

Data-Driven Prediction of Metal Fatigue Life

Chuangjie Ou, Junyu Guo

School of Mechatronic Engineering, Southwest Petroleum University, Chengdu 610500, China

Abstract

With the rapid development of additive manufacturing (AM) technology, its applications in aerospace, medical, and automotive industries are becoming increasingly widespread. This paper investigates the fatigue life prediction of additively manufactured metallic materials and proposes a parallel deep learning architecture combining Gated Recurrent Unit (GRU) and Transformer models, optimized using Hybrid Leader-Based Optimization (HLBO) to improve prediction accuracy. The model simultaneously processes the outputs of both GRU and Transformer models, leveraging their respective strengths to learn features from both local and global perspectives, thereby enhancing the accuracy of fatigue life prediction. The paper first reviews the theoretical background of GRU and Transformer, then presents the proposed parallel model architecture in detail. The core principles of the HLBO optimization algorithm are introduced, along with its application in hyperparameter optimization. Experimental results show that the proposed parallel deep learning model significantly outperforms traditional single deep learning models in terms of prediction accuracy and generalization ability. Furthermore, the performance of the HLBO optimization algorithm is evaluated using benchmark test functions such as the Ackley and Rosenbrock functions to validate its global and local search capabilities. Experimental results demonstrate that HLBO performs excellently in handling high-dimensional, multi-modal optimization problems and effectively enhances the prediction capability of the model. Finally, the model is trained and validated using the FatigueData-AM2022 dataset. The results show that the proposed model exhibits strong adaptability and high precision in predicting the fatigue life of different additively manufactured metallic materials, providing an effective tool and method for the fatigue life assessment of additively manufactured components.

Keywords

Additive Manufacturing; Fatigue Life Prediction; Hybrid Leader-Based Optimization; Deep Learning; GRU and Transformer Models.

1. Introduction

Additive manufacturing (AM) of metals has emerged as a significant manufacturing technology in recent years, attracting extensive attention and finding widespread applications in various industries. Compared with traditional subtractive manufacturing methods, AM builds components layer by layer, directly transforming digital models into physical parts without the need for cutting and machining processes required in conventional manufacturing. This approach provides greater manufacturing flexibility and design freedom. Due to its advantages, AM has demonstrated substantial potential and broad prospects in aerospace, medical, and automotive industries. Alloy materials used in additive manufacturing, such as Ti-6Al-4V, 316L, and IN718, have been widely applied in multiple fields due to their excellent properties and diverse application characteristics[1-5].

Fatigue failure is one of the primary failure mechanisms of metal components, commonly observed in aerospace, mechanical, and chemical engineering fields. Accurate fatigue life

prediction is a critical issue in metal fatigue research, holding great theoretical and practical significance in engineering and materials science. By conducting an in-depth investigation into the fatigue failure mechanisms of metals and integrating advanced life prediction methods, potential fatigue issues can be considered at the early design stage of metal components. This allows for optimized design strategies to extend the service life of critical parts. Furthermore, fatigue life prediction facilitates the early identification of fatigue defects in metallic structures, enabling the monitoring and assessment of fatigue-sensitive regions. This approach contributes to the early detection of potential fatigue-related issues, reducing maintenance costs and enhancing the reliability of metal components.

Additively manufactured alloy components are also subject to fatigue failure risks. Therefore, this study focuses on predicting the fatigue life of additively manufactured alloy materials, aiming to provide an effective method for addressing this issue.

This study introduces a parallel deep learning architecture combining Gated Recurrent Unit (GRU) and Transformer models, optimized via Hybrid Leader-Based Optimization (HLBO) to enhance prediction accuracy. Unlike traditional methods, this approach provides robust fatigue life predictions across diverse materials and loading conditions, offering theoretical and practical guidance for engineering design and lifecycle assessment.

2. Theoretical Background

This study introduces a parallel deep learning architecture combining the Gated Recurrent Unit (GRU) and Transformer models, optimized via Hybrid Leader-Based Optimization (HLBO) to enhance prediction accuracy. Unlike traditional methods, this approach provides robust fatigue life predictions across diverse materials and loading conditions, offering theoretical and practical guidance for engineering design and lifecycle assessment.

To build the fatigue life prediction model for additively manufactured metals, the following theoretical foundations are considered:

2.1. GRU

The Gated Recurrent Unit (GRU), proposed by Cho et al. in 2014, is a variant of the Recurrent Neural Network (RNN) designed to address the gradient vanishing problem in traditional RNNs while reducing the computational complexity of Long Short-Term Memory (LSTM) networks.

Unlike LSTM, which consists of three gates (forget gate, input gate, and output gate), GRU simplifies the gating mechanism by retaining only two gates:

- Update Gate: Controls how much of the previous hidden state should be retained and how much should be updated with new information.
- Reset Gate: Determines how much of the past hidden state should be forgotten.

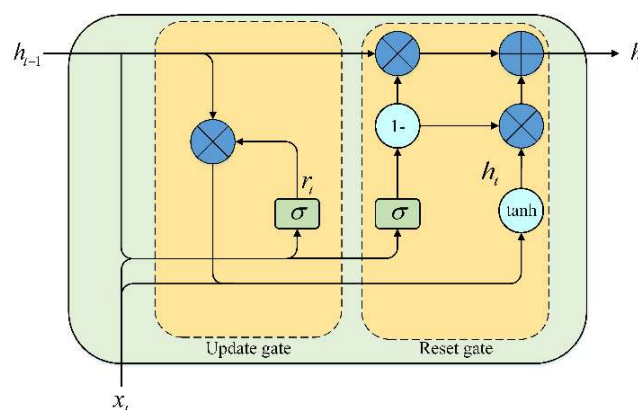


Figure 1. GRU Model Architecture Diagram

By eliminating the forget gate, GRU reduces the number of parameters, making it more computationally efficient and enabling faster convergence during training. The gating mechanism dynamically regulates information flow, allowing GRU to capture long-term dependencies in sequential data while maintaining a simpler architecture compared to LSTM. Due to its efficiency and strong sequential modeling capabilities, GRU has been widely applied in time-series prediction, natural language processing, and engineering applications, including fatigue life prediction of materials.

2.2. Transformer

The Transformer model, proposed by Google in 2017, was initially designed for natural language processing (NLP) tasks. Unlike traditional recurrent neural networks (RNNs) and long short-term memory (LSTM) networks, the Transformer model processes information based on the self-attention mechanism, which enables it to effectively capture long-range dependencies in sequential data while supporting parallel computation, significantly improving training efficiency.

In recent years, the Transformer model has demonstrated superior performance not only in NLP but also in time series forecasting, computer vision, and material fatigue life prediction. Particularly in engineering applications, the fatigue life evolution of materials often involves multi-scale and multi-factor interactions. The Transformer model's ability to capture feature dependencies on a global scale makes it an efficient deep learning solution for fatigue life prediction.

The Transformer architecture follows an encoder-decoder structure, where its core concept revolves around using the self-attention mechanism to model the global dependencies in input sequences. Compared to traditional RNNs and convolutional neural networks (CNNs), the Transformer model, equipped with multi-head self-attention mechanisms, processes input data in parallel, effectively overcoming RNN's vanishing gradient problem and CNN's local receptive field limitations, making it more suitable for long-sequence modeling.

The overall structure of the Transformer model is illustrated in Figure 2, which consists of the following key components:

- **Encoder:** Responsible for extracting features from the input sequence. It employs layer normalization and residual connections to enhance training stability and prevent gradient vanishing.
- **Decoder:** Generates the output sequence by leveraging the features extracted by the encoder. It uses an autoregressive mechanism to iteratively produce the predicted sequence.

This encoder-decoder architecture enables the Transformer model to effectively process sequential data, capturing complex dependencies and improving prediction accuracy, making it particularly suitable for fatigue life prediction and other engineering applications.

Both the encoder and decoder in the Transformer model consist of multiple sub-layers, with the most critical components being the multi-head self-attention layer and the feedforward neural network layer. Additionally, to enhance gradient flow and accelerate convergence, the Transformer model incorporates residual connections and layer normalization in each layer.

Significance of Transformer in Fatigue Life Prediction

In summary, the Transformer model primarily relies on the self-attention mechanism, enabling it to effectively model global dependencies. This capability is crucial for fatigue life prediction, as the fatigue behavior of materials and components is influenced by multiple environmental factors, which often exhibit complex interdependencies.

The multi-head attention mechanism allows the model to focus on different aspects of the sequence data from multiple perspectives, providing a more comprehensive understanding of

the fatigue process. However, due to the scarcity of data in the field of fatigue life prediction, using a Transformer model alone may lead to overfitting, especially when the model complexity is high. To mitigate this risk, additional techniques such as data augmentation, regularization, and hybrid modeling approaches are necessary to improve the model's generalization ability in fatigue life prediction tasks.

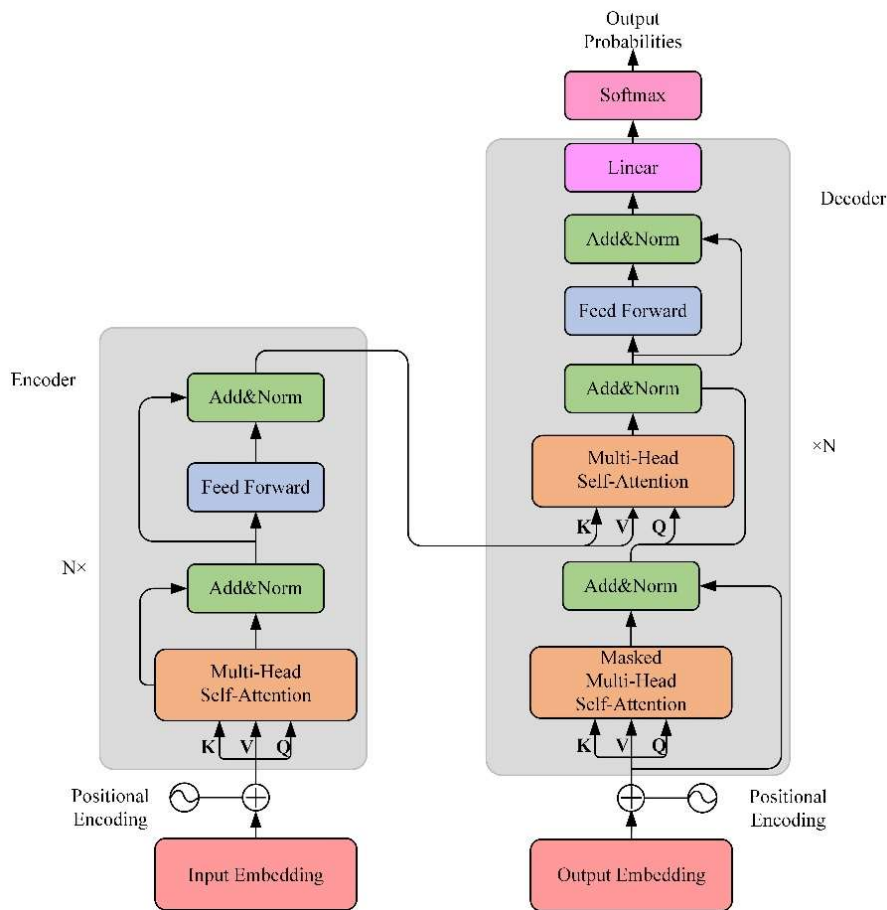


Figure 2. Schematic Diagram of Transformer Architecture Principle

2.3. The Proposed Model

To fully leverage the advantages of GRU and Transformer models and maximize the potential of the dataset, this study proposes an innovative parallel deep learning architecture that effectively combines these two models. The basic structure of the proposed model is illustrated in Figure 3. In this architecture, the input data from the dataset is simultaneously fed into both GRU and Transformer models, allowing them to perform feature extraction and prediction in parallel, thereby exploiting the unique strengths of each model.

Although GRU is traditionally used for time-series data, its ability to capture local features and perform nonlinear transformations makes it suitable for this task as well. Within this framework, GRU effectively learns local dependencies and feature interactions, even when dealing with data without explicit temporal ordering. In contrast, Transformer does not rely on time sequence information. Instead, it employs a self-attention mechanism to capture global dependencies, making it particularly efficient in modeling complex relationships among features in datasets without time-series structures.

By integrating GRU's strength in learning local dependencies and Transformer's ability to model global relationships, this parallel model architecture enhances the overall prediction accuracy and feature extraction capability, making it well-suited for fatigue life prediction tasks in complex datasets.

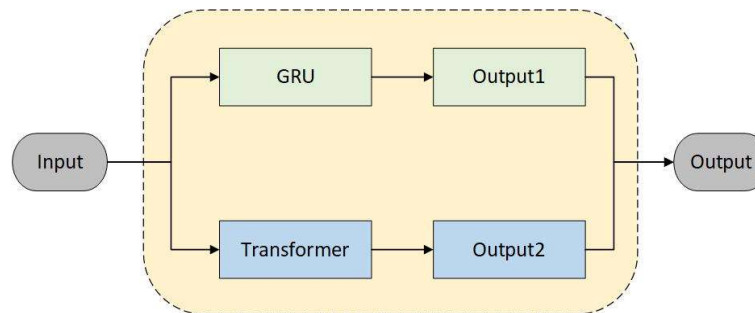


Figure 3. Schematic Diagram of Parallel Model Architecture

After separately performing computations and generating outputs from the two models, their predicted results are subsequently fed into a Fully Connected Network (FCN) for further joint processing. In this step, the FCN fuses the outputs of the GRU and Transformer models to produce a comprehensive final prediction. Consequently, the outputs of the GRU and Transformer models do not exist as standalone results; rather, they are combined through the FCN to yield more precise and robust predictions. This enables the final decision-making process to incorporate the feature representations learned by both models and leverage their respective advantages.

A key benefit of this parallel architecture lies in its ability to learn features from two different perspectives. The Gated Recurrent Unit (GRU) leverages its gating mechanisms to capture local dependencies and nonlinear relationships within the data, while the Transformer model efficiently learns global dependencies and complex interrelationships among features through its self-attention mechanism. By employing these two models in parallel, the network can simultaneously capture both local and global representations of the data, resulting in a more comprehensive understanding of multi-level features.

In addition, this parallel structure fully exploits the strengths of GRU and Transformer in modeling different types of data. GRU is highly effective when dealing with data that exhibit local dependencies, whereas the Transformer, by virtue of its self-attention mechanism, excels at capturing intricate relationships among features. Combining these two modeling approaches helps mitigate the limitations of a single model, thereby enhancing predictive accuracy and generalization capability.

In summary, this parallel deep learning architecture not only integrates the respective strengths of GRU and Transformer models but also facilitates the comprehensive exploration of data features. By effectively merging these two different types of networks, the model delivers a more holistic understanding and modeling of complex, non-temporal data. Such a network architecture provides a more effective and integrative approach to deep learning applications, improving model performance and offering new perspectives for addressing real-world complex problems.

2.4. Hybrid Leader-Based Optimization

In deep learning-based data prediction tasks, setting appropriate model hyperparameters is a critical step[7,8]. However, manually selecting these hyperparameters is often both complex and time-consuming, since it requires the adjustment and optimization of multiple parameters. Each hyperparameter can significantly affect the model's final performance, making the discovery of an optimal hyperparameter combination essential for enhancing prediction accuracy. By contrast, employing optimization algorithms to automatically tune and optimize model hyperparameters has emerged as a more efficient and commonly adopted approach. Such algorithms systematically explore various hyperparameter spaces, thereby helping models attain optimal predictive performance. A substantial body of research has demonstrated that optimization methods in hyperparameter tuning can effectively improve

model performance. Consequently, choosing a suitable optimization algorithm to perform hyperparameter optimization has become a vital strategy for boosting the predictive power of deep learning models.

Hybrid Leader-Based Optimization (HLBO)[9] is a novel intelligent optimization algorithm rooted in the concept of hybrid leadership. Its primary objective is to enhance global search capability and convergence performance by introducing multiple leaders to guide the evolution of population members. Traditional population-based optimization algorithms (such as Genetic Algorithms (GA) and Particle Swarm Optimization (PSO)) often depend on a single leader (e.g., the global best solution or a particular individual within the population), which may cause the algorithm to converge prematurely to local optima and leave large regions of the search space underexplored—especially in high-dimensional, complex optimization problems. To overcome this limitation, HLBO adopts a hybrid leadership strategy that incorporates multiple sources of information. This approach diversifies the search process, reduces reliance on a single best solution, and improves both the exploration capability of the algorithm and its ability to escape local optima.

HLBO establishes population evolution rules through a hybrid leader mechanism. Its optimization process consists of population initialization, individual quality evaluation, participation coefficient calculation, hybrid leader position calculation, individual updating, and fitness evaluation. Let the objective function of the optimization problem be:

$$\min_X f(X), \quad X = [x_1, x_2, \dots, x_m] \in \mathbb{R}^m \quad (1)$$

where X is the decision variable vector,
 m is the dimensionality of the variables,
 $f(X)$ is the objective function.

During the initialization phase, the population P consists of N solutions X_i :

$$P = X_1, X_2, \dots, X_N, \quad X_i = [x_{i1}, x_{i2}, \dots, x_{im}] \quad (2)$$

Each individual X_i is randomly initialized within the search space $[X_{\min}, X_{\max}]$. In the evolutionary process, the position of each individual is updated based on three factors: the current individual, the global best individual, and a randomly selected individual. This strategy ensures a balance between global exploration and local exploitation, enabling HLBO to more thoroughly explore the solution space and reducing the likelihood of convergence to local optima. To appropriately allocate the influence of these three components on individual updates, HLBO introduces a Participation Coefficient, which is calculated as follows.

(1) Individual Mass Calculation

HLBO employs the fitness value of each individual to gauge its mass. The mass q_i of the i -th individual is computed as follows:

$$q_i = \frac{F_i - F_{\text{worst}}}{\sum_{j=1}^N (F_j - F_{\text{worst}})}, \quad i \in \{1, 2, \dots, N\} \quad (3)$$

where F_i denotes the objective function value of the i -th individual; F_{worst} is the objective function value of the worst-performing individual in the current population; N is the population size.

By utilizing this formula, the mass of each individual is normalized to the range $[0,1]$, allowing for a quantifiable measure of each individual's contribution to the overall optimization direction.

(2) Calculation of the Participation Coefficient

Based on the computed individual mass q_i , HLBO calculates three distinct members' participation coefficients as follows:

$$PC_i = \frac{q_i}{q_i + q_{\text{best}} + q_k} \quad (4)$$

$$PC_{\text{best}} = \frac{q_{\text{best}}}{q_i + q_{\text{best}} + q_k} \quad (5)$$

$$PC_k = \frac{q_k}{q_i + q_{\text{best}} + q_k} \quad (6)$$

where q_{best} denotes the mass of the best-performing individual in the population, and q_k denotes the mass of a randomly selected individual.

(3) Hybrid Leader Mechanism

By computing the aforementioned participation coefficients, HLBO employs the following formula to generate the position of the hybrid leader:

$$HL_i = PC_i \cdot X_i + PC_{\text{best}} \cdot X_{\text{best}} + PC_k \cdot X_k \quad (7)$$

where X_i is the current individual (retaining the individual's own information), X_{best} is the global best individual (ensuring convergence toward the optimal solution), and X_k is a randomly selected individual (introducing population diversity and enhancing global exploration). Here, HL_i denotes the hybrid leader position for the i -th individual, and $PC_i, PC_{\text{best}}, PC_k$ are weighting parameters for the current individual, global best individual, and randomly selected individual, respectively, satisfying $PC_i + PC_{\text{best}} + PC_k = 1$. These parameters dynamically adjust the influence each leader has on the individual's update.

(4) Individual Updating and Fitness Evaluation

The formula for updating the position of each individual is as follows:

$$X_i^{\text{new}} = X_i + r \cdot (HL_i - X_i) \quad (8)$$

where r is a random factor introduced to enhance search stochasticity, preventing the population from prematurely converging to a local optimum. After completing the position update, HLBO calculates the new fitness value $f(X_i^{\text{new}})$. The selection criterion for retaining the updated position is given by:

$$X_i = \begin{cases} X_i^{\text{new}}, & \text{if } f(X_i^{\text{new}}) < f(X_i) \\ X_i, & \text{otherwise} \end{cases} \quad (9)$$

In other words, if the objective function value at the new position demonstrates an improvement, the new solution is adopted; otherwise, the original solution remains unchanged.

Start HLBO

1. Input the information about the optimization problem.
 2. Set parameters: N (population size) and T (maximum number of iterations).
 3. Initialize positions of the HLBO population and evaluate the objective function.
 4. For $t = 1$ to T :
 5. For $i = 1$ to N :
 6. Phase 1: Exploration
 7. Compute the mass q_i of the candidate solution using Equation (6-15).
 8. Calculate the participation coefficients PC_i , PC_k , and PC_{best} using Equations (6-16) – (6-18).
 9. Generate the hybrid leader HL_i with Equation (6-19).
 10. Determine the new position of the i -th member using Equation (6-20).
 11. Update the i -th member using Equation (6-15).
 12. Phase 2: Exploitation
 13. Determine the new position of the i -th member using Equation (6-20).
 14. Update the i -th member using Equation (6-15).
 15. End For (i)
 16. Update the best candidate solution found so far.
 17. End For (t)
 18. Output: The best candidate solution obtained via HLBO.
-
- End HLBO
-

The HLBO optimization algorithm introduces the concept of hybrid leaders, which collaboratively update multiple members within the population, including the best member, worst member, and other relevant members. These members together form a hybrid leadership group, guiding the entire search process of the algorithm. The core mechanism of this approach is to prevent excessive dependence on any single member (such as the best or worst member) by leveraging multiple guiding members. This collaborative updating strategy enhances the algorithm's flexibility and efficiency, enabling it to explore a broader search space more comprehensively.

Compared to traditional population-based optimization algorithms, HLBO demonstrates a greater ability to escape local optima. Conventional algorithms often rely on a single leader,

which may cause premature convergence or restrict the search within a local optimal region. In contrast, HLBO adopts a diversified leadership strategy, significantly improving its global search capability. By incorporating multiple guiding members, HLBO enhances exploration efficiency and integrates information from multiple sources, making the search process more diverse. This approach effectively prevents the algorithm from getting trapped in local optima, ensuring that the global optimum can be successfully located.

Thus, HLBO optimization not only improves optimization accuracy but also enhances its adaptability and robustness when dealing with complex problems. This is particularly beneficial for solving high-dimensional, multi-modal optimization problems with intricate structures, where HLBO exhibits significant advantages over conventional optimization methods.

2.5. Performance Evaluation of the Optimization Algorithm

When evaluating the performance of an optimization algorithm, it is essential to comprehensively consider its global search and local search capabilities. Global search ability refers to an algorithm's capability to explore a wide search space and locate the global optimum, while local search ability measures how well an algorithm fine-tunes solutions within a local region to refine the optimal result. Different optimization algorithms exhibit varying levels of performance in these two aspects, making a comprehensive evaluation-particularly in terms of global and local search efficiency-critical for understanding their overall effectiveness.

This section presents a brief evaluation of the HLBO algorithm's global and local search capabilities to verify its performance in these aspects. To achieve this, optimization algorithms are typically assessed using CEC benchmark functions. The CEC 2019 function set is a standardized benchmark suite specifically designed to evaluate optimization algorithms, consisting of 23 standard test functions. These functions are widely utilized in optimization performance assessments because they cover various challenges, including multi-modal landscapes, nonlinearities, and high-dimensional search spaces, effectively reflecting an algorithm's behavior under different conditions.

In this study, two standard test functions from CEC 2019 were selected to assess the HLBO algorithm's performance:

(1) Rosenbrock Function – Evaluating global search capability

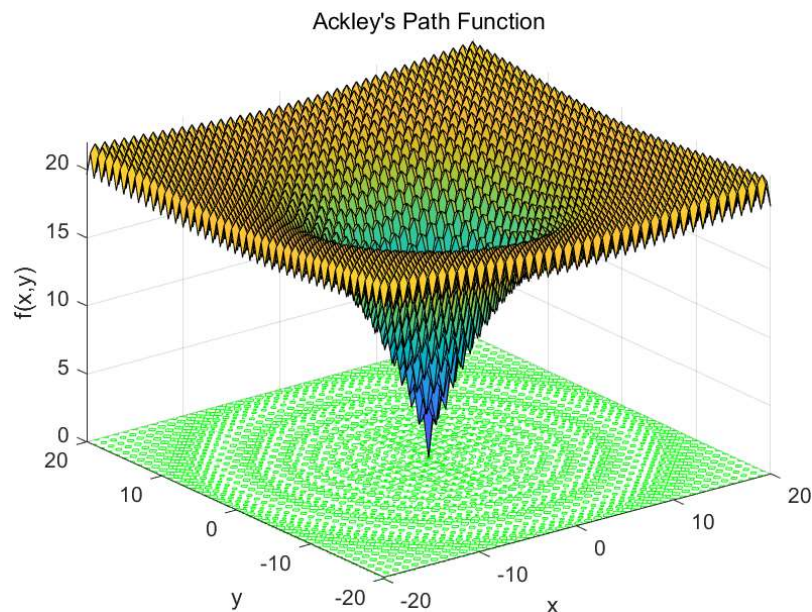
- The Rosenbrock function is frequently used to assess an algorithm's ability to perform global searches.
- It features a narrow valley with a distant global minimum, requiring the optimization algorithm to efficiently navigate from distant points and converge toward the global optimum.
- This function tests whether an algorithm can escape poor initial positions and effectively explore the broader search space before refining its solution.

(2) Ackley Function – Evaluating local search capability

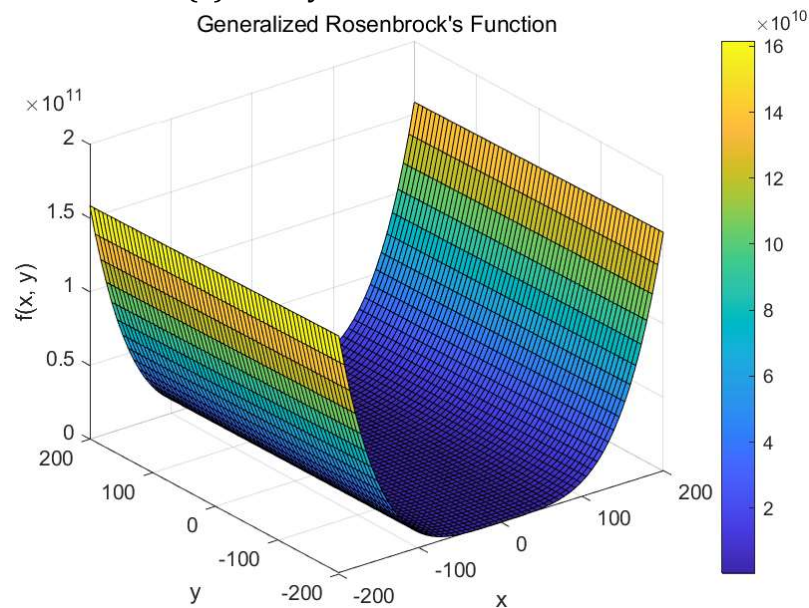
- The Ackley function is commonly used to evaluate an algorithm's ability to perform local searches.
- It has multiple local minima, requiring the optimization algorithm to distinguish the global minimum from several misleading local minima.
- This function assesses whether an algorithm can avoid premature convergence and successfully identify the best solution in a complex multi-modal landscape.

The visual representations of these two functions are shown in Figure 4, providing an intuitive understanding of their structural characteristics and the challenges they pose to optimization algorithms. By testing the HLBO algorithm on these two benchmark functions, we can further

evaluate its global and local search performance, offering a more comprehensive understanding of its potential advantages in real-world applications.



(a) Ackley Function Visualization



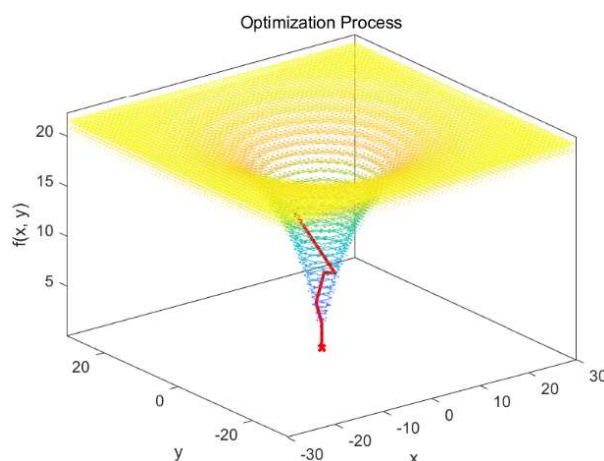
(b) Rosenbrock Function Visualization

Figure 4. Visualization of Two Benchmark Functions

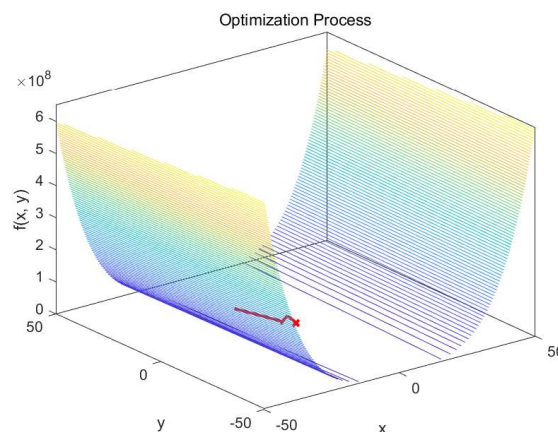
The Ackley function is a widely used multi-dimensional benchmark optimization function, commonly employed to test an optimization algorithm's performance in handling multi-modal problems, particularly those with multiple local optima. The key characteristic of the Ackley function is that, although it has a single global optimum, there exist multiple local minima around this optimal point. This multi-peak structure poses a challenge for optimization algorithms: if they lack sufficient exploration capability, they may easily become trapped in a local optimum instead of reaching the global minimum. Therefore, evaluating an optimization algorithm's performance on the Ackley function effectively assesses its global search capability, particularly its ability to avoid premature convergence to local optima.

On the other hand, the Rosenbrock function is a well-known benchmark function commonly used to evaluate an optimization algorithm's local search capability. Its shape forms a narrow valley that contains complex local minima and exhibits the characteristic of sparse gradients. The global optimum of the Rosenbrock function is located within a narrow region, requiring an optimization algorithm to efficiently search along the valley floor to ensure convergence to the global minimum. If an algorithm lacks sufficient local search ability, it may fail to escape flat regions within the valley and become trapped in suboptimal solutions. Additionally, when optimizing the Rosenbrock function, an algorithm must overcome the challenge of vanishing gradients-as the optimization process enters regions with small gradient values, the algorithm must maintain a strong search direction to prevent stagnation. Consequently, the Rosenbrock function is widely used to test an optimization algorithm's local search performance, particularly its ability to handle complex gradient landscapes and avoid suboptimal convergence.

By evaluating the HLBO algorithm on both the Ackley and Rosenbrock functions, its global and local search capabilities can be comprehensively assessed. The iteration process of the selected test functions is presented in Figure 5, demonstrating the algorithm's ability to navigate multi-modal landscapes and efficiently converge to optimal solutions.



(a) HLBO Optimization Function Iteration Process on Ackley Function



(b) HLBO Optimization Function Iteration Process on Rosenbrock Function

Figure 5. HLBO Optimization Iteration Process

Based on the iteration process plots of the selected test functions, it can be observed that the HLBO optimization algorithm performs well in both global and local search. The algorithm effectively converges to the function's minimum value during optimization, demonstrating its strong optimization capability. This ensures that HLBO can efficiently adjust hyperparameters

to enhance model performance. Therefore, HLBO not only meets the requirements for hyperparameter optimization but also maximizes prediction accuracy, making it a reliable approach for practical applications.

Hyperparameter Optimization Using HLBO

In this study, the HLBO algorithm is applied to optimize key hyperparameters in the model, focusing on the following five parameters:

(1) Learning Rate

- The learning rate determines the step size of each parameter update.
- If the learning rate is too large, the model may fail to converge or even diverge.
- If the learning rate is too small, training will be excessively slow, increasing computational resource consumption.
- Proper adjustment of the learning rate ensures efficient convergence, accelerates training, and stabilizes the model near the optimal solution.

(2) GRU Hidden State Dimension

- The hidden state dimension defines the memory capacity of GRU units.
- A small hidden state dimension may limit the model's ability to capture long-term dependencies, leading to poor performance.
- A large hidden state dimension increases computational complexity, training time, and storage demand.
- The optimal hidden dimension should be chosen based on the dataset complexity and task requirements to ensure sufficient model representation.

(3) Number of GRU Layers

- Stacking more GRU layers enhances the model's representation capacity, allowing it to learn more complex features from data.
- However, adding more layers significantly increases computational cost and training time.
- An excessive number of stacked layers may cause overfitting, especially in cases with limited data or simple tasks.
- A balance between model complexity and computational efficiency must be maintained.

(4) Number of Transformer Encoders and Decoders

- The encoder and decoder count in the Transformer model plays a crucial role in determining model complexity and representational power.
- Increasing these layers can enhance model performance, particularly when dealing with long sequences or complex features.
- However, adding more layers leads to higher computational costs, longer training times, and potential inference delays.
- The encoder-decoder structure should be fine-tuned based on task complexity and dataset size to achieve optimal performance-efficiency balance.

Impact of HLBO Optimization

By applying HLBO optimization to these key hyperparameters, the model structure and learning process can be finely adjusted, maximizing predictive performance and enhancing real-world applicability. Through comprehensive hyperparameter tuning, the proposed model achieves:

- Improved adaptability to different datasets and tasks.
- More accurate and reliable fatigue life predictions.
- Optimized computational efficiency without unnecessary complexity.

Thus, HLBO ensures the best trade-off between model accuracy, generalization ability, and computational feasibility, making it a robust optimization strategy for fatigue life prediction in additively manufactured alloys.

2.6. Prediction Performance Evaluation

To evaluate the capability of the proposed model in fatigue life prediction and further validate its advantages, this study compares its predictive performance with six other machine learning (ML) models. The effectiveness of the models is assessed using four commonly used evaluation metrics: Mean Squared Error (MSE), Mean Absolute Error (MAE), Coefficient of Determination (R^2), and Mean Absolute Percentage Error (MAPE). These evaluation metrics provide insights into model accuracy and reliability from different perspectives.

(1) Mean Squared Error (MSE) measures the average squared difference between the predicted and actual values, emphasizing large errors due to its quadratic nature. A lower MSE value indicates better predictive accuracy.

(2) Mean Absolute Error (MAE) represents the average absolute difference between predictions and actual values, providing an intuitive measure of prediction accuracy without emphasizing extreme values.

(3) Coefficient of Determination (R^2) quantifies how well the model explains the variance in the actual data. A value closer to 1 indicates that the model can effectively capture the underlying patterns in the data.

(4) Mean Absolute Percentage Error (MAPE) evaluates the relative error in predictions, making it useful for assessing model performance across different scales. A lower MAPE value suggests higher prediction reliability.

By using these four evaluation metrics, this study ensures a comprehensive assessment of the proposed model's accuracy, robustness, and generalization ability, allowing for an objective comparison with traditional machine learning methods in fatigue life prediction tasks.

When comparing predictive models using the four evaluation metrics (MSE, MAE, R^2 , and MAPE), a single metric provides limited insight into overall model performance. Additionally, when evaluating different models, some may perform well on one metric but poorly on another, making direct comparisons difficult.

To address this issue and provide a more comprehensive evaluation, this study introduces a new evaluation metric, which integrates all four performance indicators into a unified measure. The calculation formula for this new metric is given as follows:

$$new = \frac{R^2}{MAE + MSE + MAPE} \quad (10)$$

Among the four evaluation metrics, a higher R^2 value is preferred, indicating that the model better explains the variance in fatigue life. Conversely, lower values of MSE, MAE, and MAPE indicate better predictive accuracy, as they represent smaller prediction errors. Therefore, the new evaluation metric is designed such that a higher value indicates better overall model performance.

3. Dataset for Fatigue Life Prediction of Additively Manufactured Alloys

The construction of a fatigue dataset for additively manufactured alloys is crucial for predicting their fatigue life. However, publicly available datasets from authoritative research institutions remain limited, with constraints in material types and feature diversity, making them not universally applicable to all research scenarios. For instance:

- The Metallic Materials Properties Development and Standardization (MMPDS) handbook contains 213 S-N curve data for 62 metallic materials.
- The National Institute for Materials Science (NIMS) Fatigue Data Sheet includes 126 fatigue performance datasets for 59 metal materials.

While these datasets provide valuable resources for research, they also pose limitations in studying fatigue behavior under diverse materials and loading conditions.

In studies involving multiple materials and complex loading conditions, publicly available datasets often suffer from insufficient data volume and limited material variety, which may hinder the generalizability of research findings. Consequently, many researchers opt to conduct independent fatigue experiments and build their own datasets to address the issue of data scarcity. However, experimental data collection presents two major challenges:

- High experimental costs: Conducting fatigue tests is time-consuming and expensive, resulting in limited data availability, which may not comprehensively cover all materials and loading conditions.
- Data limitations: Fatigue life data obtained under specific material properties and loading conditions may not be directly applicable to other materials or conditions, limiting its generalization.

To address these challenges, this study employs the FatigueData-AM2022 dataset[10] as the foundation for model training to predict the fatigue life of additively manufactured alloys. The dataset comprises:

- 116 types of additively manufactured alloys extracted from 3,415 research papers.
- 1,610 S-N curve datasets, totaling 15,146 individual data points.

More importantly, FatigueData-AM2022 provides essential fatigue life prediction features, including basic material properties, loading conditions, and experimental methodologies.

By training the model on this dataset, this study effectively overcomes the limitations imposed by insufficient data volume that could otherwise hinder predictive performance. Additionally, since the dataset encompasses a wide range of metal types and fatigue-related features, the model exhibits strong generalization capability when predicting fatigue life across different loading conditions and material compositions. This significantly enhances the model's applicability and accuracy, making FatigueData-AM2022 a powerful tool for fatigue life prediction in additively manufactured alloys, further advancing research and applications in this domain.

3.1. Data Preprocessing

Although the FatigueData-AM2022 dataset contains extensive feature information, incorporating all features into the model does not necessarily yield the best predictive performance. Therefore, appropriate data preprocessing is crucial to ensure that the model effectively learns from the data. The specific preprocessing steps are described as follows.

First, the additive manufacturing dataset includes three types of fatigue data: stress-fatigue life (S-N) curves, strain-fatigue life (E-N) curves, and fatigue crack growth rate (da/dN) curves. This study selects S-N curve data as the primary research focus to investigate the fatigue life of metal materials under different stress conditions. Although the dataset contains 116 different additively manufactured alloys, the majority of the data is concentrated on four primary alloys: Ti-6Al-4V, IN718, 316L, and AlSi10Mg. To minimize the impact of materials with insufficient data on the analysis and reduce potential noise interference, this study removes materials with limited data, ensuring the research remains focused on these four major alloys.

Second, due to the limitations of traditional fatigue testing machines, such as low testing frequency, long experimental duration, and high testing costs, fatigue tests rarely exceed (10^7) cycles. As a result, 10^7 cycles is generally considered the upper limit in traditional

fatigue life research, and data beyond this range is scarce. Therefore, this study primarily focuses on fatigue data below 10^7 cycles, as it is more representative and abundant for analysis. Furthermore, the fatigue life values in the dataset span a broad range, from 10^2 to 10^7 cycles, creating significant scale differences between different fatigue data points. These large variations in magnitude could negatively affect the model's predictive performance. To eliminate the impact of scale differences, this study applies logarithmic transformation to fatigue life values, compressing the range to [2,7]. This transformation not only balances differences between different fatigue life scales but also ensures that the model processes the data more stably and effectively.

By implementing these preprocessing steps, this study ensures the quality and consistency of the dataset, providing clearer and more valuable information for subsequent model training. These data processing techniques help enhance the predictive accuracy and generalization ability of the model, enabling it to better capture the fatigue life characteristics of additively manufactured alloys.

3.2. Selection of Input Features

In the AM dataset, the features of fatigue life data provided in different studies vary due to differences in research focus and objectives. Therefore, for model training, it is essential to filter the features from the original dataset to ensure that the selected features effectively capture the key factors influencing fatigue life. Based on an analysis of factors affecting fatigue life, this study selected eight features from an initial pool of over 20 candidate features in the dataset as model inputs for predicting the fatigue life of additively manufactured alloys.

Metal fatigue life is primarily influenced by the physical properties of the metal material and loading conditions. In the original dataset, the heat treatment processes applied to the same metal material may vary across different studies. Even for the same type of metal, differences in its physical and mechanical properties may exist due to variations in processing techniques. The combination of different metal materials and treatment methods results in significant variations in the physical properties of metals.

This study focuses on the following four physical properties, as they have a direct impact on fatigue life:

- Ultimate Tensile Strength
- Yield Strength
- Metal Elongation
- Stress Concentration Factor

Thus, these four physical factors were first selected as key input variables in the model to ensure that the training process captures the essential influences on fatigue life prediction.

In describing loading environmental conditions, this study considers the ambient temperature in fatigue experiments. In addition to conventional room temperature conditions, some data also include high-temperature conditions. Variations in experimental temperature significantly affect metal fatigue life. High-temperature environments may lead to increased material deformation and crack propagation, resulting in a shorter fatigue life. In contrast, at low temperatures, metals exhibit higher toughness and resistance to deformation, which generally extends their fatigue life. Therefore, temperature is considered a critical factor in the loading environment and is incorporated into the model.

Additionally, the stress ratio (R) plays a crucial role in predicting fatigue performance. The stress ratio is directly related to metal fatigue life, influencing the stress distribution during the loading process. When R is close to 1, the stress concentration effect becomes more significant, potentially leading to localized stress accumulation, which accelerates crack initiation and propagation. Conversely, a lower stress ratio (smaller R value) helps reduce stress

concentration effects, thereby improving fatigue resistance. Thus, the stress ratio (R) is included as a key feature to enhance the accuracy of fatigue life prediction.

The loading frequency of the fatigue testing machine also has a significant impact on metal fatigue life behavior. At higher frequencies, certain metals may exhibit shorter fatigue lives, whereas at lower frequencies, they may have longer fatigue lives. Consequently, the loading frequency is considered an important input feature, helping the model understand how frequency affects metal fatigue behavior.

Finally, stress level is one of the most critical factors in studying metal fatigue life. Under cyclic stress loading, metals gradually accumulate damage, leading to crack initiation and propagation, ultimately causing fatigue failure. Higher stress levels generally result in shorter fatigue life, whereas lower stress levels tend to extend fatigue life. Thus, the relationship between stress level and fatigue life is among the most essential pieces of information in the dataset.

By selecting the above-mentioned physical properties and loading conditions, this study establishes a strong foundation for fatigue life prediction of additively manufactured alloys. These carefully selected input features effectively reflect how different metal materials behave under varying environmental conditions, contributing to the accuracy and generalization capability of the predictive model.

3.3. Dataset Processing Results

The figure illustrates the Kernel Density Estimation (KDE) distribution of the processed fatigue life data. After data filtering and preprocessing, the dataset includes three different additively manufactured alloys (Ti-6Al-4V, IN718, and 316L) along with eight key features, as shown in the corresponding table. Following these preprocessing steps, the final dataset consists of 1007 fatigue life records, ensuring a comprehensive representation of fatigue behavior under various loading and material conditions.

Table 1. Key Features of the Fatigue Life Dataset

Number	Feature Name	Unit
1	Fatigue Test Temperature	°C
2	Load Stress Ratio	Dimensionless
3	Load Frequency	Hz
4	Stress Concentration Factor	Dimensionless
5	Yield Strength	MPa
6	Ultimate Tensile Strength	MPa
7	Elongation	%
8	Load Stress	MPa

From Figure 6, the Kernel Density Estimation (KDE) curve (blue line) and the histogram illustrate that the processed dataset exhibits a near-normal distribution in terms of fatigue life. This distribution highlights the central tendency and overall shape of the dataset. The observed distribution pattern indicates that most of the fatigue life values are concentrated within a specific range, aligning well with the expected fatigue life distribution model. Such a distribution characteristic suggests that the dataset is well-structured and suitable for predictive modeling, providing a solid foundation for subsequent fatigue life analysis.

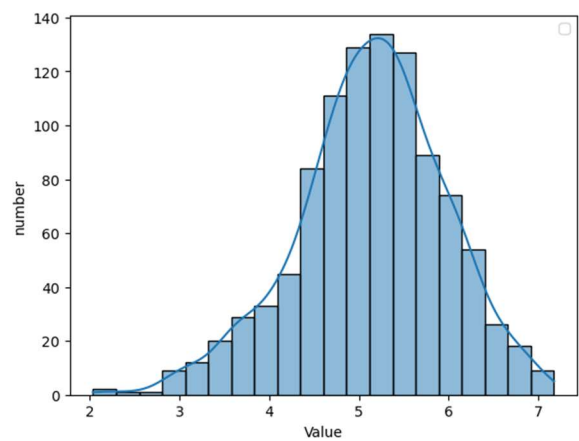


Figure 6. Kernel Density Plot of Fatigue Life in the Processed Dataset

In this study, Pearson correlation coefficient analysis was conducted on the filtered dataset to examine the linear relationships among input features and their association with fatigue life[6], as shown in Figure 7. This analysis helps to identify the degree of correlation between different features and fatigue life, providing insights into their potential influence on fatigue behavior. Understanding these correlations is essential for feature selection and model optimization, ensuring that the most relevant information is utilized for fatigue life prediction.

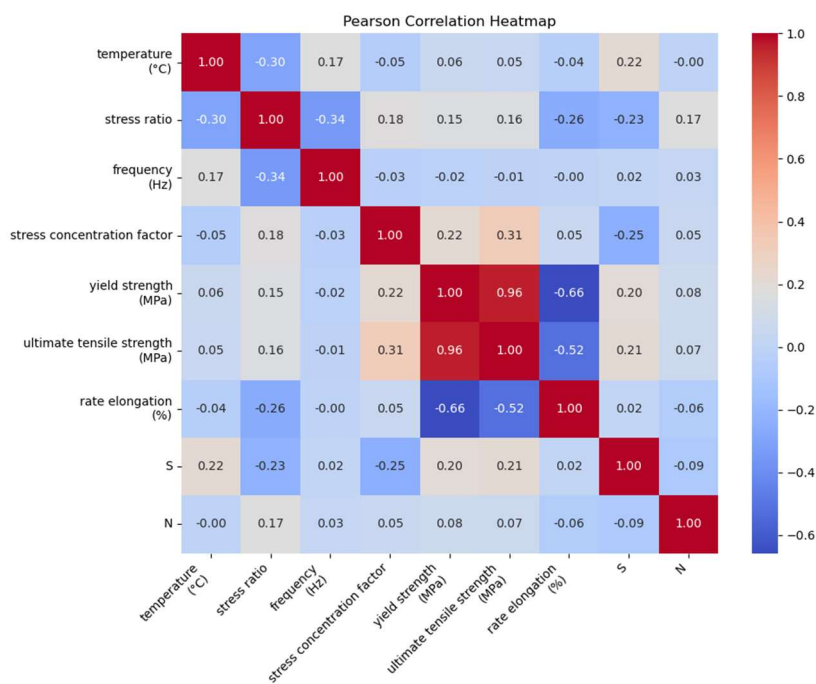


Figure 7. Fatigue Life Feature Correlation Heat Map

As shown in Figure 7, the Pearson correlation coefficients between the selected eight input features and fatigue life are generally low, indicating a weak linear relationship between them. This result suggests the presence of complex nonlinear dependencies within the data, which may be difficult to capture using simple linear regression methods. To better model these nonlinear relationships and enhance the accuracy of fatigue life prediction, this study adopts flexible deep learning models as predictive tools. Deep learning architectures are well-suited for handling nonlinear patterns within datasets, allowing for more accurate and robust fatigue life predictions. By leveraging deep learning models, the study aims

to improve the model's ability to recognize complex interactions between input features and fatigue life, ultimately leading to more precise and reliable predictive performance.

4. Deep Learning Model for Fatigue Life Prediction of Additively Manufactured Metals

In this section, the preprocessed additively manufactured metal fatigue dataset was trained with the objective of predicting the fatigue life of different metals under various loading conditions. To evaluate the training performance of the models, the newly proposed evaluation metric was employed. However, according to experimental results, even when using the same batch of test specimens, differences in microstructure and material composition during the manufacturing process lead to variations in the physical properties of metals, such as strength and hardness. These intrinsic differences contribute to the variability in fatigue life, making it challenging to achieve perfect fatigue life predictions, particularly when attempting to accurately predict the fatigue life of each specific material.

As a result, even when employing deep learning models for prediction, the model outputs often do not perfectly match the actual fatigue life values. Given this limitation, a certain tolerance standard is introduced to assess the predictive performance of the model. Specifically, the evaluation is based on the error margin of predictions, determining whether the model's predictions align with the expected level of accuracy. Even if the model cannot precisely predict the fatigue life of every individual sample, as long as it effectively captures the overall trend of the dataset and provides reasonable predictions, it can be considered to have met its expected performance to a certain degree.

Furthermore, considering that fatigue life prediction involves complex interactions among multiple influencing factors, this study integrates additional error metrics, such as Mean Absolute Error (MAE), Mean Squared Error (MSE), and the coefficient of determination (R^2), to comprehensively evaluate the model's performance across different metal materials and loading conditions. This multi-dimensional evaluation approach ensures the reliability and applicability of the model's predictions. By adopting this rigorous assessment framework, this study aims to validate the model's effectiveness and adaptability in practical applications of fatigue life prediction.

4.1. Model Architecture Design

First, several commonly used deep learning models were employed to train the dataset, including CNN, Transformer, GRU, MLP, and ResNet. The training results are presented in Figure 8. According to the bar chart in the figure, the predictive performance of these conventional deep learning models did not reach a satisfactory level, as indicated by the new evaluation metric being less than 4. However, among these models, GRU and Transformer demonstrated superior predictive performance, suggesting that these two architectures have certain advantages in handling metal fatigue data.

To further enhance prediction accuracy, this study explores a parallel combination of GRU and Transformer to leverage their respective strengths in predicting metal fatigue data, thereby improving the overall predictive performance on this dataset.

By parallelizing these two models, the approach integrates the advantages of GRU in processing sequential data with the capability of Transformer in capturing long-range dependencies. This hybrid architecture enhances the model's adaptability to complex data structures, particularly in metal fatigue prediction tasks under various material types and loading conditions.

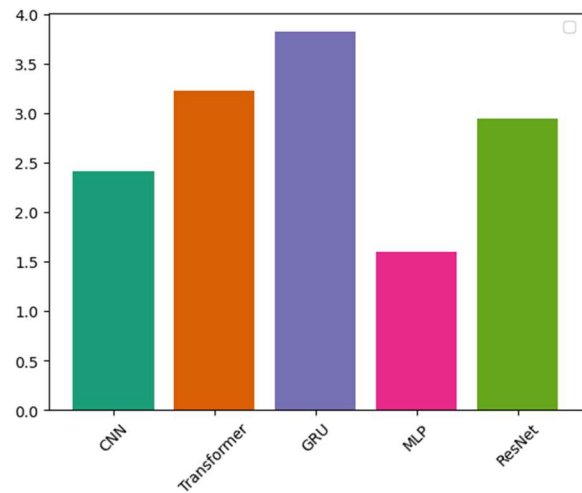


Figure 8. Comparison of Prediction Performance with New Metrics for Different Deep Learning Models

During the parameter optimization process of the proposed model, the selection of population size in the HLBO optimization method has a significant impact on the optimization effectiveness. Figure 9 presents a comparison of the model’s predictive performance under different population sizes. The results indicate that when the population size is set to 30, the optimized R² value reaches its maximum, demonstrating the best predictive accuracy. This finding suggests that a population size of 30 achieves the optimal optimization effect. Therefore, when performing parameter optimization for the model, setting the population size to 30 ensures improved predictive performance and enhanced model stability.

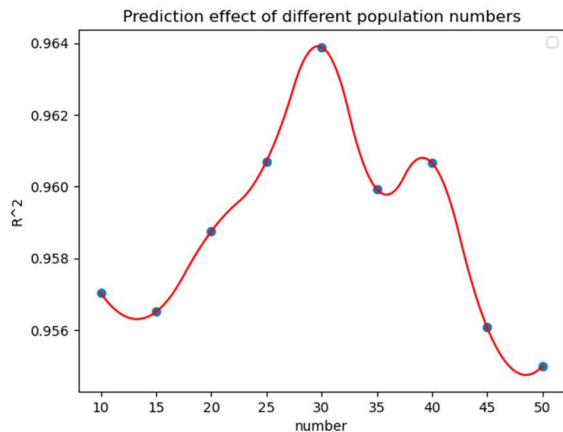


Figure 9. Impact of Different Population Sizes on Prediction Performance of the HLBO Optimization Algorithm

At the same time, by analyzing the optimization iteration process in Figure 10, it is observed that the model converged and reached its optimal state after approximately 20 iterations. This indicates that setting the number of iterations to 30 during the HLBO optimization is a reasonable choice, ensuring a stable and efficient optimization process. Therefore, in the practical application of HLBO optimization, selecting a population size of 30 and 30 iterations as parameter settings maximizes the optimization effect and enhances the prediction accuracy of the model. This configuration ensures that the optimization process sufficiently explores the solution space while avoiding unnecessary computational overhead.

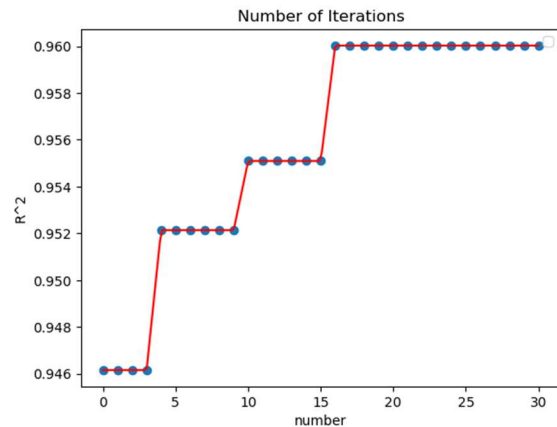


Figure 10. Iterative Process of HLBO Optimization Algorithm

The HLBO optimization algorithm was applied to select the key hyperparameters of the model, ensuring optimal performance. Through the optimization process, selecting the most suitable hyperparameters significantly enhances the predictive accuracy of the model. Table 2 presents the parameters subject to optimization along with their respective value ranges. During the optimization process, different combinations of these parameters were explored to identify the optimal solution.

Once the HLBO optimization algorithm achieved the best training performance, the optimized parameters were preserved to ensure the model's efficiency and stability in practical applications. The final optimized parameter configuration was determined, which played a crucial role in improving the model's generalization ability and accuracy. The specific optimal parameter settings, as shown in the table, include all key hyperparameter values, further enhancing the overall performance of the model.

Table 2. HLBO Optimized Model Parameter Values

Number	Name	Value Range	Final Value
1	Learning Rate	[0.0001,0.01]	0.0005
2	GRU Hidden Units	[124 , 1024]	201
3	GRU Layer	[1,6]	1
4	Number of Encoders	[1,6]	1
5	Number of Decoders	[1,6]	2

4.2. Experimental Results Analysis

After applying the proposed method, the training effectiveness of the model has been significantly improved. By implementing a parallel combination of the GRU and Transformer models, the predictive performance has shown a marked improvement compared to using either model individually. The overall predictive accuracy of the model has been significantly enhanced, demonstrating the advantage of this parallel architecture in fatigue life prediction. This structure enables more effective learning of fatigue life data patterns compared to a single deep learning model.

Furthermore, after employing Hybrid Leader-Based Optimization (HLBO) for hyperparameter tuning, the predictive performance of all models has generally improved. Notably, the proposed method, after parameter optimization, achieved significantly enhanced performance indicators, such as R², reaching a new level of accuracy. This improvement ensures that the proposed model exhibits higher reliability and predictive accuracy in metal fatigue life estimation.

To determine whether this model represents the optimal configuration, additional validation was conducted by testing other model combinations, as shown in Figure 11. For example, a sequential combination of GRU and Transformer was tested, as well as extending the parallel GRU-Transformer structure by incorporating additional models. However, experimental results confirmed that the parallel combination of GRU and Transformer yielded the best performance on the fatigue life dataset compared to other architectures. Based on evaluation metrics such as R^2 , MAE, and MSE, the parallel GRU-Transformer model demonstrated superior predictive accuracy and generalization capability. Therefore, the parallel GRU-Transformer model, optimized using HLBO, is the optimal choice in this study, achieving the best balance between prediction accuracy and model robustness.

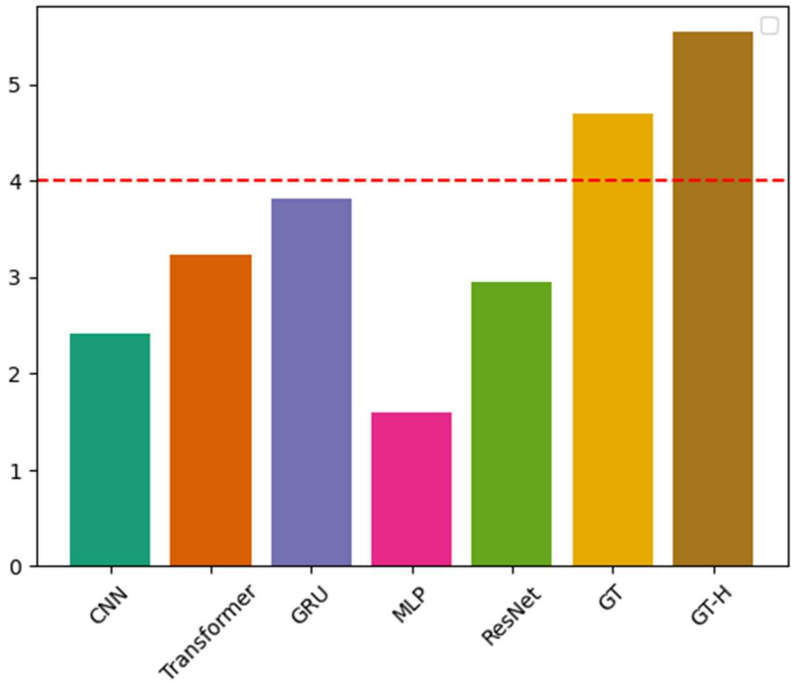


Figure 11. Comparison of Training Performance Using the New Evaluation Metric Across Different Deep Learning Models

As shown in Figure 12, the predictive performance of different models is compared using scatter plots. From the figure, it can be observed that while other prediction models achieve a certain degree of accuracy in fatigue life estimation, their performance is generally inferior to that of the proposed model. In the scatter plot of the proposed model, the prediction results exhibit higher accuracy, with data points more densely clustered. Within the prediction interval, the predicted values are more uniformly and concentrically distributed, with fewer data points exceeding the 1.5-times band. Furthermore, the test data and predicted values exhibit a clear linear relationship, indicating that the proposed model effectively fits the actual data and maintains high consistency with real-world conditions. The presence of this linear relationship further validates the model's excellent performance in fatigue life prediction, demonstrating its high prediction accuracy and stability. Consequently, the proposed model significantly outperforms conventional deep learning models in fatigue life estimation, as it better captures the underlying data patterns, thereby enhancing predictive performance.

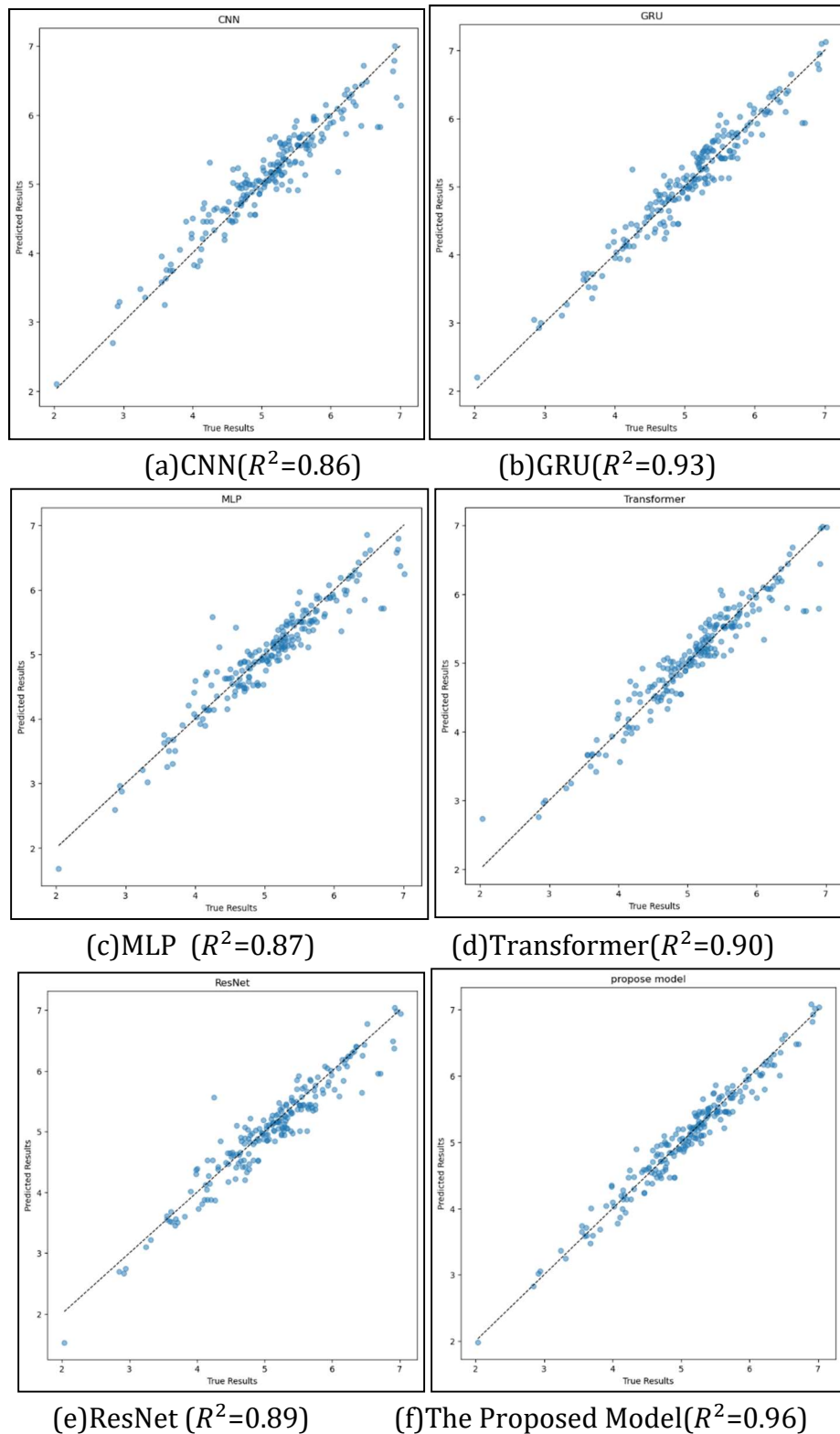


Figure 12. Comparison of Scatter Plots for Different Models

To comprehensively evaluate the differences in predictive performance among various models and validate the feasibility of the newly proposed evaluation metric, this study introduces the Squared Prediction Error (SPE) as a criterion for assessing model performance. SPE is a method used to quantify the distribution of prediction errors, effectively evaluating both the stability and accuracy of model predictions. A smaller SPE value indicates that the predicted results are closer to the actual data, signifying better predictive performance of the model.

The Squared Prediction Error (SPE) is calculated using the following formula:

$$SPE = \frac{1}{n} \sum_{i=1}^n \left(\frac{y_i - \hat{y}_i}{y_i} \right)^2 \quad (11)$$

where y_i represents the actual fatigue life data for the i -th sample,

$\hat{y}_i$ represents the predicted fatigue life data for the i -th sample,

n is the total number of samples.

By calculating the Squared Prediction Error (SPE), the error level of different model predictions can be quantitatively assessed, enabling a comparative analysis of predictive performance across models. On this basis, this study also presents box plots to provide a more intuitive visualization of the prediction results from different models. The box plot illustrates the distribution of prediction errors for each model, including the median, interquartile range, and outliers, offering a clear and comparative representation of model performance.

By integrating SPE analysis with box plots, this study conducts a comprehensive evaluation of different models in the task of fatigue life prediction. This combined approach validates the effectiveness of the newly proposed evaluation metric while further supporting the superiority of the proposed model in terms of predictive accuracy and robustness.

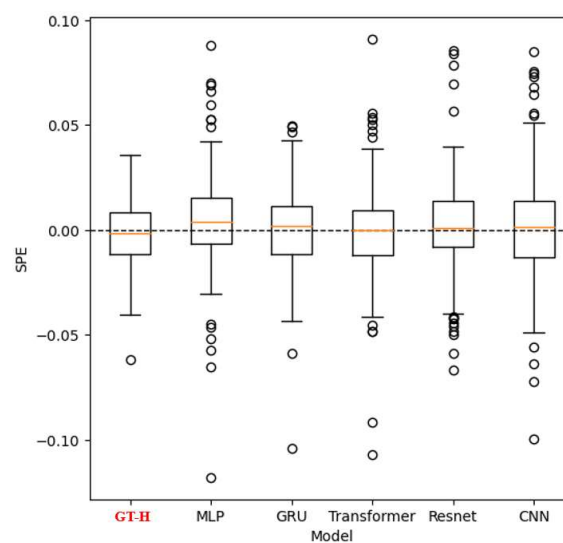


Figure 13. Box Plot Comparison of Different Models

From the box plot analysis, it is evident that the proposed model (GT-H) exhibits a shorter interquartile range (IQR) compared to other models, indicating that its predicted values are more concentrated and stable. Specifically, the MLP model's mean is significantly deviated from zero, suggesting a noticeable bias in its predictions. In contrast, the proposed model shows a mean value slightly below zero, implying that its overall predictions are more conservative, with fatigue life estimates generally lower than the actual values.

Additionally, the proposed model displays fewer outliers, and the magnitude of these outliers is relatively lower, further demonstrating its superiority in metal fatigue life prediction. Compared to other models, the proposed approach not only produces more concentrated and stable predictions but also effectively suppresses the occurrence of extreme outliers, indicating strong robustness in predictive performance.

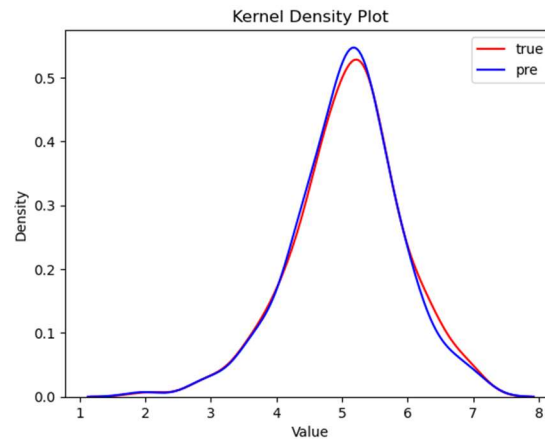


Figure 14. Comparison of Kernel Density Estimation (KDE) Between Predicted and Actual Fatigue Life Values

The conclusions drawn from the scatter plots and box plots are consistent with the findings obtained earlier, further validating the effectiveness of the proposed evaluation metric and demonstrating the model's excellent predictive capability. By comparing the kernel density estimation (KDE) distributions of actual and predicted values, it is evident that the predicted values closely align with the actual values. Particularly, in the fatigue life range of 10^6 to 10^7 cycles, the predicted values tend to be slightly lower than the actual values, as indicated by the fact that most data points in the scatter plot are distributed below the standard reference line. This suggests that the proposed prediction model exhibits a conservative tendency in predicting high-cycle fatigue life. Conversely, for fatigue life around 10^5 cycles, the predicted values are slightly larger than the actual values, with the scatter plot data points predominantly located above the standard reference line. This observation is consistent with the actual distributions shown in the box plots and scatter plots, further verifying the model's predictive performance across different fatigue life ranges.

By employing the HLBO optimization algorithm for hyperparameter tuning, the proposed model surpasses several traditional machine learning models in predictive accuracy. Unlike conventional empirical formulas that rely on extensive experimental data and are typically tailored for specific materials or particular loading conditions, the proposed model effectively predicts fatigue life across multiple metal materials and various loading scenarios.

The parallel structure enhances the predictive performance of the proposed model compared to single-model architectures, further validating the advantages of parallel network structures in establishing the relationship between loading conditions and fatigue life. The parallel model is more adept at capturing complex feature information within the data, adapting to the diverse properties of additively manufactured alloys and varying loading conditions, thereby achieving more precise predictions.

In summary, the proposed model demonstrates superior predictive performance in fatigue life estimation compared to other models. Notably, when applied to different materials and environmental conditions, its predictive capability remains relatively stable, indicating that the model is well-suited for complex real-world engineering applications. Its strong generalization ability and robustness make it a promising tool for practical fatigue life assessment in additively manufactured components.

5. Discussion

The study validates the superiority of the proposed model in predicting the fatigue life of additively manufactured metals, particularly when dealing with various metal materials and loading conditions, where the model demonstrates accurate predictive performance.

This study introduces a parallel model integrating the Gated Recurrent Unit (GRU) and Transformer architectures, specifically designed for fatigue life prediction of additively manufactured metallic materials. To further enhance predictive performance, the Hybrid Leader-Based Optimization (HLBO) algorithm is employed for hyperparameter optimization. The model is trained and evaluated using preprocessed data from the FatigueData-AM2022 dataset. Subsequently, to further verify its effectiveness, the proposed model is compared against five other deep learning models, and the results indicate that the proposed model achieves higher prediction accuracy.

The main conclusions are as follows:

(1)Advantages of the Parallel Model: The proposed parallel model effectively processes fatigue data of additively manufactured metals while leveraging the strengths of both the GRU and Transformer models. By capturing data features from multiple perspectives, the model can extract potential relationships within the dataset, leading to improved prediction accuracy.

(2)HLBO-based Parameter Optimization: The HLBO algorithm is utilized for automatic hyperparameter optimization, reducing manual tuning efforts. Experimental validation confirms that HLBO exhibits excellent global and local search capabilities, further enhancing the predictive performance of the model.

(3)Model Superiority and Generalization Ability: Compared to other deep learning models, the proposed parallel model demonstrates a significant improvement in predictive accuracy. The dataset used in this study includes fatigue life data for three different additively manufactured alloy materials under various loading conditions. Notably, different material types were not explicitly included as input features, yet the model successfully learned the physical properties of different materials and loading conditions, enabling accurate fatigue life predictions. Furthermore, the model exhibits strong generalization ability in predicting fatigue life across welded structures and varying load conditions, showcasing its adaptability and robustness.

Based on these findings, the study confirms the effectiveness of the proposed model in predicting the fatigue life of additively manufactured metals. The model's ability to handle diverse metal materials and loading conditions suggests its broad application potential, providing reliable predictive support for the field of additively manufactured alloys.

References

- [1] Abroug F, Monnier A, Arnaud L, Balcaen Y, Dalverny O. High cycle fatigue strength of additively manufactured AISI 316L Stainless Steel parts joined by laser welding. *Eng Fract Mech* 2022;275:108865.
- [2] Yue P, Ma J, Dai CP, Zhang JF, Du W. Probabilistic framework for reliability analysis of gas turbine blades under combined loading conditions. *Structures* 2023; 55:1437–46.
- [3] Wu H, Sun C, Xu W, Chen X, Wu X. A novel evaluation method for high cycle and very high cycle fatigue strength. *Eng Fract Mech* 2023;109482.
- [4] Li H, Huang CG, Guedes SC. A real-time inspection and opportunistic maintenance strategies for floating offshore wind turbines. *Ocean Eng* 2022;256:111433.
- [5] Xu S, Zhu SP, Hao YZ, Liao D, Qian G. A new critical plane-energy model for multiaxial fatigue life prediction of turbine disc alloys. *Eng Fail Anal* 2018;93: 55–63.

- [6] Zhang W, Wu C, Zhong H, Li Y, Wang L. Prediction of undrained shear strength using extreme gradient boosting and random forest based on Bayesian optimization. *Geosci Front* 2021;12:469–77.
- [7] Agrawal A, Deshpande PD, Cecen A, Basavarsu GP, Choudhary AN, Kalidindi SR. Exploration of data science techniques to predict fatigue strength of steel from composition and processing parameters. *Integr Mater Manuf Innov* 2014;3:90–108.
- [8] Sameen MI, Pradhan B, Lee S. Self-Learning Random Forests Model for Mapping Groundwater Yield in Data-Scarce Areas. *Nat Resour Res* 2019;28:757–75.
- [9] Dehghani, M., Trojovský, P. Hybrid leader based optimization: a new stochastic optimization algorithm for solving optimization applications. *Sci Rep* 12, 5549 (2022).
- [10] Zhang, Z., Xu, Z. Fatigue database of additively manufactured alloys. *Sci Data* 10, 249 (2023).