



Optimizing energy consumption in heating systems through advanced machine learning and hybrid optimization for precision heating load prediction

Chengcheng Cai¹ · Na Feng¹ · Qianqian Liu¹

Received: 4 November 2024 / Accepted: 19 March 2025 / Published online: 29 March 2025
© Indian National Science Academy 2025

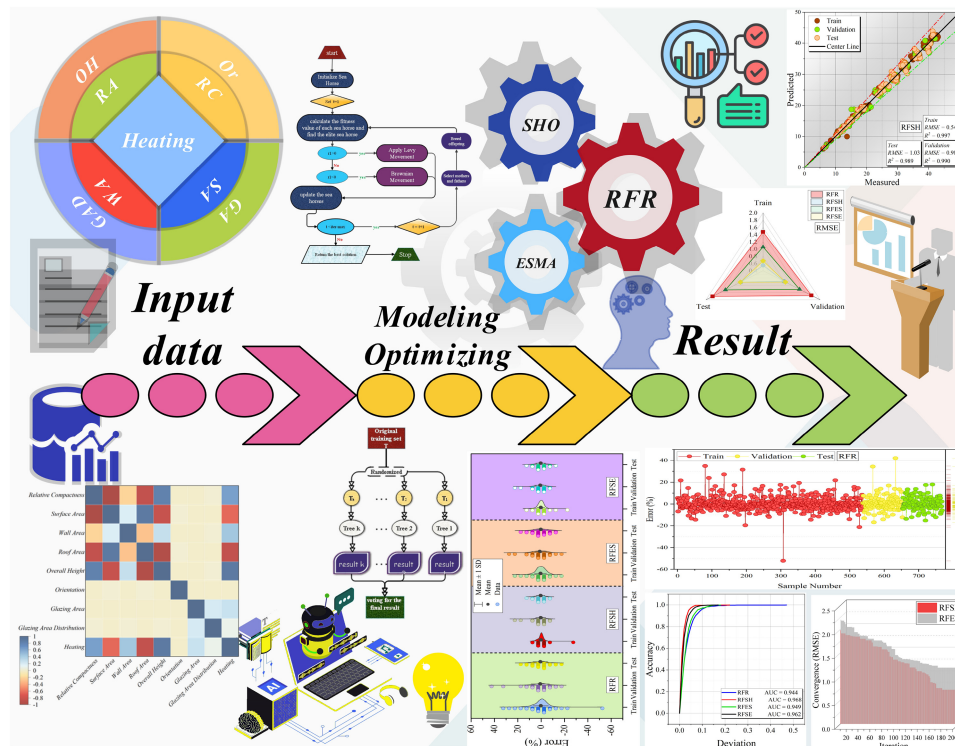
Abstract

Precise heating load prediction would be fundamental in optimizing energy consumption and assuring efficient operation of heating systems. Advanced machine learning methods were explicitly employed in this respect, using a random forest regressor that has shown high precision in heating load prediction. Then, the developed model was optimized by using two different methodologies: Sea Horse Optimizer and Equilibrium Slime Mould Algorithm. The approach of optimization ensembles, when put together, is termed RFSE, which brings out unique strengths and further enhances the overall accuracy of predictions. Of special note was the synergy between the RF model and the SHO, the interaction of which had the RFSH model showcase extraordinary predictive prowess. The RFSH model had a very impressive R-squared of 0.997, showing that this model fits the data very well, and the RMSE value was as low as 0.550, hence showing minimal error in predictions. This, therefore, makes the RFSH model very reliable while estimating the heating load. In this work, the strategic integration of ML with specialized optimization techniques underpins immense potential in solving complex real-world challenges, such as heating load prediction. This paper, therefore, constitutes a very important milestone in the quest toward optimization of energy efficiency and heating systems: precise and effective heating load forecasting. These advances ultimately contribute to energy conservation and cost reduction in heating applications.

✉ Na Feng
qzyuchen83@163.com

¹ Weifang Engineering Vocational College, Weifang 262500, Shandong, China

Graphical abstract



Keywords Building energy · Heating load · Random forest · Sea-horse optimizer · Equilibrium slime mould algorithm

Introduction

Machine learning and artificial intelligence are revolutionary technologies that allow computers to learn from experience and make intelligent decisions based on data input (Haykin 2009; Gordon, et al. 2022). L algorithms identify patterns and, therefore, play an important role in healthcare, finance, and autonomous vehicles (Jordan and Mitchell 1979). Artificial intelligence, on the other hand, includes reasoning, problem-solving, and a number of other cognitive processes similar in nature to those of human thought (Zhou 2021; Behnam Sedaghat et al. 2023; Sadaghat et al. 2023). AI goes beyond machine learning into other areas, including chatbots, virtual assistants, and robotics. AI and ML are of great importance in the precise forecast of heating energy consumption in buildings, thereby improving energy efficiency (Zhong et al. 2024). Such methods apply historical data, weather forecasting, and the characteristics of a building to make accurate forecasts of heating requirements. Artificial Intelligence-driven prediction is at the heart of sustaining a cost-effective building operation through the optimization of heating systems and the reduction of energy losses. This has shifted the focus of modern research toward building

energy efficiency improvements due to increased concerns about energy wastage and its considerable environmental impact. Buildings are, in fact, one of the largest sectors regarding energy consumption and greenhouse gas emissions; therefore, greater effectiveness by professionals is pursued to reduce these environmental impacts (Zhang et al. 2020; Guelpa et al. 2019; Shamshirband et al. 2015).

In this pursuit, much attention is directed towards multifaceted strategies for effective energy management (Han et al. 2024). Among the recent focuses of interest is the accurate forecasting of heating and energy consumption, which provides the backbone for these strategies. Precise energy consumption forecasts are important for any targeted efforts in effective energy conservation (Jaffal et al. 2009; Catalina et al. 2008; Xue et al. 2020; Nieto et al. 2019). Effective monitoring of building energy usage variation will establish operational inefficiencies and possible opportunities for energy conservation. This will be important in reducing wastage, increasing energy efficiency, and ensuring the optimum operation of the energy-saving techniques (Le et al. 2019; Taki and Rohani 2022). A number of research works have established the game-changing effect even from relatively small improvements in building energy consumption prediction. Such improvements can drastically cut down

overall energy consumption and are thus considered a major step towards sustainable energy use (Huang et al. 2024).

Besides, with proper forecasts, decision-makers can know when to optimize their energy usage to the advantage of building managers and residents. Such proactive measures may include optimizing lighting systems in line with anticipated insights into heating as well as energy consumption, which shall be put up with the occupant and daylight patterns accordingly (Ding et al. 2018). This lowers energy waste and improves residents' quality of life. It can also be utilized to program HVAC settings to meet forecasted demand patterns and prevent unnecessary heating loads for huge amounts of energy savings (Biswas et al. 2016; Lim et al. 2012).

Accurate consumption forecasts make the long-term expense of installing energy-efficient equipment extremely economical. For example, high-efficiency appliances, different types of insulation, and contemporary building automation systems can all lead to significant energy savings. In any regard, energy conservation would rely heavily on modifications to occupant behavior spurred With incentives and awareness-raising initiatives backed by predictive models. It will align the construction occupant behavior with energy-saving goals and create a sustainable construction culture. The cumulative effect of such initiatives will result in less waste, heightened efficiency, as well as a greener future. These valuable insights will be provided to all stakeholders so that they can make informed decisions. Energy-efficient buildings with less environmental impact will be developed through multi-faceted technology, equipment, and human aspects. This domain constitutes a prime area of research and innovation for a sustainable future in the face of global energy challenges (Roy et al. 2018).

On the other hand, Afzal et al. (Afzal et al. 2023) worked on an analysis of a challenging task building energy consumption forecasting for a heating/cooling application. They introduced the integration of multi-layer perceptron neural networks with eight meta-heuristic algorithms to make extra cautious tuning and hyperparameter optimization of the MLP model. Their elaborate analysis produced impressive results. The model MLP-PSOGWO achieved the highest overall R^2 values in most cases, with a maximum of 0.998 in heating load prediction. All in all, the best model shown had higher precision, accuracy, and efficiency in computations. Ahmad (Ahmad et al. 2017) estimated the accuracy of decision-based models in predicting building energy demand against those using neural networks. Zhang applied the support vector regression method to solve a structural energy consumption time-ordered data prediction issue (Zhang et al. 2016). Moradzadeh et al. (Moradzadeh et al. 2020) applied the Support Vector Regression and Multi-Layer Perceptron models. While it focused on Heating Load prediction, the MLP model performed outstandingly with an outstanding R-value of 0.9993.

The AI behind this breakthrough study is tuned to a novel synthesis that yields unprecedented predictive accuracy and optimization. Specifically, this seeks to address how ingeniously hybridization techniques can be married to the RF model as a means of extending the frontier of predictability reliability to new heights. These pioneering hybrid models represent a paradigm shift in the ways of conventional approaches that set new benchmarks in predictive analytics. Full-domain evaluations were performed in autonomous and hybrid modes to test their prowess with criticality. Performance metrics established and recorded with some degree of analogy in the past, such as R^2 , RMSE, MSE, SI, MNB, $n10_index$, and MARE, are used. This pragmatic, multi-faceted approach tends to avoid model bias while providing rich information on model efficiency. Further, the Seahorse optimizer and Equilibrium Slime Mould Algorithm were selected for the hybrid models due to their different strengths. An ensemble of SHO and ESMA with the RF model is visible in a prominent location. This strategic decision leverages the unique strengths of each optimizer and synergistically enhances the performance of predictive systems. This not only pushes the frontier of predictive analytics but also lays a strong foundation for the future in heating load predictions. This indeed stands as a testimony to the unwavering commitment to excellence in predictive modeling and promises a bright trajectory for advancements in this critical domain.

Resources and procedures

Information congregation

This work develops the efficiency and robustness analysis of intelligent models for the forecast of building heating demands. In this respect, a broad dataset of heating load production is used, wisely split into three subsets: the Train Set makes up 70% of the total; the Validation Set makes up 15%; and the Test Set makes up 15% as well. These subsets bear different responsibilities across different phases in model development and testing. Where the Training Set acts as the basis for the training of the predictive model, the Validation Set assists in fine-tuning the model parameters with a view to avoiding overfitting. Finally, the performance of the model on unseen data is rigorously evaluated by the Testing Set. The kernel of the predictive model will be constituted based on eight critical input variables that represent a great deal of impact on the heating behavior of the building. These variables epitomize some crucial characteristics of the building, which are evidenced by heat exchange and insulation properties. Tables 1, 2, 3 outline the numerical features of contribution and production variables for Training, Validation,



Table 1 The arithmetical properties of the contribution and production parameters for training phase

Parameters	Indicators				
	Category	Min	Utmost	Avg	St. Dev
RC	Input	0.620	0.980	0.7640	0.1060
SA (m ²)	Input	514.50	808.50	672.0650	88.1970
WA (m ²)	Input	245.0	416.50	318.1360	43.8810
RA (m)	Input	110.250	220.50	176.9650	45.1080
OH (m)	Input	3.50	7.0	5.2370	1.7500
Or	Input	2.0	5.0	3.5190	1.1280
GA (%)	Input	0.0	0.40	0.2400	0.1340
GAD	Input	0.0	5.0	2.7960	1.5570
Heating (kW)	Output	6.010	43.10	22.3260	10.1420

Table 2 The arithmetical properties of the contribution and production parameters for validating phase

Parameters	Indicators				
	Category	Min	Utmost	Avg	St. Dev
RC	Input	0.620	0.980	0.7560	0.1020
SA (m ²)	Input	514.50	808.50	677.4780	85.2410
WA (m ²)	Input	245.0	416.50	321.6960	44.4800
RA (m)	Input	110.250	220.50	177.8910	44.8960
OH (m)	Input	3.50	7.0	5.2040	1.7490
Or	Input	2.0	5.0	3.4090	1.0870
GA (%)	Input	0.0	0.40	0.2020	0.1280
GAD	Input	0	5	2.809	1.571
Heating (kW)	Output	6.77	42.11	21.464	9.843

Table 3 The arithmetical properties of the contribution and production parameters for testing phase

Parameters	Indicators				
	Category	Min	Utmost	Avg	St. Dev
RC	Input	0.62	0.98	0.773	0.108
SA (m ²)	Input	514.5	808.5	664.270	89.467
WA (m ²)	Input	245	416.5	317.009	41.160
RA (m)	Input	110.25	220.5	173.630	45.385
OH (m)	Input	3.5	7	5.357	1.747
Or	Input	2	5	3.504	1.098
GA (%)	Input	0	0.4	0.243	0.128
GAD	Input	0	5	2.896	1.494
Heating (kW)	Output	6.37	41.73	23.061	9.982

and Testing Phases, respectively. The list includes relative compactness, wall area, surface area, roof area, overall height, orientation, glazing area (%), and glazing area

distribution, along with their range and average values accordingly (Shariah et al. 1997).

The present data is an excellent basis on which predictive policies for heating demand in buildings can be realized. In this regard, it also forms an extension of previous studies, many of whose statistical analyses were required to understand characteristics and the distribution of interest. The study evaluates intelligent models using a partitioned dataset for the production of heating loads and places a strong emphasis on Training, Validation, and Testing Sets in model development. These key inputs drive heating behavior as per eight important input variables. It provides proof of the robustness and foundation that underpin this dataset for predictive strategies based on reliance upon prior research and statistical analysis. Plotting the correlation between input and output variables is displayed in Fig. 1.

Random forest (RF)

Principle of RF

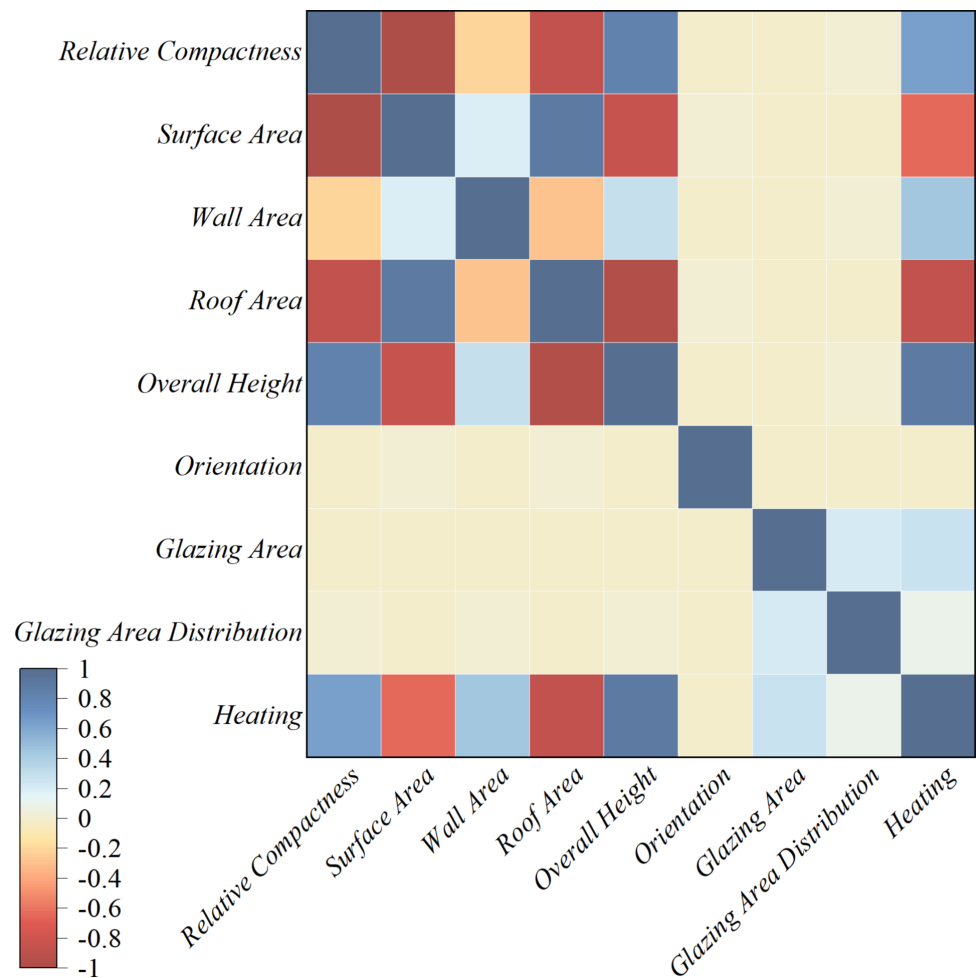
An RF system comprises a collection of tree-modeled categorizers displayed as $\{b(x, \mathfrak{N}_i), q = 1, \dots\}$, with apiece sapling making an individual ballot to ascertain the prevailing category for a provided contribution x . In this context, the $\{\mathfrak{N}_i\}$ signify random courses with identical properties that are individually generated.

An RF comprises a variety of arboreal configurations classifiers Crafted through the utilization of a training dataset and a stochastic element $\{\mathfrak{N}_i\}$, for the q -th tree within Breiman's framework (Biau and Scornet 2016). The stochastic elements are uncorrelated and uniform across any pair of trees, culminating in the establishment of a classification system $b(x, \mathfrak{N}_i)$, with 'x' denoting the input vector, when executing the procedure 'l' aeras, an array of categorizers is produced $\{b_1(x), b_2(x), \dots, b_q(x)\}$. This results in the formation of a resource for crafting numerous organizational frameworks. A traditional consensus vote determines the system's outcome, and the result process is computed accordingly.

$$B_{(x)} = \operatorname{argmax}_p \sum_{i=1}^q F(b_i(x) = V) \quad (1)$$

Every tree possesses the privilege to endorse the most optimal categorization outcome for an assumed contribution factor, and the aggregate of these distinct arboreal decision models is signified as $B(x)$. The outcome mutable is labeled as 'V,' and the pointer purpose is symbolized by $F(\cdot)$ (Sarica et al. 2017). Figure 2 shows the process for identifying the best possible classification outcome.

Fig. 1 Input and output variable correlation plot



Types of RF

The boundary metric, utilized in RF to measure how much the mean votes for the right session at X, V exceeds those for the wrong class, can be articulated as follows:

$$mc(X, V) = av_l F(b_l(X) = V) - \max_{j \neq V} av_l F(b_l(X) = j) \quad (2)$$

A higher number indicates greater precision as well as trust in the classification prediction, which is measured by the error rate function (Khajavi and Rastgoo 2023). The error in prediction beyond the training data for the categorization model may be characterized as:

$$OS^* = O_{X,V}(mc(X, F) < 0) \quad (3)$$

The stochastic element was put forth by Leo Breiman, where $b_q(X) = b(x, \mathfrak{N}_q)$, Adhering to the Rule of Big Numbers when there are enough decision trees (Lin et al. 2017). With the augmentation in the number of decision trees, OS^* gradually approaches a fixed value for nearly all orders of $\mathfrak{N}_1$. Breiman further showcased that RF is

resistant to overfitting and can ultimately provide a constant value for the prediction mistake.

$$O_{x,v}(O_\theta(b_l(x, \theta) = V) - \max_{j \neq V} O_\theta(b_l(x, \theta) = j) < 0) \quad (4)$$

Leo Breiman's work also established an upper bound for the generalization error.

$$OS^* \leq \bar{\beta}(1 - z^2)/z^2 \quad (5)$$

The generalization error in (RF) is impacted by two variables: the tree strength, indicated by 'z,' and the tree correlation, symbolized by the average correlation value ' $\bar{\beta}$ '. Reduced interconnectedness between the trees is shown by a declining correlation value, which improves RF performance.

Sea-horse optimizer (SHO)

Zhao introduced the SHO algorithm as a new meta-heuristic technique in 2022. The three main components of the SHO



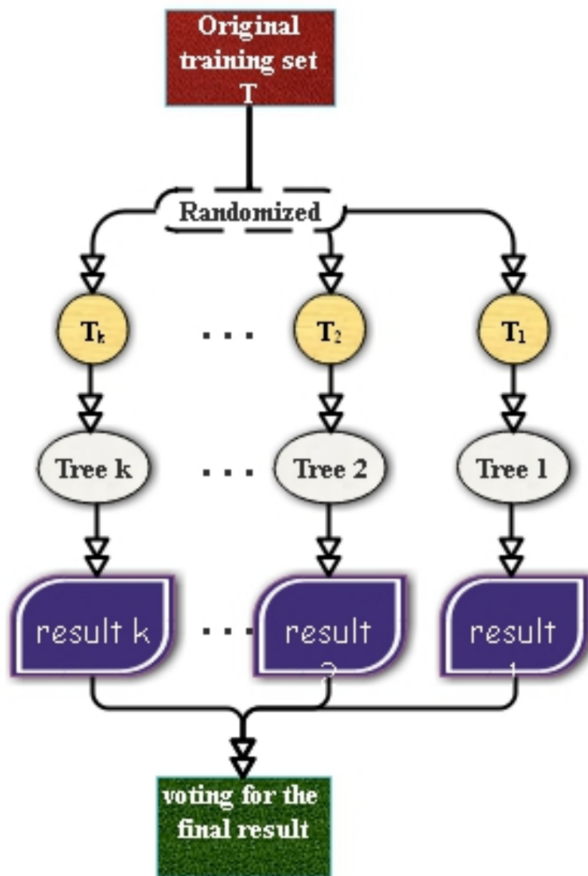


Fig. 2 Schematic of RF

algorithm—movement, hunting, and reproduction—are population-based meta-heuristics that simulate the social behavior of sea horses. To balance the advantages of exploration and exploitation, the algorithm employs both local and global search functions. The act of reproduction balances the mobility behavior; this is intended for searching locally, whereas the pursuit behavior is intended for worldwide search (Aribowo 2023). First, a population of possible solutions is generated via the *SHO* algorithm.

$$SH = \begin{bmatrix} x_1^1 & \dots & x_1^{dim} \\ \vdots & \dots & \vdots \\ x_{pop}^1 & \dots & x_{pop}^{dim} \end{bmatrix} \quad (6)$$

where *pop* indicates the population size utilized in the *SHO* technique and *dim* indicates the number of dimensions in the searching region, each sea horse in the population represents a potential solution inside the problem's search space. The elite person in a minimizing optimization issue is identified as the individual with the least favorable fitness score, indicated by X_{elite} . X_{elite} can be obtained using Eq. (7).

$$X_{elite} = \operatorname{argmin}(f(X_i)) \quad (7)$$

The function $f(\cdot)$ for a given problem represents the cost function, which assesses the appropriateness of potential solutions in the search space. *Levyflight* and Brownian motion are the two states involved in sea horse motion behavior. While *Levyflying* imitates the size of a sea horse's stride, enabling them to migrate as well as investigate diverse places to avoid excessive local exploitation, Brownian motion allows for greater exploration in the search space. In order to ascertain the sea horse's updated location in iteration t , these two possibilities may be articulated as follows:

$$X_{new}^2(t+1) = \begin{cases} X_i(t) + \operatorname{Levy}(\lambda) (X_{elite}(t) - X_i(t)) \times x \times y \times z \times X_{elite}(t) & r_1 > 0 \\ X_i(t) + \operatorname{rand} * l * \beta_t * X_{elite}(t) & r_1 < 0 \end{cases} \quad (8)$$

The metric λ of the Lévy flight distribution formula is picked at random from the interval $[0,2]$, defining *Levy*. The coordinates x, y , and z indicate the spiral movement component of *SHO*. The step size of the Lévy flight is controlled by the constant coefficient l ; the random walk coefficient for Brownian motion is denoted by βt . The normal random number r_1 is utilized to make the Brownian motion component stochastic (Zhao et al. 2023).

Sea horses' hunting strategies could conclude in either success or failure. When a sea horse moves more quickly, it may grab its prey; when it fails, it must examine the search field more thoroughly. Mathematically, this hunting behavior may be expressed as:

$$X_{new}^2(t+1) = \begin{cases} b * (X_{elite} - \operatorname{rand} * X_{new}^1(t)) + (1-b) * X_{elite}(t) & r_2 > 0.1 \\ (1-b) * (X_{new}^1(t) - \operatorname{rand} * X_{elite}) + b * X_{new}^1(t) & r_2 < 0.1 \end{cases} \quad (9)$$

where, after hunting at iteration t , the sea horse's new position is shown by $X_{new}^1(t)$, r_2 is the number inside $[0,1]$ that is created at random b is a straight-line lowering parameter that adjusts the seahorse's searching step length. Based on their fitness levels, sea horses' reproductive behavior separates the population into male and female groups, with males handling reproduction.

$$\text{mothers} = X_{sort}^2 \left(\frac{pop}{2} + 1 : pop \right) \quad (10)$$

while X_{sort}^2 indicates all X_{sort}^2 organized in an ascending sequence of their respective fitness scores, and the parents stand for the male and female populations in general, correspondingly. To generate a fresh batch of potential solutions, the algorithm chooses half of the population's most suitable members. The i -th offspring's expression is as follows:

$$X_i^{offspring} = r_3 X_i^{father} + (1-r_3) X_i^{mother} \quad (11)$$

Among the populations of men and women, X_i^{father} and X_i^{mother} individuals are chosen at random, where the random number between $[0,1]$ is denoted by r_3 .

The SHO approach, which is especially made to address optimization problems has displayed promising results in numerous applications employing continuous search spaces. Figure 3 shows the flowchart for the suggested SHO method.

With its efficacy and efficiency, the SHO algorithm is a promising method for various applications, and it provides a novel viewpoint on tackling optimization issues.

Equilibrium slime mould algorithm (ESMA)

Developing successful and efficient optimization techniques may be inspired by slime mold's foraging habits. Each slime

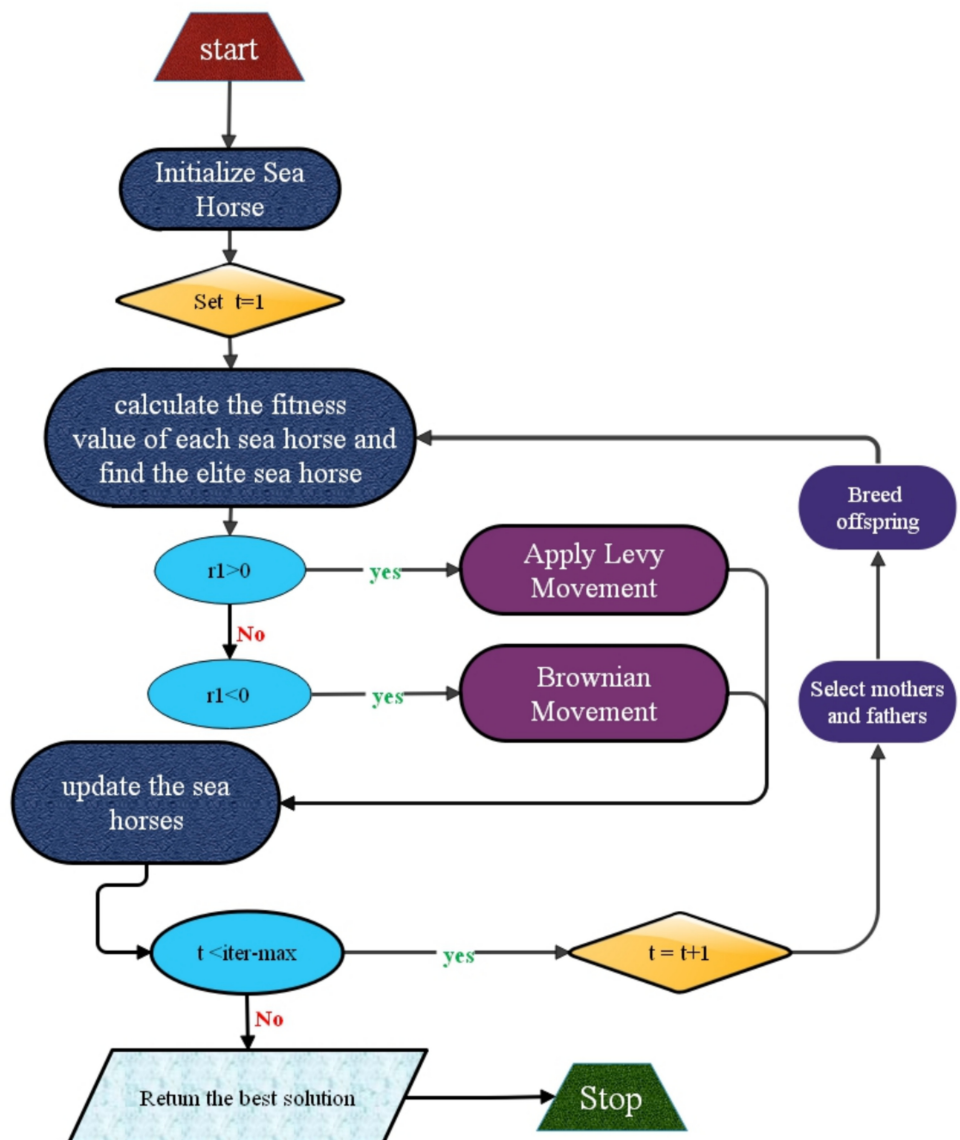
mold's initial location vector is determined randomly via a randomization procedure (Yin et al. 2022).

$$\vec{X}_i(t=1) = r_1 \cdot (UB - LB) + LB, i = 1, 2, \dots, N \quad (12)$$

Following the next repetition $(t+1)$, the usual location for the i -th goo mold, denoted as $X_i(j=1, 2, \dots, N)$, is determined using SMA in the manner described below:

$$\vec{X}_i(t+1) = \begin{cases} r_1 \cdot (UB - LB) + LB & r_1 < z \\ \vec{X}_{Gbest} + \text{step}_b \left(\vec{U} \cdot \frac{\vec{X}_i - \vec{X}_D}{C} \right) & r_2 < P_i(t) \text{ and } r_1 \geq z \\ \text{step}_b \cdot \vec{X}_i(t) & r_2 \geq P_i(t) \text{ and } r_1 \geq z \end{cases} \quad (13)$$

Fig. 3 The flowchart of the proposed SHO algorithm



$\vec{X}_{Gbest}$ denotes the price of the world-class optimal health achieved across repetition 1 toward t . Additionally, the variables r_1 and r_2 represent arbitrary values that fall between $[0, 1]$. A probability represented by z is used to eliminate and spread the slime mold. In the framework of this investigation, the constant value of z is 0.03 (Naik et al. 2021). Equation (14) The physical fitness values are arranged in increasing direction.

$$[sortf, sortIndex] = sort(f), \text{ where } f = \{f_1, f_2, \dots, f_N\} \quad (14)$$

Equation (15) is employed to compute $\vec{U}$.

$$\vec{U}(sortIndex(j)) = \begin{cases} 1 + r_3 \cdot \log\left(\frac{f_{Lbest} - sortf(j)}{f_{Lbest} - f_{Lworst}} + 1\right) & 1 \leq j \leq \frac{N}{2} \\ 1 - r_3 \cdot \log\left(\frac{f_{Lbest} - sortf(j)}{f_{Lbest} - f_{Lworst}} + 1\right) & \frac{N}{2} < j \leq N \end{cases} \quad (15)$$

A randomly generated arbitrary number, r_3 , in the interval $[0, 1]$, is used. f_{Lworst} and f_{Lbest} , respectively, display the local fitness values acquired throughout the current iteration and reflect the least favorable and highest values. Equations (16–17) to determine these fitness scores.

$$f_{Lbest} = sortf(1) \quad (16)$$

and

$$f_{Lworst} = sortf(N) \quad (17)$$

The parameter, P_i , which denotes the probability of chosen the route for the i -th gunk mold, is defined by the equation shown that follows:

$$P_i = \tanh\left|f(X_i) - f_{Gbest}\right| \quad (18)$$

For each $i = 1, 2, \dots, N$, the fitness value of the i -th slime mold in X_i is determined by $f(X_i)$. The supreme performance level recorded at the onset of the repetition up toward the present restatement is represented by f_{Gbest} . An even distribution from $-a$ to a determines the amount of steps in scale, which is represented by $(step)_a$. In a similar manner, a uniform dispersion from $-b$ to b determines the step's size, denoted by $(step)_b$. Equation (19), a function of the current iteration t and the maximum iteration T , yields the values of a and b :

$$a = \operatorname{arctanh}\left(-\left(\frac{t}{T}\right) + 1\right) \quad (19)$$

and

$$b = 1 - \frac{t}{T} \quad (20)$$

Equation (20) shows that even with the SMA's encouraging findings, the search procedure still has to be improved. Keep in mind that adding random slime molds might cause the search to go in a different direction. The effectiveness of the search procedure for choosing individuals X_D and X_C from an inventory of N slime molds may be limited by the local minimums. This section presents the EOSM, a novel optimization method. This algorithm replaces the position vector $\vec{X}_A$ array within deduced characteristics after a symmetry pond of 4 larger location courses. After that, the average position of this selection is computed utilizing the Equilibrium Enhancer (EO) concept as a reference. Equation (21) identifies the balanced pool's constituent parts with precision.

$$\begin{aligned} \vec{X}_{eq(1)} &= X(sortIndex(1)) \\ \vec{X}_{eq(2)} &= X(sortIndex(2)) \\ \vec{X}_{eq(3)} &= X(sortIndex(3)) \\ \vec{X}_{eq(4)} &= X(sortIndex(4)) \\ \vec{X}_{ave} &= \frac{\vec{X}_{eq(1)} + \vec{X}_{eq(2)} + \vec{X}_{eq(3)} + \vec{X}_{eq(4)}}{4} \end{aligned} \quad (21)$$

The equilibrium pool is built using a collection of five-position vectors, represented by $\vec{X}_{eq, pool}$

$$\vec{X}_{eq, pool} = \left\{ \vec{X}_{eq(1)}, \vec{X}_{eq(2)}, \vec{X}_{eq(3)}, \vec{X}_{eq(4)}, \vec{X}_{ave} \right\} \quad (22)$$

In ESMA, the place course for the i -th goocast, $X_i(j = 1, 2, \dots, N)$, during the novel repetition $(t + 1)$ is represented by the following Eqs. (23a), (23b), and (23c):

$$\vec{X}_i(t + 1) = r_1 \cdot (UB - LB) + LB, \text{ when } r_1 < z \quad (23a)$$

$$\vec{X}_i(t + 1) = \vec{X}_{Gbest} + \vec{step}_a \cdot \left(\vec{U} \cdot \vec{X}_{eq} - \vec{X}_D \right), \text{ when } r_2 < p_i(t) \text{ and } r_1 \geq z \quad (23b)$$

$$\vec{X}_i(t + 1) = \vec{step}_b \cdot \vec{X}_i(t), \text{ when } r_2 \geq p_i(t) \text{ and } r_1 \geq z \quad (23c)$$

The position vector $\vec{X}_{eq}$ is acquired by choosing a vector at random from the equilibrium pool. The algorithmic instrument z guarantees the effectiveness of ESMA by restricting restricted local occurrence and promoting

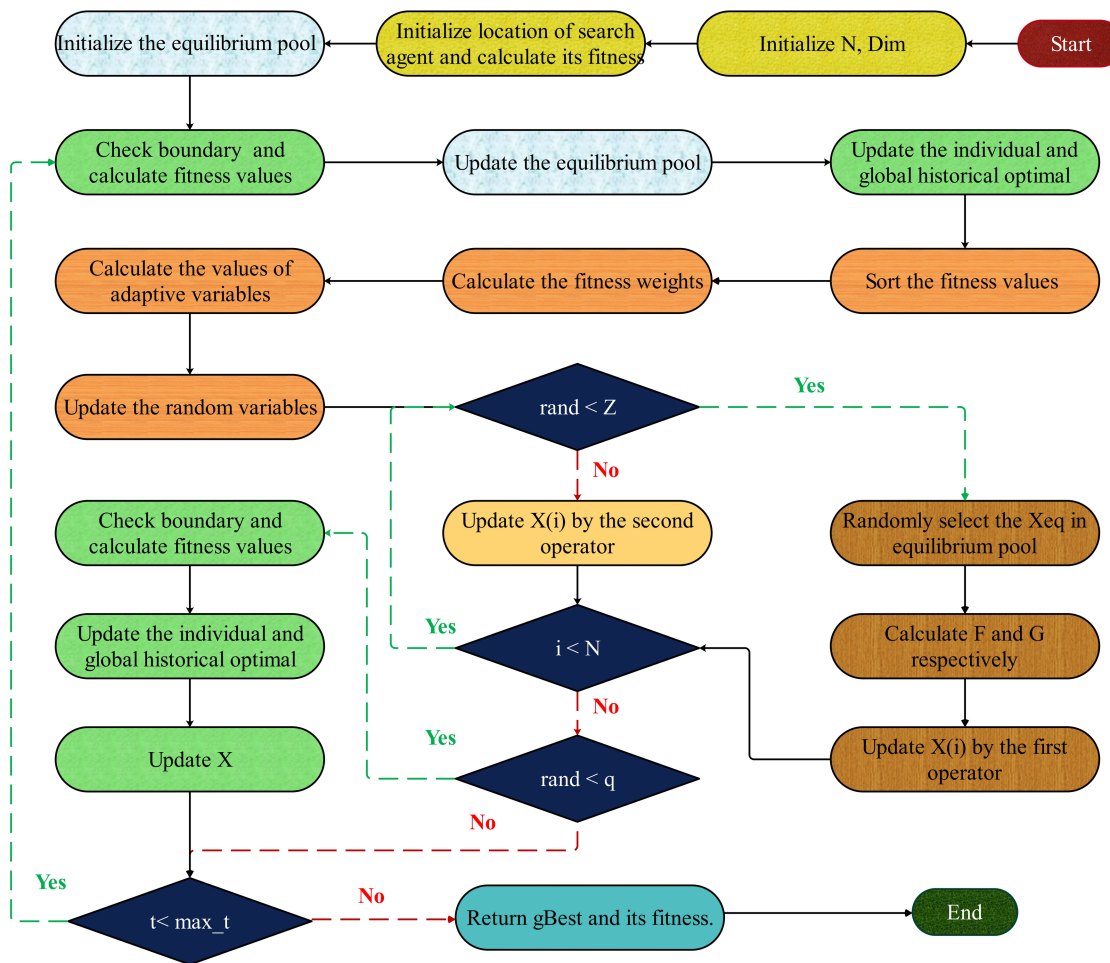


Fig. 4 The ESMA flowchart

exploration throughout the search phase. This is accomplished using a threshold value of 0.03 that was determined through testing. Notably, the *ESMA* method modifies the place itinerary in the next cycle by fusing a randomly generated path with the local best spot and the global ideal

position obtained from the current highest managing storehouse. This approach enables a fair trade-off between usage and examination. Detailed in Algorithm 1 is the suggested *ESMA*. In Fig. 4, an *ESMA* flowchart is shown.]



Begin**Inputs:** N, D, T, UB, LB , and z .

Initialization: $x_i = \{x_i^1, x_i^2, \dots, x_i^D\}$ for $i = 1, 2, \dots, N$ at a random position within the at the outset of iteration $t = 1$, establish a search range $[LB, UB]$ for a D – dimensional space to create $X = [\vec{X}_1, \vec{X}_2, \dots, \vec{X}_i, \dots, \vec{X}_N]'$.

While ($t \leq T$)

- Evaluate the fitness $f(x)$ of N slime mould.
- Sort the fitness value in ascending order (for minimization problem) using Eq
- Control the allowance issue $\vec{U}$ by means of Reckoning. (21).
- Revise the $\vec{X}_{Gbest}$.
- Regulate the a and b .

Aimed at $i = 1$: amount of gunk mould (N)

- Generate a random value r_1 .

Doubt $r_1 < z$

- Revise the location course $\vec{X}_i(t+1) = r_1 \cdot (UB - LB) + LB$.

Else If $r_1 \geq z$

- Update the probability p_i , $\overrightarrow{step_a}$ and $\overrightarrow{step_b}$.
- Arbitrarily select a single positional course $\vec{X}_{eq}$ after the s
- Make a chance worth r_2 .

If $r_2 < p_i$

- Choice a chance location $\vec{X}_D$ a course after X .
- Inform the location path $\vec{X}_i(t+1) = \vec{X}_{Gbest} + \overrightarrow{step_a} \cdot (\vec{U} \cdot \vec{X}_{eq} - \vec{X}_D)$.

Else If $r_2 \geq p_i$

- Update the position vector $\vec{X}_i(t+1) = \overrightarrow{step_b} \cdot \vec{X}_i(t)$.

End If**End If****End For**→ $t = t + 1$ **End while****Output:** $\vec{X}_{Gbest}$ **End****Performance evaluation methods**

A variety of metrics, including the previously mentioned Mean Square Error (MSE), Mean Normalized Bias (MNB), Correlation Coefficient (R^2), Scatter index (SI), Root Mean Square Error (RMSE), Mean Absolute Relative Error (MARE), and n10_index are used in this article to assess the models. An algorithm that has a high R^2 score has done well across the stages of testing, validation, and training. On the other side, smaller MSE and RMSE levels are preferable as they show fewer model failures. Equations (24–30) are used to calculate these measures.]

$$\bullet \text{ Coefficient of Correlation } R^2 = \left(\frac{\sum_{i=1}^W (h_i - \bar{h})(l_i - \bar{l})}{\sqrt{\left[\sum_{i=1}^W (h_i - \bar{h})^2 \right] \left[\sum_{i=1}^W (l_i - \bar{l})^2 \right]}} \right)^2 \quad (24)$$

$$\bullet \text{ Root Mean Square Error RMSE} = \sqrt{\frac{1}{W} \sum_{i=1}^W (l_i - h_i)^2} \quad (25)$$

$$\bullet \text{ Mean Square Error MSE} = \frac{1}{W} \sum_{i=1}^W (l_i - h_i)^2 \quad (26)$$

$$\bullet \text{ Mean Normalized Bias MNB} = \frac{1}{W} \sum_i^w \frac{l_i - h_i}{l_i - \bar{l}} \quad (27)$$

$$\bullet \text{ Scatter Index SI} = \frac{RMSE}{meam(l_i)} \quad (28)$$

$$\bullet \text{ N10_index } n10_{index} = \frac{w10}{w} \quad (29)$$

$$\bullet \text{ Mean Absolute Relative Error MARE} = \frac{1}{W} \sum_j^w \frac{|l_i - h_i|}{|\bar{l} - \bar{h}|} \quad (30)$$

In these equations, h_i and l_i allude to the forecasted and experimental principles correspondingly. The mean values of the experimental samples and predicted are denoted by $\bar{h}$ and $\bar{l}$, alternatively. W signifies the quantity of data points under examination.

Hybridization and ensembling

Hybridization and ensembling are probably the most meaningful strategies in predictive modeling, where each brings forth a different advantage in the case of complex tasks related to heating load prediction. In general, hybridization refers to the putting together of two or more techniques or methodologies into one single better technique. In the context provided herein, this is illustrated by the integration of state-of-the-art machine learning techniques with specialized optimization algorithms. The combination of the RFR model and SHO with ESMA allowed researchers to leverage the strengths of ML and optimization in improving heating load predictions. On the other hand, ensembling involves aggregating the predictions of several models into one final prediction, which usually outperforms individual models. This means that for the RFSE approach, assembly is done by combining the random forest model's predictions with the predictions of the SHO and ESMA optimization algorithms. This ensembling of outputs from the various techniques gives better performance compared to any single model or optimization algorithm. The synergy between hybridization and ensembling is realized in the RFSH model that emerges as the top performer in the

heating load prediction, while the integration of the Random Forest model with the Sea-Horse Optimizer will be done in a strategic way to enable researchers to leverage its strengths through complementary capabilities of both methods toward delivering an exceptionally accurate and reliable predictive tool. In any case, this work's successful application of hybridization and ensembling underlines their potential concerning efficient solution finding in complex real-world challenges. Combining various methodologies and exploiting the collective intelligence stemming from multiple models allows researchers to attain breakthroughs in predictive accuracy and performance, as shown in the context of heating load prediction. Such strategies represent an interesting tool for enhancing predictive modeling and encouraging innovation across various fields.

K-fold cross validation

During the k-fold cross-validation process, the dataset is partitioned into 'k' equal subsets (folds). The model is then trained iteratively 'k' times, where in each iteration, one fold is designated as the test set, while the remaining k–1 folds are utilized for training. This cycle continues until each fold has been used exactly once for testing, ensuring a robust evaluation of the model's performance across all data segments. In this study, a five-fold cross-validation approach was adopted, dividing the dataset into five distinct subsets. The model underwent five training cycles, with each iteration utilizing four folds for training and the remaining fold for validation. This systematic evaluation framework provides a comprehensive performance assessment, mitigating overfitting risks and enhancing the model's generalizability. Table 4 presents the R^2 and RMSE values obtained from the fivefold cross-validation process. The results indicate that the 5th fold yielded the highest R^2 value of 0.9777 and the lowest RMSE of 1.5088, demonstrating superior predictive performance in that specific iteration.

Results and discussion

Investigate the convergence curve

Figure 5 gives the convergence patterns in detail of the 3D hybrid models that tap the synergy between different

Table 4 The results of the K-Fold cross-validation

Model	Indicator	Number of K-fold				
		K1	K2	K3	K4	K5
RFR	RMSE	1.6520	1.5491	1.5523	1.6437	1.5088
	R^2	0.9734	0.9764	0.9763	0.9734	0.9777



Fig. 5 A three-dimensional bar plot illustrating the convergence of hybrid models

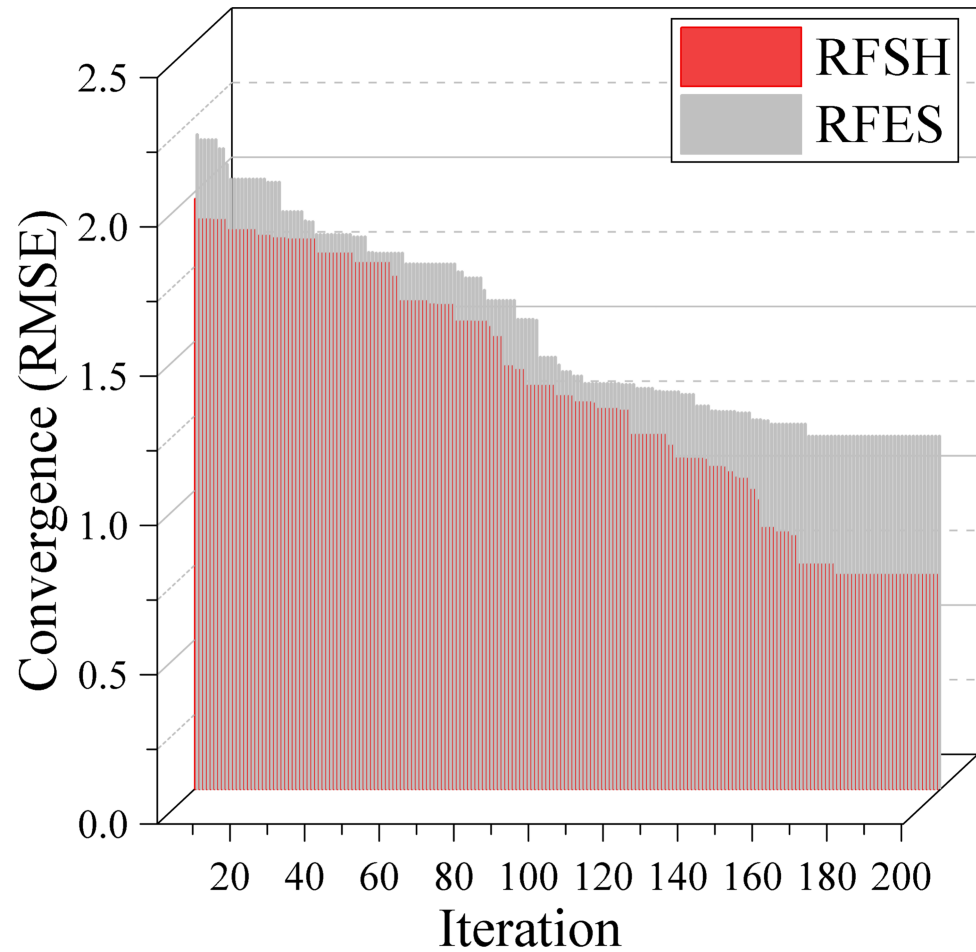


Table 5 The result of developed models for RF

Model	Phase	Index values						
		RMSE	R ²	MSE	SI	MNB	n10_index	MARE
RF	Train	1.469	0.979	2.159	0.066	− 0.003	0.913	0.043
	Validation	1.761	0.969	3.101	0.082	− 0.008	0.870	0.051
	Test	1.815	0.968	3.296	0.079	− 0.007	0.904	0.047
	All	1.572	0.976	2.470	0.070	− 0.005	0.905	0.045
RFESH	Train	0.550	0.997	0.302	0.025	− 0.001	0.998	0.016
	Validation	0.986	0.990	0.972	0.046	− 0.004	0.974	0.028
	Test	1.030	0.989	1.061	0.045	− 0.005	1.000	0.027
	All	0.718	0.995	0.516	0.032	− 0.002	0.995	0.020
RFES	Train	1.054	0.989	1.111	0.047	− 0.005	0.955	0.036
	Validation	1.427	0.979	2.037	0.066	− 0.012	0.930	0.050
	Test	1.458	0.979	2.127	0.063	− 0.003	0.870	0.049
	All	1.184	0.986	1.402	0.053	− 0.006	0.939	0.040
RFSE	Train	0.670	0.996	0.449	0.030	− 0.003	0.985	0.022
	Validation	1.027	0.989	1.055	0.048	− 0.008	0.965	0.034
	Test	1.035	0.989	1.072	0.045	− 0.004	0.974	0.031
	All	0.796	0.994	0.633	0.036	− 0.004	0.980	0.025

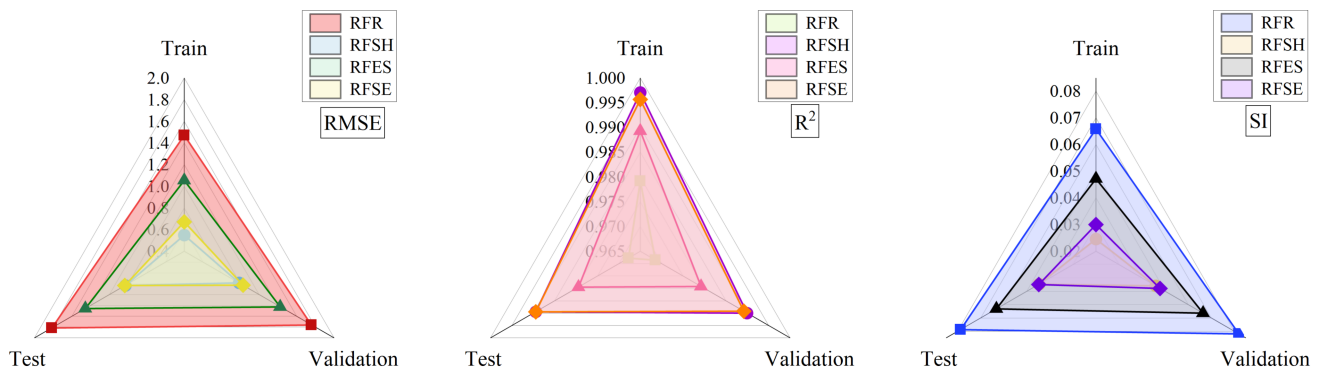


Fig. 6 Comparison between developed models based on RMSE, R^2 , and SI

machine learning algorithms. The root mean square error, probably one of the most important indications of model performance due to its relation with the prediction accuracy, is plotted against the y-axis of this graph while the x-axis plots the number of iterations, which follows the development of model refinement. Each of the hybrid models, labeled RFSH and RFES, respectively, was positioned differently along the z-axis. From plotting, as the iterations increased, the RMSE kept falling for all models. This shows improved predictive performance due to the incorporation of additional training data. The model RFSH, particularly, always remains better than RFES and other models, which is justified by its continuous lower RMSE. Among these, the RFSH model confirms its remarkable accuracy in heating load forecasting during the last iteration, with the very low value of RMSE amounting to 0.550. In conclusion, Fig. 5 highlights the exceptional predictive ability of the RFSH model, emphasizing its superiority over RFES models and confirming its accuracy in heating load forecasting.

Evaluation of developed models

The investigation's primary objective was to forecast heating load utilizing four distinct models: RFR, RFSH, RFSE, and RFES. In order to ensure an unbiased assessment of these models' performance, the research divided the experimental data into three phases: training (70%), validation (15%), and testing (15%). In order to conduct a thorough assessment and direct comparative analysis of these algorithms, the research employed five statistical metrics, namely R^2 , RMSE, MNB, SI, MSE, $n10_index$, and MARE, as displayed in Table 5. Figure 6, on the other hand, visually presented the assessment of developed models, with a focus on RMSE, R^2 , and SI for comparison.

- The central focus of the model evaluation was directed toward the R^2 values, which gauge how well the independent variable accounts for changes in the dependent

variable. During the training phase, the RFSH model demonstrated exceptional predictive accuracy, achieving the highest R^2 value of 0.997 among all models, surpassing the others. Conversely, the RFR model exhibited slightly lower R^2 values of 0.979 during the training phase. □ □ □

- Furthermore, the study investigated other error indicators, such as RMSE, which ranged from 0.550 to 1.469. Interestingly, the RFR model exhibited the greatest error rates, whilst the RFSH model exhibited the lowest mistakes.
- Concerning MSE, the RFR model yielded the highest and least favorable values during the training phase (MSE: 2.159), while the RFSH model recorded the lowest and most favorable values of 0.302 during the training phase.
- All things considered, the data point to the RFSH model performing better in certain phases than the RFR, RFSE, and RFES models. However, it is important to take other factors like model complexity, computing efficiency, and simplicity of implementation into account when choosing a model for real-world applications.†

Figure 7 presents a scatter plot of hybrid models' performance in the training, validation, and testing stages, using R^2 and RMSE criteria. The RFSH model shows exceptional accuracy and has the core line surrounded by densely packed data points, indicating minimal dispersion and high agreement. The RFES, RFSE, and RFR models show comparable performance levels, but broader dispersion implies higher error and slightly reduced accuracy compared to the RFSH model. The study highlights the importance of R^2 and RMSE in assessing model performance.

In Fig. 8, a line symbol plot effectively visualizes the models' related proportions of errors developed during the three phases of train, validation, and testing. In the training phase, RFSH stands out with the lowest error rate, at a mere 28%. This is underscored by most of its error values clustering around the 10% mark, indicating remarkable accuracy



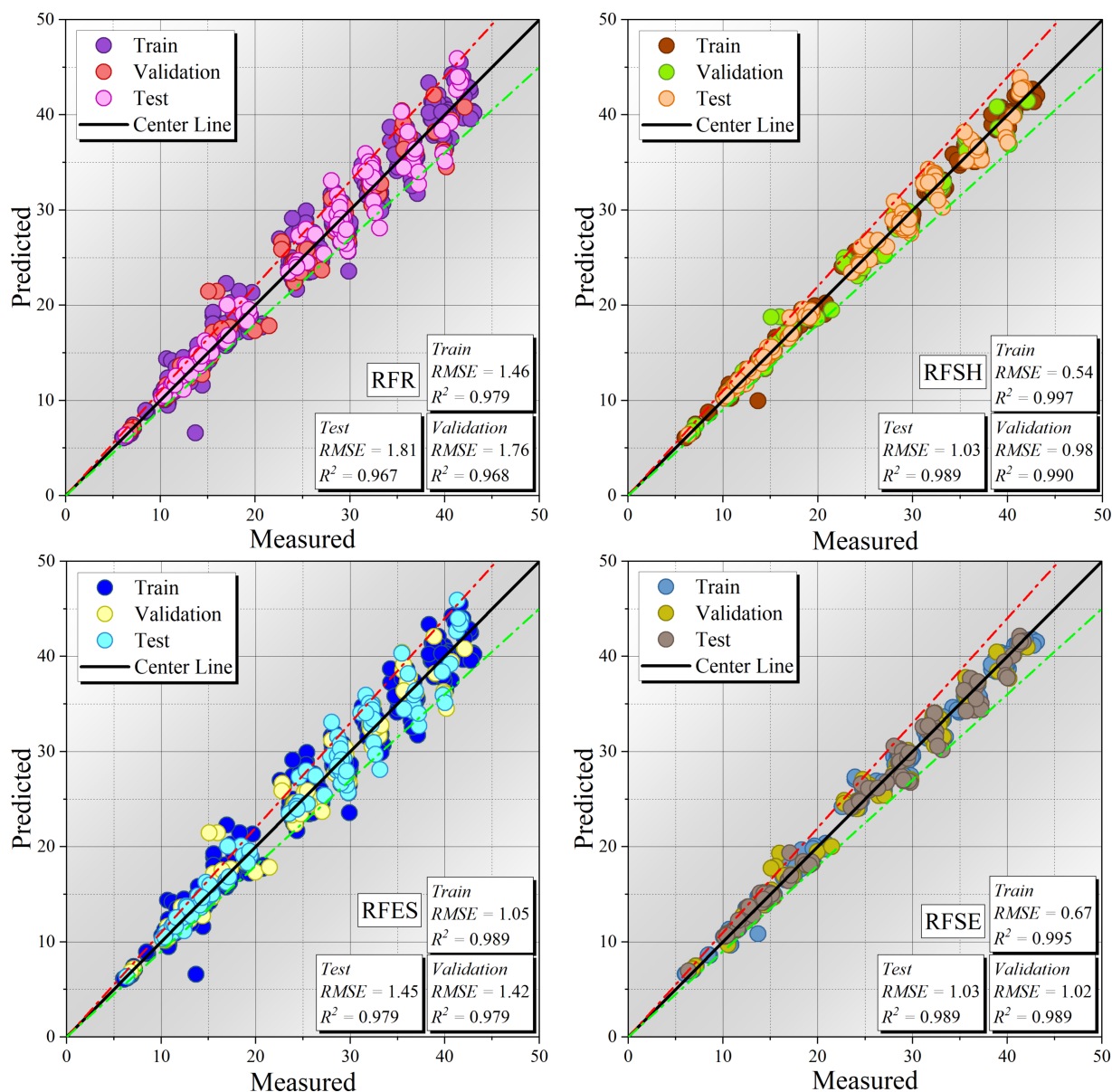


Fig. 7 The scatter plot for the developed models

and consistency in this phase. Conversely, a different picture emerges when turning attention to the RFR, RFSE, and RFES models. These models exhibit a more extensive distribution of error percentages. In the case of RFR, there is a notable concentration of error values exceeding 59% during the training phase, suggesting significant challenges in achieving accuracy during this stage. For the same period, that is, during the validation phase, RFES gives error values clustered around 32%, revealing its inability to hold on to performance. This graph shows striking differences in the level of performance between the models at different stages, all the while attesting to the accuracy and reliability during training of RFSH but signaling problems for the other RFR,

RFSE, and RFES, most especially during training and validation phases.

A half-violin diagram in Fig. 9 shows the error percentages for the models this research looked at. During the training phase, the mean error rate was an incredible 0%, the RFSH model distinguished itself. Its error distribution showed a well-formed normal curve with little dispersion, continuously staying below the 10% cutoff, indicating positive outcomes. The RFR model, on the other hand, displayed deviation in every phase and a normal curve that was symmetrical and equally distributed. However, this model maintained an error rate below 22%. Of all four, the differences between RFES and RFSE were highly marked and variable.

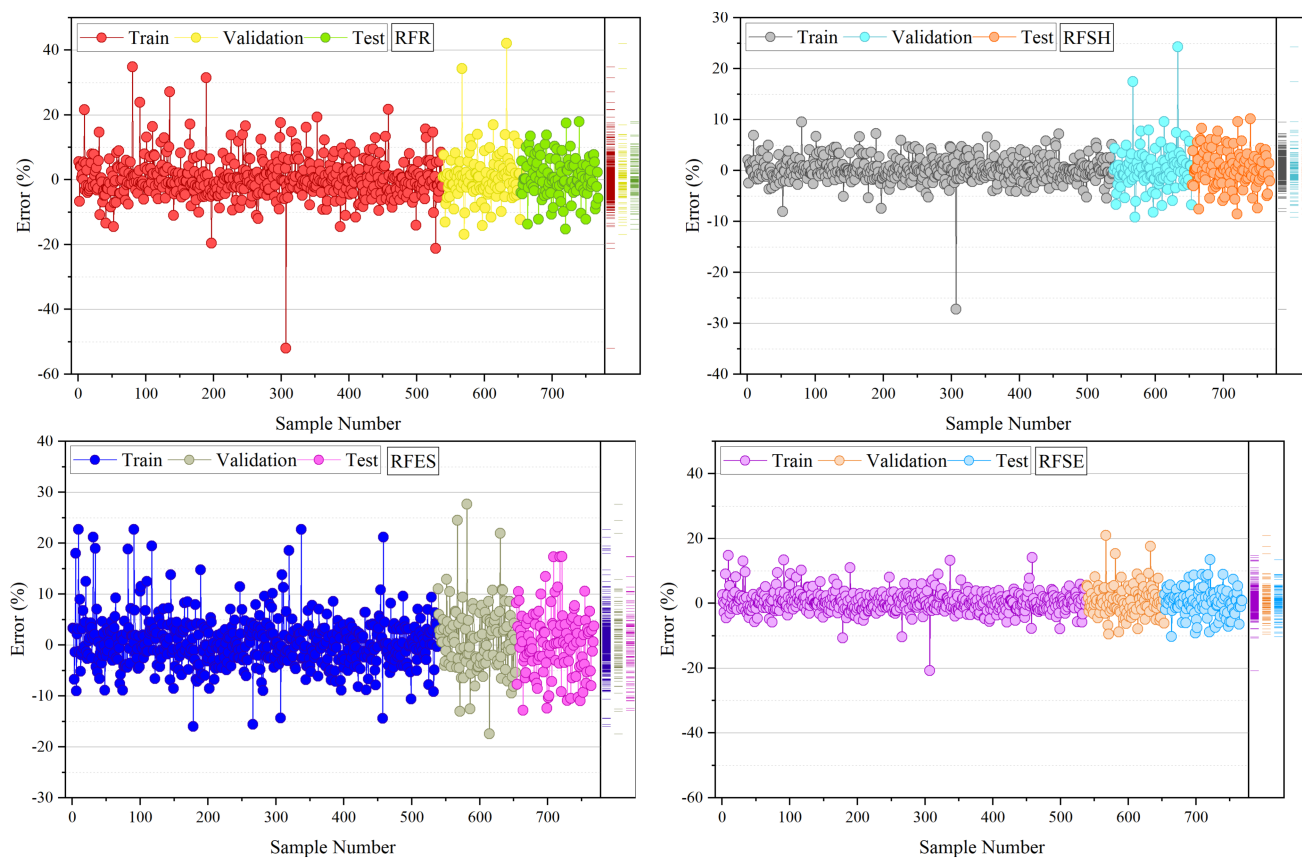


Fig. 8 Mistake fraction aimed at the cross replicas in accordance with the Line Symbol plot

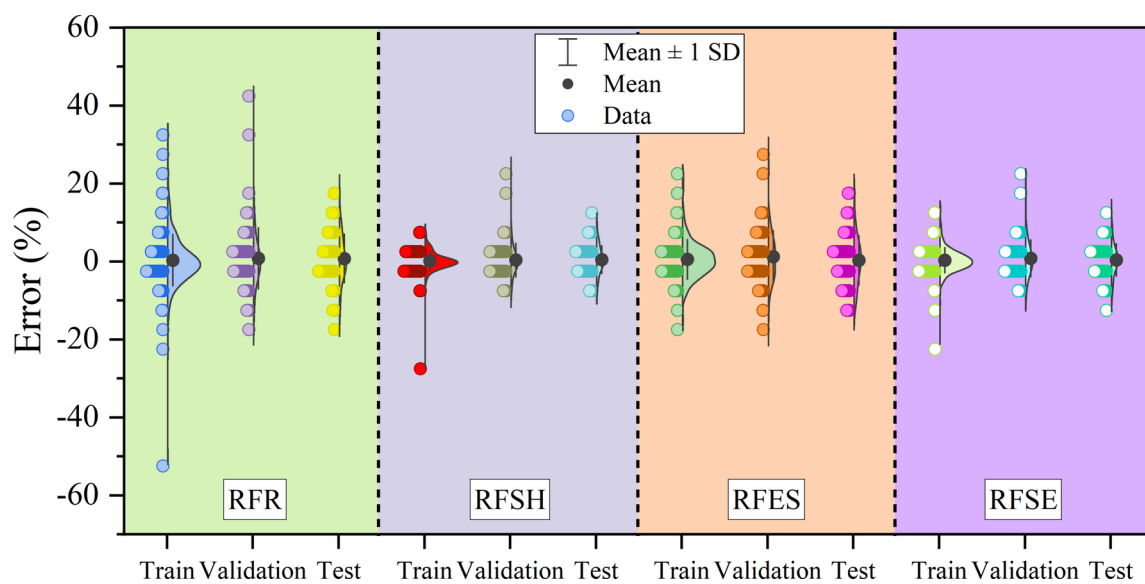


Fig. 9 The developed models' half violin of errors

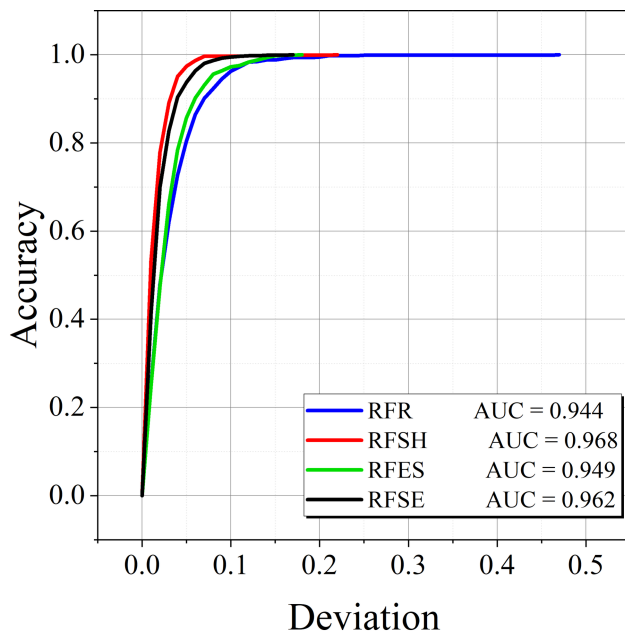


Fig. 10 Producing a REC curve to evaluate the effectiveness of the top-performing hybrid models

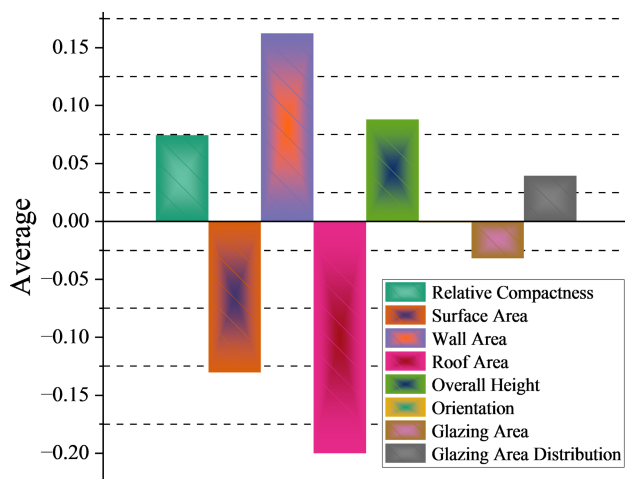


Fig. 11 Results of SHAP sensitivity analysis for the impact of inputs on the output values

An unexpected feature in the validation process was an outlier data point consisting of more than 18% of the data, a rare feature when carrying out statistical analyses.]

Figure 10 shows the REC curve, comparing the prediction performance of four different ML models against heating load prediction. The x-axis gives the deviation from the predicted value while the y-axis depicts the accuracy, fully informing about the model performance for various error threshold levels. The vertical dotted line at zero provides the reference point for the error threshold. There are four marked curves on the graph, each belonging to one of the model's

predictive performances. The RFSH is quite exceptionally high in terms of accuracy; this reflects the high predictive power for heating load. RFES and RFSE follow with quite good performances, whereas the curve associated with the RFR model is fairly below in terms of accuracy. The area under the curve, as defined by AUC-REC, is a key metric used to predict overall model performance. In this figure, the respective metric value is given next to each corresponding curve where values of 0.944 for RFR, 0.949 for RFES, 0.962 for RFSE, and an impressive value of 0.968 for RFSH. While the higher value of AUC-REC corresponds to better model performance, once again, the RFSH model qualifies as a dominant and accurate one in the prediction of heating load.

Figure 10 speaks volumes about the excellence of RFSH, integrating the RF model with two optimization techniques. This is to maximize the synergy between the strengths of Machine Learning and optimization methodologies. The proposed predictive model also shows better performance in handling various intricacies of heating load prediction with the highest accuracy and robust performance. The RFSH model follows a natural process and stands as the best option for providing better enlightenment in securing better energy efficiency and optimizing heating systems in real-world applications.

Sensitivity analysis based on SHAP

Understanding the influence of individual features on model predictions is crucial for interpretability, and Shapley Additive Explanations (SHAP) provide a mathematically rigorous approach to achieve this. SHAP values decompose a model's predictions, quantifying the contribution of each feature in a transparent and interpretable manner. As illustrated in Fig. 11, the analysis indicates that Wall Area and Overall Height have the most significant positive impact on HL values. Conversely, Surface Area and Roof Area exhibit the strongest negative influence, suggesting their relatively lower importance in shaping the model's predictions. By leveraging SHAP, this study ensures a comprehensive understanding of feature importance, ultimately enhancing the trustworthiness and explainability of the machine learning model.

Conclusion

The paper proposes a novel method for improving the accuracy of heating load forecasting by applying machine learning techniques, particularly focusing on random forest regression algorithms. The proposed methodology will be cost-effective and will simplify the process of heating load estimation considerably. At the heart of this effort lies an

advanced ML model based on the algorithmic foundation of RFR. The four models obtained by the research are RFR, RFSH, RFSE, and RFES. These are skillfully combined models of SHO and ESMA meta-heuristic algorithms with the ensembling to take advantage of the strengths of the optimizer technique in skillful fusion for high accuracy and minimizing errors. These models obtained here are fully validated through extensive training, efficient model validation, and testing using laboratory samples from the published literature. The performances of the models were evaluated in terms of several metrics, such as RMSE, MSE, R^2 , SI, MNB, n10_index, and MARE. All these metrics together give an in-depth understanding of the forecast performance of general models as well as their efficiency in estimating heating load. This is a great contribution to knowledge about the elements that influence heating load and the building energy area through the application of machine learning methods. It, therefore, provides the basis for a more realistic and harmonized estimation of heat load in a wide range of technical applications. The analysis showed that the RFR model had the lowest R^2 value, which was 1.6% marginal variance, while the highest R^2 values were consistently contributed by the RFSH models. Error indicators also pointed out that the RFSH models outperformed RFR, RFSE, and RFES models in performance by constantly showing fairly lower error levels. The RFSH models recorded the lowest RMSE values in each phase on a continuous basis; these differences were 71%, 38%, and 44% greater than those of the RFR, RFSE, and RFES models, respectively. This highlights the outstanding ability of these models to provide precise and trustworthy heating load estimates.†

A crucial aspect of model development involves balancing complexity and accuracy, as more sophisticated models often enhance predictive performance but may also introduce computational challenges and reduced interpretability. In this study, while the RFSH models demonstrated superior accuracy, consistently achieving the highest R^2 values and the lowest RMSE, their complexity, stemming from the integration of SHO and ESMA meta-heuristic algorithms, must be considered. In contrast, the RFR model, though less accurate, offers a simpler and computationally efficient alternative. The trade-off highlights that while advanced ensembling and optimization techniques significantly improve prediction accuracy and minimize errors, they also require greater computational resources and may be more challenging to implement in real-world applications. Future research will explore strategies to optimize this balance, ensuring that enhanced model performance does not come at the expense of practicality and efficiency.

Author contributions All authors contributed to the study's conception and design. Data collection, simulation and analysis were performed

by CC, NF, QL. The first draft of the manuscript was written by NF and all authors commented on previous versions of the manuscript. All authors have read and approved the manuscript.

Data availability The authors do not have permission to share data.

Declarations

Conflicts of interest The authors declare no competing interests.

References

- Afzal, S., Ziapour, B.M., Shokri, A., Shakibi, H., Sobhani, B.: Building energy consumption prediction using multi-layer perceptron neural network-assisted models; comparison of different optimization algorithms. *Energy* **282**, 128446 (2023)
- Ahmad, M.W., Mourshed, M., Rezgui, Y.: Trees vs neurons: comparison between random forest and ANN for high-resolution prediction of building energy consumption. *Energy Build.* **147**, 77–89 (2017)
- Aribowo, W.: A novel improved sea-horse optimizer for tuning parameter power system stabilizer. *J. Robot. Control (JRC)* **4**(1), 12–22 (2023)
- Biau, G., Scornet, E.: A random forest guided tour. *TEST* **25**, 197–227 (2016)
- Biswas, M.A.R., Robinson, M.D., Fumo, N.: Prediction of residential building energy consumption: a neural network approach. *Energy* **117**, 84–92 (2016). <https://doi.org/10.1016/j.energy.2016.10.066>
- Catalina, T., Virgone, J., Blanco, E.: Development and validation of regression models to predict monthly heating demand for residential buildings. *Energy Build.* **40**(10), 1825–1832 (2008). <https://doi.org/10.1016/j.enbuild.2008.04.001>
- Ding, Y., Zhang, Q., Yuan, T., Yang, K.: Model input selection for building heating load prediction: a case study for an office building in Tianjin. *Energy Build* **159**, 254–270 (2018)
- Gordon, M. L. et al.: Jury learning: Integrating dissenting voices into machine learning models. In: *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*, pp. 1–19, (2022)
- Guelpa, E., Marincioni, L., Capone, M., Deputato, S., Verda, V.: Thermal load prediction in district heating systems. *Energy* **176**, 693–703 (2019)
- Han, Y., Xi, W., Yang, S., Wang, W., Wu, J.: Robust design based on cost-quality model in micro-manufacturing. *Eksplotacja i Niezawodność Mainten. Reliab.* (2024). <https://doi.org/10.17531/ein/190380>
- Haykin, S.: *Neural networks and learning machines*, 3/E. Pearson Education India (2009)
- Huang, M., Yu, W., Yang, F.: Analysis of remaining useful life of slope based on nonlinear wiener process. *Eksplotacja i Niezawodność Mainten. Reliab.* (2024). <https://doi.org/10.17531/ein/187160>
- Jaffal, I., Inard, C., Ghiaus, C.: Fast method to predict building heating demand based on the design of experiments. *Energy Build* **41**(6), 669–677 (2009)
- Jordan, M.I., Mitchell, T.M.: Machine learning: trends, perspectives, and prospects. *Science* **349**(6245), 255–260 (2015)
- Khajavi, H., Rastgoo, A.: Predicting the carbon dioxide emission caused by road transport using a random forest (RF) model combined by meta-heuristic algorithms. *Sustain. Cities Soc.* **93**, 104503 (2023)
- Le, L.T., Nguyen, H., Dou, J., Zhou, J.: A comparative study of PSO-ANN, GA-ANN, ICA-ANN, and ABC-ANN in estimating the



- heating load of buildings' energy efficiency for smart city planning. *Appl. Sci.* **9**(13), 2630 (2019)
- Lim, T.S., Schaefer, L., Kim, J.T., Kim, G.: Energy benefit of the underfloor air distribution system for reducing air-conditioning and heating loads in buildings. *Indoor Built Environ.* **21**(1), 62–70 (2012)
- Lin, W., Wu, Z., Lin, L., Wen, A., Li, J.: An ensemble random forest algorithm for insurance big data analysis. *IEEE Access* **5**, 16568–16575 (2017)
- Moradzadeh, A., Mansour-Saatloo, A., Mohammadi-Ivatloo, B., Anvari-Moghaddam, A.: Performance evaluation of two machine learning techniques in heating and cooling loads forecasting of residential buildings. *Appl. Sci.* **10**(11), 3829 (2020)
- Naik, M.K., Panda, R., Abraham, A.: An entropy minimization based multilevel colour thresholding technique for analysis of breast thermograms using equilibrium slime mould algorithm. *Appl. Soft Comput.* **113**, 107955 (2021)
- Nieto, P.J.G., García-Gonzalo, E., Lasheras, F.S., Paredes-Sánchez, J.P., Fernández, P.R.: Forecast of the higher heating value in biomass torrefaction by means of machine learning techniques. *J. Comput. Appl. Math. Comput. Appl. Math.* **357**, 284–301 (2019)
- Roy, S.S., Roy, R., Balas, V.E.: Estimating heating load in buildings using multivariate adaptive regression splines, extreme learning machine, a hybrid model of MARS and ELM. *Renew. Sustain. Energy Rev.* **82**, 4256–4268 (2018)
- Sadaghat, B., Javadzade Khiavi, A., Naeim, B., Khajavi, E., Sadaghat, H., Taghavi Khanghah, A.R.: The utilization of a naïve bayes model for predicting the energy consumption of buildings. *J. Artif. Intell. Syst. Model.* **1**(01), 72–89 (2023). <https://doi.org/10.22034/JAISM.2023.422292.1003>
- Sarica, A., Cerasa, A., Quattrone, A.: Random forest algorithm for the classification of neuroimaging data in Alzheimer's disease: a systematic review. *Front Aging Neurosci.* **9**, 329 (2017)
- Shamshirband, S., et al.: Heat load prediction in district heating systems with adaptive neuro-fuzzy method. *Renew. Sustain. Energy Rev.* **48**, 760–767 (2015)
- Shariah, A., Tashtoush, B., Rousan, A.: Cooling and heating loads in residential buildings in Jordan. *Energy Build* **26**(2), 137–143 (1997)
- Taki, M., Rohani, A.: Machine learning models for prediction the higher heating value (HHV) of municipal solid waste (MSW) for waste-to-energy evaluation. *Case Stud. Therm. Eng.* **31**, 101823 (2022)
- Tejani, G.G., Sadaghat, B., Kumar, S.: Predict the maximum dry density of soil based on individual and hybrid methods of machine learning. *Adv. Eng. Intell. Syst.* **2**(03), 95–106 (2023)
- Xue, G., Qi, C., Li, H., Kong, X., Song, J.: Heating load prediction based on attention long short term memory: a case study of Xingtai. *Energy* **203**, 117846 (2020)
- Yin, S., Luo, Q., Zhou, Y.: EOSMA: an equilibrium optimizer slime mould algorithm for engineering design problems. *Arab. J. Sci. Eng.* **47**(8), 10115–10146 (2022)
- Zhang, F., Deb, C., Lee, S.E., Yang, J., Shah, K.W.: Time series forecasting for building energy consumption using weighted Support Vector Regression with differential evolution optimization technique. *Energy Build* **126**, 94–103 (2016)
- Zhang, Q., Tian, Z., Ma, Z., Li, G., Lu, Y., Niu, J.: Development of the heating load prediction model for the residential building of district heating based on model calibration. *Energy* **205**, 117949 (2020)
- Zhao, S., Zhang, T., Ma, S., Wang, M.: Sea-horse optimizer: a novel nature-inspired meta-heuristic for global optimization problems. *Appl. Intell.* **53**(10), 11833–11860 (2023)
- Zhong, Y., Li, H., Zhuang, X.: A resilience-driven two-stage operational chain optimization model for unmanned weapon system-of-systems under limited resource environments. *Eksploatacja i Niezawodność Maintn. Reliab.* (2024). <https://doi.org/10.17531/ein/188198>
- Zhou, Z.-H.: *Machine learning*. Springer Nature, Singapore (2021)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.