

Optimizing Random Forest Models for Early Detection of Defects in Steel

Tri Surawan¹, Adhitio Satyo Bayangkari Karno², Widi Hastomo^{3,*}, Reza Fitriansyah³, Ahmad Eko Saputro⁴, Indra Bakti³

¹ Department of Mechanical Engineering, Faculty of Technology Industry, Jayabaya University, Indonesia; e-mail: tri.surawan@gmail.com

² Department of Information System, Faculty of Engineering, Gunadarma University, Depok; Indonesia; e-mail: adh1t10.2@gmail.com

³ Department of Information Technology, Ahmad Dahlan Institute of Technology and Business, Jakarta; Indonesia; e-mail: widie.has@gmail.com; rezafsyh@gmail.com; indra.chufran@gmail.com

⁴ Department of Management, Ahmad Dahlan Institute of Technology and Business, Jakarta; Indonesia; e-mail: ahmadeko23@gmail.com

*Correspondence: e-mail: widie.has@gmail.com

Diterima: 8 Juni 2024; Review: 28 Juni 2024; Disetujui: 29 Juni 2024; Diterbitkan: 30 Juni 2024

Abstract

In the manufacturing sector, steel plate defects are a severe issue that may result in significant losses for a company's finances and image. The purpose of this study is to evaluate how well three machine learning algorithms detect steel plate flaws. The accuracy, area under the ROC curve (ROC-AUC), and Log-Loss of the method were used to assess its performance using a dataset that was downloaded from www.kaggle.com. Based on the findings, the Random-Forest algorithm performed best overall, having the lowest Log-Loss of 0.9327, an accuracy of 0.6722, and an AUC value of 0.9222. Research using other algorithms is still very open to be carried out to get better results. Research utilizing other algorithms is still very much open to be conducted in order to get better outcomes.

Keywords: Random Forest Models, Early Detection, Defects in Steel.

1. INTRODUCTION

Steel, a key material in many different areas of production, has to overcome challenges regarding defects on plates during manufacturing processes. To overcome the problem, modern techniques including machine learning have been developed to accurately predict these defects. Different researchers have adopted fresh strategies ranging from deep learning-oriented defect detection approaches to time-series-based neural networks predicting steel characteristics or fusion systems integrating handcrafted and abstract features, toward identifying shortcomings (Akhyar et al., 2023).

Steel plate defect prediction is an important industrial consideration in the domain of manufacturing and construction. For example, a lightweight YOLO-ACG methodology has been developed to improve the detection of defects in steel surfaces with accuracy and speed balance (Yang et al., 2023). Similarly, recent studies have shown that time-series neural networks based on LSTM offer improved predictions regarding features of steel, such as yield strength and ultimate

tensile strength (Yang et al., 2023). Manufacturing processes can be supported by implementing predictive models that can help identify scrap parts without increasing physical inspections thereby enhancing product quality while minimizing cost (Shang et al., 2023). These innovations are beneficial because they lower risks associated with defects in steel plates ensuring structural integrity, operational efficiency, and safety standards in industrial operations.

This research is designed to improve defect prediction in sheets of iron by using methods of machine learning Random Forest (Ghasemkhani et al., 2023). The investigation uses ensemble learning techniques to handle complexities in the detection of defects as a means to enhance the quality evaluation approaches (Knott et al., 2023). The study helps develop defect prediction methodologies in the steel industry thereby assisting industrial stakeholders in improving quality control practices and operational efficiency in their industries (Akhyar et al., 2023).

Complete analyses of ensemble learning methods Random Forest in raising the accuracy of steel plate defect prediction are lacking in the current literature (Shang et al., 2023). To fill this vacuum, recent studies have presented innovative techniques for identifying surface defects in steel. These techniques include the use of convolutional neural networks in a multiscale local and global feature fusion mechanism (Zhang et al., 2023), the logistic model tree forest approach for fault prediction in stainless steel plates (Ghasemkhani et al., 2023), and an adaptive weighting multimodal fusion classification system for precise defect identification (Tu et al., 2023). With precision rates as high as 99.0% in certain instances, these developments have demonstrated notable improvements in defect detection accuracy (Wu et al., 2023). Through the integration of these inventive techniques, scholars can elevate the current benchmark in evaluating steel quality and create more efficient methods for identifying flaws in industrial applications.

The main objective of this research is to determine how well the Random Forest machine learning algorithm can improve the precision and reliability of steel plate defect prediction. Ultimately, it is hoped that the results of this research will expand defect prediction methods in the steel sector and provide useful suggestions for practitioners and stakeholders who wish to improve operational effectiveness and quality control procedures in industrial environments.

2. METHOD

The steps used throughout the preprocessing process will be covered in this section.

2.1. Dataset

By downloading a CSV file ("train.csv") containing 19,219 rows and 33 columns of data, the dataset was acquired from www.kaggle.com (Walter Reade, 2024). There are 33 columns in the "train.csv" dataset: 7 columns are dedicated to goals and 26 columns are dedicated to features. Seven kinds of steel plate defects—Pastry, ZScratch, KScratch, Stains, Dirtiness, Bumps, and

Other Faults—are included in the target data. The dimensions and form of the steel plate defect are fully explained by these 26 features, which are as follows: Three coordinates—Pixels_Areas, MinLuminosity, MaxLuminosity, and YMin, YMax—represent the defect area, while four coordinates represent the defect location. TypeOfSteelA300, TypeOfSteelA400, SteelPlateThickness, EdgesIndex, BlankIndex, SquareIndex, OutsideXIndex, EdgesXIndex, EdgesYIndex, OutsideGlobalIndex are the ten material features and indices. Three logarithmic features (LogArea, Log_XIndex, Log_YIndex) and three statistical features (OrientationIndex, LuminosityIndex, SigmoidAreas).

2.2. Preprocessing

This section will describe the stages carried out in the preprocessing process.

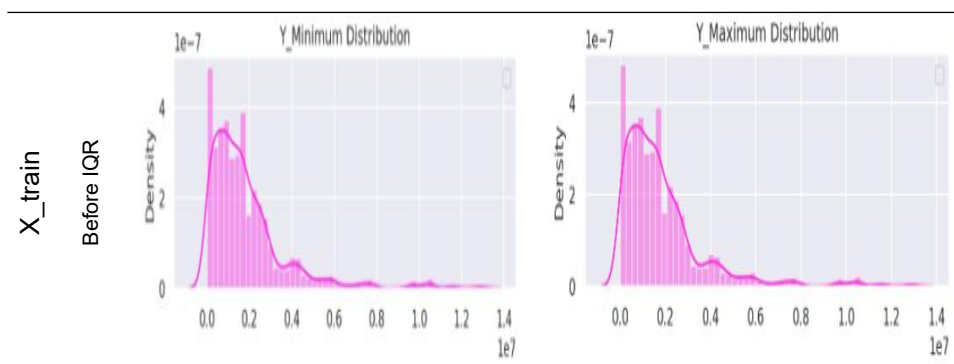
Splitting

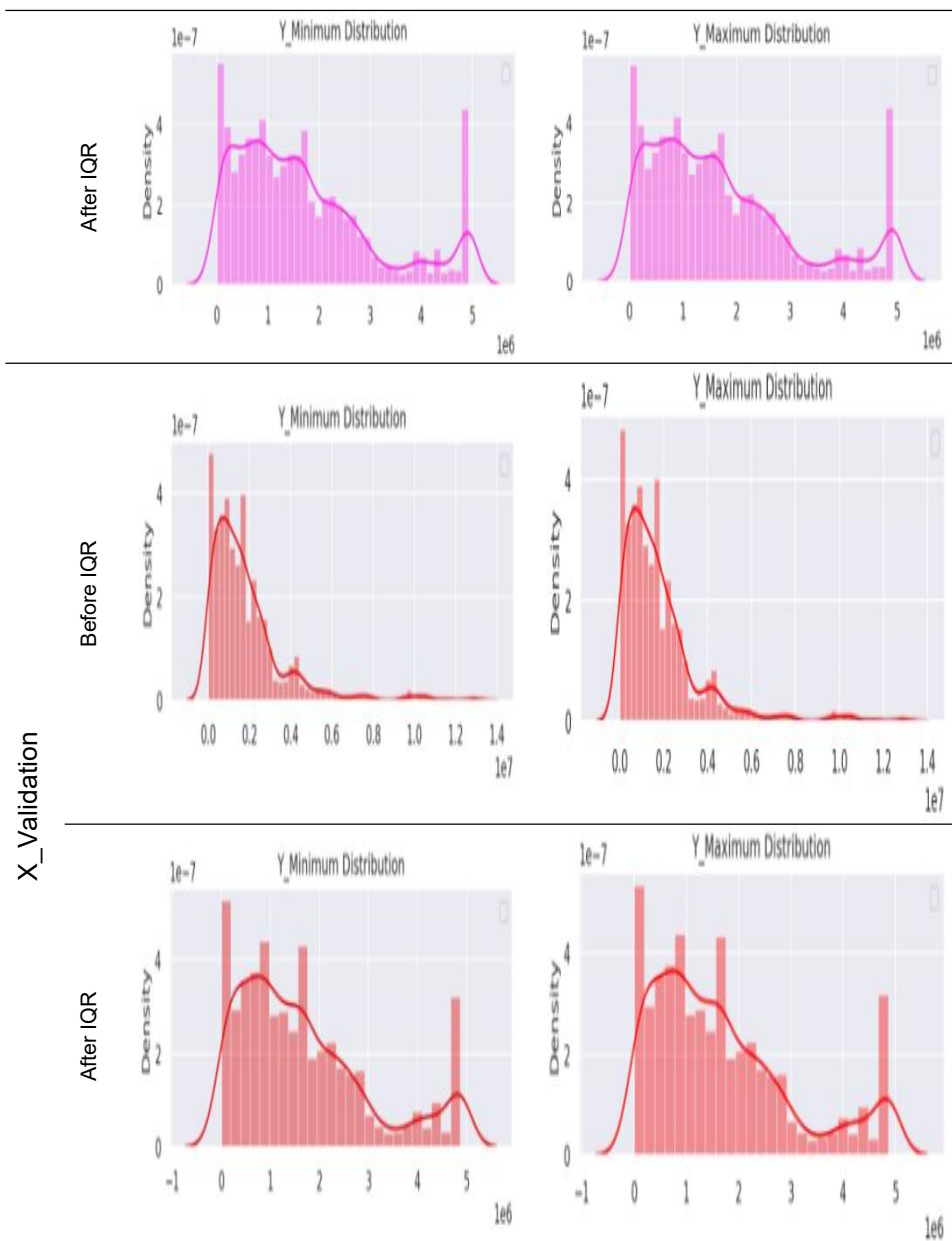
The target data was shrunk from seven columns to one, bringing the total down to twenty-seven columns. The 27 data columns were split into 26 columns for feature data (x) and 1 column for target data (y) in order to prepare for the training process. The original 19,219 rows of data are separated into 14,414 rows of training data (75%) and 4,805 rows of validation data (25%) for these two categories of data (x and y).

Outlier

When processing data, one well-liked and effective method for identifying and handling outliers is the IQR (Interquartile Range) Method. Since this non-parametric approach is less likely to produce extreme findings or outliers, it is particularly useful when applied to datasets with heavy-tailed or non-normal distributions (Yulianto et al., 2023).

Figure 1: Distribution plot with IQR applied before and after





Source : Research Results (2022)

Figure 1 is the IQR is the difference between the 75th and 25th percentiles, or the third and first quartiles, of the data distribution. Since Q1 and Q3 represent the first and third quartiles, respectively, all data that fall outside of the range between the $-1.5 * IQR + Q1$ limit and the $1.5 * IQR + Q3$ limit are considered outliers (Hoaglin et al., 1986). Due to page limitations, this example only shows a scatter plot of many data_train and data_validation features that show the presence of outliers, as well as a scatter plot of features that follow the IQR stage.

Imbalanced data

A class imbalance arises when the data set used to identify steel plate defects contains more samples from the non-defect class (majority class) than from the defect class (minority class) (Hastomo et al., 2021). This discrepancy in data has the ability to tamper with machine learning algorithms and impair minority class categorization accuracy. To solve the issue in this study, the Synthetic Minority Over-sampling (SMOTE) technique (Kuhn, M., & Johnson, 2013) was applied to the data set. Using an oversampling approach called SMOTE, an artificial minority sample is created by interpolating from pre-existing minority samples for the minority class. This process is continued until the minority class acquires the necessary quantity of synthetic samples. SMOTE is applied to a data set of steel plates that suffer from class imbalance between defects and non-defects. By using the value $k = 5$ to find the nearest neighbors in the SMOTE process. The number of synthetic samples generated was adjusted to achieve a more balanced ratio between defective and non-defective classes (Karno et al., 2023).

Scaling

In machine learning, feature scaling is a crucial step in the pre-processing of data. Feature scaling aims to place all feature values on a single scale, within a predetermined range. This is important because differences in feature size can affect the performance of several machine-learning techniques (Kuhn, M., & Johnson, 2013). Among the most used feature scaling methods is StandardScaler. A scaling method called StandardScaler modifies each characteristic such that its average (mean) is equal to zero and its variance is equal to one (Géron, 2022). Using this method, the average and standard deviation of each feature value are divided (Bishop, 2006). StandardScaler is readily built-in in Python's scikit-learn package by utilizing the StandardScaler class from the preprocessing module (Tao et al., 2022). Furthermore, several theoretical and empirical studies have demonstrated that machine learning models perform better when StandardScaler is used (D.K. et al., 2019).

2.3. Model Building

Random Forest Classifier

Random Forest is a versatile and widely-used ensemble learning method that extends the idea of Bagging to multiple decision trees. Introduced by (Breiman, 2001), Random Forest not only reduces overfitting by averaging multiple decision trees but also improves the model's accuracy by introducing randomness in the feature selection process during tree construction. Random Forest builds upon the Bagging method by incorporating an additional layer of randomness. In addition to creating bootstrap samples of the training data, Random Forest also selects a random subset of features at each split in the decision tree construction. This random feature selection helps to de-

correlate the individual trees, leading to a more diverse and robust ensemble. Each tree in the Random Forest is grown to the maximum depth without pruning, and the final prediction is obtained by averaging the predictions of all individual trees (for regression tasks) or by majority voting (for classification tasks). This dual randomness—bootstrap sampling of data and random feature selection—enables Random Forest to achieve high accuracy while maintaining generalization.

2.4. Performance parameter

Confusion Matrix

In the field of machine learning, evaluating the performance of classification models is crucial for understanding their effectiveness and areas of improvement (Widi Hastomo et al., 2022). The confusion matrix is a fundamental tool for this purpose, providing a detailed breakdown of a model's predictions. It allows practitioners to visualize the performance of an algorithm by comparing the predicted classifications with the actual classes. A confusion matrix is a square matrix that summarizes the performance of a classification (Hastomo, W., Karno, A. S. B., Kalbuana, N., Nisfiani, E., & Lussiana, 2021). For a binary classification problem, the matrix is typically structured as follows:

Figure 2. Each cell in the matrix represents a different type of prediction

	Predicted Positive	Predicted Negative
Actual Positive	True Positive (TP)	False Negative (FN)
Actual Negative	False Positive (FP)	True Negative (TN)

Source : Research Results (2022)

Figure 2 is the matrix represents a different type of prediction, each cell in the matrix represents a different type of prediction:

- **True Positives (TP):** The number of correctly predicted positive cases.
- **False Positives (FP):** The number of negative cases incorrectly predicted as positive.
- **True Negatives (TN):** The number of correctly predicted negative cases.
- **False Negatives (FN):** The number of positive cases incorrectly predicted as negative.

Several important performance metrics can be derived from the confusion matrix:

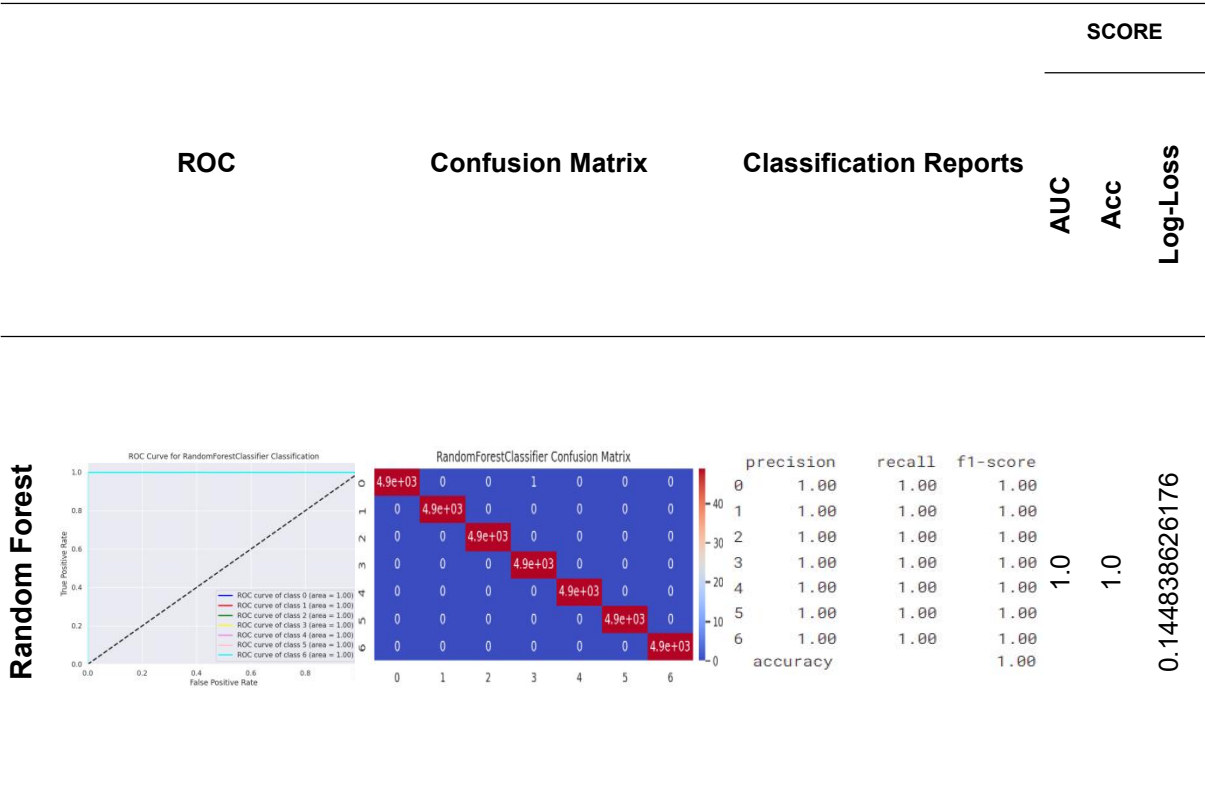
- **Accuracy:** The proportion of total correct predictions (TP + TN) out of all predictions.
- **Precision:** The proportion of true positive predictions out of all positive predictions.
- **Recall (Sensitivity):** The proportion of actual positives correctly identified.
- **Specificity:** The proportion of actual negatives correctly identified.
- **F1 Score:** The harmonic mean of precision and recall.

Log loss.

Logarithmic Loss, commonly known as Log Loss, is a fundamental metric used to evaluate the performance of classification models, particularly in probabilistic frameworks. It provides a measure of how well the predicted probabilities of a model match the actual class labels. Lower log loss values indicate better performance, with a log loss of zero representing perfect predictions. Log Loss is commonly used in fields such as finance, healthcare, and marketing, where understanding the uncertainty of predictions is critical. It is also a standard evaluation metric in many machines learning competitions, including those hosted by Kaggle. Libraries such as Scikit-learn in Python provide easy-to-use functions to calculate Log Loss, facilitating its widespread use.

3. RESULTS AND DISCUSSION

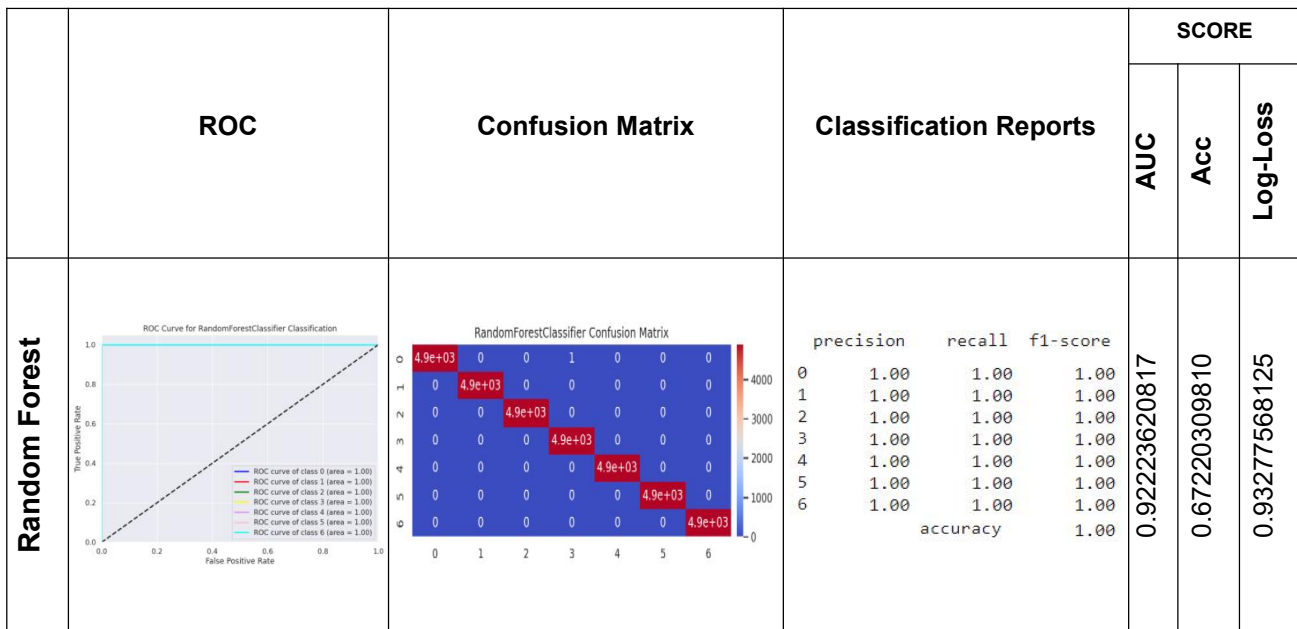
Figure 3. Training Results



Source : Research Results (2022)

Figure 3 is illustrate the use of the random forest algorithm in the training. ROC, Confusion Matrix, precision, recall, f1-score, AUC, accuracy, and log loss are used to assess each type of algorithm. The training results show that the Random-Forest technique has a low log loss value of around 0.14 and AUC and Acc values greater than 1.0.

Figure 4. Validation Results



Source : Research Results (2022)

Figure 4 is illustrate the use of the random forest algorithm in the validation stages. However, when the training model was tested using evaluation data, the Random-Forest method with a Log Lost value reached 0.93277568, while the lowest Log Lost value was 0.93277568125.

4. CONCLUSION

Tests show that machine learning can detect weaknesses in steel plates more precisely. This study assesses that machine learning techniques can be used to find weaknesses in steel plates. This study finds the best solution to the error detection problem. Based on the data used in this research, Random-Forest is an efficient algorithm for detecting weaknesses in steel plates. The Random-Forest Log-Loss value needs to be increased to get better results even though it is the smallest of the two other algorithms used in this research. There is still great hope for further research to use different algorithms to obtain even higher accuracy and Log-Lost values.

REFERENCES

- Akhyar, F., Liu, Y., Hsu, C.-Y., Shih, T. K., & Lin, C.-Y. (2023). FDD: a deep learning–based steel defect detectors. *The International Journal of Advanced Manufacturing Technology*, 126(3), 1093–1107. <https://doi.org/10.1007/s00170-023-11087-9>
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer google scholar.
- Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32.

- <https://doi.org/10.1023/A:1010933404324>
- D.K., T., B.G. P., & Xiong, F. (2019). Auto-detection of epileptic seizure events using deep neural network with different feature scaling techniques. *Pattern Recognition Letters*, 128, 544–550. <https://doi.org/https://doi.org/10.1016/j.patrec.2019.10.029>
- Géron, A. (2022). *Hands-on machine learning with Scikit-Learn, Keras, and TensorFlow*. O'Reilly Media, Inc.
- Ghasemkhani, B., Yilmaz, R., Birant, D., & Kut, R. A. (2023). Logistic Model Tree Forest for Steel Plates Faults Prediction. *Machines*, 11(7). <https://doi.org/10.3390/machines11070679>
- Hastomo, W., Karno, A. S. B., Kalbuana, N., Nisfiani, E., & Lussiana, E. T. P. (2021). Deep Learning Optimization for Stock Predictions during the Covid-19 Pandemic. *JEPIN (Jurnal Edukasi Dan Penelitian Informatika)*, 7(2), 133–140. <https://jurnal.untan.ac.id/index.php/jepin/article/view/47411>
- Hastomo, W., Bayangkari Karno, A. S., Kalbuana, N., Meiriki, A., & Sutarno. (2021). Characteristic Parameters of Epoch Deep Learning to Predict Covid-19 Data in Indonesia. *Journal of Physics: Conference Series*, 1933(1), 012050. <https://doi.org/10.1088/1742-6596/1933/1/012050>
- Hoaglin, D. C., Iglewicz, B., & Tukey, J. W. (1986). Performance of Some Resistant Rules for Outlier Labeling. *Journal of the American Statistical Association*, 81(396), 991–999. <https://doi.org/10.1080/01621459.1986.10478363>
- Karno, A. S. B., Hastomo, W., Surawan, T., Lamandasa, S. R., Usuli, S., Kapuy, H. R., & Digdoyo, A. (2023). Classification of cervical spine fractures using 8 variants EfficientNet with transfer learning. *International Journal of Electrical and Computer Engineering*, 13(6), 7065–7077. <https://doi.org/10.11591/ijece.v13i6.pp7065-7077>
- Knott, A.-L., Stauder, L., Ruan, X., Schmitt, R. H., & Bergs, T. (2023). Potential of prediction in manufacturing process and inspection sequences for scrap reduction. *CIRP Journal of Manufacturing Science and Technology*, 44, 55–69. <https://doi.org/https://doi.org/10.1016/j.cirpj.2023.04.012>
- Kuhn, M., & Johnson, K. (2013). *Applied predictive modeling*. New York: springer.
- Shang, Z., Li, B., Chen, L., & Zhang, L. (2023). Defects Prediction Method for Radiographic Images Based on Random PSO Using Regional Fluctuation Sensitivity. *Sensors*, 23(12). <https://doi.org/10.3390/s23125679>
- Tao, H., Awadh, S. M., Salih, S. Q., Shafik, S. S., & Yaseen, Z. M. (2022). Integration of extreme gradient boosting feature selection approach with machine learning models: application of weather relative humidity prediction. *Neural Computing and Applications*, 34(1), 515–533. <https://doi.org/10.1007/s00521-021-06362-3>
- Tu, F. M., Wei, M. H., Liu, J., & Liao, L. L. (2023). An adaptive weighting multimodal fusion classification system for steel plate surface defect. *Journal of Intelligent & Fuzzy Systems*, 45,

3501–3512. <https://doi.org/10.3233/JIFS-230170>

Walter Reade, A. C. (2024). *Steel Plate Defect Prediction*. <https://www.kaggle.com/competitions/playground-series-s4e3>

Widi Hastomo, Nur Aini, Adhitio Satyo Bayangkari Karno, & L.M. Rasdi Rere. (2022). Metode Pembelajaran Mesin untuk Memprediksi Emisi Manure Management. *Jurnal Nasional Teknik Elektro Dan Teknologi Informasi*, 11(2), 131–139. <https://doi.org/10.22146/jnteti.v11i2.2586>

Wu, S., Chu, H., & Cheng, C. (2023). Deep Network for Steel Surface Defect Detection Based on Attention Mechanism. *2023 2nd International Conference on Innovations and Development of Information Technologies and Robotics (IDITR)*, 89–93. <https://doi.org/10.1109/IDITR57726.2023.10145981>

Yang, Z., Wang, Y., Xu, F., Li, X., Yang, K., Xia, W., Cai, J., Xie, Q., & Xu, Q. (2023). Online Prediction of Mechanical Properties of the Hot Rolled Steel Plate Using Time-series Deep Neural Network. *ISIJ International*, 63(4), 746–757. <https://doi.org/10.2355/isijinternational.ISIJINT-2022-383>

Yulianto, R., Rusli, M. S., Satyo, A., Karno, B., Hastomo, W., & Kardian, A. R. (2023). *Innovative UNET-Based Steel Defect Detection Using 5 Pretrained Models Innovative UNET-Based Steel Defect Detection Using 5 Pretrained Models*. 10(4), 2365–2378.

Zhang, L., Fu, Z., Guo, H., Sun, Y., Li, X., & Xu, M. (2023). Multiscale Local and Global Feature Fusion for the Detection of Steel Surface Defects. *Electronics*, 12(14). <https://doi.org/10.3390/electronics12143090>