



Research Article

Machine learning in flow boiling: predicting bubble lift-off diameter despite data limitations

Atta Heydarpour TABRIZI¹, Mousa MOHAMMADPOURFARD^{2,*},
Mostafa MOHAMMADPOURFARD³

¹Faculty of Mechanical Engineering, University of Tabriz, 51656 Tabriz, Iran

²Department of Energy Systems Engineering, Izmir Institute of Technology, 35430 Izmir, Türkiye

³Renewable Energy Program, Texas Tech University, Lubbock, TX 79409, USA

ARTICLE INFO

Article history

Received: 28 May 2024

Revised: 29 September 2024

Accepted: 04 October 2024

Keywords:

Bubble Size; Machine Learning
Techniques; Predictive Model;
Regression Model

ABSTRACT

This study concentrates on applying machine learning techniques to flow boiling in order to predict the bubble lift-off diameter. This prediction is critical because the diameter plays a key role in understanding boiling dynamics and calculating heat transfer rates. Additionally, accurately predicting this diameter is essential for optimizing thermal systems and enhancing energy efficiency. By evaluating the performance of three different machine learning algorithms: M5 tree, multilinear regression, and random forest, we aimed to assess their effectiveness in providing reliable predictions even with limited experimental data. This research is essential as it demonstrates the potential of machine learning to enhance predictive accuracy in scenarios where obtaining extensive datasets is challenging. Our findings show that these machine-learning techniques are effective for accurate predictions. The results show that the coefficient of determination ranged from 0.64 to 0.94, indicating how well the models fit the data. The root mean square error was between 0.009 and 0.02, and the mean absolute error ranged from 0.0004 to 0.02. Based on the findings, it can be inferred that the machine learning algorithms used in this study are reliable for predicting bubble lift-off diameter. This reliability also extends to other experimental parameters, suggesting that these techniques can be effectively applied in various contexts with limited data. This study demonstrates the potential of machine learning to predict experimental parameters and advances previous research by identifying key factors that influence bubble lift-off diameter.

Cite this article as: Tabrizi AH, Mohammadpourfard M, Mohammadpourfard M. Machine learning in flow boiling: predicting bubble lift-off diameter despite data limitations. J Ther Eng 2025;11(4):1051–1063.

INTRODUCTION

Subcooled flow boiling is a highly efficient way to remove heat from a source using the extra latent heat of a fluid. To

understand how flow boiling works, many researchers have conducted experiments or developed theories to explain the heat transfer process. They have developed models and

*Corresponding author.

*E-mail address: Mousafard@iyte.edu.tr

This paper was recommended for publication in revised form by
Editor-in-Chief Ahmet Selim Dalkılıç



correlations to explore the relationship between wall temperature and heat flux, with bubble diameter being a critical factor in these models. Bubble size is typically characterized using two main types of diameters: the bubble lift-off diameter (BLOD) and the bubble departure diameter (BDD) [1]. In boiling, as a bubble grows, it usually expands until it detaches from the nucleation site; at this point, its size is called the BDD. The bubble may then slide along the wall while continuing to grow until it eventually lifts off or lifts off directly into the liquid without sliding. In both cases, BLOD refers to the size of the bubble as it lifts off from the wall and enters the bulk liquid.

Several studies have focused on predicting the critical size at which a bubble initiates lift-off. These investigations often use force equilibrium principles to analyze growing bubbles, incorporating dimensionless parameters such as the Jacob and Prandtl numbers. These models have demonstrated good agreement with experimental data [2]. Furthermore, force balance approaches have been widely used to predict both bubble lift-off and departure diameters, with some models precisely estimating evaporation heat flux under low heat flux and flow velocity conditions [3].

The determination of bubble departure points has been explored by considering the balance of forces on a bubble at its nucleation site alongside the temperature distribution in the liquid near the heated wall [4]. Different analyses have considered forces such as buoyancy, drag, and surface tension, with models differing in their assumptions about bubble shape and behavior [5].

In saturated horizontal flow boiling, simplified force balance equations have been used to model BDD and BLOD by neglecting specific forces deemed negligible due to the minimal contact area between the wall and the bubble [6]. Optimization studies on two-phase closed thermosyphons have highlighted the benefits of using fluids with superior surface tension properties, significantly improving power input and fluid performance [7]. Additionally, super-hydrophobic coatings have been shown to accelerate boiling in systems using fluid mixtures compared to water alone [8]. Beyond force analysis, semi-empirical correlations have been employed to calculate BLOD, with new correlations being developed from experimental data in flow boiling scenarios [9]. Comprehensive databases compiled from these studies have been used to assess and refine predictive models, leading to more accurate correlations that combine parameters such as bubble nucleation frequency and BLOD [10], [11].

With the rise of machine learning (ML) and data mining, researchers are now using these techniques to predict experimental outcomes better. These methods bring new insights into understanding different aspects of boiling phenomena. For example, probabilistic ML models were used in micro-structured surfaces to predict pool boiling heat transfer, achieving up to a 30% improvement in prediction accuracy compared to traditional correlations. The research highlighted the boiling Reynolds number as the

most significant parameter and noted that these models provided better uncertainty estimates than deterministic approaches [12].

In another investigation, Bard et al. [13] explored the prediction of heat transfer coefficients in mini/micro-channels during saturated flow boiling through machine learning techniques. Their findings indicate that while machine learning proved highly effective in forecasting the heat transfer coefficient across diverse fluids, it encountered challenges in accurately predicting exceptionally high outlier data, particularly when water served as the working fluid. Qiu et al. [14] reported that optimized ML models performed better than highly reliable generalized pressure drop correlations and performed well across individual datasets, channel configurations, and flow regimes. He et al. [15] explored subcooled flow boiling to predict the bubble departure frequency (BDF). They curated a comprehensive dataset encompassing BDF across four working fluids in subcooled flow boiling. The study extensively examined nine regression models based on machine learning techniques. Moreover, it scrutinized various input parameters, including dimensionless and geometric factors, to determine the most effective approach. Overall, the XGBoost model emerged as the top performer in predicting BDF, surpassing even highly dependable generalized prediction correlations.

Zhang et al. [16] employed ML techniques to forecast critical heat transfer on downward-facing surfaces. In order to enhance the applicability of their research, they incorporated machine learning after compiling the most readily available critical heat flux (CHF) data from pool boiling on such surfaces. Given the limited availability of experimental data, they supplemented their dataset by generating pseudo data by fitting existing experimental records. Cabarcos et al. [17] investigated the use of ML algorithms to predict temperature in nucleate flow boiling. They evaluated the effectiveness of different algorithms, such as random forest, artificial neural networks, XGBoost, AdaBoost, and support vector machine. The critical heat flux (CHF) classification findings reveal that the support vector machine algorithm outperforms the others, whereas the boosting methods (XGBoost and AdaBoost) show overfitting.

As we know in subcooled flow boiling BLOD plays a significant role; therefore, many efforts have been made to predict this parameter under various conditions. However, a review of the available studies reveals that investigations of BLOD are limited, and existing models need further evaluation. Furthermore, the potential of ML techniques such as random forest, the M5 tree, and multiple linear regression has yet to be explored. Hence, this study evaluates the efficacy of the M5 tree, random forest, and multiple linear regression algorithms for predicting BLOD. This paper makes the following contributions:

1. Predict BLOD in the presence of a magnetic field using ML methods and available data.
2. Despite data limitation, assessing the predictive performance of ML models (M5 tree and random forest).

3. Performing feature importance analyses for assessment of the influence and effect of different parameters on target variables (BLOD).

BUBBLE LIFT-OFF DIAMETER MODELING

Researchers typically consider the forces acting on a single bubble to model BLOD. Since the primary focus of this study is utilizing ML algorithms to predict BLOD, the following section will provide a concise overview of the significant forces acting on bubbles. This summary aims to enhance the readers' comprehension of the approach adopted in this study.

Force Analysis

The fundamental concept of bubble lift-off can be summarized as follows: initially, a bubble is formed at the nucleation site and undergoes gradual growth. Once it reaches a specific

size, it detaches from the nucleation site and may slide along the heating surface. Subsequently, vaporization transpires at the inner surface of the bubble, while condensation occurs at the outer surface of the bubble's tip extends beyond the superheated layer [2]. The fate of the bubble, whether it continues to grow or undergoes condensation, is determined by the combined influence of these two processes. Nevertheless, after a certain distance downstream from the nucleation site, the bubble ultimately detaches from the surface of the heater. Figure 1 illustrates a schematic representation of an active nucleation site in upward subcooled flow boiling.

Force Balance for a Single Bubble

Numerous forces influence BLOD, as illustrated in Figure 1. Considering their impact on bubbles at the nucleation site is crucial for accurately predicting BLOD under different conditions. These forces can be decomposed and projected into the x- and y-directions and their respective values are provided as follows [18]:

$$\sum F_x = F_{sx} + F_{sl} + F_{dux} + F_{mg} \quad (1)$$

$$\sum F_y = F_{sy} + F_{duy} + F_p + F_g + F_{qs} \quad (2)$$

where F_{sx} , F_{sl} , F_{dux} and F_{mg} are the surface tension force, shear lift force, unsteady drag force (growth force), and magnetic force in the x-direction, respectively. F_{sy} , F_{duy} , F_p , F_{gs} , F_g the surface tension, the unsteady drag force, the pressure force, the quasi-steady force, and the gravity force in the y-direction, respectively. These forces play a crucial role in determining the dynamics of the bubble in a given system and contribute to the overall motion and equilibrium of the bubble.

Evaluation of Available Correlations for Blod

As mentioned in the introduction, various models have been proposed to estimate the BLOD, employing force analysis or empirical correlations. Some of these models and their experimental conditions are summarized in Table 1 and Table 2.

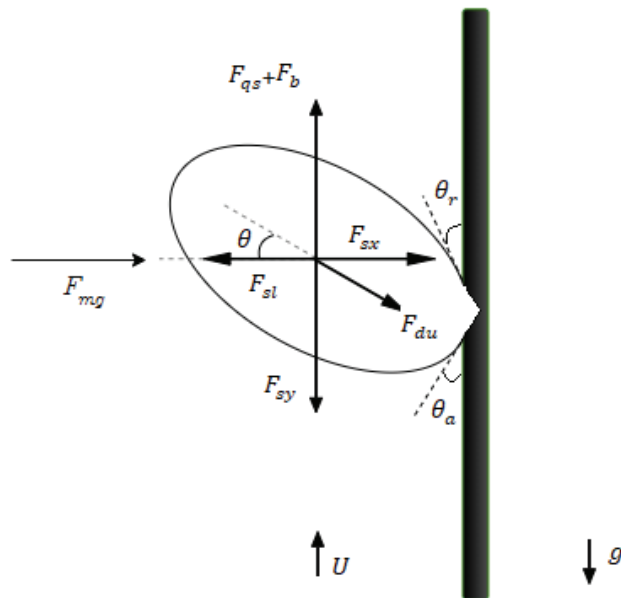


Figure 1. Force balance on a single bubble at the nucleation site.

Table 1. Available correlations for predicting BLOD

Authors	Correlation	Parameters
Situ et al. [2]	$D_{lo}^* = \frac{4\sqrt{22/3} b^2}{\pi} Ja_e^2 Pr^{-1}$ $Ja_e = \frac{\rho_l c_{pl} S(T_w - T_{sat})}{\rho_g h_{fg}}$	$P = 0.101 \text{ MPa}$ $U_l = 0.487 - 0.939 \text{ m/s}$ $T_{in} = 80 - 98.5 \text{ }^\circ\text{C}$
Prodanovic et al. [9]	$D_{e,jc}^+ = 440.98 Ja^{-0.708} \theta^{-1.112} (\rho_l/\rho_v)^{1.747} Bo^{0.124}$ $\theta = \frac{T_w - T_l}{T_w - T_{sat}}, Bo = \frac{q_w}{gh_{fg}}$	$P = 0.105 - 0.3 \text{ MPa}$ $U_l = 0.08 - 0.84 \text{ m/s}$ $\Delta T_{sub} = 10 - 30 \text{ K}$
Chu et al. [10]	$D_{lo}^+ = 12788.5 Ja^{-0.28} \theta^{-1.07} (\rho_l/\rho_v)^{1.36} Bo^{0.35}$	

Table 2. Specific experimental parameters

Parameters	Situ et al.[2]	Prodanovic et al.[9]	Chu et al.[10]
Direction	Vertical	Vertical	Vertical
Channel	Annulus	Annulus	Annulus
Fluid	Water	Water	Water
Pressure (Mpa)	0.101	0.105-0.3	0.145
Mass Flow rate (Kg/m ² s)	466.75–899.96	74.54–804.43	301–702
Subcooling (°C)	3–20	10–60	3.4–22.6
Data point	90	54	14

In order to thoroughly evaluate the effectiveness of the theoretical and empirical models, a quantitative comparison between the corresponding experimental data and the predictive models has been done. For this purpose, experimental data sets provided by Situ et al. [2], Zeng et al. [6], Prodanovic et al. [9], Chu et al. [10], Okawa et al. [19], and Basu et al.[20] have been used in this study.

The mean relative error was employed as a metric to assess the accuracy of the models. This parameter is calculated by taking the average of the absolute differences between the values predicted by the models and the values measured in the experiments, divided by the number of data points in the database.

$$\varepsilon = \frac{1}{n} \sum_{i=1}^n \frac{|D_{lo,measured} - D_{lo,predicted}|}{D_{lo,measured}} \quad (3)$$

The variable n denotes the number of data points available in the database. Table 3 summarizes the findings of this comparison, presenting the mean relative errors for each of the four predictive models. For instance, Situ's model predicts Zeng's experimental data with a relative error of 84.67%, while the corresponding parameter for Prodanovic's data is 736.53% (Fig. 2).

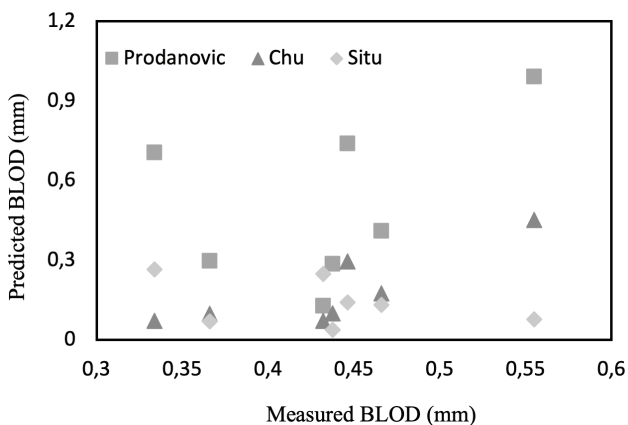
The relative errors of existing correlations in the preceding section underscore the necessity for further investigations in predicting BLOD. The primary objective of this study is to employ ML approaches to predict BLOD, aiming to align with experimental findings closely.

Machine Learning Approach For Predicting Experimental Results

ML algorithms can be applied in experimental studies across various scientific disciplines. These algorithms can analyze complex data, uncover patterns, and make predictions that traditional statistical methods might struggle with. For example, ML algorithms can classify experimental data into categories or perform regression tasks to predict a continuous outcome based on input variables. This study will consider the M5 model, random forest regression (RF), and multiple linear regression (MLR) for predicting BLOD.

The M5 Model Tree

The M5 model tree, presented by Quinlan[21], is a data-driven approach that modifies the traditional decision tree concept [22]. This nonlinear model, which links input and output variables, is based on the 'separate-and-conquer'

**Figure 2.** The assessment of available experimental data vs predictive models.**Table 3.** Relative error between experimental data and available models

Models	Experiments	Situ et al. [2]	Zeng et al. [6]	Prodanovic et al. [9]	Chu et al. [10]	Okawa et al. [19]
	Situ et al. [2]	40.22%	84.67%	736.53%	129.24%	318.66%
	Prodanovic et al. [9]	237.65%	85.53%	73.57%	20.19%	41.11%
	Basu et al.[20]	102.53%	69.54%	44.63%	21.71%	36.77%

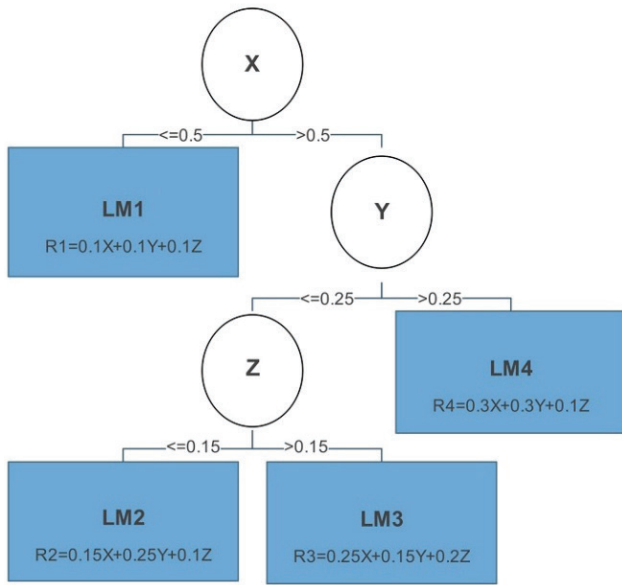


Figure 3. Visual presentation of M5 tree model.

approach. In this method, the range of input variables is divided into distinct subspaces, each with its own dedicated model [23]. M5 model trees are designed to handle classification and regression tasks, making them adaptable for various data analysis applications. Figure 3 visually represents a basic M5 model tree [24].

The splitting algorithm operates recursively, measuring variability at each node by assessing standard deviation. This helps determine the reduction in error [25]. The key idea is to identify the attribute that offers the most significant reduction in error. That is achieved by quantifying the error at each node using standard deviation reduction (SDR), which is the main criterion for making splits[26]:

$$SDR = sd(T) - \sum \left(\frac{|T_i|}{|T|} sd(T_i) \right) \quad (4)$$

Here, T represents the count of instances arriving at the node, T_i signifies the subset of instances exhibiting the i th potential test outcome, and sd corresponds to the standard deviation of the observed values.

The M5 model is a good choice when dealing with a limited amount of data for the following reasons [27]:

- The model's structure is easy to interpret, which is beneficial when working with a small dataset. This interpretability helps understand how the model functions and makes decisions, which is essential for identifying patterns in the data.
- The algorithm includes pruning techniques to prevent overfitting, essential with limited data. Pruning simplifies the model and reduces the risk of capturing noise, leading to more reliable predictions.

- The M5 model tree algorithm can handle missing values through surrogate splits, effectively using incomplete datasets.

Random Forest Regression

The random forest regression (RF) approach was initially introduced by Breiman [28]. RF is a robust and versatile ML technique widely used for predictive modeling and regression tasks. To gain a clearer understanding of the RF (Random Forest) approach, it is important first to become familiar with bootstrap aggregation, also known as bagging. This technique serves as a means to reduce the variance of an estimated prediction function [29–31].

Random Forest (RF) has advanced the concept of bagging by aggregating a large ensemble of decorrelated trees and averaging their predictions. RF constructs numerous decision trees during the training phase, making it well-suited for classification and regression tasks. In regression tasks, the final prediction is the average of each tree's predictions. RF typically outperforms single decision trees by reducing the risk of overfitting to the training data [32].

Consider an average of B identically and independently distributed (i.i.d.) random variables, each with variance σ^2 , resulting in a variance $\frac{1}{B}\sigma^2$. However, when the variables are identically distributed (i.d.) but not necessarily independent, and they exhibit positive pairwise correlation, the variance of the average is [31]:

$$\rho\sigma^2 + \left(\frac{1-\rho}{B} \right) \sigma^2 \quad (5)$$

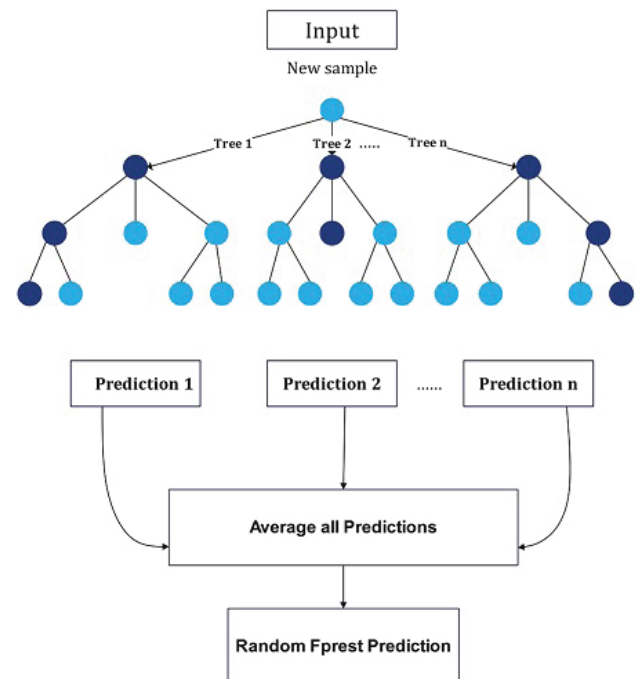


Figure 4. Visual presentation of random forest algorithm.

The benefits of averaging are limited by the degree of correlation between pairs of bagged trees. As B increases, the influence of the second term diminishes while the first term remains. RF aims to enhance the variance reduction of bagging without significantly increasing variance by reducing the correlation between trees. That is achieved during the tree-growing process by randomly selecting input variables. Specifically, when constructing a tree on a bootstrapped dataset, randomly select $m \leq p$ input variables as candidates for splitting.

In many instances, the value of m is typically $p/3$ or even 1. After growing $\{T(x; \Theta_b)\}_{b=1}^B$ trees, the random forest regression predictor is as follows [24]:

$$\hat{f}_{rf}^B(x) = \frac{1}{B} \sum_{b=1}^B T(x; \Theta_b) \quad (6)$$

Lowering the value of m intuitively reduces the correlation between any two trees in the ensemble, thereby decreasing the variance of the average using Eq 6. Here, Θ_b represents the split variables, cut points at each node, and terminal-node values of the b th random forest tree. Figure 4 illustrates this algorithm schematically.

Multiple Linear Regression

MLR, also known simply as multiple regression, is a statistical method that utilizes multiple explanatory variables to forecast the outcome of a response variable. Its primary objective is establishing the linear association between the explanatory (independent) variables and the response (dependent) variable. Essentially, multiple regression extends the principles of ordinary least squares (OLS) regression by accommodating more than one explanatory variable [33]. It helps to determine how changes in multiple independent factors affect the dependent variable. The model finds the best-fitting linear relationship between the independent variables and the dependent variable by minimizing the differences between observed and predicted values [34]:

$$y_i = \beta_0 + \beta_1 x_{i,1} + \beta_2 x_{i,2} + \dots + \beta_p x_{i,p} + \varepsilon_i \quad (7)$$

where y_i is the scalar dependent variable y and $x_{i,k}$ ($k = 1 \dots P$) is the k th independent in the i th observation. The term ε_i denotes the disturbance and signifies the portion of y_i that remains unexplained.

It should be noted that MLR has $p + 1$ parameters, including β_0 as intercept and p slope coefficients, each of which corresponds to an independent variable. The intercept β_0 signifies the baseline value of y when both independent variables are at zero, that is when $x_1 = x_2 = 0$. It represents the slope coefficient linked with the first independent variable, x_1 , indicating how y changes as x_1 increases by one unit while keeping β_0 constant. β_2 , on the other hand, is the slope coefficient associated with β_0 , illustrating the change in y for every one unit increase in β_0 while controlling for

the influence of x_1 , and the term ε_i represents the disturbance and has a variance of σ_2 .

In the following sections, we will compare the performance of three popular predictive modeling methods: the M5 tree, Random Forest (RF), and MLR. Each method is used to build models that learn from past data and make predictions for new data. While they all aim to predict outcomes, their techniques, assumptions, and how well they work for different types of problems differ.

Performance Evaluation

In order to assess the performance of ML methods, several metrics must be considered. The following statistical criteria can be considered for this purpose.

Coefficient of Determination

The correlation coefficient provides insight into the intensity of the linear connection between two variables. Furthermore, it indicates whether the linearity is substantial enough to warrant the application of a model to the dataset. This parameter can be formulated as follows [35]:

$$R^2 = \frac{\sum_{i=1}^n (DO_i - \langle DO \rangle)(DS_i - \langle DS \rangle)}{\sqrt{\sum_{i=1}^n (DO_i - \langle DO \rangle)^2} \sqrt{\sum_{i=1}^n (DS_i - \langle DS \rangle)^2}} \quad (8)$$

In this formula, DO_i and DS_i are observed and simulated data, respectively, and $\langle DO \rangle$ is the mean of observed data, and $\langle DS \rangle$ is the mean of simulated data.

Root Mean Square Error

Root Mean Square Error (*RMSE*) signifies the sample's standard deviation of the disparities between actual values and predicted. While *RMSE* is an effective gauge of precision, it is best suited for contrasting predictive discrepancies among various models concerning specific variables due to its reliance on the scale.

It gauges the overall effectiveness spanning the complete dataset spectrum and offers a robust evaluation of the model. A flawless model would yield an *RMSE* of 0 in an ideal scenario. It can be defined as [35]:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (DS_i - DO_i)^2} \quad (9)$$

Mean Absolute Error

Mean Absolute Error (*MAE*) is a statistical metric used to quantify the average magnitude of errors between predicted values and actual values. It measures the absolute difference between these values, disregarding whether the prediction overestimates or underestimates the actual value.

This parameter is defined as:

$$MAE = \frac{1}{n} \sum_{i=1}^n |DS_i - DO_i| \quad (10)$$

RESULTS AND DISCUSSION

Predicting BLOD as an experimental parameter is valuable, especially when conducting numerous experiments under various conditions is impractical. Machine learning algorithms like M5, RF, and MLR can help predict this parameter across different scenarios. Therefore, evaluating and comparing these predictions with actual experimental

values is crucial. This section uses these algorithms to predict BLOD and compare the predictions with experimental results. We evaluated the performance of the machine learning methods using three established statistical measures: the coefficient of determination (R^2), $RMSE$, and MAE .

Higher R^2 and lower $RMSE$ and MAE values indicate better accuracy in the models' predictions. Experimental

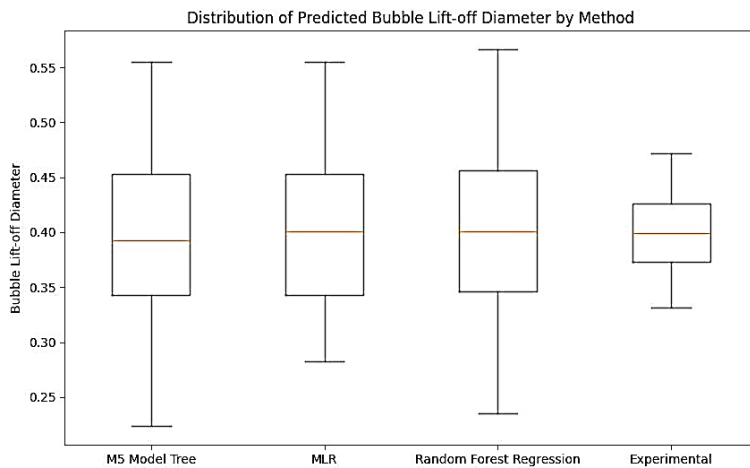


Figure 5. Distribution of BLOD predicted by different ML models.

Table 4. Predicting BLOD by using M5 and RF methods

Measured Bubble lift-off diameter (Tabrizi et al. [36])	M5 Model Tree	MLR	Random forest regression (RF)
0.4878	0.5419	0.4995	0.4301
0.3923	0.3705	0.4226	0.3977
0.3336	0.3119	0.3453	0.3781
0.5419	0.5202	0.5330	0.4672
0.4463	0.4246	0.4562	0.4283
0.3877	0.4006	0.3788	0.4009
0.5549	0.5419	0.5665	0.4721
0.4593	0.4463	0.4897	0.4341
0.4006	0.3877	0.4124	0.4082
0.4661	0.5202	0.4445	0.4220
0.3705	0.3923	0.3677	0.3891
0.3119	0.3336	0.2904	0.3720
0.5202	0.5331	0.4781	0.4554
0.4246	0.4376	0.4012	0.4161
0.3659	0.3789	0.3239	0.3924
0.5331	0.5549	0.5116	0.4625
0.4376	0.4593	0.4347	0.4245
0.3789	0.3659	0.3574	0.4012
0.3779	0.4320	0.3896	0.3743
0.2823	0.3364	0.3127	0.3438
0.2237	0.2823	0.2354	0.3313
0.4320	0.4449	0.4231	0.3928
0.3364	0.3494	0.3462	0.3611
0.2778	0.3364	0.2689	0.3423
0.4449	0.4320	0.4566	0.3991
0.3494	0.3364	0.3798	0.3687
0.2907	0.3364	0.3024	0.3501

data of Tabrizi et al. [36] were used to generate predictions with the M5, RF, and MLR methods. The results are presented in Tables 4 and 5. To better understand the models' performance, the distribution of these predictions is shown in Figure 5.

To interpret the results, defining the parameters in Table 5 is essential. R^2 , a coefficient of determination, is a statistical measure that indicates how well the variability in the dependent variable can be predicted based on the independent variables in a regression model. It simply shows how well the regression model fits the observed data points [37].

The value of R^2 falls between 0 and 1; $R^2 = 0$ means that the model does not explain any of the variability in the dependent variable around its mean, indicating that the model does not capture any discernible patterns in the data. Conversely, $R^2 = 1$ means that the model perfectly explains all the variations in the dependent variable. When $0 < R^2 < 1$, it indicates the proportion of the variance in the dependent variable explained by the independent variables included in the model. For example, an R^2 value of 0.75 means that the independent variables in the model explain 75% of the

Table 5. Accuracy and effectiveness of predictive models in comparison to actual observed values

	R^2	RMSE	MAE
M5 model tree	0.6405	0.0295	0.0263
RF	0.9872	0.0095	0.0072
MLR	0.9416	0.0204	0.0004

variability in the dependent variable. In comparison, the remaining 25% is unexplained and may be due to other factors or randomness.

RMSE shows the average size of errors in a model's predictions. It helps measure how close the predictions are to the actual values. By taking the square root, RMSE is expressed in the same units as the target variable, making it easier to understand [38]. MAE measures the average error size between actual values and predicted. It provides a straightforward way to assess how much a model's predictions deviate from the actual values without considering

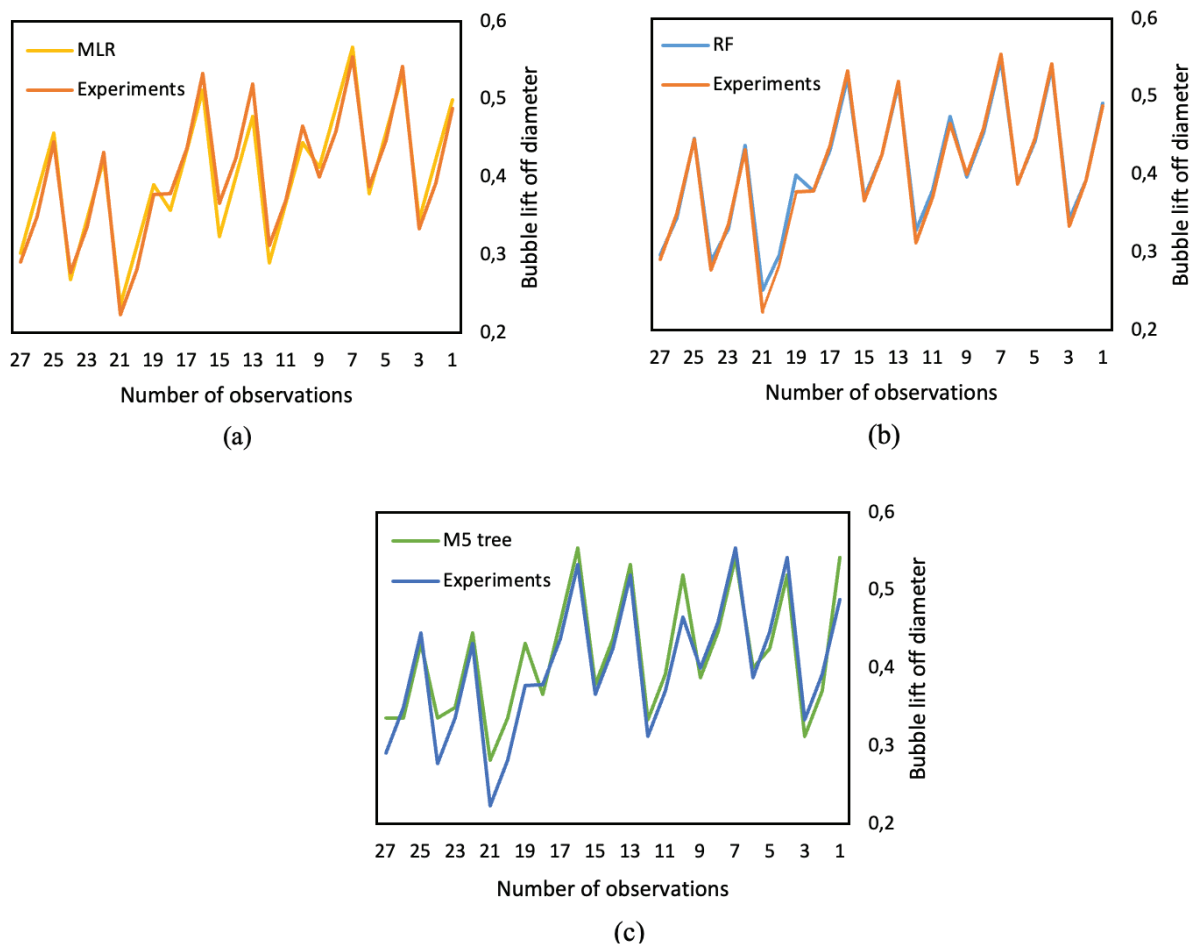


Figure 6. Predicted BLOD by using ML models including a) multi-linear regression, b) random forest, c) m5 tree, and comparison with experimentally measured values.

the direction of the errors [39]. To evaluate the comparative fluctuations of the applied modeling methodologies against the experimental data, Figure 6 shows a graph of the relationship between the number of observations and the BLOD.

Cross-Validation

In machine learning and data mining, a common challenge is dealing with a limited number of datasets, often just a few hundred observations. Splitting this data into modeling and testing sets can be problematic. The testing set might be too small to provide reliable results, or the modeling set may need more data to build an accurate predictive model [40].

A practical solution is cross-validation, which helps balance the challenges of limited data. Cross-validation allows us to evaluate a model's performance on various subsets of the data, giving insights into its predictive accuracy.

Given the limited studies on BLOD, especially in a magnetic field, cross-validation is essential when applying machine learning techniques to predict BLOD under these conditions. The results of the cross-validation are shown in Table 6, providing insights into how well the M5, RF, and MLR models predict the experimental data.

Table 6. Average Mean Square error

ML algorithm	AMSE
M5 tree	0.0010
RF	0.0011
MLR	0.0010

The RF (Random Forest) algorithm shows limited accuracy in predicting experimental results. That may be due to the complex nature of the model and the need for precise hyperparameter tuning. Hyperparameter tuning involves finding the optimal values for the model's hyperparameters, which are set before the learning process and control how the algorithm behaves.

In the M5 Algorithm, three critical hyperparameters stand out [21, 41]:

- Number of Instances Per Leaf (N): This parameter dictates the minimum number of instances needed to form a leaf node in the regression tree. A higher value of N may lead to simpler trees, mitigating the risk of overfitting while potentially sacrificing model flexibility.
- Pruning Method: Various pruning methods, such as error-based or cost-complexity pruning, are employed to prevent overfitting by eliminating tree branches that do not significantly enhance predictive accuracy.
- Minimum Significance Level (α): This parameter determines the significance level utilized during pruning, determining whether the improvement achieved by

splitting a node is statistically significant enough to warrant the split.

On the other hand, RF features four distinct hyperparameters:

- Number of Trees (n estimators): This parameter specifies the quantity of trees in the forest. A higher number of trees may enhance generalization performance but could increase computational demands.
- Maximum Depth of Trees: This parameter sets the maximum depth of each decision tree in the forest. Deeper trees can capture more intricate relationships in the data but may elevate the risk of overfitting.
- Number of Features Considered for Split: RF randomly selects a subset of features for each split in a decision tree. This hyperparameter controls the size of this subset, with a smaller subset introducing more randomness and potentially preventing overfitting.
- Minimum Samples per Leaf: Similar to M5, this parameter specifies the minimum number of samples needed to form a leaf node in each decision tree, aiding in controlling the complexity of individual trees and mitigating overfitting.

Overall, the hyperparameter tuning process for M5 and RF involves experimenting with different values for their respective hyperparameters and selecting the combination that yields the best performance on a validation set or through cross-validation. The specific hyperparameters and tuning strategies may vary depending on the dataset and the desired balance between model complexity, interpretability, and predictive accuracy. The main point is that if the dataset is small or if cross-validation is performed on a limited number of folds, the variability in the performance metrics may be higher, making it difficult to distinguish between models.

Bootstrap Validation

Cross-validation results show the importance of investigating the performance of the models more. Therefore, in this section, we performed bootstrap validation.

In summary, bootstrapping methods are used to assess how well a sample's parameter value estimates the larger population's true value [42]. The theory behind this approach has been published previously [43-46]. For bootstrap validation, we organized our data into arrays based on variables such as subcooling, pressure, mass flux, and BLOD and then randomly resampled to create bootstrap datasets.

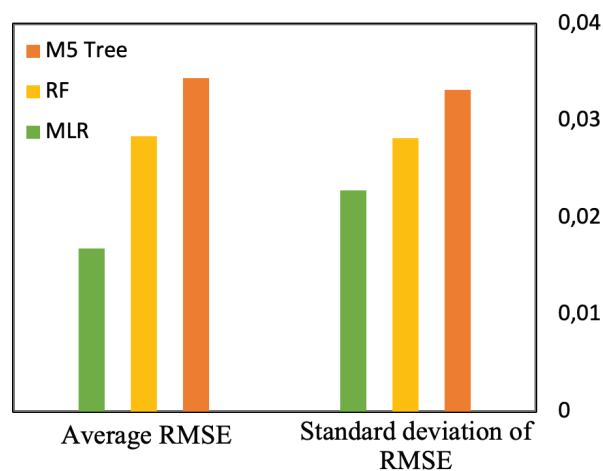
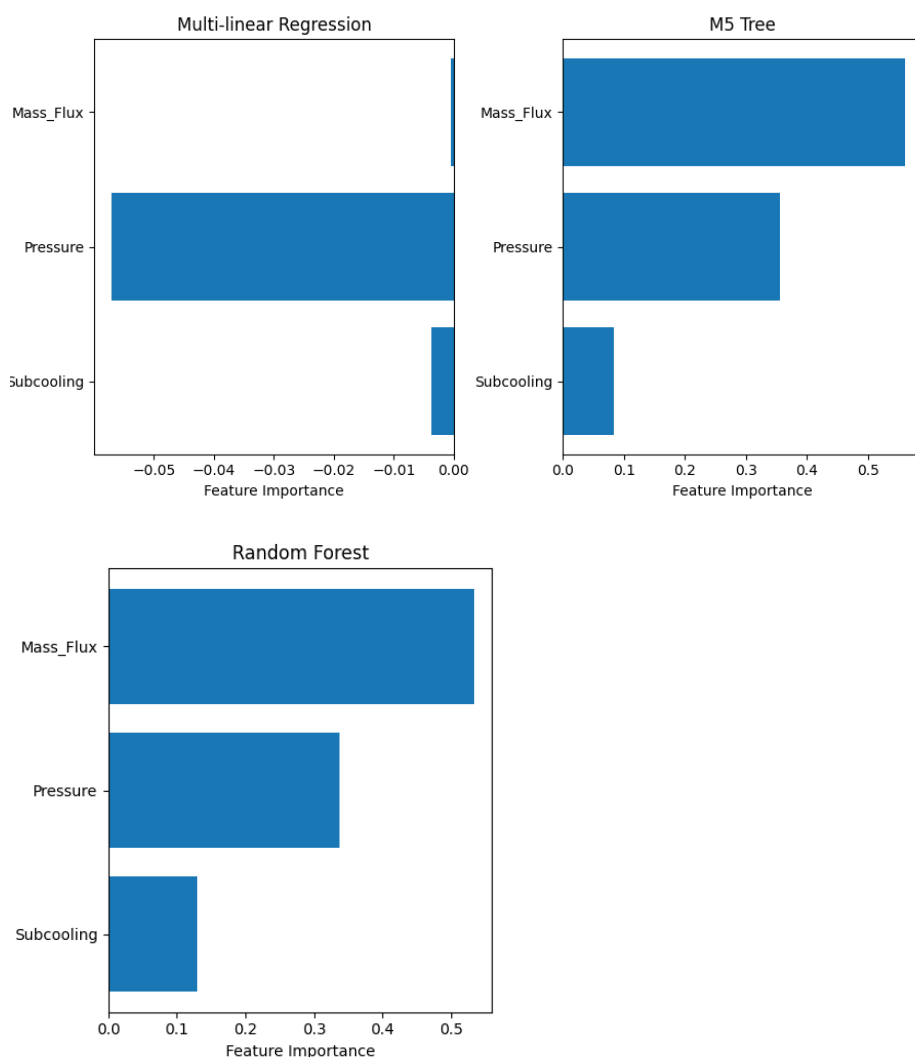
We trained our model on that specific subset of data for each bootstrap sample. Each iteration of the bootstrap validation process involves training the model on a slightly different version of the original dataset. For each validation set, we calculated performance metrics including MAE, RMSE, and MSE. To obtain stable estimates of these metrics, we repeated this process for 1,000 iterations. The results are summarized in Table 7.

Table 7. Calculated Average RMSE and Standard Deviation of RMSE to evaluate the model's performance

Model	Average RMSE	Standard Deviation of RMSE
M5 Tree	0.0332	0.0344
RF	0.0282	0.0284
MLR	0.0228	0.0168

The results show that MLR model perform better with the lowest average *RMSE* about 0.0228 this value for MLR suggests that its performance demonstrates more uniformity across diverse iterations of the bootstrap validation in comparison to the other models which means this model provides reliable predictions.

For a better understanding of the performance of the models, the obtained results from the validation process are depicted in Figure 7.

**Figure 7.** Average *RMSE* and standard deviation of *RMSE* comparison between three different ML models.**Figure 8.** Feature importance for evaluating the effect of different parameters on the target variable prediction.

Effect of Different Parameters on Blod (Feature Importance)

By performing feature importance analyses, valuable insights into the complex dynamics of bubble motion and optimized process conditions can be gained accordingly. Figure 8 illustrates the importance of different parameters in predicting BLOD using the ML method. This figure shows that operation condition pressure and mass flux significantly impact the prediction of the target value [47].

Figure 8 shows that the feature importance in the MLR model is negative, unlike the two other models. This difference arises from the models' approaches. In MLR, the feature coefficient shows the change in the target value (in our study, BLOD) per unit change of that feature (while holding other features constant). Therefore, negative coefficients indicate that an increase in the corresponding feature leads to a decrease in the target variable. In this model, feature importance does not necessarily show the importance of the features in predicting the target value. It simply indicates the magnitude and direction of each feature in prediction.

On the other hand, feature importance in tree-based models like RF and M5 tree models measures how much each feature contributes to decreasing the impurity (e.g., variance) in the prediction. Therefore, this parameter is typically positive and presents each feature's relative importance in making predictions. Higher feature importance values indicate that the feature is more critical for making accurate predictions.

Several recent studies have emphasized the critical role of geometry in determining heat transfer performance in boiling and evaporation phenomena. Dalkılıç [48] provides a comprehensive review of flow boiling in mini and micro-channels with enhanced geometries, underscoring how geometric modifications significantly affect thermal behavior. Koca et al. [49] simulate boiling heat transfer in rectangular milli-channels, showing how even minor geometric changes can lead to notable differences in heat transfer efficiency. Similarly, Nakhjavani and Zadeh [50] investigate annular heat exchangers using nanofluids, where geometry directly influences flow dynamics and thermal characteristics. Basnet et al. [51] further extend this discussion by modeling droplet evaporation and heating processes, where surface shape and boundary conditions play essential roles. These findings collectively suggest that geometry is not merely a physical parameter but a potential high-impact feature. Therefore, incorporating geometric descriptors in machine learning-based feature importance analyses could be a promising direction for future research, potentially enabling more accurate and generalizable models for complex thermal systems.

CONCLUSION

Although the experimental results are precious for investigating a physical phenomenon, there are limitations in the number of tests performed and the results obtained.

Therefore, finding a way to predict experimental results is significant. In this study, the ML techniques was used to predict the BLOD in flow boiling, achieving results that closely matched experimental data. Through rigorous calculations and assessments, including cross and bootstrap validation of algorithms, we confirmed that models like Random Forest, M5 tree, and Multiple linear regression can accurately predict BLOD. Among these, the Multi-linear regression model had the best accuracy.

We also found that mass flux and pressure are the most influential factors affecting the performance of these ML models. That highlights the need for precise control of these parameters in future studies. Our findings underscore the potential of ML techniques to provide valuable insights, especially when experimental data is limited. This study can pave the way for more extensive research utilizing machine learning to predict experimental outcomes in similar contexts. Expanding the dataset and exploring additional machine learning algorithms could further enhance predictive accuracy. Additionally, examining these models' long-term stability and performance under different operational conditions will be crucial for ensuring their reliability. By integrating ML with experimental research, it is possible to enhance our understanding and efficiency in predicting critical parameters in complex phenomena like flow boiling.

NOMENCLATURE

d_w	Bubble contact diameter on the heater surface (m)
f	Bubble departure frequency (Hz)
k	Thermal conductivity ($W/m^{\circ}C$)
T_{sat}	Saturation temperature ($^{\circ}C$)
u	velocity (m/s)

Greek Symbols

α	Thermal diffusivity (m^2/s)
ΔT_{sat}	Wall superheat ($^{\circ}C$)
ΔT_{sub}	Liquid subcooling ($^{\circ}C$)
ρ	Density (kg/m^3)

AUTHORSHIP CONTRIBUTIONS

Authors equally contributed to this work.

DATA AVAILABILITY STATEMENT

The authors confirm that the data that supports the findings of this study are available within the article. Raw data that support the finding of this study are available from the corresponding author, upon reasonable request.

CONFLICT OF INTEREST

The author declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

ETHICS

There are no ethical issues with the publication of this manuscript.

STATEMENT ON THE USE OF ARTIFICIAL INTELLIGENCE

Artificial intelligence was not used in the preparation of the article.

REFERENCES

- [1] Hoang NH, Chu IC, Euh DJ, Song CH. A mechanistic model for predicting the maximum diameter of vapor bubbles in a subcooled boiling flow. *Intern J Heat Mass Transf* 2016;94:174–179. [\[CrossRef\]](#)
- [2] Situ R, Hibiki T, Ishii M, Mori M. Bubble lift-off size in forced convective subcooled boiling flow. *Intern J Heat Mass Transf* 2005;48:5536–5548. [\[CrossRef\]](#)
- [3] Cho YJ, Yum SB, Lee JH, Park GC. Development of bubble departure and lift-off diameter models in low heat flux and low flow velocity conditions. *Intern J Heat Mass Transf* 2011;54:3234–3244. [\[CrossRef\]](#)
- [4] Levy S. Forced convection subcooled boiling—prediction of vapor volumetric fraction. *Intern J Heat Mass Transf* 1967;10:951–965. [\[CrossRef\]](#)
- [5] Staub FW. The void fraction in subcooled boiling—prediction of the initial point of net vapor generation. *J Heat Transf* 1968;90:151–157. [\[CrossRef\]](#)
- [6] Zeng LZ, Klausner JF, Bernhard DM, Mei R. A unified model for the prediction of bubble detachment diameters in boiling systems—II. Flow boiling. *Intern J Heat Mass Transf* 1993;36:2271–2279. [\[CrossRef\]](#)
- [7] Hosseinzadeh K, Moghaddam MA, Hatami M, Ganji DD, Ommi F. Experimental and numerical study for the effect of aqueous solution on heat transfer characteristics of two phase close thermosyphon. *Intern Commun Heat Mass Transf* 2022;135:106129. [\[CrossRef\]](#)
- [8] Hosseinzadeh K, Ganji DD, Ommi F. Effect of SiO₂ super-hydrophobic coating and self-rewetting fluid on two phase closed thermosyphon heat transfer characteristics: An experimental and numerical study. *J Mol Liq* 2020;315:113748. [\[CrossRef\]](#)
- [9] Prodanovic V, Fraser D, Salcudean M. Bubble behavior in subcooled flow boiling of water at low pressures and low flow rates. *Intern J Multip Flow* 2002;28:1–19. [\[CrossRef\]](#)
- [10] Chu IC, No HC, Song CH. Bubble lift-off diameter and nucleation frequency in vertical subcooled boiling flow. *J Nuc Sci Technol* 2011;48:936–949. [\[CrossRef\]](#)
- [11] Ünal HC. Maximum bubble diameter, maximum bubble-growth time and bubble-growth rate during the subcooled nucleate flow boiling of water up to 17.7 MN/m². *Intern J Heat Mass Transf* 1976;19:643–649. [\[CrossRef\]](#)
- [12] Mehdi S, Borumand M, Hwang G. Accurate and robust predictions of pool boiling heat transfer with micro-structured surfaces using probabilistic machine learning models. *Intern J Heat Mass Transf* 2024;226:125487. [\[CrossRef\]](#)
- [13] Bard A, Qiu Y, Kharangate CR, French R. Consolidated modeling and prediction of heat transfer coefficients for saturated flow boiling in mini/micro-channels using machine learning methods. *App Therm Eng* 2022;210:118305. [\[CrossRef\]](#)
- [14] Qiu Y, Garg D, Kim SM, Mudawar I, Kharangate CR. Machine learning algorithms to predict flow boiling pressure drop in mini/micro-channels based on universal consolidated data. *Intern J Heat Mass Transf* 2021;178:121607. [\[CrossRef\]](#)
- [15] He Y, Hu C, Li H, Hu X, Tang D. Reliable predictions of bubble departure frequency in subcooled flow boiling: A machine learning-based approach. *Intern J Heat Mass Transf* 2022;195:123217. [\[CrossRef\]](#)
- [16] Zhang J, Zhong D, Shi H, Meng J, Chen L. Machine learning prediction of critical heat flux on downward facing surfaces. *Intern J Heat Mass Transf* 2022;191:122857. [\[CrossRef\]](#)
- [17] Cabarcos A, Paz C, Suarez E, Vence J. Application of supervised learning algorithms for temperature prediction in nucleate flow boiling. *App Therm Eng* 2024;240:122155. [\[CrossRef\]](#)
- [18] Klausner JF, Mei R, Bernhard DM, Zeng LZ. Vapor bubble departure in forced convection boiling. *Intern J Heat Mass Transf* 1993;36:651–662. [\[CrossRef\]](#)
- [19] Okawa T, Kaiho K, Sakamoto S, Enoki K. Observation and modelling of bubble dynamics in isolated bubble regime in subcooled flow boiling. *Nuc Eng Des* 2018;335:400–408. [\[CrossRef\]](#)
- [20] Basu N. Modeling and experiments for wall heat flux partitioning during subcooled flow boiling of water at low pressures. (ProQuest Dissertations & Theses). Los Angeles: University of California; 2003. Report No: 3081174. [\[CrossRef\]](#)
- [21] Quinlan JR. Learning With Continuous Classes. *Proceedings of Australian Joint Conference on Artificial Intelligence*, 16-18 November 1992. Hobart: Scientific Research; 1992. s. 343–348.
- [22] Feng Y, Hao W, Li H, Cui N, Gong D, Gao D. Machine learning models to quantify and map daily global solar radiation and photovoltaic power. *Renew Sustain Energy Review* 2020;118:109393. [\[CrossRef\]](#)
- [23] Fan J, Ma X, Wu L, Zhang F, Yu X, Zeng W. Light gradient boosting machine: An efficient soft computing model for estimating daily reference evapotranspiration with local and external meteorological data. *Agricult Water Manag* 2019;225:105758. [\[CrossRef\]](#)
- [24] Nieto PJG, Gonzalo EG, Menendez L, Prado LA. Predicting the critical superconducting temperature using the random forest, MLP neural network, M5 model tree and multivariate linear regression. *Alex Eng J* 2024;86:144–156. [\[CrossRef\]](#)

- [25] Fan J, Wang X, Zhang F, Ma X, Wu L. Predicting daily diffuse horizontal solar radiation in various climatic regions of China using support vector machine and tree-based soft computing models with local and extrinsic climatic data. *J Clean Product* 2020;248:119264. [\[CrossRef\]](#)
- [26] Adnan RM, Liang Z, Trajkovic S, Kermani MZ, Li B, Kisi O. Daily streamflow prediction using optimally pruned extreme learning machine. *J Hydro* 2019;577:123981. [\[CrossRef\]](#)
- [27] Makridakis S, Spiliotis E, Assimakopoulos V, Chen Z, Gaba A, Tsetlin I, et al. The M5 uncertainty competition: Results, findings and conclusions. *Intern J Forecast* 2022;38:1365–1385. [\[CrossRef\]](#)
- [28] Das KS, Samui P, Sabat AK. Prediction of field hydraulic conductivity of clay liners using an artificial neural network and support vector machine. *Intern J Geomech* 2012;12:606–611. [\[CrossRef\]](#)
- [29] Deisenroth MP, Faisal AA, Ong CS. *Mathematics for machine learning*. 1st ed. Cambridge: Cambridge University Press; 2020. p. 398. [\[CrossRef\]](#)
- [30] Hastie T, Tibshirani R, Friedman J. *The elements of statistical learning. data mining, inference, and prediction*. 2nd ed. New York: Springer; 2009. p. 200. [\[CrossRef\]](#)
- [31] Bishop CM. *Pattern recognition and machine learning*. 2nd ed. Cambridge: Springer; 2006. p. 645–678.
- [32] Xia Y. Correlation and association analyses in microbiome study integrating multiomics in health and disease. *Prog Mol Biol Transl Sci* 2020;171:309–491. [\[CrossRef\]](#)
- [33] Berry WD. Probit/Logit and Other Binary Models. In: Leonard KK, editors. *Encyclopedia of Social Measurement*. New York: Elsevier; 2005. p. 161–169. [\[CrossRef\]](#)
- [34] Huang S. Linear regression analysis. In: Tierney RJ, Rizvi F, Ercikan K, editors. *International Encyclopedia of Education*. 4th ed. Oxford: Elsevier; 2023. p. 548–557. [\[CrossRef\]](#)
- [35] Chai T, Draxler RR. Root mean square error (RMSE) or mean absolute error (MAE)? – Arguments against avoiding RMSE in the literature. *Geosci Model Dev* 2014;7:1247–1250. [\[CrossRef\]](#)
- [36] Tabrizi AH, Aminfar H, Mohammadpourfard M, Zonouzi SA. Bubble lift-off diameter and frequency in ferrofluid subcooled flow boiling. *Heat Transf Eng* 2023;44:512–529. [\[CrossRef\]](#)
- [37] Ross SM. *Linear Regression*. In: Ross SM, editor. *Introductory statistics*. 3rd ed. Cambridge: Academic Press; 2010. p. 537–604. [\[CrossRef\]](#)
- [38] Tyagi K. Regression analysis. In: Pandey R, Khatri SK, Singh NK, Verma P, editors. *Artificial intelligence and machine learning for edge computing*. 1st ed. Cambridge: Academic Press; 2022. p. 53–63. [\[CrossRef\]](#)
- [39] Schneider P, Xhafa F. *Anomaly detection: Concepts and methods*. In: Schneider P, Xhafa F, editors. *Anomaly detection and complex event processing over IoT data streams*. 1st ed. Cambridge: Academic Press; 2022. p. 49–66. [\[CrossRef\]](#)
- [40] Refaeilzadeh P, Tang L, Liu H. Cross-Validation. In: Liu L, Özsu MT, editors. *Encyclopedia of database systems*. 1st ed. Boston: Springer US; 2009. p. 532–538. [\[CrossRef\]](#)
- [41] Holmes G, Hall M, Prank E. Generating Rule Sets from Model Trees. In: Foo N, editor. *Advanced Topics in Artificial Intelligence*. 1st ed. Berlin: Springer; 1999. p. 1–12. [\[CrossRef\]](#)
- [42] Parke J, Holford NHG, Charles BG. A procedure for generating bootstrap samples for the validation of nonlinear mixed-effects population models. *Comput Method Prog Biomed* 1999;59:19–29. [\[CrossRef\]](#)
- [43] Efron B. Bootstrap methods: Another look at the jackknife. In: Kotz S, Johnson NL, editors. *Breakthroughs in statistics: Methodology and distribution*. 1st ed. New York: Springer; 1992. p. 569–593. [\[CrossRef\]](#)
- [44] Efron B, Gong G. A leisurely look at the bootstrap, the jackknife, and cross-validation. *Am Stat* 1983;37:36–48. [\[CrossRef\]](#)
- [45] Efron B, Tibshirani R. Bootstrap methods for standard errors, confidence intervals, and other measures of statistical accuracy. *Statist Sci* 1986;1:54–75. [\[CrossRef\]](#)
- [46] Efron B, Tibshirani RJ. *An introduction to the bootstrap*. 1st ed. New York: Chapman and Hall/CRC; 1994. p. 456. [\[CrossRef\]](#)
- [47] Saarela M, Jauhiainen S. Comparison of feature importance measures as explanations for classification models. *SN App Sci* 2021;3:272. [\[CrossRef\]](#)
- [48] Dalkılıç AS. A review of flow boiling in mini and microchannel for enhanced geometries. *J Therm Eng* 2018;4:2037–2074. [\[CrossRef\]](#)
- [49] Koca A, Khalaji MN, Sepahyar S. Boiling heat transfer simulation in rectangular micro-channels. *J Therm Eng* 2021;7:1432–1447. [\[CrossRef\]](#)
- [50] Basnet M, Senthilkumar D, Yuvaraj R. Mathematical modelling of evaporation rate and heating of biodiesel blends of a single-component droplets. *J Therm Eng* 2023;9:1572–1584. [\[CrossRef\]](#)
- [51] Nakhjavani SH, Zadeh MAA. Flow boiling heat transfer characteristics of titanium oxide/water nanofluid (tio2/di water) in an annular heat exchanger. *J Therm Eng* 2020;6:592–603. [\[CrossRef\]](#)