

基于ICSA优化WELM的不平衡条件下财务风险预测研究

郭宇新, 刘媛华

上海理工大学管理学院, 上海

收稿日期: 2025年9月29日; 录用日期: 2025年11月24日; 发布日期: 2025年12月1日

摘要

企业财务风险预测对维护市场稳定和投资者权益具有关键作用,但在实际数据中存在类别分布不平衡的挑战,传统预测模型往往对少数类(如ST企业)的识别能力有限。为提高不平衡数据下的预测性能,本文提出一种利用改进乌鸦搜索算法(ICSA)优化加权极限学习机(WELM)的方法。通过动态调整参数和引入莱维飞行机制、多个体加权学习策略和差分进化算法提升算法搜索效率,并结合加权极限学习机对少数类样本进行识别。以2022~2024年A股上市公司为研究对象,采取交叉验证的方法对模型进行性能评估。实验结果表明,本文所提出模型在准确性、召回率、G-mean等关键指标上优于对比模型,能够比较有效地识别ST企业,为财务风险预测提供了更可靠的技术手段。

关键词

财务风险预测, 改进乌鸦搜索算法, 加权极限学习机, 不平衡数据

Research on Financial Risk Prediction under Imbalanced Conditions Based on ICSA-Optimized WELM

Yuxin Guo, Yuanhua Liu

Business School, University of Shanghai for Science and Technology, Shanghai

Received: September 29, 2025; accepted: November 24, 2025; published: December 1, 2025

Abstract

The prediction of corporate financial risk plays a critical role in safeguarding market stability and

protecting investor interests. However, real-world datasets often present challenges related to class imbalance, wherein traditional prediction models tend to exhibit limited capability in identifying minority classes, such as Special Treatment (ST) enterprises. To enhance predictive performance under imbalanced data conditions, this paper proposes a method that optimizes the Weighted Extreme Learning Machine (WELM) using an Improved Crow Search Algorithm (ICSA). By dynamically adjusting parameters and incorporating mechanisms such as Lévy flight, a multi-individual weighted learning strategy, and a differential evolution algorithm, the search efficiency of the algorithm is significantly improved. This optimized approach is integrated with WELM to enhance the identification of minority class samples. Using A-share listed companies from 2022 to 2024 as the research sample, cross-validation is employed to evaluate model performance. Experimental results demonstrate that the proposed model outperforms baseline models in key metrics including accuracy, recall, and G-mean, showing effective identification of ST enterprises and providing a more reliable technical solution for financial risk prediction.

Keywords

Financial Risk Forecast, Improved Crow Search Algorithm, Weighted Extreme Learning Machine, Unbalanced Data

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. 引言

在复杂多变的经济环境中,企业的财务健康状况是衡量其生存能力与发展潜力的重要依据。作为企业管理与投资决策中的核心环节,财务风险预测工作的主要目的就在于通过科学的量化分析,提前识别企业潜在的财务困境,从而为风险防控提供合理的理论支撑和实践指导。在资本市场中,有效的财务风险预警不仅关系到企业的可持续发展,更直接影响着投资者的资产安全和整体市场资源的合理分配。然而,财务风险的形成往往具有多因素交织、非线性演变的特征,导致传统预测方法在准确性和时效性上面临严峻挑战[1]。

从国内外研究来看,财务风险预测模型经历了从传统统计方法到机器学习方法的演进。早期的研究主要依赖统计学方法,例如逻辑回归[2]和多变量分析[3]等模型,这些方法的优点是计算简单、解释性强,但难以处理财务指标与风险状态之间复杂的非线性关系。近些年来,随着机器学习成为主流研究方法,越来越多的学者将群智能优化算法[4] [5]、神经网络[6]、支持向量机[7]、深度学习[8]等方法应用到财务风险预测领域。在机器学习方法中,群智能优化算法结合神经网络的预测模型显示出较好的应用潜力[9]。群智能优化算法通过模拟自然界中群体的智能行为实现参数寻优,并将结果用于优化神经网络、支持向量机等预测模型的关键参数,从而提高了分类精度[10]。然而,很多研究者在预测过程中会避免或降低财务数据的不平衡比例,因此在实际应用中对少数类(ST企业)的识别依旧存在问题。

目前对财务风险预测的研究主要面临以下两点困境:首先,财务数据具有典型的高维特性,各类财务指标间存在复杂的相关关系,传统统计方法难以有效捕捉其非线性规律[11]。其次,风险企业(如ST公司)在样本中的占比通常不足5%,导致训练数据呈现严重的类别不平衡分布[12]。这种数据分布特点会使常规的机器学习模型过度拟合多数类样本,从而降低对风险企业的识别灵敏度。

针对上述问题,本研究提出将改进乌鸦搜索算法与加权极限学习机结合的思路处理在不平衡条件下的财务风险预测问题。乌鸦搜索算法(Crow Search Algorithm, CSA)因其参数简单、收敛速度快、全局搜索

能力强等特点, 被广泛应用于特征选择和模型优化[12]。但 CSA 容易陷入局部最优或搜索效率偏低等问题[13], 因此本文借鉴文献[14]对该算法提出了改进方法, 以提高算法的优化效果。而极限学习机(Extreme Learning Machine, ELM)作为单隐层前馈神经网络的快速学习算法, 因其训练速度快、泛化性能好等特点, 在财务风险预测中获得广泛应用[15]。加权极限学习机(Weighted Extreme Learning Machine, WELM)在 ELM 的基础上引入了权重, 通过给少数类样本更高的权重, 强制模型更加关注这些关键样本, 从而提高对 ST 企业的预测能力。二者结合改进的策略不仅能克服单一算法的性能局限, 更能形成协同优化效应, 为处理财务风险预测中的高维、非线性、不平衡数据提供新的方法论支持。

2. 基本原理介绍

2.1. 乌鸦搜索算法原理

乌鸦搜索算法(Crow Search Algorithm, CSA)是受自然界乌鸦觅食行为启发提出的一种群智能优化算法[16]。研究发现, 乌鸦具有出色的空间记忆能力, 能将多余食物储存在不同位置并长期记忆这些藏匿点。与此同时, 它们还会通过观察同伴行为来获取食物藏匿信息, 并伺机窃取其他乌鸦储存的食物。基于这种智能行为机制, CSA 算法构建了包含位置更新与记忆更新两个核心模块的数学模型。

在算法中, 优化问题的解空间通过 d 维搜索空间进行表征, 其中每只乌鸦所处的位置代表一个潜在解。假设乌鸦的种群规模为 N , 最大迭代次数为 MIT 。第 i 只乌鸦在 d 维空间的位置可以表示为一个向量, 代表当前的一个解:

$$x_i^{gen} = [x_{i,1}^{gen}, x_{i,2}^{gen}, \dots, x_{i,d}^{gen}], \quad i = 1, 2, \dots, N; gen = 1, \dots, MIT \quad (1)$$

在每次迭代过程中, 每只乌鸦会记忆其已知的最佳藏食位置 m , 该位置代表了当前个体所发现的最优解。为寻求更优解, 乌鸦在搜索空间中通过跟踪同伴的方式探索潜在的更好的位置。具体来说就是, 算法在每一轮迭代中, 将执行以下位置更新机制:

1) 跟随阶段: 在这个阶段, 乌鸦 i 将随机选择种群中另一只乌鸦 j 进行跟踪, 试图找到乌鸦 j 的食物藏匿点。位置更新公式如下:

$$x_i^{gen+1} = \begin{cases} x_i^{gen} + r_i \times FL_i^{gen} \times (m_j^{gen} - x_i^{gen}), & \text{if } r_j \geq AP_j^{gen} \\ \text{a random position,} & \text{otherwise} \end{cases} \quad (2)$$

其中, x_i^{gen} 是乌鸦 i 在第 gen 次迭代时的位置; x_i^{gen+1} 是乌鸦 i 在第 $gen+1$ 次迭代时的新位置; m_j^{gen} 是乌鸦 j 在第 gen 次迭代时已知的最佳位置(即其藏匿点); r_i 和 r_j 是范围在 $[0, 1]$ 内的均匀分布的随机数; FL_i^{gen} 是乌鸦 i 在第 gen 次迭代时的飞行长度, 这个参数控制着搜索的步长; AP_j^{gen} 是乌鸦 j 在第 gen 次迭代时的感知概率, 它决定了被跟踪的乌鸦 j 是否能够意识到自己被跟踪, 从而采取反制措施(即让跟踪者飞到一个随机位置)。

2) 记忆更新阶段: 在乌鸦移动到新位置后, 需要更新其记忆(即已知的最佳藏匿点)。记忆更新公式如下:

$$m_i^{gen+1} = \begin{cases} x_i^{gen+1}, & \text{if } fit(x_i^{gen+1}) \text{ is better than } fit(m_i^{gen}) \\ m_i^{gen}, & \text{otherwise} \end{cases} \quad (3)$$

其中, m_i^{gen} 是乌鸦 i 在第 gen 次迭代时的记忆(历史最佳位置); $fit(\bullet)$ 是适应度函数。对于最小化问题, 如果新位置的适应度值更小(更好), 则更新记忆。

2.2. 极限学习机原理

极限学习机(Extreme Learning Machine, ELM)作为一种高效的单隐层前馈神经网络学习算法, 以其极

快的训练速度和良好的泛化能力受到广泛关注[17]。其核心思想在于随机初始化输入层与隐层之间的连接权重和阈值, 且在训练过程中无需迭代调整这些参数, 仅通过最小二乘法求解输出层权重, 从而大幅降低计算复杂度。ELM 网络结构如图 1 所示。

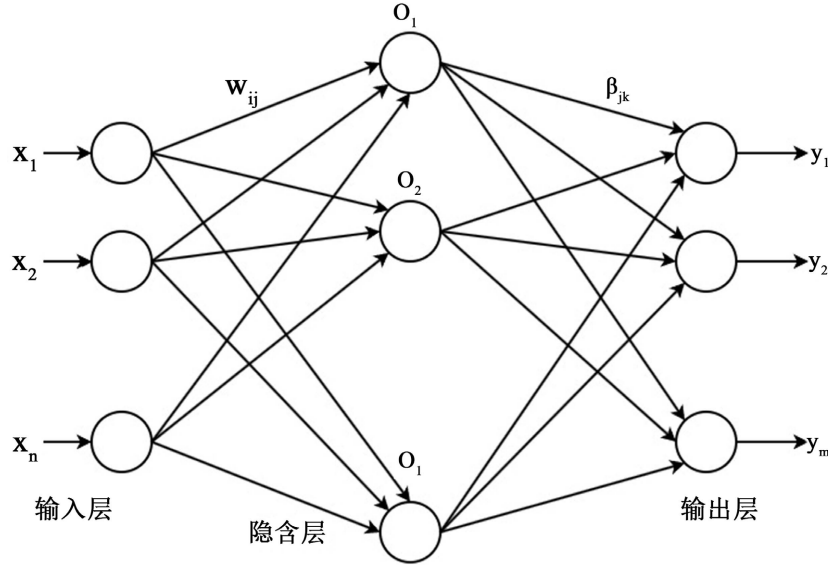


Figure 1. ELM network structure diagram

图 1. ELM 网络结构图

对于一个具有 N 个样本的回归问题, ELM 的数学模型可以表示为:

$$\sum_{i=1}^L \beta_i g(w_i \cdot x_j + b_i) = \hat{y}_j, \quad j=1, \dots, N \quad (4)$$

其中, x_j 是输入样本, y_j 是对应的输出, β_i 是输出权重, w_i 是输入权重, b_i 是阈值, $g(\bullet)$ 是激活函数, L 是隐层节点数。

为了方便矩阵运算也可以写成:

$$H\beta = Y \quad (5)$$

其中, $X = [x_1, x_2, \dots, x_N]^T \in \mathbb{R}^{N \times n_{in}}$ 是输入样本矩阵; $Y = [y_1, y_2, \dots, y_N]^T \in \mathbb{R}^{N \times n_{out}}$ 是目标输出矩阵;

$W = [w_1, w_2, \dots, w_L]^T \in \mathbb{R}^{n_{in} \times L}$ 是输入层到隐层之间的权重矩阵; $b = [b_1, b_2, \dots, b_L]^T \in \mathbb{R}^{1 \times L}$ 是隐层偏置向量; $g(\bullet)$ 是隐层激活函数; $H = g(XW + b)$ 是隐层输出矩阵; $H \in \mathbb{R}^{N \times L}$ 。

$\beta = [\beta_1, \beta_2, \dots, \beta_L]^T \in \mathbb{R}^{L \times n_{out}}$ 是隐层到输出层之间的权重矩阵。

ELM 的目标是最小化训练误差, 通常是通过最小化均方误差(MSE)来完成:

$$\min_{\beta} \|H\beta - Y\|^2 \quad (6)$$

2.3. 加权极限学习机原理

传统 ELM 在处理分类问题时, 通常假设训练数据集是平衡的, 即各类别样本数量大致相当。当面对现实世界中普遍存在的不平衡数据时, 例如财务风险预警中 ST 公司(少数类)与非 ST 公司(多数类)样本数量差异巨大的情况, 传统 ELM 的分类决策边界会倾向于多数类, 导致对少数类样本的识别率(如召回

率)下降。加权极限学习机(WELM)在 ELM 的基础上引入了样本权重,使得模型在训练过程中能够更关注某些样本。这可以通过在最小化损失函数时引入一个样本权重矩阵来实现。WELM 的核心思想是将 ELM 中的均方误差最小化问题转化为加权均方误差最小化问题。则改动后的数学表达式如下:

$$\min_{\beta} \sum_{i=1}^N w_i \|h_i \beta - y_i\|^2 \quad (7)$$

其中, $w_i \geq 0$ 是第 i 个样本的权重; h_i 是第 i 个样本在隐层输出矩阵 H 中的第 i 行, 即 $H \in \mathbb{R}^{N \times L}$ 。 y_i 是第 i 个样本的目标输出。

若将加权均方误差最小化问题改写为矩阵形式。引入一个对角加权矩阵 $W_{diag} \in \mathbb{R}^{N \times N}$, 其对角线元素为样本权重 w_i , 其余元素为 0, 则 WELM 的矩阵形式如下:

$$\min_{\beta} Tr\left((H\beta - Y)^T W_{diag} (H\beta - Y)\right) \quad (8)$$

其中, $Tr(\bullet)$ 表示矩阵的迹。

3. 乌鸦搜索算法改进策略

相较于遗传算法、粒子群等优化算法, CSA 仅需设置飞行长度(flight length, FL)和感知概率(Awareness Probability, AP)两个参数, 但也因此, 标准 CSA 存在两个固有缺陷: 一是固定 AP 值设置难以平衡全局探索与局部开发的关系, 在优化高维财务数据时易陷入局部最优; 二是传统位置更新策略对种群多样性保护不足, 在处理非凸、多峰的财务风险预测问题时搜索效率偏低[5]。这些局限性导致算法在优化不平衡数据分类器时, 难以有效提升对少数类样本的识别精度。综上, 本文参考相关文献[14]对乌鸦搜索算法提出了以下的改进:

3.1. 基于参数动态调整策略的乌鸦位置更新优化方法

在标准 CSA 算法基础上, 引入感知概率(AP)动态调整机制。通过构建 AP 值动态递变函数, 使算法在迭代前期保持较大 AP 值以增强全局探索能力, 随着迭代进程逐渐减小 AP 值以强化局部开发能力。该方法的数学表达如下:

$$AP = AP_{\max} - \left(AP_{\max} \times \frac{gen^3}{gen_{\max}^3} \right) \quad (9)$$

其中, gen 是迭代次数。

在全局搜索阶段引入莱维飞行机制以优化搜索行为。位置迭代过程中, 当乌鸦 j 感知到被跟踪时(即在 $r < AP$ 状态下), 传统 CSA 算法中乌鸦 j 会随机选择飞行方向, 即执行全局盲目搜索。采用莱维飞行策略对搜索路径进行引导, 通过其特有的短步长与偶尔长距离跳跃相结合的搜索特性, 既保持全局探索能力, 又有效提高收敛效率。相应数学表达如下:

$$x_i^{gen+1} = x_i^{gen} \times \left(1 + \alpha \times r_{\alpha} \times \sigma^{\gamma^{1/s}} \right), \quad \text{if } r_j < AP_j^{gen} \quad (10)$$

$$\sigma = \sqrt{\frac{\Gamma(1+s) \times \sin(\pi s/2)}{\Gamma(1+s/2) \times s \times 2^{(s-1)/2}}} \quad (11)$$

其中, α 是步长缩放因子, 控制乌鸦 i 搜索范围; r_{α} 是区间(0, 1)区间内的随机数; γ 、 σ 服从标准正态分布; $\Gamma(x) = (x-1)!$; s 是取值范围在[1] [2]之间的常数, x_i^{gen+1} 为乌鸦 i 在第 $gen + 1$ 次迭代下的个体最优藏食位置。

在局部搜索阶段采用多个体变因子加权的协同学习机制。在位置更新过程中, 当乌鸦 j 未感知到被跟踪时(即在 $r > AP$ 状态下), 传统算法中乌鸦 i 仅单向跟随乌鸦 j 进行局部开发, 加入多个体变因子加权学习方法, 使个体能够同时向种群中多个优质个体学习, 有效增强种群多样性并抑制早熟收敛。该方法的数学表达如下:

$$x_i^{gen+1} = x_i^{gen} + r_i(1, d) \times FL_i^{gen} \times (\lambda^{gen} m_j^{gen} + (1 - \lambda^{gen}) b^{gen-1} - x_i^{gen}), \quad \text{if } r_j \geq AP_j^{gen} \quad (12)$$

$$\lambda = \lambda_{\max} - \frac{gen \times (\lambda_{\max} - \lambda_{\min})}{gen_{\max}} \quad (13)$$

其中, $r_i(1, d)$ 是区间 $[0, 1]$ 之间的 d 维随机变量, λ^{gen} 是第 gen 次迭代时的加权学习因子, b^{gen-1} 是第 $gen-1$ 代种群的最佳藏食位置, r_j 为乌鸦 j 在区间 $(0, 1)$ 之间的随机数。

3.2. 基于差分进化算法的个体最优藏食位置更新方法

在位置更新阶段, 采用 CSA 与差分进化算法(DE)的混合优化策略。该方法先通过 CSA 的全局搜索和局部开发完成初步位置更新, 然后引入 DE 算法对当前种群进行变异、交叉操作, 最后通过选择机制保留解空间中的最优个体。这种混合策略有效扩展了种群的搜索潜力, 显著增强了算法的整体优化性能。

差分进化算法[18]作为一类基于群体智能的元启发式优化方法, 其求解过程首先在解空间内随机初始化种群, 随后通过循环执行变异、交叉与选择三种核心操作驱动种群进化, 直至满足终止条件或获得满意解。标准差分进化算法中各操作的数学表达如下:

$$v_i^{gen} = x_{r_1}^{gen} + F \times (x_{r_2}^{gen} - x_{r_3}^{gen}) \quad (14)$$

其中, v 是变异得到的解, x 是种群中的解, $r_1 - r_3$ 是 $1 - NP$ 中的 3 个互不相等的随机整数, gen 是迭代次数, F 是变异概率。

$$z_{i,d}^{gen} = \begin{cases} v_{i,d}^{gen}, & \text{if } rand \leq CR \\ x_{i,d}^{gen}, & \text{otherwise} \end{cases} \quad (15)$$

其中, $rand$ 是 $(0, 1)$ 的随机数, CR 是交叉概率。

$$x_i^{gen+1} = \begin{cases} z_i^{gen}, & \text{if } f(z_i^{gen}) \leq f(x_i^{gen}) \\ x_i^{gen}, & \text{otherwise} \end{cases} \quad (16)$$

其中, $f(z)$ 是种群个体解 z 的适应度函数值。

3.3. 小结

改进乌鸦搜索算法(ICSA)与加权极限学习机(WELM)融合模型的核心在于构建一个高效的协同优化框架, 旨在通过参数联合调优提升模型在不平衡财务数据下的预测精度。该融合机制主要包含两个层面: 一是利用 ICSA 强大的全局与局部搜索能力对 WELM 的关键参数进行优化; 二是通过 ICSA 自身的改进策略确保优化过程能够有效应对不平衡数据带来的挑战。

ICSA 通过 AP 值动态调整机制, 在迭代过程中自适应地平衡全局探索与局部开发能力。这一特性使其能够有效地在复杂的参数空间中为 WELM 寻找最优的输入权重和阈值。在优化过程中, ICSA 将 WELM 在验证集上的性能指标作为其适应度函数, 直接以提升对少数类(ST 公司)的识别能力为目标进行搜索。莱维飞行策略的引入增强了算法跳出局部最优的能力, 有助于发现更优的参数组合, 而多个体变因子加

权学习与差分进化算法促进了种群信息的交流与继承, 加速收敛并维持种群的多样性, 防止早熟。这种协同机制共同确保了优化过程的鲁棒性和高效性。

这种协同优化流程的理论依据在于, WELM 的性能对其参数设置极为敏感, 尤其是在处理高维、非线性的财务指标数据时。传统的网格搜索或随机搜索方法计算成本高昂且难以保证在不平衡数据上获得最佳泛化性能。ICSA 作为一种高效的元启发式算法, 通过模拟乌鸦的智能觅食行为, 能够以较高的效率逼近全局最优解。ICSA 输出的最优参数配置使 WELM 能够构建一个决策边界更有利于少数类样本的分类器, 从而在不平衡的企业财务数据上实现更高的预测精度和更强的泛化能力。该融合模型有效地将智能优化算法的搜索优势与加权极限学习机的快速学习能力相结合, 为处理财务风险预测中的类别不平衡问题提供了一种有效的解决方案。

4. ICSA-WELM 财务风险预测模型构建

企业财务风险预测模型求解流程如图 2 所示。

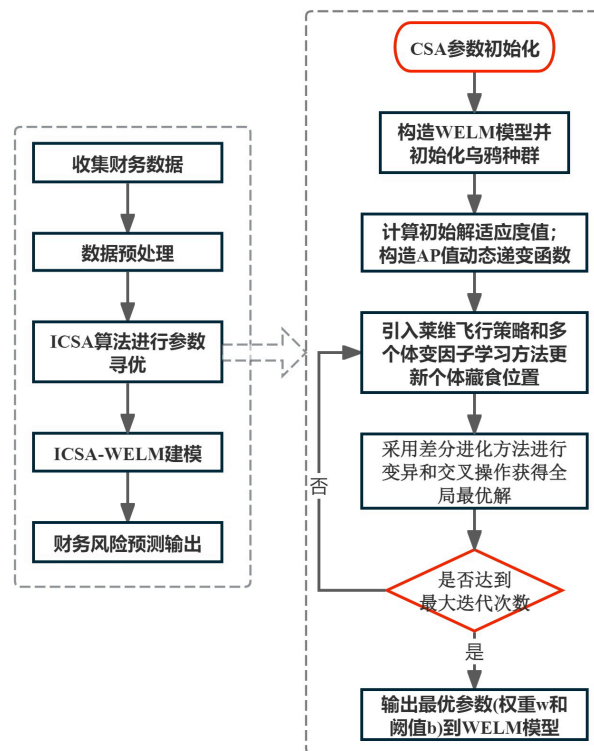


Figure 2. Flowchart for solving the financial risk prediction model based on ICSA-WELM

图 2. 基于 ICSA-WELM 的财务风险预测模型求解流程图

由图 2 可知, 本文提出的 ICSA-WELM 财务风险预测模型的核心在于利用改进乌鸦搜索算法(ICSA)对加权极限学习机(WELM)的输入权重 w 与阈值 b 进行优化, 将二者视为 ICSA 中的“藏食位置”, 实现参数全局最优取值。模型的主要步骤如下:

随机生成 N 个初始解(即乌鸦初始位置), 在训练集上构建 WELM 模型进行训练, 计算各解的初始适应度值与最优位置, 并依据递变规则构造 AP 值动态递变函数, 实现 ICSA 参数的迭代更新。

引入莱维飞行搜索策略与多个个体变因子加权学习方法, 使个体在跟随群体获取最优藏食位置的同时,

结合父代信息生成自身候选位置, 进而计算当前种群适应度, 检验新位置可行性, 并确定当前迭代最优解。

采用差分进化方法对当前种群进行变异与交叉操作, 通过选择机制保留解空间中的优秀个体, 从而获得全局最优位置。

当算法达到预设最大迭代次数时, 将所获得的最优藏食位置解码为 WELM 的输入权重与阈值, 用于最终训练 WELM 模型, 并在测试集上完成财务风险预测任务。

5. 实证分析

5.1. 数据选择与预处理

5.1.1. 样本选取与数据来源

本文选取 2022 年至 2024 年中国上市 A 股公司作为研究对象, 财务指标数据主要来源于 CSMAR 数据库。研究样本涵盖沪深两市所有 A 股上市公司, 但剔除了金融类企业和数据缺失严重的公司, 以确保数据的可比性和完整性。

根据以往学者的研究, 在划分企业的财务状况时, 通常根据企业是否被证监会特别处理(Special Treatment, ST)分为财务危机组和财务健康组。根据《上海证券交易所股票上市规则新规》, 被标记为 ST 或 *ST 的企业称为危机组。如此当企业出现轻度财务风险时, 可以及时采取干预措施, 避免发生危机。根据我国上市公司信息披露监管要求, 企业是否在 t 年度被实施特别处理, 取决于其在 $t-1$ 年度披露的财务报告所反映的经营与财务状况, 且以连续两年亏损而被标为 *ST 作为上市公司出现退市风险的标志。也就是说, 当企业在第 t 年被标记为 ST 或 *ST 时, 企业的财务状况在 $t-1$ 或 $t-2$ 年就出现了问题。因此, 本文将 2022~2024 年之间, 上市时间超过 5 年且首次被标记为 ST 或 *ST 的企业都视为企业出现财务风险的情况, 因此选择所有财务指标均采用 $t-3$ 年的数据, 即预测 ST 或 *ST 状态前三年的财务数据, 更符合财务风险具有滞后性的特点。

综上, 本文共筛选出存在财务风险的企业 151 家, 财务健康企业 4925 家, 不平衡比例约为 1:32。

5.1.2. 指标筛选

财务指标的选取遵循全面性、代表性和可获取性原则。参考国内外经典财务预警研究, 主要从五个维度构建指标体系: 偿债能力、盈利能力、营运能力、发展能力和现金流量, 筛选出包括流动比率、资产负债率、应收账款周转率、存货周转率、总资产收益率、净资产收益率、营业收入增长率、经营活动现金流量比率等 27 项核心指标。这些指标能够从不同角度反映企业的财务状况, 为风险预测提供多维度的信息支持。但其中部分指标可能存在相关性, 因此首先需要对数据指标进行相关性分析以减少多重共线的可能性。本文运用 MATLAB 对指标进行筛选, 得出有关的相关性系数如表 1 所示。

Table 1. Correlation analysis among financial indicators

表 1. 财务指标之间的相关性分析

	流动比率	速动比率	利息保障倍数	资产负债率	长期借款与总资产比	权益乘数
流动比率	1	0.941**	-0.114**	-0.849**	-0.498**	-0.836**
速动比率	0.941**	1	-0.127**	-0.835**	-0.499**	-0.822**
利息保障倍数	-0.114**	-0.127**	1	0.098**	0.087**	0.103**
资产负债率	-0.849**	-0.835**	0.098**	1	0.535**	0.986**

续表

长期借款与总资产比	-0.498**	-0.499**	0.087**	0.535**	1	0.531**
权益乘数	-0.836**	-0.822**	0.103**	0.986**	0.531**	1
	资产报酬率	总资产净利润率	净资产收益率	营业毛利率	营业利润率	销售费用率
资产报酬率	1	0.980**	0.794**	0.325**	0.662**	-0.012*
总资产净利润率	0.980**	1	0.791**	0.350**	0.681**	0.021**
净资产收益率	0.794**	0.791**	1	0.258**	0.596**	-0.041**
营业毛利率	0.325**	0.350**	0.258**	1	0.633**	0.543**
营业利润率	0.662**	0.681**	0.596**	0.633**	1	0.043**
销售费用率	-0.012*	0.021**	-0.041**	0.543**	0.043**	1
	应收账款周转率	存货周转率	营运资本周转率	流动资产周转率	固定资产周转率	总资产周转率
应收账款周转率	1	0.229**	0.068**	0.426**	-0.001	0.264**
存货周转率	0.229**	1	0.130**	0.480**	0.027**	0.325**
营运资本周转率	0.068**	0.130**	1	0.322**	0.138**	0.384**
流动资产周转率	0.426**	0.480**	0.322**	1	0.046**	0.817**
固定资产周转率	-0.001	0.027**	0.138**	0.046**	1	0.338**
总资产周转率	0.264**	0.325**	0.384**	0.817**	0.338**	1
	资本积累率	总资产增长率	净利润增长率	营业收入增长率	营业收入现金净含量	全部现金回收率
资本积累率	1	0.591**	0.286**	0.291**	0.133**	0.197**
总资产增长率	0.591**	1	0.218**	0.372**	0.139**	0.202**
净利润增长率	0.286**	0.218**	1	0.407**	0.140**	0.207**
营业收入增长率	0.291**	0.372**	0.407**	1	0.021**	0.094**
营业收入现金净含量	0.133**	0.139**	0.140**	0.021**	1	0.847**
全部现金回收率	0.197**	0.202**	0.207**	0.094**	0.847**	1

注: **表示在 0.01 级别(双尾), 相关性显著; *表示在 0.05 级别(双尾), 相关性显著。

在判断变量间的相关性时, 一般认为当相关系数绝对值 $|r| \geq 0.8$ 时, 两个变量高度相关; $0.5 \leq |r| < 0.8$ 为中度相关; $0.3 \leq |r| < 0.5$ 为低度相关; $|r| < 0.3$ 则认为两个变量基本不相关。因此, 本文剔除了相关性分析表中高度相关的指标后, 构建了如表 2 所示的财务风险预测指标体系。

Table 2. Financial risk prediction indicator system

表 2. 财务风险预测指标体系

类型	指标	计算公式
偿债能力	流动比率(X_1)	流动资产 ÷ 流动负债
	利息保障倍数(X_2)	(利润总额 + 利息费用) ÷ 利息费用
	资产负债率(X_3)	总负债 ÷ 总资产 × 100%
	长期借款与总资产比(X_4)	长期借款 ÷ 总资产 × 100%

续表

盈利能力	总资产净利润率(X_5)	净利润 $\div$ 平均总资产 $\times 100\%$
	净资产收益率(X_6)	净利润 $\div$ 平均净资产 $\times 100\%$
	营业毛利率(X_7)	(营业收入 - 营业成本) $\div$ 营业收入 $\times 100\%$
	营业利润率(X_8)	营业利润 $\div$ 营业收入 $\times 100\%$
	销售费用率(X_9)	销售费用 $\div$ 营业收入 $\times 100\%$
营运能力	应收账款周转率(X_{10})	营业收入 $\div$ 平均应收账款余额
	存货周转率(X_{11})	营业成本 $\div$ 平均存货余额
	营运资本周转率(X_{12})	营业收入 $\div$ 平均营运资本
	固定资产周转率(X_{13})	营业收入 $\div$ 平均固定资产净值
	总资产周转率(X_{14})	营业收入 $\div$ 平均总资产
发展能力	资本积累率(X_{15})	(期末所有者权益 - 期初所有者权益) $\div$ 期初所有者权益
	总资产增长率(X_{16})	(期末总资产 - 期初总资产) $\div$ 期初总资产 $\times 100\%$
	净利润增长率(X_{17})	(本期净利润 - 上期净利润) $\div$ 上期净利润 $\times 100\%$
	营业收入增长率(X_{18})	(本期营业收入 - 上期营业收入) $\div$ 上期营业收入 $\times 100\%$
现金流量	营业收入现金净含量(X_{19})	销售商品、提供劳务收到的现金 $\div$ 营业收入 $\times 100\%$
	每股经营活动产生的现金流量净额(X_{20})	经营活动现金流量净额 $\div$ 总股本
	每股现金净流量(X_{21})	现金及现金等价物净增加额 $\div$ 总股本

5.1.3. 数据预处理与模型参数设置

由于不同财务指标的量纲和量级差异较大，在进行实验分析之前，先将数据归一化到[0, 1]范围内。对原始数据集进行标准化处理能削弱各个指标的值之间存在的大小差异，并能减小计算过程的复杂程度。归一化处理方式为：

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}} \tag{17}$$

本文采用分层抽样方法将数据集按 8:2 的比例划分为训练集与测试集，确保两个子集中正负类样本的分布比例与总体数据集保持一致，避免因数据划分不当导致的评估偏差。具体而言，将总样本的 80% 作为训练集，用于模型构建和参数优化；20%作为测试集，用于验证模型最终的泛化能力。为了降低因数数据划分随机性所带来的评估波动，引入了交叉验证的方法，对训练集进行多次重复训练，取测试集预测精度的平均值作为适应度值。以上实验均使用 MATLAB R2024a 进行编程。模型的具体参数参考表 3。

Table 3. Model structural parameters
表 3. 模型结构参数

参数	最大迭代次数	种群规模	飞行长度	最大感知概率	加权学习因子	步长缩放因子	变异概率	交叉概率	隐层节点数
符号	MIT	N	FL	$AP_{\max}$	λ	α	F	CR	H
值	50	20	2	0.5	[0.05, 0.95]	0.01	0.5	0.7	[80, 120]

5.2. 实验结果与分析

本节将运用改进乌鸦搜索算法优化的加权极限学习机分类器对样本进行分类处理。对于高度不平衡的二分类问题(ST 企业占少数), 仅使用总体准确率(Accuracy)会带来严重的误导, 因为一个简单地将所有样本预测为“非 ST”的模型也能获得很高的准确率, 但这完全丧失了风险预警的意义。因此, 本研究除准确率外, 还加入了灵敏度(Sensitivity, 也称为召回率 Recall)、特异度(Specificity)、精确率(Precision)、F1 值(F1-Score)、受试者工作特征曲线下面积(Area Under the ROC Curve, AUC)以及 G-mean 指标。

灵敏度衡量了模型正确识别出的 ST 公司占有所有真实 ST 公司的比例, 是风险预测中最为关键的指标, 值越高, 说明漏报的财务危机企业越少。特异度衡量了模型正确识别出的非 ST 公司占有所有真实非 ST 公司的比例。精确率则衡量了在被模型预测为 ST 的公司中, 真正是 ST 公司的比例。F1 值是精确率和召回率的调和平均数, 能够综合平衡这两者的表现, 是评估不平衡分类模型的一个非常核心的指标。AUC 值则从整体上衡量了模型分类能力的优劣, 值越接近 1, 说明模型的分类性能越好。G-mean 根据正负类精度乘积的开方能够反映算法的总体性能。G-mean 的取值越大, 代表算法的分类性能越好, 识别水平越高, 表达式如下:

$$\text{G-mean} = \text{Sensitivity} \times \text{Specificity} \quad (18)$$

本文选取了基础 ELM 模型、WELM 和基于遗传算法 GA、粒子群算法 PSO、基础 CSA 优化后的 ELM 进行性能对比。实验结果如表 4 所示。

Table 4. Experimental results on corporate data
表 4. 企业数据实验结果

算法	Accuracy	Sensitivity (Recall)	Specificity	Precision	F1-score	AUC	G-mean
ELM	0.9695	0.0667	0.9970	0.4000	0.1143	0.7383	0.2578
WELM	0.7389	0.6333	0.7421	0.0696	0.1254	0.7541	0.6856
GA-WELM	0.7271	0.7333	0.7269	0.0756	0.1371	0.6888	0.7301
PSO-WELM	0.7222	0.6777	0.7320	0.0435	0.0784	0.6132	0.5411
CSA-WELM	0.7754	0.7000	0.7777	0.0875	0.1556	0.7425	0.7378
ICSA-WELM	0.8315	0.7667	0.8335	0.1230	0.2120	0.7941	0.7994

观察表中的数据, 可以发现, 基础的极限学习机模型尽管准确率达到 96.95%, 但召回率低于 6.67%。也就是说, 该模型的高准确率是牺牲了对少数类(ST 企业)识别所得出来的结果, 说明该模型无法准确识别出企业的财务风险状况。

通过对不同模型进行交叉验证得出的结果进行比较, 分析如下:

在整体准确率方面, 所有模型的表现均只在 70%~85%之间, 但对少数类的识别远远高于基础 ELM 模型, 说明对 ELM 的权重进行调整, 可以有效的避免模型的分类决策边界倾向于多数类, 使模型在训练过程中能够更关注到 ST 企业, 增强对 ST 企业的识别率。

对比其他模型, 本文提出的 ICSA-WELM 模型对企业财务风险预测的准确率、召回率、AUC 值都较高, 说明该模型的预测性能较好。模型的 G-mean 值达到了 0.7994, 可以认为模型对 ST 企业 and 非 ST 企业的识别水平都相对较高。

综上, 经改进算法优化后的 WELM 模型能够比较有效地从海量的正常公司中挖掘出潜在的财务风险

信号,降低困境企业的漏报率。同时,由于其引入了有效的加权机制和算法优化,精确率也能保持在合理水平,避免了误报率的大幅上升。但是整体预测结果受少数类的影响,准确率和识别力还有待提高,后续需对模型进行进一步改进或调整。

6. 结论与展望

本文构建了基于改进乌鸦搜索算法结合加权极限学习机的财务风险预测模型,对上市 A 股企业进行了实证分析,通过严谨的实验和针对不平衡数据设计的评估指标体系,验证了 IC-SA-WELM 模型在不平衡条件下的财务风险预测中的性能。实验结果表明,该模型的整体表现尚可,尤其在识别财务危机企业上具有较高的敏感性和可靠性,为企业和投资者进行有效的风险预测提供了更为强大的工具。但模型精确率仅为 83.15%,表明存在一定程度的误报现象。这种高召回率与低精确率的组合特征,实际上反映了财务风险预测领域普遍面临的类别不平衡问题。

从统计学角度看,模型性能指标的矛盾性揭示了风险预警系统的本质特征。高召回率确保了对潜在风险企业的广泛覆盖,而相对较低的精确率则意味着需要投入额外资源进行二次验证。在实际应用中,可以通过调整分类决策阈值来平衡这两个指标。当监管机构需要全面筛查风险企业时,可适当降低阈值以提高召回率;而金融机构在制定信贷政策时,则可提高阈值以确保精确性。这种动态调整机制使模型能够适应不同应用场景的需求。

模型的黑箱特性是另一个需要正视的局限性。虽然深度学习模型具有强大的特征提取能力,但其决策过程缺乏透明性,可能影响监管机构和企业的信任度。未来研究可考虑引入可解释性 AI 技术,如 LIME 或 SHAP 方法,来揭示模型的关键决策因素。这种改进不仅能提升模型的可信度,还能帮助识别影响财务风险的核心指标,为风险防控提供更直接的指导。

从应用价值来看,该模型在实际场景中可为不同主体提供多维度风险管控支持。对企业而言,模型能够提前三年基于财务指标输出风险概率,帮助管理层识别潜在财务隐患,企业可据此调整资产负债结构、优化现金流管理或加强内部控制,从而防范风险实质性发生。对投资者而言,模型输出可作为投资决策的参考依据,辅助识别高风险标的,降低投资组合的潜在损失。但应注意的是,模型仍有可以改进的空间,企业应当综合治理结构、行业周期等因素综合评判风险状况。另外,财务数据具有时效性特征,模型参数需要定期重新训练以保持预测准确性。同时,可以考虑引入增量学习技术,使模型能够持续适应市场环境变化。针对不同行业的企业,可以探索建立细分领域的专用模型,以进一步提高预测精度。

参考文献

- [1] 李光荣, 李风强. 基于几种神经网络方法的公司财务风险判别研究[J]. 经济经纬, 2017, 34(2): 122-127.
- [2] Li, Q., Cai, D. and Wang, H. (2012) Study on Network Finance Risk on the Basis of Logit Model. *Advances in Intelligent Systems and Computing*, **136**, 213-220. https://doi.org/10.1007/978-3-642-27711-5_29
- [3] Canbas, S., Cabuk, A. and Kilic, S.B. (2005) Prediction of Commercial Bank Failure via Multivariate Statistical Analysis of Financial Structures: The Turkish Case. *European Journal of Operational Research*, **166**, 528-546. <https://doi.org/10.1016/j.ejor.2004.03.023>
- [4] 王文雯, 姚欣. 注意力机制结合卷积神经网络的医院财务风险预警方法研究[J]. 现代科学仪器, 2023, 40(2): 166-173.
- [5] 张亚男, 刘人境, 陈慧灵. 基于粒子群算法和核极限学习机的财务危机预测模型[J]. 统计与决策, 2019, 35(9): 67-71.
- [6] 严莉红. 基于 LSTM 神经网络的饲料企业财务风险预警模型构建[J]. 饲料研究, 2023, 46(3): 130-134.
- [7] 李祥飞, 张再生, 刘珊珊. 改进布谷鸟搜索 SVM 在财务风险评估中的应用[J]. 计算机工程与应用, 2015, 51(23): 218-225.
- [8] 向有涛, 王明, 曹琳. 基于多目标深度学习模型的财务风险预测方法[J]. 统计与决策, 2022, 38(10): 184-188.

-
- [9] 邱福祿. 基于 SSA 优化 BP 神经网络上市公司财务危机预警研究[J]. 知识经济, 2024(24): 3-7.
- [10] Taabodi, M.H., Niknam, T., Sharifhosseini, S.M., Aghajari, H.A., Bornapour, S.M., Sheybani, E., *et al.* (2025) Optimizing Rural Mg's Performance: A Scenario-Based Approach Using an Improved Multi-Objective Crow Search Algorithm Considering Uncertainty. *Energies*, **18**, Article 294. <https://doi.org/10.3390/en18020294>
- [11] 林敏. 基于大数据的财务风险预警模型探究——评《大数据财务分析》[J]. 中国科技论文, 2022, 17(2): 233.
- [12] 任婷婷, 鲁统宇, 崔俊. 基于改进 AdaBoost 算法的动态不平衡财务预警模型[J]. 数量经济技术经济研究, 2021, 38(11): 182-197.
- [13] Das, S., Sahu, T.P. and Janghel, R.R. (2022) Stock Market Forecasting Using Intrinsic Time-Scale Decomposition in Fusion with Cluster Based Modified CSA Optimized Elm. *Journal of King Saud University-Computer and Information Sciences*, **34**, 8777-8793. <https://doi.org/10.1016/j.jksuci.2021.10.004>
- [14] 霍闪闪, 苏兵, 王章权, 等. 基于改进乌鸦搜索算法的极限学习机分类方法[J]. 计算机仿真, 2023, 40(8): 370-375, 487.
- [15] Wang, J., Lu, S., Wang, S. and Zhang, Y. (2021) A Review on Extreme Learning Machine. *Multimedia Tools and Applications*, **81**, 41611-41660. <https://doi.org/10.1007/s11042-021-11007-7>
- [16] Askarzadeh, A. (2016) A Novel Metaheuristic Method for Solving Constrained Engineering Optimization Problems: Crow Search Algorithm. *Computers & Structures*, **169**, 1-12. <https://doi.org/10.1016/j.compstruc.2016.03.001>
- [17] Huang, G., Zhu, Q. and Siew, C. (2006) Extreme Learning Machine: Theory and Applications. *Neurocomputing*, **70**, 489-501. <https://doi.org/10.1016/j.neucom.2005.12.126>
- [18] Price, K. (1996) Differential Evolution: A Fast and Simple Numerical Optimizer. *Proceedings of North American Fuzzy Information Processing*, Berkeley, 19-22 June 1996, 524-527.