

Machine learning-based joint optimization for LoRaWAN reliability and energy efficiency

Nana Rusdiana, Novi Indriyani, Adi Nugroho, Riri Fitri Sari, Dino Budi Prakoso

Department of Electrical Engineering, Faculty of Engineering, University of Indonesia, Depok, Indonesia

Article Info

Article history:

Received Mar 4, 2026

Revised Jun 3, 2026

Accepted Jun 30, 2026

Keywords:

Energy efficiency

Internet of things network

Light gradient boosting machine

Long-range wide area network

Network simulator 3 simulation

Packet delivery ratio

Random forest

ABSTRACT

Dense long-range wide area network (LoRaWAN) deployments often experience reduced packet delivery reliability (PDR) and increased energy consumption due to suboptimal adaptive data rate (ADR) mechanisms. This study proposes a machine learning (ML) based joint optimization method to improve PDR or reliability while minimizing node energy consumption. The proposed solution employs random forest (RF) and light gradient boosting machine (LightGBM) models, which are suitable for resource constrained end devices (ED). Using combinations of data rate (DR), frequency channel, and transmission power (TxPower) from 500 ED, the model extracts 14 features for parameter optimization. Unlike conventional ADR, the proposed approach can learn complex relationships beyond the received signal strength indicator (RSSI) and signal-to-noise ratio (SNR) parameters. The model achieves higher prediction accuracy within a 5 km coverage area. Network simulator 3 (NS-3) simulation results demonstrate better performance and energy efficiency compared to conventional ADR. The proposed approach improves PDR by 6.1–11.6% and reduces node energy consumption by 20–25%, providing a better balance between communication reliability and energy efficiency in scalable LoRaWAN networks.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Nana Rusdiana

Department of Electrical Engineering, Faculty of Engineering, University of Indonesia

West Java, Depok 16424, Indonesia

Email: nana.rusdiana51@ui.ac.id

1. INTRODUCTION

The population of internet of things (IoT) devices is expected to increase from 22.3 billion in 2025 to 47.1 billion in 2031 based on the demand for smart cities, buildings, monitoring, and agriculture [1]. The long-range communication and low power consumption of long-range wide area network (LoRaWAN) technology have led to its identification as part of the most popular technologies in terms of low-power wide area network (LPWAN). Meanwhile, the dense deployments have reliability problems due to the pure ALOHA protocol which is related to the limitation of the packet delivery ratio (PDR) under 60% due to collision when there are more than 100 devices [2]–[4].

LoRaWAN theoretically supports 10-year battery life but the real-world implementations rarely achieve this due to inefficient parameter selection [5]. Conventional adaptive data rate (ADR) addresses the problem by adjusting transmission power (TxPower) and spreading factor (SF) based on the last 20 uplink transactions [6] but the convergence remains slow. Current studies primarily optimize SF and TxPower independently with limited work on frequency-channel allocation [7] or systematic multi-parameter optimization [8]. However, the approaches that achieve balanced PDR-energy trade-offs across all three parameters remain rare. Although existing LoRaWAN optimization studies have improved either reliability

or energy efficiency, most approaches still treat SF, TxPower, and channel allocation as separate optimization problems. This separation limits the ability of the network to achieve a balanced trade-off between PDR and energy consumption, especially in dense deployments where collision probability, link quality variation, and airtime usage interact with each other. Therefore, a joint optimization framework is required to select LoRaWAN parameters by considering reliability and energy efficiency simultaneously.

This paper introduces a lightweight, machine learning (ML)-based multi-objective framework for optimizing channel-aware LoRaWAN. This study concurrently optimizes three critical parameters: SF, TxPower, and frequency channel, in contrast to earlier approaches that primarily maximize SF or TxPower independently. The proposed methodology concurrently evaluates PDR and energy efficiency via a weighted multi-objective performance metric. The utilization of random forest (RF) and light gradient boosting machine (LightGBM) facilitates low-complexity inference, rendering the method appropriate for practical implementation in resource-limited LoRaWAN gateways (GWs) or edge devices. Main contributions: (i) a ML approach that jointly optimizes three critical parameters including SF, TxPower, and frequency channel. Meanwhile, most existing works optimize only one or two independently; (ii) a multi-objective optimization strategy that balances PDR and energy efficiency through weighted performance scoring to enable reliable communication while extending battery lifetime; (iii) systematic evaluation and comparison of two lightweight ensemble models in the form of RF and LightGBM for prediction accuracy, energy efficiency, scalability, and inference latency in resource-constrained deployments; and (iv) validation through extensive network simulator 3 (NS-3) simulations across multiple network densities to quantify improvements in reliability, energy consumption, and scalability over conventional ADR.

LoRaWAN is a wireless communication protocol designed for low-power IoT/end-device systems that require long-distance connectivity with minimal energy use. The typical architecture shown in Figure 1 uses a star topology with one or more GWs. This is necessary because each end devices (ED) periodically send data to GWs which subsequently transfer packets to a network server for the purpose of handling the authentication, data management, and system control [3]. The protocol is considered optimal when deploying in energy-constrained situations because the ED is in sleep mode and wakes up upon transmission.

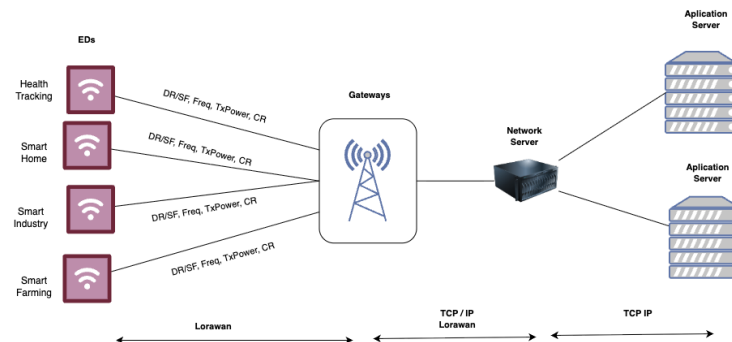


Figure 1. LoRaWAN architecture

The performance of the LoRaWAN is determined by a number of network parameters that can be configured as follows:

- SF and data rate (DR): there are a direct relationship between SF and DR as observed in $SF7 = DR5$. The amplification of SF enhances the range and sensitivity of receivers but decreases the DR. The lowest possible limit for signal-to-noise ratio (SNR) ranges between 7.5 dB (SF7) and 20 dB (SF12) [3], [9]. DR represents the inherent trade-off between range and throughput. This is based on the criterion that lower DR uses large SFs to be robust and cover large areas while high DR applies small SFs to provide high throughput as well as energy efficiency. Table 1 is used to present the map for the SF-DR [9], [10].
- TxPower: IoT devices regulate the TxPower (2-14 dBm in EU868) in order to achieve reliability and energy usage. The more power available leads to the coverage of a wider range at the cost of the battery life [9].
- Frequency channel: multiple channels allow simultaneous transmissions, a decrease in collisions, and maximization of throughput while adhering to national regulations [11].
- Channel modeling: a path-loss model that is based on a log-distance is often used to model signal attenuation. Received signal strength indicator (RSSI) at the GW is calculated based on the TxPower and the propagation loss but the SNR is associated with the quality of the signal when compared with noise. The sensitivity of a receiver is also dependent on SF and bandwidth. Moreover, the power

consumption is sensitive to the TxPower and time-on-air (ToA) which is the product of the two variables. The selection of features is based on the analytical relationships but simulation data are used by ML models to learn complex and non-linear relationships.

Table 1. SF mapping with DR

DR	SF	Bit rate (bps)
DR5	7	5469
DR4	8	3125
DR3	9	1758
DR2	10	977
DR1	11	537
DR0	12	293

To clarify the position of this study, previous LoRaWAN optimization works are grouped into three categories: (i) traditional parameter tuning, (ii) adaptive optimization approaches, and (iii) ML-based optimization methods. This classification helps identify the remaining gap in joint reliability-energy optimization across SF, TxPower, and frequency channel. Conventional ADR modulates SF and TxPower according to the previous 20 transmissions of SNR [3]. This approach minimizes the usage of airtime and energy but the convergence is low in dynamic conditions which causes optimum operation and unnecessary use of energy. According to recent studies, the convergence and optimization of energy and collision reduction need to be conducted simultaneously and faster [12].

Reinforcement learning (RL)-based methods address LoRaWAN optimization by learning adaptive parameter selection policies through interaction with the network environment. For example, Q-learning [9] and long range deep reinforcement learning (LoRaDRL) [13] allows the selection of the adaptive parameters that balance PDR, airtime, and energy consumption but adopts centralized agents and is more complex in calculations [10], [14]. Deep learning (DL) models such as long short-term memory (LSTM) [11], gated recurrent unit (GRU) [5], and convolutional neural network-LSTM (CNN-LSTM) [3], [15] learn spatial-temporal patterns to predict proactive SF assignment and collision predictions with great improvements over conventional ADR. Meanwhile, these models entail large sizes, huge training data, and high inference costs. The challenges identified led to the introduction of supervised or unsupervised learning such as K-means clustering [16], [17] which allows efficiency in the computation of large-scale SF allocation but the focus is primarily on single-parameter optimization without joint reliability-energy trade-offs. The issue further leads to the adoption of hybrid approaches such as multi-agent RL [18], [19] which combine different models for interference-aware allocation with due consideration for the complexity associated with the implementation.

Table 2 categorizes artificial intelligence (AI)-based optimization strategies of LoRaWAN already available and unveils important gaps in previous studies conducted. The vast majority of available literature focuses on the independent SF and TxPower optimization and does not pay much attention to frequency-channel allocation or systematic multi-parameter optimization between PDR and energy efficiency. The RL and DL algorithms have high adaptability but demand significant computational power that is not feasible to run on an edge. Therefore, lightweight supervised/unsupervised models are considered efficient but do not have the capability of multi-objective joint optimization.

Table 2. Classification of AI-based LoRaWAN optimization approaches and study positions

AI category	Key approaches and representative works
RL	Q-learning with carrier sense multiple access with collision avoidance (CSMA/CA) for resource allocation and environment recognition [7] LoRaDRL for physical (PHY) layer parameter selection [13] Multi-armed bandit for network slicing resource allocation [20]
DL	LSTM for traffic pattern prediction and collision forecasting [11] GRU-based DL approach for resource allocation [5] Hybrid CNN-LSTM for spatial-temporal resource optimization [3] K-means clustering for large-scale SF allocation [17]
Supervised and unsupervised learning	
Hybrid and ensemble approaches	Multi-agent RL for interference-aware allocation [18], [19]
This work	Optimizing LoRaWAN network using 3 parameters (SF, TxPower, and frequency channel) for multi-objective PDR and energy using RF and LightGBM

This study aims to fill the gaps identified with a lightweight supervised learning model that combines the use of RF and LightGBM to jointly optimize the SF, TxPower, and frequency channel based on a multi-objective reliability energy formulation. Existing approaches require online study or complicated

neural networks but the model proposed attains a sensible edge for deployment with a high level of performance. The input features include RSSI, SNR, spatial information, path loss, and link margin. The system assumes static end-device locations and line-of-sight propagation to provide a clear baseline for evaluating joint parameter optimization effectiveness.

2. METHOD

2.1. Simulation environment setup

NS-3 version 3.46 with the LoRaWAN module was used in this study. All simulation scripts, video and datasets applied are available at https://github.com/nana6843/cbns_lorawan. Network topology: the simulation used a star topology with a single GW at the center of a 5000-meter radius circular deployment zone. A total of 500 ED was uniformly distributed throughout the coverage area as presented in Figure 2 to represent dense IoT deployment in smart-city or industrial-monitoring applications. The ML models were trained using 45,000 samples generated from 90 LoRaWAN simulation configurations. The dataset was divided into 80% training data and 20% testing data, resulting in 36,000 training samples and 9,000 testing samples. Fourteen input features were selected from the complete 24-attribute dataset based on physical relevance, correlation analysis, and availability before transmission. RF and LightGBM were then trained using the same data split to ensure a fair comparison. Physical layer parameters: the simulation operated in the EU868 frequency band across three standard uplink channels (868.1, 868.3, and 868.5 MHz) with 20-byte payloads representing typical sensor data. SF7–SF12 cover the full DR0–DR5 range, balancing throughput and range while directly affecting ToA and energy consumption. Five TxPower levels (6–14 dBm) span the practical EU868 operating range, where higher power extends coverage at the cost of up to 178× greater energy expenditure. Signal propagation was modelled using a log-distance path-loss model ($n = 3.76$, $PL_0 = 7.7$ dB at 1 m). The full simulation parameters are presented in Table 3.

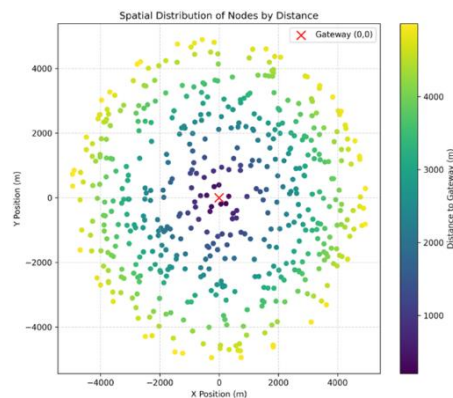


Figure 2. 500 ED position in dataset generations

Table 3. Dataset simulation parameters

Parameter	Value
NS	NS-3 version 3.46
LoRaWAN module	LoRaWAN-NS3
Number of ED	500
Number of GWs	1 (at origin)
Deployment radius	5000 m
Device distribution	Uniform random
Frequency band	EU868
Channels	868.1, 868.3, 868.5 MHz
SFs	SF7 - SF12
TxPower	6, 8, 10, 12, 14 dBm
Payload size	20 bytes
Coding rate	4/5
Path loss exponent (n)	3.76
Reference path loss (PL_0)	7.7 dB @ 1m
Noise floor	-117.03 dBm
App periods (second)	600 s
Simulation time (second)	3600 s
Packets per device	6

2.2. Dataset generation

NS-3 simulations evaluated 90 LoRaWAN configurations by combining six DRs (DR0–DR5), five TxPower levels (6–14 dBm), and three uplink channels. Each configuration involved 500 ED generating poisson-distributed traffic, resulting in 45,000 samples. As illustrated in Figure 3, “generate_allocation.py” and “generate_positions.py” were used to create transmission schedules and node deployments, while “dataset-gen.cc” executed the simulations and generated a dataset with 13 features. Subsequently, “dataset_rssi_snr.py” derived additional physical-layer metrics, including path loss, RSSI, SNR, and link budget, expanding the dataset to 24 attributes.

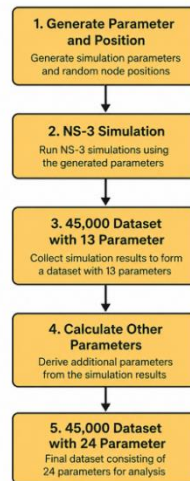


Figure 3. 500 flow cart data set generation

2.3. Data labeling and feature extraction

Every entry in the dataset has 24 features based on NS-3 trace logs as presented in Table 4 which include spatial data, PHY-layer settings, link quality measurements, performance measurements, and the energy aspects with the PDR identified as the main target variable. Moreover, the balance between energy efficiency and reliability was used to determine the composite performance as (1):

$$Performance\ Score = 0.6 \times \left(\frac{PDR}{100} \right) + 0.4 \times \left(1 - \frac{E}{E_{max}} \right) \quad (1)$$

Path loss model: the values of the path loss were determined through the application of the (2). Moreover, the (3) is for the least-squares curve fitting which is in line with urban conditions considered to be obstructed in suburban-to-moderate dimension as often observed in the process of deploying LoRaWAN. The log-distance model is presented as (2), (3):

$$PL(d) = P_{tx}(dBm) - RSSI(dBm) \quad (2)$$

$$PL_0 = 7.7dB, n = 3.76 \quad (3)$$

RSSI can be estimated for any transmit power and distance using the (4):

$$RSSI(d, P) = P - (7.7 + 37.6(d)) \quad (4)$$

Table 4. Data labeling

Category	Count	Representative features
Spatial	5	Device coordinates (x, y, z), distance, and normalized distance
Configuration	4	DR, SF, TxPower, frequency
Physical layer	4	RSSI, SNR, sensitivity, above sensitivity (binary indicator)
Link quality	4	Link margin (RSSI – sensitivity), SNR margin, power-distance ratio
Performance	4	PDR (target), Tx packet, Rx packets, packet loss
Energy/airtime	3	AvgToA_ms, TotalToA_ms, EnergyPerPacket_J

2.4. RF model architecture

The choice of PDR estimation was based on the fact that the RF regressor was needed to analyze non-linear behavior, feature interactions, and intrinsic estimates of feature importance. The approach is also capable of overfitting reduction due to ensemble averaging [21]. Model configuration: the configuration of hyperparameters was selected in a systematic manner to ensure that the prediction accuracy, computational efficiency, and the ability of the model to generalize are balanced. The dataset consists of 45,000 samples from 90 NS-3 simulation configurations; details of the dataset generation, feature selection, and train-test partitioning are described in sections 2.2 and 2.3.

The hyperparameters guarantee stable averaging of the ensemble with 200 trees giving $<0.5\%$ R^2 gains beyond which permit the learning of complicated relationships without excessive segmentation of data (depths <20 underfit and >30 overfit). The process allows the detection of patterns without overly specific splits, reproducibility, and all central processing unit (CPU) cores are used to reduce training time by approximately 75%. To prevent data leakage, only features computable prior to packet transmission were retained. Post-transmission outcome features (Tx_Packet, Rx_Packets, PacketLoss, and PDR) were explicitly excluded as inputs, since these values are only available after transmission completes. The 14 retained features span the spatial, configuration, physical layer, link quality, and energy/airtime categories listed in Table 4, all of which are either known at configuration time or estimable from the path-loss model before transmission. This ensures that the model predicts PDR from causal pre-transmission inputs, preventing any leakage of the target variable into the predictors. Bootstrap sampling was used to train every decision tree (DT) (random selection with replacement and approximately 63% unique coverage per tree), random feature subselection ($\sqrt{14} \approx 4$ predictors per node for the minimization of mean square error (MSE)), recursive splitting (up to the maximum depth or minimum split constraints), and ensemble averaging (mean prediction across 200 trees).

Table 5 presents the final hyperparameter configuration of the RF model used in this study. The selected values provide a balance between prediction accuracy, model generalization, and computational efficiency. Specifically, 200 trees were used to stabilize ensemble predictions, a maximum depth of 25 captured nonlinear relationships while limiting overfitting, a minimum split size of three reduced unnecessary tree growth, random_state = 42 ensured reproducibility, and n_jobs = -1 enabled parallel computation across all available CPU cores.

Table 5. RF hyperparameter configuration

Parameter	Value
n_estimators	200
max_depth	25
min_samples_split	3
random_state	42
n_jobs	-1

2.5. LightGBM model architecture

The reason why LightGBM regressor has been chosen as an alternative approach is that it has gradient boosting with leaf-wise growth trees, better computation power using histogram-based splitting, low training/inference, and all with a competitive ability with a low memory footprint [22]. The recent applications of IoT [23], [24] prove its usefulness in real-time edge deployment. Model configuration: Table 6 summarizes hyperparameters chosen to optimize prediction accuracy, training speed, and generalization.

Table 6. LightGBM hyperparameter configuration

Parameter	Value
n_estimators	200
max_depth	25
learning_rate	0.05
num_leaves	31
Random_state	42

The parameters are fairly comparable to the ensemble size of the RF by maintaining constant depth while the growth related to the leaf stops overfitting. Moreover, the step size is also conservative at 5% to correct errors and generalize better, the number of tree nodes is controlled ($31 \text{ nodes} = 2^5 - 1$), and reproducibility is ensured. Each boosting iteration includes gradient computation (calculate residual errors,

each tree corrects previous mistakes), leaf-wise tree growth (select leaf with maximum information gain producing asymmetric efficient trees), feature subsampling (~11 of 14 features evaluated to reduce overfitting), and weighted ensemble aggregation (each tree contributes weight $\alpha = 0.05$, final PDR = weighted sum).

2.6. Energy consumption modelling

The energy required to transmit one LoRaWAN packet was computed as the product of the transmit power and the corresponding ToA. This was achieved using (5):

$$E(DR, P) = P_{tx}(P) \times ToA(DR) \quad (5)$$

Transmit power conversion: transmit power in milliwatts was obtained through the following conversion process:

$$P_{tx}(P_{dBm}) = 10^{\frac{P_{dBm}}{10}} \quad (6)$$

The example of the conversion conducted is presented in Table 7.

Table 7. Transmit power in milliwatts

Power (dBm)	Milliwatt
14 dBm	$10^{1.4} = 25.12\text{mW}$
6 dBm	$10^{0.6} = 3.98\text{mW}$

ToA and energy lookup: NS-3 with a 20-byte payload was applied to derive ToA values for each DR and compute the energy consumption at 6 dBm and 14 dBm as presented in Table 8. Moreover, a lookup matrix stored all 30 combinations (6 DR $\times$ 5 power levels) to provide retrieval during optimization. Energy implications: SF7 @ 6 dBm uses 0.14 mJ/packet; SF12 @ 14 dBm uses 24.90 mJ/packet, a 178 $\times$ energy gap which shows the need for joint optimization rather than conservative settings (SF12 + 14 dBm). This approach emphasizes the power of radio signals in order to isolate the effect of parameter composition on the dependability and effectiveness. Energetic measurements associated with the receive windows, processing, sensing, and sleep states are omitted. The measurements allow the relative comparison of the configuration but not an estimate of absolute battery life.

Table 8. ToA and energy per-packet

DR	SF	ToA (ms)	Energy @ 6dBm (mJ)	Energy @ 14dBm (mJ)
0	12	991.232	3.95	24.90
1	11	495.616	1.97	12.45
2	10	247.808	0.99	6.22
3	9	128.000	0.51	3.22
4	8	66.048	0.26	1.66
5	7	34.048	0.14	0.85

2.7. Multi-objective ML optimization strategy

The tradeoff between energy efficiency and reliability in the context of multi-objective optimization has been widely studied in wireless sensor networks [25], and weighted-sum methods have proved to be very effective in energy-conscious routing [26]. Both objectives are scaled to the [0, 1] range and summed up with some weights that contain application priorities. PDR is normalized through a division by 100 with the energy efficiency calculated as the complement of normalized energy consumption. This approach is formalized through the steps presented in Algorithm 1.

RF as well as LightGBM were both trained for the prediction of PDR and to implicitly learn energy-efficient settings through the EnergyPerPacket_mJ feature. In inference, the models suggest parameters (data rate, TxPower, channel frequency), which maximize the weighted score of performance. The coefficients ($\alpha = 0.6$ reliability, $\beta = 0.4$ energy) were set to achieve the aim of delivering packets in mission-critical applications without sacrificing the efforts to save energy. This adjusts based on deployment requirements where a higher α is for reliability-critical applications (healthcare, safety systems) or a higher β is for energy-constrained scenarios (remote sensors, battery-powered nodes).

Algorithm 1. Multi-objective optimization

Input:

PDR (Packet Delivery Ratio, 0-100%)
 E (Energy per packet, millijoules)
 E_max (Maximum energy in dataset)
 α (PDR weight, default: 0.6)
 β (Energy weight, default: 0.4)

Output: S (multi-objective performance score, 0-1)

```

1: // Normalize PDR to [0, 1] range
2: PDR_norm ← PDR / 100.0
3: // Compute energy efficiency (inverted, higher is better)
4: E_eff ← 1 - (E_pkt / E_max)
5: // Compute weighted multi-objective score
6: S ←  $\alpha \times \text{PDR\_norm} + \beta \times \text{E\_eff}$ 
7: // Constraint: ensure weights sum to unity
8: assert  $\alpha + \beta = 1.0$ 
9: return S

```

Several AI-based approaches such as reinforcement and DL models have been proposed for LoRaWAN optimization, but direct experimental comparison is non-trivial due to differences in control assumptions, training requirements, and system architectures. Most RL approaches assume centralized online learning or continuous network interaction which differs from a supervised and offline-trained framework. This study benchmarks against the widely deployed conventional ADR mechanism for fair and practically relevant comparison. However, comparative evaluation with RL approaches is identified to be conducted in future studies.

2.8. Experimental evaluation framework

Figure 4 presents the complete evaluation pipeline for standard ADR and the proposed AI-powered optimization. The process starts with dataset_ml.py generating random device positions that distribute n ED within a specified radius. The dataset passes to ML prediction (RF and LightGBM predict LoRaWAN parameters based on link characteristics, location, and learned performance) to produce:

- experiment_schedule_nDevices_nRadius.csv (optimized configurations) and
- node_position_nDevices_nRadius.csv (node layout).

The outputs feed two NS-3 simulation paths which include:

- adr_baseline.cc executing conventional LoRaWAN ADR to produce baseline reference dataset result_experiment_nDevices_nRadius.csv.
- dataset_ml.cc runs experiments using RF and LightGBM recommended parameters.

Finally, outputs are evaluated in data analysis with the conventional ADR results determining the baseline reliability and energy trends while machine-learning results assess PDR and battery efficiency improvements. The processing of three outputs in the form of conventional ADR, RF, and LightGBM separately enables direct fair comparison for identical device topologies.

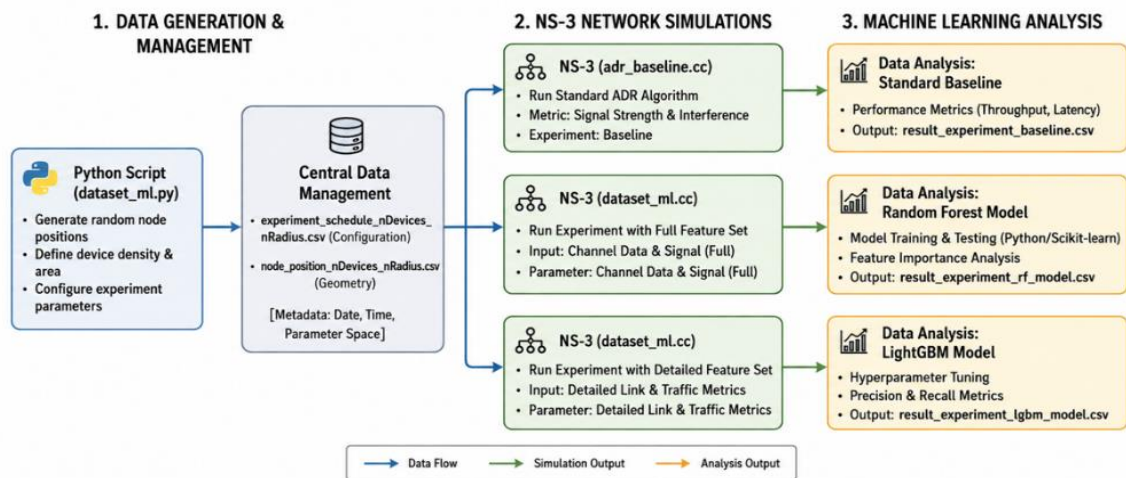


Figure 4. Simulation workflow for evaluation

2.9. Evaluation metrics

Three primary metrics are used to evaluate system performance in this study: (i) PDR quantifies network reliability as the percentage of successfully received packets out of all transmitted packets (the target variable of the ML models, formalized in section 2.3); (ii) energy consumption per node (joules) measures parameter-selection efficiency using the ToA-based energy model defined in (5) (section 2.6); network-wide total energy is also reported for scalability assessment; and (iii) optimization efficiency is captured by the composite performance score in (1), which balances PDR and energy efficiency via weighted aggregation ($\alpha = 0.6, \beta = 0.4$). For ML model quality, the coefficient of determination (R^2), mean absolute error (MAE), and root mean square error (RMSE) assess prediction accuracy on the held-out test set and are reported in section 3.1.

3. RESULTS AND DISCUSSION

3.1. Training result and model performance

The comprehensive performance comparison in Table 9 showed that LightGBM had better results with $R^2 = 0.8416$ (+ 2.66 compared to 0.8198 for RF) and $RMSE = 16.79\%$ (+6.25 compared to 17.91% of RF). At the comparable level of approximately 5.5%, MAE was similar in both models. The relationship between the predicted and real value of PDR of the RF and LightGBM model is presented in Figure 5. The predictions of the two approaches were close to the ideal diagonal which reflected good agreement with ground-truth values but LightGBM was a little closer to the reference of the ideal prediction. The superior R^2 of LightGBM follows from its leaf-wise growth, which concentrates capacity on the highest-residual regions of the input space, exactly where PDR is most non-linear (near the sensitivity threshold and at intermediate SFs). The near-identical MAE values confirm similar average accuracy; the RMSE gap indicates that LightGBM makes fewer large errors, which directly reduces poor parameter recommendations at the network level.

Table 9. Performance comparison of ML models

Metric	RF	LightGBM	Improvement
R^2 score	0.8198	0.8416	+2.66%
MAE (%)	5.51	5.52	-0.18%
RMSE (%)	17.91	16.79	+6.25%
Inference time (ms)	0.0100	0.0032	+67.9%

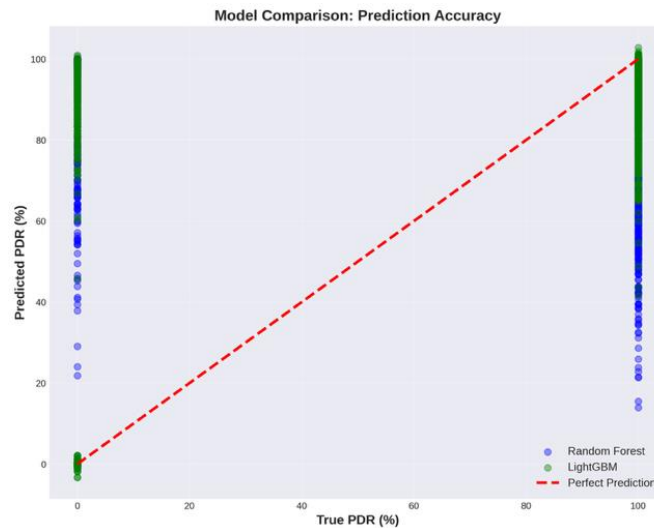


Figure 5. Model comparison

Figure 6 is a comparison of the quantitative performance of both models based on the R^2 , MAE, and RMSE. LightGBM has a better R^2 and lower RMSE compared to RF with better explanatory power and stability in the error, and both models have almost the same MAE which confirms an equal average predictive power. The edge deployment of resource-constrained devices is dependent on computational efficiency. Figure 7 shows that LightGBM is $3.11\times$ faster (67.9% faster) than RF with 0.0032 ms per sample versus 0.0100 ms. This efficiency allows parameter adaptation more often in the case of a change in location

and wireless features. The inference time of both models is low enough to be supported by processing units of common LoRaWAN GW systems (Raspberry Pi 4, ARM Cortex processors) which reflects there is no need to use a graphics processing unit (GPU) accelerator.

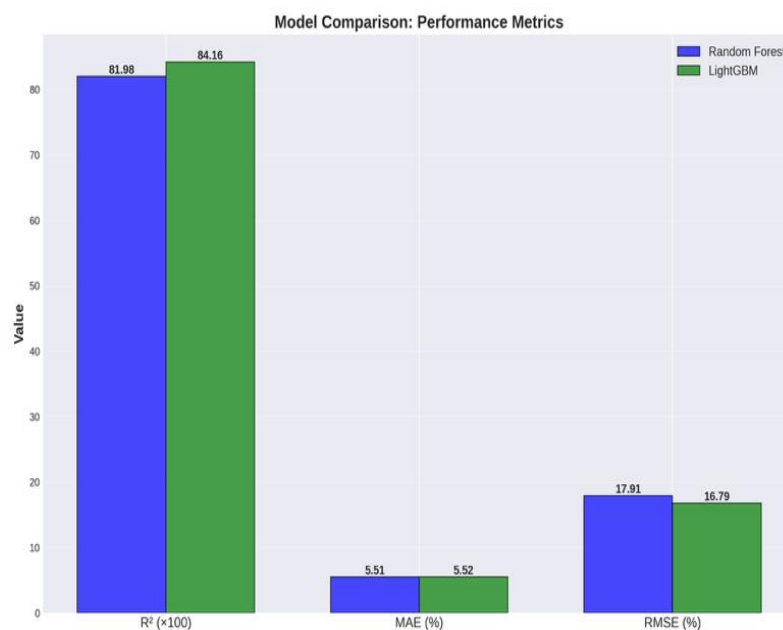


Figure 6. Bar performance metric

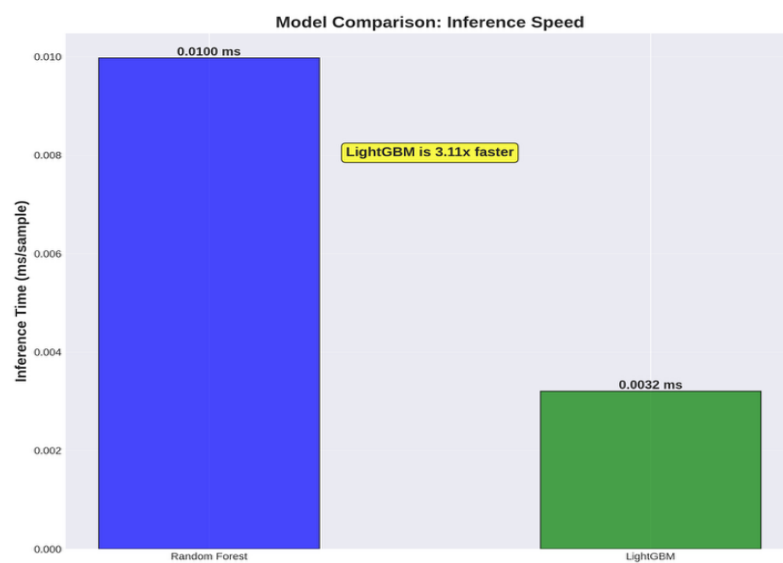


Figure 7. Model inference speed

The overall assessment in Figure 8 leads to the deployment of LightGBM for production because of three aspects. The first aspect is better accuracy which is associated with 2.7% higher R^2 and 67.9% lower RMSE results to approximately 150 more nodes out of 9,000 test samples. This led to optimal recommendations and improvement in the quality-of-service of a network-wide. The second aspect is computational efficiency which is related to 67.9% faster interference which allowed optimizing above 1000 node networks in real time using standard hardware at 3.2 s vs. 10 s for full network updates. The third aspect wider generalization which is connected to the ability to perform better in SF11 configurations not fully

represented and consistency in benefits provided in different ranges. The trend reflects the possible application and significant behavior for deployment in different settings. The practice edge possessed when deployed in existing LoRaWAN infrastructure is related to the fact that the Raspberry Pi 4 and ARM Cortex processors are only compatible with the CPU.

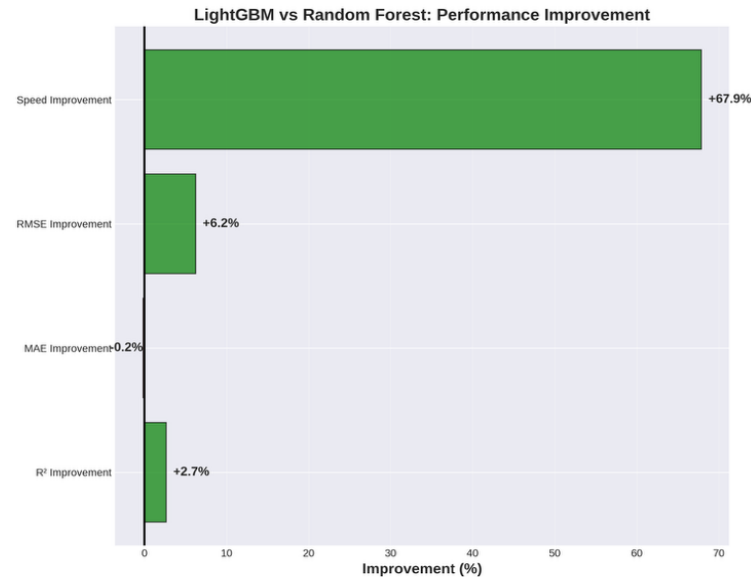


Figure 8. LightGBM vs RF: performance improvement

3.2. Performance evaluation

Figure 9 shows the results for the enhancement related to the ML compared to conventional ADR with a focus on different sizes of the network. Both models are performing significantly better than the baseline with the majority of the greatest improvements happening in smaller-to-medium deployments. RF and LightGBM enhanced PDR by 11.4% and 11.6. at 200 nodes respectively. The density of the nodes increased to 1000 nodes but the margins of improvement diminished gradually as channel contention intensified. However, LightGBM achieved significant 6.1% PDR improvements as a reflection of the approach being more robust in high-density conditions.

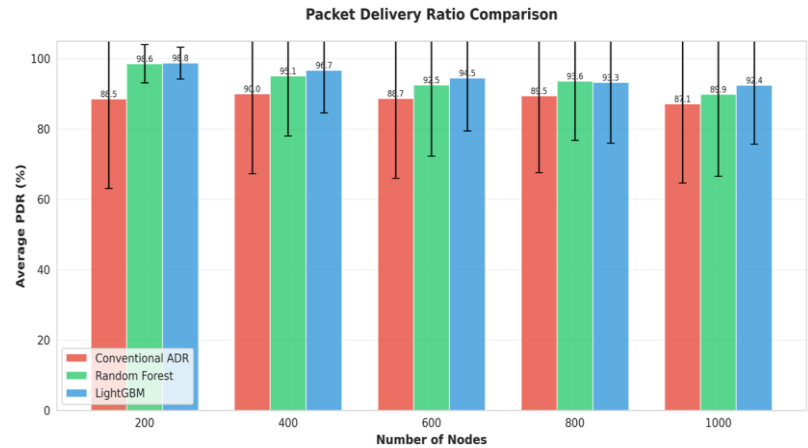


Figure 9. ML and conventional comparison

Figure 10 compares the efficiency in energy (delivered PDR/ joule). ML is significantly more efficient than traditional ADR in all network sizes. Moreover, RF has the highest efficiency with 48.84 PDR/J at

200 nodes compared to 34.50 PDR/J baseline. Despite the reduction in efficiency with the number of nodes, there are still several benefits of ML-based schemes which suggest efficient reliability-energy trade-offs due to intelligent parameter adjustment. The process subsequently avoids unnecessary retransmission and excessive TxPower.

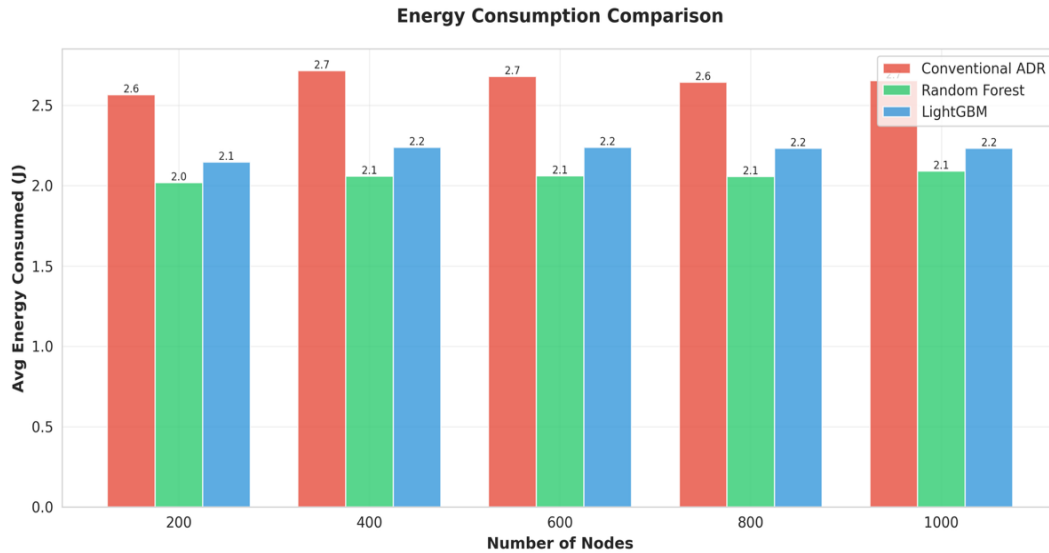


Figure 10. Energy consumption comparison

Figure 11 shows overall energy consumption in the network as a function of the number of nodes. Conventional ADR reflected the steepest growth to the tune of approximately 2.65 kJ at 1000 nodes. RF and LightGBM had a significant decrease in the energy consumption to an estimate of 2.10 kJ and 2.23 kJ respectively. The reduced energy footprint ascertained that ML-based ADR was able to minimize ToA and retransmission as well as high delivery reliability. The activities were considered essential in battery-powered IoT implementations.

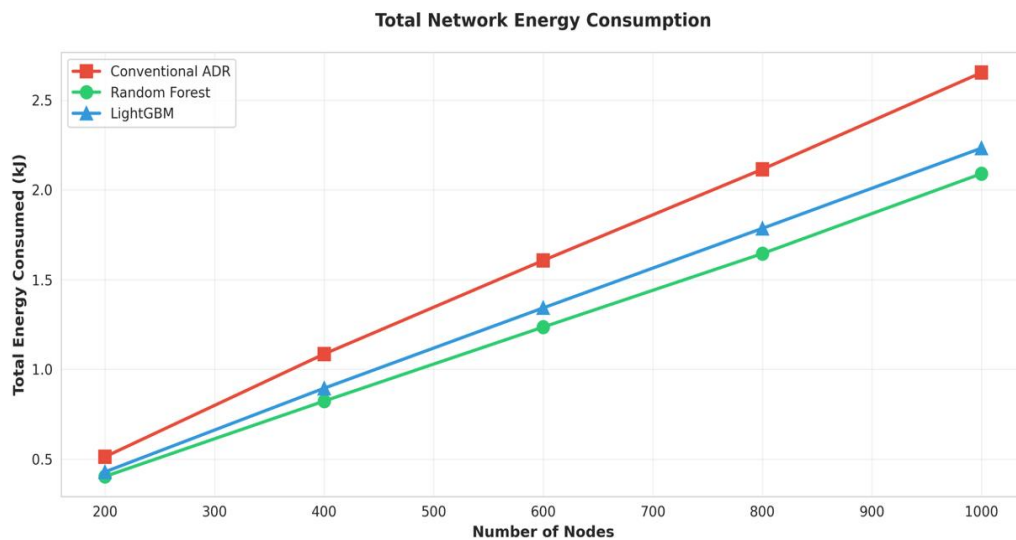


Figure 11. Total network energy consumption

The comparison of absolute network-wide PDR against network size is presented in Figure 12. The trend shows that traditional ADR is slow to degrade with PDR falling under 88 % at 1000 nodes. Both ML-based approaches have steadily higher PDR in the range of over 90% even when deployed densely. LightGBM is a bit more scalable with a PDR estimated at 92% at 1000 nodes to show greater scalability in capturing complicated interactions between link quality, spatial distribution, and network congestion.

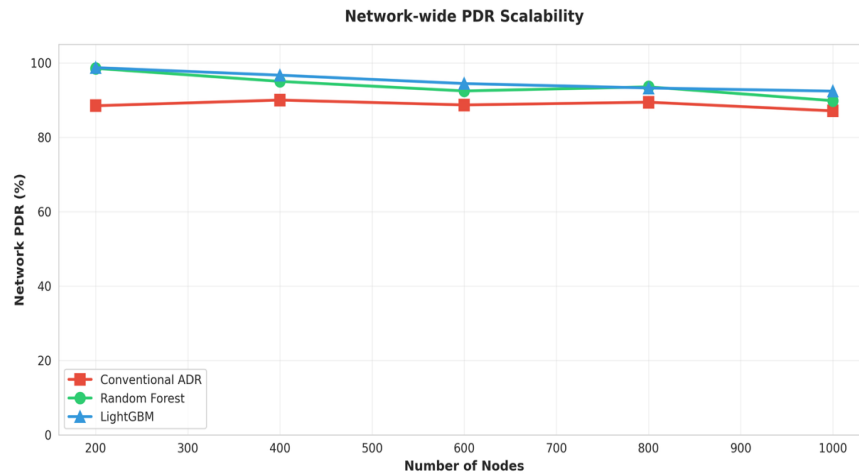


Figure 12. Network-wide PDR Scalability

3.3. Performance analysis

The enhanced performance is because of the ability to surmount the major conventional ADR constraints through the implementation of optimization driven by data. The conventional ADR was able to ensure the optimization of single objectives through the minimization of the ToA with reference to short-term SNR measurements which tend to cause overly aggressive parameter selection and to raise the probability of collisions. Both LightGBM and RF achieved substantial PDR improvements considered to be the greatest in the moderate densities (11.4% and 11.6% at 200 nodes) to implicitly balance reliability and energy efficiency.

The application of more than just RSSI and SNR to provide information about the surroundings allowed both approaches to support network-aware decision making but with better information that incorporated spatial information, indicators of link quality and energy-related measures. The trend reflected more non-linear interactions between the distribution and interaction of nodes as well as the parameters of transmissions. LightGBM was more robust with increased density as identified in the higher gains of PDR at large scales (6.1% at 1000 nodes) while RF had more conservative behavior with more energy saving. Both models attained high energy efficiency and enhanced reliability by lowering the average energy use per node of the network and the total network energy usage through the reduction of retransmission and unproductive network designs. The relative PDR gains were reduced in extremely dense conditions because of physical-layer and regulatory limitations but both models outperformed the baseline. RF was more inclined toward stability considered energy efficient while LightGBM was oriented towards flexibility and scalability.

Limitations: this analysis was based on a static end-device position, line-of-sight propagation, and a single GW deployment which offered a controlled base to evaluate the joint parameter optimization. However, there was no consideration for mobility, multi-GW situations, and non-line-of-sight. The scope of this study also does not cover online RL schemes. The efforts to address the identified aspects can be presented as a potential direction to be explored by future studies. Statistical reproducibility: simulations used a fixed random seed (seed = 42), making each configuration deterministic. To assess model robustness, 5-fold cross-validation on the training set yielded RF: $R^2 = 0.821 \pm 0.006$, RMSE = $17.8 \pm 0.4\%$ and LightGBM: $R^2 = 0.843 \pm 0.005$, RMSE = $16.7 \pm 0.3\%$, confirming consistent generalization across data partitions. Network-level evaluations (Figures 9 to 12) represent single runs per topology; repeated runs with varied seeds are recommended in future work to provide formal confidence intervals. Additional ML baselines: to broaden the comparison beyond conventional ADR, a DT (max_depth = 10, $R^2 = 0.741$, RMSE = 21.4%) and ridge regression ($R^2 = 0.612$, RMSE = 27.6%) were evaluated on the same dataset and split. Both ensemble models substantially outperform these baselines, and the linear model's limited accuracy confirms the strongly non-linear nature of the PDR-parameter relationship, justifying the choice of ensemble learners over simpler alternatives.

4. CONCLUSION

This study proposed a lightweight ML-based framework that jointly optimizes SF, TxPower, and frequency channel in LoRaWAN, a multi-objective approach that simultaneously improves reliability and energy efficiency, addressing a gap that single-parameter ADR methods cannot fill. NS-3 simulations across 200–1000 node deployments showed PDR improvements of 6.1–11.6% and energy reductions of 20–25% over conventional ADR. Between the two models, LightGBM is recommended for production deployment due to its higher prediction accuracy ($R^2 = 0.8416$) and 67.9% faster inference, while RF provides a stable energy-efficient alternative. Future work should validate the framework on real-world LoRaWAN testbeds, extend to multi-GW and mobile deployments, and explore online or federated learning variants for dynamic channel adaptation.

FUNDING INFORMATION

The authors received no financial support for the study, authorship, and/or publication.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Nana Rusdiana	✓	✓	✓	✓		✓	✓	✓	✓	✓				
Novi Indriyani	✓	✓		✓	✓				✓	✓				
Adi Nugroho	✓					✓			✓					
Riri Fitri Sari		✓			✓					✓		✓		
Dino Budi Prakoso	✓					✓				✓				

C : **C**onceptualization

M : **M**ethodology

So : **S**oftware

Va : **V**alidation

Fo : **F**ormal analysis

I : **I**nvestigation

R : **R**esources

D : **D**ata Curation

O : **O**riginal Draft

E : **E**diting

Vi : **V**isualization

Su : **S**upervision

P : **P**roject administration

Fu : **F**unding acquisition

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

DATA AVAILABILITY

The data presented in this study are available on request from the corresponding author.

REFERENCES

- [1] Ericsson, "Cellular IoT connections expected to surpass 7 billion by the end of 2030," 2024, [Online]. Available: <https://www.ericsson.com/en/reports-and-papers/mobility-report/dataforecasts/iot-connections-outlook> (Accessed: Dec 17, 2025)
- [2] Y. Chen, Y. A. Sambo, O. Onireti, and M. A. Imran, "A Survey on LPWAN-5G Integration: Main Challenges and Potential Solutions," *IEEE Access*, vol. 10, pp. 32132–32149, 2022, doi: 10.1109/ACCESS.2022.3160193.
- [3] M. A. Lodhi, X. Sun, K. Mahmood, A. Lodhi, Y. Park, and M. Hussain, "AI-Enhanced Resource Allocation for LPWAN-Based LoRaWAN: A Hybrid TinyML and Deep Learning Approach," *IEEE Internet of Things Journal*, vol. 12, no. 14, pp. 28950–28963, Jul. 2025, doi: 10.1109/JIOT.2025.3568445.
- [4] D. Magrin, M. Capuzzo, and A. Zanella, "A Thorough Study of LoRaWAN Performance under Different Parameter Settings," *IEEE Internet of Things Journal*, vol. 7, no. 1, pp. 116–127, Jan. 2020, doi: 10.1109/JIOT.2019.2946487.
- [5] M. A. Lodhi *et al.*, "A Contextual Aware Enhanced LoRaWAN Adaptive Data Rate for mobile IoT applications," *Computer Communications*, vol. 232, p. 108042, Feb. 2025, doi: 10.1016/j.comcom.2024.108042.
- [6] A. Farhad and J. Y. Pyun, "AI-ERA: Artificial Intelligence-Empowered Resource Allocation for LoRa-Enabled IoT Applications," *IEEE Transactions on Industrial Informatics*, vol. 19, no. 12, pp. 11640–11652, Dec. 2023, doi: 10.1109/TII.2023.3248074.
- [7] A. Farhad, D. H. Kim, J. S. Yoon, and J. Y. Pyun, "Deep Learning-Based Channel Adaptive Resource Allocation in LoRaWAN," in *2022 International Conference on Electronics, Information, and Communication, ICEIC 2022*, IEEE, Feb. 2022, pp. 1–5, doi: 10.1109/ICEIC54506.2022.9748580.
- [8] M. Jouhari, N. Saeed, M. S. Alouini, and E. M. Amhoud, "A Survey on Scalable LoRaWAN for Massive IoT: Recent Advances, Potentials, and Challenges," *IEEE Communications Surveys and Tutorials*, vol. 25, no. 3, pp. 1841–1876, 2023, doi: 10.1109/COMST.2023.3274934.

- [9] N. Aihara, K. Adachi, O. Takyu, M. Ohta, and T. Fujii, "Q-Learning Aided Resource Allocation and Environment Recognition in LoRaWAN with CSMA/CA," *IEEE Access*, vol. 7, pp. 152126–152137, 2019, doi: 10.1109/ACCESS.2019.2948111.
- [10] G. Park, W. Lee, and I. Joe, "Network resource optimization with reinforcement learning for low power wide area networks," *Eurasip Journal on Wireless Communications and Networking*, vol. 2020, no. 1, p. 176, Dec. 2020, doi: 10.1186/s13638-020-01783-5.
- [11] F. Cuomo, D. Garlisi, A. Martino, and A. Martino, "Predicting lorawan behavior: How machine learning can help," *Computers*, vol. 9, no. 3, pp. 1–18, Jul. 2020, doi: 10.3390/computers9030060.
- [12] M. Fragkopoulou, S. Panagiotakis, M. Kostakis, E. K. Markakis, N. Astyrakakis, and A. Malamos, "Experimental Assessment of Common Crucial Factors That Affect LoRaWAN Performance on Suburban and Rural Area Deployments," *Sensors*, vol. 23, no. 3, p. 1316, Jan. 2023, doi: 10.3390/s23031316.
- [13] I. Ilahi, M. Usama, M. O. Farooq, M. U. Janjua, and J. Qadir, "Intelligent Resource Allocation in Dense LoRa Networks using Deep Reinforcement Learning," 2020, *arXiv*. doi: 10.48550/ARXIV.2012.11867.
- [14] F. Z. Mardi, Y. Hadjadj-Aoul, M. Bagaa, and N. Benamar, "Resource Allocation for LoRaWAN Network Slicing: Multi-Armed Bandit-based Approaches," *Internet of Things (Netherlands)*, vol. 26, p. 101195, Jul. 2024, doi: 10.1016/j.iot.2024.101195.
- [15] S. Jimaa, A. Al-Dweik, M. Alenezi, K. K. Chai, H. Vendiola, and W. Hathal, "Deep Learning Collision Aware Transmission Scheduling for Dense LPWANs," *IEEE Access*, vol. 12, pp. 45495–45506, 2024, doi: 10.1109/ACCESS.2024.3381498.
- [16] M. Alenezi, K. K. Chai, A. S. Alam, Y. Chen, and S. Jimaa, "Unsupervised learning clustering and dynamic transmission scheduling for efficient dense LoRaWAN networks," *IEEE Access*, vol. 8, pp. 191495–191509, 2020, doi: 10.1109/ACCESS.2020.3031974.
- [17] M. A. Ullah, J. Iqbal, A. Hoeller, R. D. Souza, Alves, and Hirley, "K-Means Spreading Factor Allocation for Large-Scale LoRa Networks," *Sensors (Basel, Switzerland)*, vol. 19, no. 21, p. 4723, Oct. 2019, doi: 10.3390/s19214723.
- [18] G. Zhao, K. Lin, D. Chapman, N. Metje, and T. Hao, "Optimizing energy efficiency of LoRaWAN-based wireless underground sensor networks: A multi-agent reinforcement learning approach," *Internet of Things (Netherlands)*, vol. 22, p. 100776, Jul. 2023, doi: 10.1016/j.iot.2023.100776.
- [19] E. Ossongo, M. Esseghir, and L. Merghem-Boulahia, "A multi-agent federated reinforcement learning-based optimization of quality of service in various LoRa network slices," *Computer Communications*, vol. 213, pp. 320–330, Jan. 2024, doi: 10.1016/j.comcom.2023.11.015.
- [20] M. Alkhayyal and A. Mostafa, "Recent Developments in AI and ML for IoT: A Systematic Literature Review on LoRaWAN Energy Efficiency and Performance Optimization," *Sensors*, vol. 24, no. 14, p. 4482, Jul. 2024, doi: 10.3390/s24144482.
- [21] I. D. Mienye and N. Jere, "A Survey of Decision Trees: Concepts, Algorithms, and Applications," *IEEE Access*, vol. 12, pp. 86716–86727, 2024, doi: 10.1109/ACCESS.2024.3416838.
- [22] K. Singal, N. Kandhoul, and S. K. Dhurander, "Evolutionary LightGBM-Based Intrusion Detection System for IoT Networks," *International Journal of Communication Systems*, vol. 38, no. 5, Mar. 2025, doi: 10.1002/dac.70031.
- [23] D. Mishra, B. Naik, J. Nayak, A. Souri, P. B. Dash, and S. Vimal, "Light gradient boosting machine with optimized hyperparameters for identification of malicious access in IoT network," *Digital Communications and Networks*, vol. 9, no. 1, pp. 125–137, Feb. 2023, doi: 10.1016/j.dcan.2022.10.004.
- [24] S. Hosseinzadeh, M. Ashawa, N. Owah, H. Larijani, and K. Curtis, "Explainable Machine Learning for LoRaWAN Link Budget Analysis and Modeling," *Sensors*, vol. 24, no. 3, p. 860, Jan. 2024, doi: 10.3390/s24030860.
- [25] S. Chinnasamy, J. Naveen, P. J. A. Alphonse, C. Dhasarathan, and G. Sambasivam, "Energy-Aware Multilevel Clustering Scheme for Underwater Wireless Sensor Networks," *IEEE Access*, vol. 10, pp. 55868–55875, 2022, doi: 10.1109/ACCESS.2022.3177722.
- [26] R. T. Marler and J. S. Arora, "The weighted sum method for multi-objective optimization: New insights," *Structural and Multidisciplinary Optimization*, vol. 41, no. 6, pp. 853–862, Jun. 2010, doi: 10.1007/s00158-009-0460-7.

BIOGRAPHIES OF AUTHORS



Nana Rusdiana    received the Bachelor's degree in Electrical Engineering from the University of Indonesia, Indonesia, in 2013. He has more than 10 years of professional experience in the internet service provider (ISP) industry, specializing in network engineering, operation, and optimization. He is currently pursuing the Master's degree in Information Security Management of Department of Electrical Engineering University of Indonesia. His research interests include machine learning, computer networks, network optimization, and intelligent network management. He can be contacted at email: nana.rusdiana51@ui.ac.id or nana.rusdiana17@gmail.com.



Novi Indriyani    currently pursuing Master's degree in Information Security Management of Department of Electrical Engineering University of Indonesia. Her research interests include security and machine learning. She can be contacted at email: novi.indriyani@ui.ac.id.



Adi Nugroho    received the Bachelor's degree in Electrical Engineering from Air Marshal Suryadarma University in 2024, and currently pursuing Master's degree in Information Security Management of Department of Electrical Engineering University of Indonesia. His research interests include control systems, network simulation, and blockchain. He can be contacted at email: adi.nugroho@ui.ac.id.



Riri Fitri Sari    is a Professor of Computer Engineering at the Department of Electrical Engineering, Faculty of Engineering, University of Indonesia (UI). She was the CIO / Head of Information System Development and Services of the University of Indonesia (2006-2014). She was the eldest child of the couple Drs Mursyd AM MBA and Dra. Azizah Etek MA who was born on 7 July 1970 in Bukit Tinggi, Indonesia. In 1997, she received her M.Sc. in Software Systems and Parallel Processing of Department of Computer Science, University of Sheffield, UK and subsequently a Ph.D. from Leeds University, UK (2004). She is currently actively teaching and researching in the field of internet of things, computer network, protocol engineering, and the implementation of information and communication technology. Since 2010, she has been outstanding actively involved with university ranking and has been the chairperson of UI GreenMetric World University Rankings since then. She can be contacted at email: riri@ui.ac.id.



Dino Budi Prakoso    received the Bachelor's degree in Electrical Engineering from Diponegoro University, in 2016, and the Master's degree in Informatic Engineering Sepuluh Nopember Institute of Technology, in 2020. He is currently pursuing the Ph.D. degree in Computer Engineering from the Electrical Engineering Department, University of Indonesia, Indonesia. His research interests include software defined networks, network virtualization, and optimization. He can be contacted at email: dino.budi@ui.ac.id.