性增强

通过引入 SHAP 值、LIME 等解释性工具，提升模型的透明度和可解释性，满足金融监管和实际应用的需求。

## 参考文献

- [1] 韩国栋. 民间借贷正规化路径的选择模型研究[J]. 财政研究, 2012, 28(6): 16-19.
- [2] 段辉娜, 王雪梅, 孙敬怡. 互联网消费金融对居民消费行为的影响研究[J]. 商业经济研究, 2020(7): 48-52.
- [3] Chen, T. and Guestrin, C. (2016) XGBoost. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco, 13-17 August 2016, 785-794.  
<https://doi.org/10.1145/2939672.2939785>
- [4] 庞素琳, 巩吉璋. C5.0 分类算法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.
- [5] 杨海江, 魏秋萍, 张景肖. 基于改进的 AdaBoost 算法的信用评分模型[J]. 统计与信息论坛, 2011, 26(2): 27-31.
- [6] 萧超武, 蔡文学, 黄晓宇, 等. 基于随机森林的个人信用评估模型研究及实证分析[J]. 管理现代化, 2014(6): 111-113.
- [7] 李学锋. 基于 XGBoost 的个人信贷违约预测研究[J]. 电脑知识与技术, 2019, 15(33): 192-194.
- [8] 崔佳旭, 杨博. 贝叶斯优化方法和应用综述[J]. 软件学报, 2018, 29(10): 3068-3090.
- [9] 李雪静. 国外 P2P 网络借贷平台的监管及对我国的启示[J]. 金融理论与实践, 2013(7): 101-104.