

多源数据融合的张家口地区芸豆产量预测研究

孙伟健¹ 武丽媛¹ 赵喜清^{1*} 尚启兵²

SUN Weijian WU Liyuan ZHAO Xiqing SHANG Qibing

摘要

芸豆是我国农业生产中重要的作物之一，在北方地区被广泛种植，建立一个可靠、准确的产量预测模型可以为河北张家口地区芸豆的种植管理提供决策支持。基于张家口地区的气象数据（温度、大气湿度、降水量、光照）和8项作物农艺性状（株高、主茎节数、主茎分支、单株结荚、荚长、单荚粒数、单株粒重、百粒重），通过改进传统预测方法，以达到芸豆的准确预测，使用决策树（DT）、自适应提升算法（Adaboost）、极端梯度提升（XGBoost）、支持向量机（SVM）和K最邻近（KNN）5种机器学习方法构建芸豆产量预测模型，并比较模型之间的预测准确度差异，其中XGBoost模型预测性能最佳， R^2 、RMSE、MAPE分别为0.760%、6.934%、6.149%，之后使用可解释方法SHAP对5种机器学习模型的特征变量进行重要度排序，其中单株粒重、月均光照时间、单荚粒数、百粒重、月平均温度、主茎分支、主茎节数、单株结荚8个特征占据全部特征重要性的89.81%，对SHAP筛选后特征进行组合，选出最优的6个特征后将其输入到XGBoost模型中并引入PSO算法，在SHAP与PSO组合下，芸豆产量的预测准确度要优于传统XGBoost模型，其 R^2 、RMSE、MAPE分别为0.897%、4.784%、3.652%。文章所使用的SHAP-PSO-XGBoost模型通过结合气象数据和作物产量数据对张家口地区芸豆的产量预测有着较好的准确度，通过对产量的预测为该地区芸豆的种植决策提供一定的科学依据。

关键词

芸豆；产量预测；XGBoost；SHAP；PSO

doi: 10.3969/j.issn.1672-9528.2025.11.033

0 引言

芸豆，又称四季豆、扁豆，属于豆科菜豆属植物，在我国北方地区广泛种植，是当地农业生产中的重要作物^[1]。河北省张家口市地处华北平原与内蒙古高原的过渡地带，气候冷凉且光照充足，具备芸豆生长的适宜环境条件，尤其在张家口坝上地区，芸豆种植面积已达十几万亩。

芸豆兼具食用与药用价值，其蛋白质为全价蛋白，营养结构优良^[2]。加之地理环境的先天优势，使其成为张家口地区的核心粮食作物之一。因此，构建科学精准的芸豆产量预测模型，不仅可为当地种植户及种植企业提供决策支持，优化种植管理方案，还能为区域粮食安全保障提供技术支撑。

目前，关于各类粮食作物的预测模型有很多，常见的是使用一些传统的回归模型，例如使用线性回归^[3]、多元回归^[4]等。但这些回归模型在各类粮食作物预测方面所考虑的

特征因素较少，通常只是单一地选择气象因素、作物因素、土壤因素等其中的一项作为特征因素进行预测^[5]。XGBoost（extreme gradient boosting）是一种基于梯度提升框架的高性能机器学习算法，通过极端的梯度提升，不断拟合上一轮模型的训练残差，从而达到提升整个模型预测精度的目的。使用能够整合多源数据，包括气象数据、土壤数据、田间管理数据、作物信息等的模型，可以使其在各类粮食产量预测方面有着较好的准确度和可信度^[6]。但XGBoost在模型输出特征排序方面缺乏一定的直观解释，不利于学者的理解，无法有效地从中获取影响产量的特征因素^[7]。

SHAP（shapley additive explanations），是一种基于博弈论的可解释机器学习方法，其主要作用是通过计算模型中每个特征对模型预测的边际贡献值来解释模型的预测结果，通过合作博弈理论当中的公平分配原则，使得模型当中的每个特征对预测结果产生的影响得到公平评估^[8]。SHAP的主要优势源于它不仅可以对模型中的所有特征进行排序，还可以对某个单一特征做出直观解释，从而展示各个特征是如何影响最终的预测结果^[9]。使用SHAP模型搭配XGBoost，不仅可以验证模型是否符合农业的基本常识理论，还可以发现一些常被忽视的变量之间的交互作用，这种交互作用可

1. 河北北方学院 河北张家口 075000

2. 张家口市农业科学院 河北张家口 075000

[项目基金] 河北北方学院校级科研项目（XJ2024016）；国家食用豆产业技术体系（CARS-08-Z02）

以达到“1+1>2”的效果，通过挖掘这些变量之间的关系，将其转换为实际的灌溉策略和田间管理策略等^[10]。

本文研究主要通过张家口农业综合试验站的田间数据，以芸豆的亩产作为因变量，气象因子、作物农艺性状数据作为特征因素，通过使用决策树（DT）、自适应提升算法（Adaboost）、极端梯度提升（XGBoost）、支持向量机（SVM）和 K 最邻近（KNN）5 种机器学习算法构建芸豆的预测模型，之后引入 SHAP 事后可解释方法对特征进行排序，筛选出在产量预测模型上表现最优的 6 个特征变量，将其输入到预测效果最佳的模型中构建组合预测模型，以达到为张家口地区芸豆的种植产业提供一定的管理建议和种植决策的理论依据。

1 材料与方法

1.1 数据来源

研究区域位于河北省张家口市张北试验基地（41°09'31.38"N, 114°43'11.53"E），海拔 1 467 m，张北县属于温带大陆性季风气候，同时具有较明显的高原气候特征。

田间试验采用随机区组的排列方式，重复 3 次，单个小区面积为 10 m²，单株作物间距 14 cm，行距 0.5 m，试验地类为水浇地，土质为沙壤土，使用桂湖复合肥（18-18-18）20 公斤/亩（0.03 kg/m²）。芸豆的整个生长发育期内及时进行病虫害的处理，保证芸豆的正常生长，待芸豆成熟后，对各小区内芸豆进行采摘测量其各项农艺性状（株高、主茎节数、主茎分支、单株结荚、荚长、单荚粒数、单株粒重、百粒重）及亩产量。

气象数据主要选取芸豆生育期（4 月—9 月）各月份的月最高温度、月最低温度、月平均温度、月均降雨量、月均大气湿度、月均光照时长（数据来源：<https://www.heywhale.com>）。

将芸豆的产量作为输出，芸豆的农艺性状和气象数据作为输入，使用 MATLAB 进行模型建立，构建决策树（DT）、自适应提升算法（Adaboost）、极端梯度提升（XGBoost）、支持向量机（SVM）和 K 最邻近（KNN）5 种机器学习模型。

1.2 数据预处理

首先对农艺性状数据进行数据清洗，检查并处理异常值，针对该项数据，采用 3σ 方法对所有数据进行离群点识别，当出现离群数据时，则视为录入错误并剔除，若为真实数据则保留并进行标注。气象数据主要捕获极端值，因其具有非线性和强耦合性等特点，所以采用卡尔曼滤波法对数据进行去噪。各项数据的缺失值处理采用高斯过程回归（GPR）时空协同插值算法进行处理，通过结合周期性核函数与径向基核函数构建复合协方差模型，高斯过程回归能够准确描述温湿度、光照等气象要素的时空演变特征，从而生成最优插补值。

1.3 模型构建

选择使用决策树（DT）、自适应提升算法（Adaboost）、极端梯度提升（XGBoost）、支持向量机（SVM）和 K 最邻近（KNN）5 种机器学习方法构建芸豆产量预测模型，其中决策树是一种常用的监督学习算法，广泛应用于分类和预测任务。通过构建树状结构对数据进行分割，模拟人类决策过程。自适应提升算法（Adaboost）和极端梯度提升（XGBoost）都属于集成学习（Ensemble Learning）中的 Boosting 方法，其核心思想是通过组合多个弱学习器（如决策树）来构建一个强学习器。支持向量机（SVM）是一种经典的监督学习算法，主要是通过创建一个最优平面使不同类别数据之间的间隔最大化，从而提高模型的泛化能力。K 最邻近（KNN）是一种简单的非参数监督学习算法，广泛应用于各种分类和回归任务。因其在最后预测结果时才会计算样本与训练集的距离并进行决策，所以 K 最邻近的训练速度较快，但在最后阶段的预测需要的计算资源较高，主要通过局部邻居的投票实现预测，是一种适合轻量级任务的模型。

将提取好的 8 项农艺性状和 6 项气象参数进行归一化处理后，将其作为模型的特征变量输入到 5 种机器学习模型中，构建芸豆的产量预测模型，按 7:3 的比例划分训练集和测试集，并进行 K 折叠交叉验证，对模型进行评估。

1.4 模型性能评价

为比较 5 种芸豆产量预测模型的性能，使用决定系数（coefficient of determination, R^2 ）、均方根误差（root mean squared error, RMSE）和平均百分比误差（mean absolute percentage error, MAPE）作为模型评估的主要指标。其中决定系数（ R^2 ）是用来评估模型拟合效果的一项常用指标，决定系数越大，说明该模型拟合效果越好，但过大则可能意味着该模型存在过拟合的情况。均方根误差 RMSE 是用来评估模型预测误差大小的指标，RMSE 值越小，意味着模型的预测准确度更高。平均百分比误差（MAPE）通过使用百分比大小可以更直观地反映模型的误差大小，其范围为 [0, +∞)，当 MAPE 为 0 时，表示该模型为完美模型，当 MAPE 值大于 100 时，表示该模型没有任何实际意义，为低质模型。通过使用上述三种模型评估指标对 5 种芸豆产量预测模型进行评估，可以对预测模型进行更综合性的分析。

$$R^2 = 1 - \frac{\sum_i (\hat{y}_i - y_i)^2}{\sum_i (y_i - \bar{y})^2} \quad (1)$$

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2} \quad (2)$$

$$MAPE = \frac{100\%}{n} \sum_{i=1}^n \left| \frac{\hat{y}_i - y_i}{y_i} \right| \quad (3)$$

式中： $\hat{y}_i$ 为芸豆产量的预测值； y_i 为芸豆的实际产量； $\bar{y}$ 为 y_i

的平均值； n 为样本的数量。

1.5 SHAP 特征筛选

使用 5 种机器学习模型后，通过对 5 种模型性能进行评估后从中选出最优预测模型，后引入 SHAP 事后解释模型对模型当中的各个特征进行贡献度排序，从而筛选出最佳的特征变量，使模型达到最优的产量预测效果。

SHAP 是一种对模型的预测结果拆分成各个特征对预测结果的贡献度，保证公平、一致的分配各特征重要性，通过数据可视化来展示所有特征是如何影响预测结果的，不仅可以用来做出单个或多个特征的解释（局部解释），也可以对整体模型的全部特征进行解释（全局解释）。SHAP 的核心便是沙普利值，沙普利值的计算公式可表示为：

$$\phi_i(f, x) = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|!(|F| - |S| - 1)!}{|F|!} (f(S \cup \{i\}) - f(S)) \quad (4)$$

式中： ϕ_i 为模型中各特征的对预测结果的贡献度； F 为模型的所有特征的集合； S 为 F 当中的一个子集； $f(S)$ 为表格使用特征子集 S 时的模型预测值。对于高维度的数据 SHAP 模型无法有效地对其进行精确计算，故而通常采用近似的方法进行计算，即使用加权线性回归来近似沙普利值，具体表示为：

$$\min_{\phi_0, \phi} \sum_{k=1}^K \left(f(x_{S_k}) - \left(\phi_0 + \sum_{i \in S_k} \phi_i \right) \right)^2 \cdot \pi(S_k) \quad (5)$$

式中： K 为随机生成的特征子集，对于每个 S_k 生成输入数据 X_{sk} ，并计算预测值。

1.6 粒子群优化算法

粒子群优化算法（PSO）是一种模拟群体智能行为的优化方法，通过粒子在搜索空间中跟踪个体历史最优和群体全局最优来迭代更新位置与速度，逐步逼近最优解，本文中使用时 PSO 算法在特征组合后对模型进行优化，其流程图如图 1 所示。

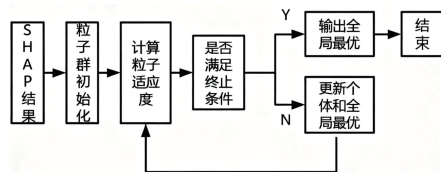


图 1 PSO 算法流程图

2 结果与分析

2.1 不同机器学习产量预测模型效果对比

5 种机器学习产量预测模型在预测芸豆产量的效果表现如表 1 所示在训练集中，自适应提升算法（Adaboost）在决定系数 R^2 上表现最优，达到了 0.776，但在其他部分性能上

要弱于 XGBoost 和 SVM 模型，其中 XGBoost 虽然在决定系数要弱于自适应提升算法（Adaboost），但其在均方根误差 RMSE 和百分比误差 MAPE 上表现最优，分别为 6.156%、5.611%。在测试集中，XGBoost 模型展现出了优越的性能，其在决定系数 R^2 上表现最优，为 0.760%，但在其他性能上与训练集有所不同，其中在均方根误差 RMSE 上表现最优的是 XGBoost 模型，为 6.934%，在百分比误差 MAPE 表现最优的为 SVM，为 6.073%。由于在训练集的 R^2 上 Adaboost 和 XGBoost 虽然有所差异，但总体差距不大，且对比 RMSE 和 MAPE 方面，XGBoost 要优于 Adaboost，综合全部评价指标，XGBoost 要优于其他模型，基于此，后续选择使用 XGBoost 模型进行下一步的分析和研究。

表 1 不同机器学习算法的芸豆产量预测模型性能对比

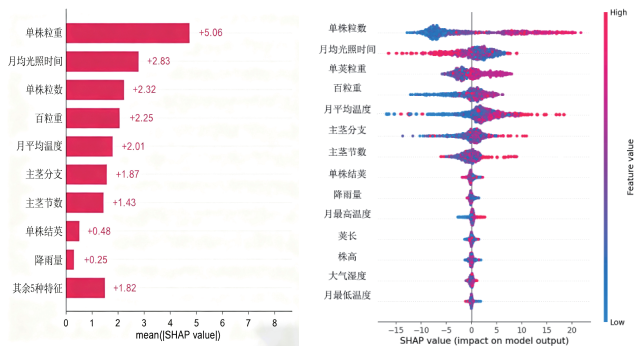
模型	训练集			测试集		
	R^2	RMSE	MAPE/%	R^2	RMSE	MAPE/%
DT	0.761	7.953	6.121	0.754	8.623	6.648
Adaboost	0.776	7.642	5.781	0.757	7.639	6.811
XGBoost	0.773	6.156	5.611	0.760	6.934	6.149
SVM	0.769	7.213	5.784	0.757	8.021	6.073
KNN	0.763	8.291	6.528	0.749	8.981	6.795

2.2 特征变量的筛选

为进一步获取模型中各个特征变量对预测模型结果的贡献度，通过使用 SHAP 解释模型对 XGBoost 模型当中的各特征变量进行分析，从图 1 可以发现，各个特征虽然对模型的预测均存在着一定的贡献度，但大部分特征所产生的贡献度并不多。其中对芸豆产量预测产生主要影响的前两个特征是单株粒重和光照时长，单株粒重是评价作物产量的核心因素之一，通过对作物的单粒大小、形状、重量来判断品种的优劣，是帮助农户筛选出优良作物品种的最常见也是最传统的方法。而单株粒重易受到环境因素的影响，例如光照、温度、水肥等，其中光照对单株粒重的影响最大，与特征筛选结果相互印证。而芸豆是一种喜光作物，尤其是可见光中的红光和蓝光，能够有效地增加叶片的光合速率，促进芸豆中碳水化合物的积累，为其籽粒的灌浆提供基础。虽然芸豆作为一种喜光植物，但芸豆对于光周期较为敏感，虽然现在培育的最新品种对光周期无过高要求，但其仍然影响着芸豆的开花时间，当光照时间过长时会导致芸豆开花期延长。

通过蜂群图 2（b）可以很好地观察各个特征变量的 SHAP 值分布情况，蜂群图是对 SHAP 解释模型所对 XGBoost 模型各特征之间关系的最直观展示。图 2 结果可以发现，当芸豆的单株粒重对芸豆产量呈正相关，当单株粒重越大时，其亩产越高。同时发现当光照时间取值过高时，会

对芸豆的产量出现负影响，这与芸豆的生理特性有关。



(a) SHAP 特征重要性结果图 (b) SHAP 蜂巢图

图 2 SHAP 特征解释结果

根据特征解释结果，计算模型当中各特征的贡献度，从而筛选出最优的特征变量，通过表 2 可以发现前 8 个特征变量对模型的贡献度为 89.81%，接近 90%，因此，前 8 个特征变量在芸豆产量预测模型中发挥着主要作用，故此选择单株粒重 (X_1)、月均光照时间 (X_2)、单株粒数 (X_3)、百粒重 (X_4)、月平均温度 (X_5)、主茎分支 (X_6)、主茎节数 (X_7)、单株结荚 (X_8) 作为主要特征输入到 XGBoost 模型中，同时使用不同特征组合构建模型，从而使得模型性能最优。

表 2 模型特征贡献度排序

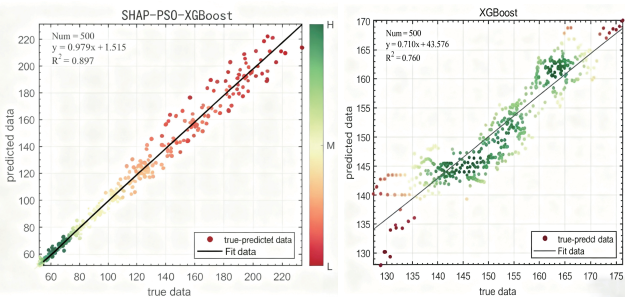
特征	贡献度	累计贡献度
单株粒重	5.06	24.90
月均光照时间	2.83	38.82
单荚粒数	2.32	50.24
百粒重	2.25	61.31
月平均温度	2.01	71.24
主茎分支	1.87	80.41
主茎节数	1.43	87.45
单株结荚	0.48	89.81

2.3 特征变量的组合

为进一步提升 XGBoost 模型的预测性能，将使用 SHAP 方法筛选出的特征变量输入到模型中，并逐步剔除贡献度最低的变量，从单株结荚 (X_8) 依次开始剔除，从而得到最优模型，从表 3 可以看出，通过逐渐剔除贡献度最低的特征模型整体性能呈现先上升后下降的趋势，当特征数量为 6 时各模型性能指标达到最佳，其 R^2 、RMSE、MAPE 分别为 0.834%、5.661%、4.103%，之后引入 PSO 优化算法，构建 SHAP-PSO-XGBoost 模型，其 R^2 、RMSE、MAPE 分别为 0.897%、4.784%、3.652%。图 3 为原 XGBoost 模型和 SHAP-PSO-XGBoost 模型的真实值和预测值的相关性散点图对比，可以发现，使用了 SHAP-PSO-XGBoost 的芸豆产量预测模型在预测结果上要更接近芸豆的实际产量，且误差较小。

表 3 不同特征组合预测结果

特征数量	R^2	RMSE	MAPE/%
$X_1+X_2+X_3+X_4+X_5+X_6+X_7+X_8$	0.783	6.547	5.317
$X_1+X_2+X_3+X_4+X_5+X_6+X_7$	0.792	6.236	4.924
$X_1+X_2+X_3+X_4+X_5+X_6$	0.834	5.661	4.103
$X_1+X_2+X_3+X_4+X_5$	0.801	5.948	4.872
$X_1+X_2+X_3+X_4$	0.761	6.667	5.498
$X_1+X_2+X_3$	0.754	7.035	6.799
X_1+X_2	0.731	7.843	7.263
X_1	0.658	8.362	8.014



(a) SHAP-PSO-XGBOOST (b) XGboost

图 3 XGBoost 与 SHAP-PSO-XGBoost 模型预测散点对比图

3 研究结果

本文所使用的决策树 (DT) 其是通过递归地从全部特征中选择最优特征进行数据分割，每次分割旨在最大化子集的纯度 (如基尼指数或信息增益)，最终生成由根节点、内部节点和叶节点，从而输出结果^[11]。决策树直观易解释，能够处理数值和类别特征，但容易过拟合，常需剪枝或集成方法 (如 XGBoost、Adaboost) 优化性能，故此在本文中所表现的性能较为普通。适应提升算法 (Adaboost) 是一种经典的集成学习 (Boosting) 方法，通过迭代训练多个弱分类器 (如决策树桩) 并对错误的特征给予其更高的权重比，使后续模型更重视此特征，最终将所有弱分类器的预测结果加权投票组合成强分类器^[12]。针对气象、农艺性状等多源特征，Adaboost 通过每轮分类精度从而动态地调整样本权重和模型权重，达到提升整体性能的目的，因此在本文中处理复杂气象问题时有着良好的表现。XGBoost 是一种基于梯度提升的机器学习算法，通过迭代的训练弱学习器并优化损失函数来提高模型性能，在结构化的数据上通常要优于 Adaboost^[13]，且为了防止模型出现过拟合的情况可以对模型使用正则化、早停等策略，但其超参数过度，训练速度较慢，在本文中表现最佳，且通过结合 SHAP 解释模型构建的预测模型在芸豆产量预测方面有着较好的准确性和解释性。支持向量机 (SVM) 是一种监督学习算法，通过核函数隐式映射到高维空间，无需显式计算高维特征，但 SVM 对缺失值和噪声较为敏感^[14]，

所以使用其构建的预测模型表现平庸。K 邻近 (KNN) 是一种惰性学习算法,适用于非线性的问题,且无需训练就可以加入新的数据,但是面对高维数据时会产生距离计算失败的情况,需要进行降维处理^[15],因此其构建的预测模型表现不佳。

通过使用 SHAP 解释模型对所有特征进行解释,可以从发现其中对产量影响最大的农艺性状是单株粒重和单株粒数,通过 SHAP 解释可以发现这两种农艺性状当数量过多或粒重过大时会对产量产生负向影响,因为粒数过多时会导致资源的竞争,例如光照、水分、肥力等,从而使粒重下降,而粒重过大时则可能导致作物倒伏,与赵丹等人^[16]的研究结果一致。而影响着产量的主要气象因素是光照,因芸豆对光照较为敏感,且光照时间是影响单株粒重和单株粒数的主要因素,因光照强度和芸豆叶片的光合速率呈高度正相关,直接影响着籽粒灌浆物质的合成,与王桂梅等人^[17]对芸豆落蕾、落花、落荚探究的影响因素一致。

4 结论

在产量预测模型上,XGBoost 相较于决策树、自适应提升算法、支持向量机、K 邻近算法有着较好的效果。之后使用 SHAP 对 XGBoost 构建的芸豆产量预测模型各特征进行解释,并展示所有特征在模型当中的贡献度,从而筛选出影响模型预测的主要特征,通过逐渐剔除贡献度最低的特征最终确定最佳的特征组合为 6 个,将其作为变量输入到 PSO-XGBoost 模型中构建 SHAP-PSO-XGBoost 模型,相较于原 XGBoost 模型, R^2 提高了 0.137%,RMSE 和 MAPE 分别降低了 2.15%、2.997%,说明使用 SHAP-PSO-XGBoost 构建的芸豆产量模型有着更好的准确性和可解释性。

参考文献:

- [1] 张紫帆, 邱思思, 刘春秀, 等. 芸豆蛋白功能特性及改性的研究进展 [J]. 食品与发酵工业, 2024, 50(5): 357-366.
- [2] 孔春丽, 段彩苹, 芦晶, 等. 芸豆中的大分子功能组分及其应用研究进展 [J]. 食品与生物技术学报, 2022, 41(12): 8-18.
- [3] 周永生, 黄润生. 基于线性回归的广西粮食产量预测模型探讨 [J]. 现代商贸工业, 2011, 23(5): 98.
- [4] 田秀芹. 基于多元线性回归的粮食产量预测 [J]. 科技创新与应用, 2017(16): 3-4.
- [5] 李松涛. 水稻作物产量预测模型的研究: 基于多源数据回归模型 [J]. 农机化研究, 2024, 46(10): 27-31.
- [6] 朱志畅, 葛焱, 臧晶荣, 等. 基于无人机图像和 SHAP 特征筛选的小麦田间产量预测方法研究 [J]. 麦类作物学报, 2025, 45(2): 264-274.

- [7] 付金鹏, 王哲. 基于 GA-XGBoost 算法的河南省粮食产量预测研究 [J]. 现代计算机, 2024, 30(6): 107-110.
- [8] XIANG K, CHEN H L, YAO C Q, et al. Data-driven machine learning with SHAP for predicting multiple biochar properties from agricultural and forestry waste [J]. Biomass and bioenergy, 2025, 200: 107958.
- [9] 严格齐, 焦洪超, 林海, 等. 基于 XGBoost-SHAP 的奶牛热应激预测与可解释性研究 [J]. 农业机械学报, 2025, 56(4): 408-414.
- [10] 王元, 王玥, 周宇琛, 等. 应用 SHAP 可解释机器学习模型估测森林蓄积量 [J]. 东北林业大学学报, 2025, 53(5): 66-73.
- [11] DU Z H, YANG L, ZHANG D X, et al. Corn variable-rate seeding decision based on gradient boosting decision tree model [J]. Computers and electronics in agriculture, 2022, 198: 107025.
- [12] KODURI S B, GUNISETTI L, RAMESH C R, et al. Prediction of crop production using adaboost regression method [J]. Journal of physics: conference series, 2019, 1228: 012005.
- [13] SHANKAR K, MOORTHY M. An efficient IoT-based crop yield prediction framework using optimal ensemble learning and hybridized optimization model [J]. Earth science informatics, 2025, 18: 156.
- [14] LI J H, BAI J L, YU B Z. C-Mn yield forecasting model based on SVM [J]. IOP conference series: earth and environmental science, 2021, 781: 022024.
- [15] SITIENEI M, OTIENO A, ANAPAPA A. An application of K-nearest-neighbor regression in maize yield prediction [J]. Asian journal of probability and statistics, 2023, 24(4): 1-10.
- [16] 赵丹, 王荣芳, 张荣华. 芸豆新品种在高海拔地区适应性鉴定初报 [J]. 新农业, 2023(18): 27-28.
- [17] 王桂梅, 邢宝龙, 张旭丽, 等. 芸豆落蕾落花落荚的原因及防治措施 [J]. 中国农业信息, 2016(18): 93-96.

【作者简介】

孙伟健 (2000—), 男, 黑龙江齐齐哈尔人, 硕士研究生, 研究方向: 农业信息化。

武丽媛 (2000—), 女, 河北张家口人, 硕士研究生, 研究方向: 农业信息化。

赵喜清 (1970—), 通信作者 (email: 149883913@qq.com), 男, 河北张北人, 硕士, 教授、硕士研究生导师, 研究方向: 智能教育、机器证明和云数据生态安全。

(收稿日期: 2025-06-10 修回日期: 2025-11-04)