

LoRa-LQE: An Energy-Efficient Link Quality Estimation Mechanism for Mobile LoRaWAN Networks

Geraldo A. Sarmiento Neto   [Universidade Federal do Piauí | geraldosarmiento@ufpi.edu.br]

Thiago A. Ribeiro da Silva  [Instituto Federal do Maranhão | thiago.allisson@ufpi.edu.br]

Fernando J. V. Santos  [Instituto Federal do Maranhão | fernando.vieira@ufpi.edu.br]

Pedro F. Ferreira de Abreu  [Universidade Federal do Piauí | pedroffda@ufpi.edu.br]

Luis H. de O. Mendes  [Universidade Federal do Piauí | luishenriqueom@ufpi.edu.br]

J. Valdemir dos Reis Jr  [Universidade Federal do Piauí | valdemirreis@ufpi.edu.br]

 Universidade Federal do Piauí, Campus Universitário Ministro Petrônio Portella, Ininga, Teresina, PI, 64049-550, Brazil.

Received: 23 December 2025 • **Accepted:** 10 July 2026 • **Published:** 31 July 2026

Abstract Smart infrastructures increasingly rely on mobile IoT applications, such as fleet tracking and asset monitoring, which require robust long-range connectivity. Although LoRaWAN is a prominent technology in this context, its standard Adaptive Data Rate mechanism suffers significant performance degradation in dynamic scenarios due to delayed and outdated channel estimation. To address this limitation, this work proposes LoRa-LQE, a link quality estimation scheme specifically designed for mobile environments. By jointly leveraging SNR history and channel stability indicators, the proposed heuristic optimizes the allocation of Spreading Factor and Transmission Power on a per-device basis. Extensive simulations conducted in NS-3 evaluate the approach under varying node densities and mobility profiles. Across the evaluated scenarios, the results show that LoRa-LQE outperforms state-of-the-art solutions, achieving packet delivery ratio improvements of up to 12.1% and energy efficiency improvements of up to 46.6% compared to Mobile ADR, the second-best mobility-aware scheme considered, thereby enabling more reliable communication for mobile IoT assets.

Keywords: Adaptive Data Rate, Energy Efficiency, LoRaWAN, Link Quality Estimation, Mobility

1 Introduction

The Internet of Things (IoT) has become a fundamental paradigm for large-scale connectivity in smart applications. In this context, wireless sensor networks enable services such as smart lighting, automated waste management, and remote utility metering [Acosta-Garcia *et al.*, 2025]. These applications require years-long battery operation while transmitting small payloads over long distances, often in electromagnetically dense environments. Consequently, network solutions must prioritize energy efficiency and coverage over high data rates, which challenges traditional cellular and short-range technologies [Loubany *et al.*, 2023].

To meet these requirements, Low Power Wide Area Networks (LPWANs) have emerged as the preferred connectivity solution, with Long Range Wide Area Network (LoRaWAN) standing out due to its use of unlicensed sub-gigahertz bands and a robust physical layer based on Long Range (LoRa) modulation [Almuhaya *et al.*, 2022]. Its star-of-stars topology enables large-scale deployments with low infrastructure cost and simplified network management [Soy, 2023]. The openness and flexibility of the standard have driven widespread adoption, enabling interoperable deployments across vendors and regions.

The performance of the LoRa physical layer is determined by the configuration of four primary transmission parameters: Spreading Factor (SF), Bandwidth (BW), Coding Rate (CR),

and Transmission Power (TP). The SF determines the number of chips per symbol, where higher values increase receiver sensitivity and range but significantly extend the airtime. BW influences the data rate, with wider bandwidths enabling faster transmissions at the expense of noise robustness. The CR introduces redundancy into the data stream, improving resilience to interference. Finally, TP controls the signal power, directly affecting both the link budget and the energy consumption of the end device [Jouhari *et al.*, 2023].

To dynamically manage these parameters, the LoRaWAN standard employs the Adaptive Data Rate (ADR) mechanism, which enables the network server to adjust the data rate of end devices based on uplink quality estimation. By analyzing the Signal-to-Noise Ratio (SNR) of received packets, the server determines the minimum resources required to sustain a reliable link [Soy, 2023]. Under favorable channel conditions, ADR reduces the SF and TP, lowering airtime and energy consumption. This approach is particularly effective for stationary devices with relatively stable channel conditions.

While many IoT applications increasingly rely on the continuous monitoring of mobile assets [Li *et al.*, 2024; Ren *et al.*, 2022], the effectiveness of the standard ADR mechanism degrades significantly in environments characterized by node mobility. As devices move, the radio channel undergoes rapid fluctuations caused by shadowing and multipath fading, which often render channel state information outdated before the network can react [Acosta-Garcia *et al.*, 2025]. Since the standard

ADR relies on SNR averaging over a packet history, it frequently fails to adapt quickly to these variations, leading to temporary connectivity losses or overly conservative parameter selections. As a result, unnecessary energy consumption and reduced network capacity are commonly observed [Farhad *et al.*, 2020].

To mitigate these limitations, precise Link Quality Estimation (LQE) is essential for evaluating wireless channel states through physical layer indicators, providing the statistical foundation for informed adaptation. This study proposes LoRa-LQE, a heuristic mechanism that utilizes LQE-derived intelligence to dynamically adjust transmission parameters. By integrating normalized SNR indicators, instability penalties, and efficiency boosts, the scheme simultaneously captures instantaneous fluctuations and short-to-medium-term channel stability trends [Iqbal *et al.*, 2023]. Consequently, this approach reduces energy waste by identifying reliable transmission windows and preventing the premature high-power configurations often found in standard backoff strategies [Liu *et al.*, 2021].

The performance of the proposed scheme was evaluated through extensive simulations using the Network Simulator 3 (NS-3), covering scenarios with network densities ranging from 200 to 1000 end devices. The analysis included various mobility profiles, categorized into specific speed classes to represent diverse urban dynamics [Teymuri *et al.*, 2023]. The results demonstrate that LoRa-LQE outperforms standard approaches across the evaluated scenarios, achieving up to 12.1% higher Packet Delivery Ratio (PDR) and up to 46.6% greater energy efficiency than Mobile ADR, the second-best mobility-aware scheme analyzed. These findings validate the ability of the mechanism to maintain robust connectivity without compromising the battery life of mobile sensor nodes.

The main contributions of this paper can be summarized as follows:

- Development of LoRa-LQE, a heuristic-based link quality estimation mechanism tailored for dynamic LoRaWAN environments;
- Formulation of an analytical-heuristic model that integrates SNR history and channel stability metrics to optimize the data rate by adjusting transmission parameters;
- Implementation and extensive performance evaluation of the proposed scheme within the NS-3 simulator under varying node densities and speed classes;
- Demonstration of significant improvements in PDR and energy consumption compared to the standard ADR and other state-of-the-art schemes designed specifically for mobile environments.

The remainder of this paper is organized as follows. Section 2 reviews related work on link quality estimation and parameter adaptation. Section 3 provides an overview of LoRaWAN and its ADR mechanism. Section 4 details the proposed scheme. Section 5 describes the methodology and simulation setup. Section 6 discusses the results, analyzing the performance of LoRa-LQE. Finally, Section 7 presents the conclusions and directions for future research.

2 Related Work

Link Quality Estimation is an essential component in the design of reliable wireless networks, influencing decisions ranging from routing to transmission parameter adjustment. In the literature, LQE techniques generally fall into data-driven approaches that utilize machine learning, and model-based or heuristic approaches focusing on specific network constraints and parameter optimization.

A significant portion of recent research employs machine learning to classify link states with high accuracy. Pham *et al.* [2021] proposed a method based on Support Vector Machines (SVM) specifically for LoRa networks. By collecting Received Signal Strength Indication (RSSI), SNR, and Packet Reception Rate (PRR) data, the authors developed a model to classify the connection level between nodes, demonstrating that learning-based classifiers can effectively capture the non-linear relationship between hardware metrics and packet delivery. Similarly, in the broader context of Wireless Sensor Networks (WSN), Liu *et al.* [2021] introduced an LQE method based on an Improved Weighted Extreme Learning Machine (IWELM). Their approach utilizes Particle Swarm Optimization (PSO) to adjust hidden nodes and weights, enhancing prediction accuracy compared to traditional Extreme Learning Machine (ELM). In more complex propagation environments, Tesfaw *et al.* [2023] applied deep learning, specifically Gated Recurrent Units (GRU), to estimate link quality in Reconfigurable Intelligent Surface (RIS)-assisted Unmanned Aerial Vehicle (UAV) communications. Their work highlights the necessity of handling complex channel characteristics in three-dimensional mobile deployments. While these data-driven methods achieve high accuracy, they typically demand substantial computational resources and extensive training datasets, which may not always be feasible for resource-constrained end devices operating in dynamic environments requiring immediate adaptation.

To address the resource limitations inherent in low-power networks, lightweight and metric-fusion estimators have been explored. Liu *et al.* [2020] proposed a fluctuation-insensitive estimator for WSNs that fuses SNR and Link Quality Indication (LQI) using weighted Euclidean distance after Kalman filtering. This approach aims to balance accuracy and overhead without relying on computationally intensive learning models. In the domain of Wireless Body Sensor Networks (WBSN), Iqbal *et al.* [2023] developed a multi-hop Quality of Service (QoS)-aware routing protocol that integrates a Predicting Link Quality Estimation (PLQE) mechanism. The scheme uses an Exponential Weighted Moving Average (EWMA) to smooth short-term variations in the packet forwarding ratio and capture link reliability more accurately. Combined with its probabilistic prediction model, this approach handles the fast changes caused by body movements, ensuring more reliable data transmission than methods based only on instantaneous measurements.

The interaction between LQE and network configuration, particularly routing [Isyaku *et al.*, 2021], is another focal point in the state-of-the-art. Shahid *et al.* [2024] introduced the Link-Quality based Energy-Efficient Routing (LQEER) protocol for WSN within IoT applications. By integrating

link quality estimation with residual energy metrics into a route cost function, LQEER selects reliable forwarder nodes to minimize packet loss and latency while extending network lifetime. In mobile ad-hoc networks (MANET), Dev *et al.* [2025] investigated energy-efficient routing algorithms, emphasizing the impact of node mobility and residual energy on network lifetime.

In the specific domain of LoRa transmission parameter adjustment, Moreira *et al.* [2025] proposed LoRaBB, an algorithm tailored for optimizing communication in challenging environments like the Amazon rainforest. The authors designed a binary search methodology to balance the trade-off between convergence speed and the selection of robust physical parameters, aiming to mitigate signal attenuation caused by dense foliage. By characterizing the deployment scenario, LoRaBB adapts settings to ensure link reliability. However, this approach relies on a search-based heuristic that iteratively probes configurations, contrasting with analytical models that seek to directly derive optimal settings from channel state information, a distinction that becomes significant in highly mobile scenarios.

Existing solutions predominantly focus either on static WSN scenarios or employ complex learning models that may induce latency in real-time adaptation. Furthermore, while mobility is addressed in MANET and UAV studies, analytical models specifically tailored for LoRaWAN link quality estimation under mobility constraints remain underexplored. The reliance on heavy data processing or heuristic search limits the applicability of such methods in scenarios where energy efficiency and rapid reaction to channel changes are paramount. The proposed LoRa-LQE scheme aims to bridge this gap by providing a solution for dynamic transmission parameter adjustment, specifically designed to handle the challenges imposed by mobility in LoRaWAN environments. Table 1 summarizes the related works referenced in this section, highlighting the methodologies applied and the metrics adopted in each study.

3 Background

This section reviews the main LoRaWAN transmission parameters and their impact on communication performance. It then describes the standard ADR mechanism, which dynamically adjusts transmission parameters to improve network efficiency and reduce energy consumption. Finally, it discusses the limitations of ADR under dynamic channel conditions, motivating the proposed LoRa-LQE scheme.

3.1 LoRaWAN and Transmission Parameters

LoRaWAN is a low-power wide-area networking protocol designed for long-range communications and large-scale IoT deployments. It operates over the LoRa physical layer, which employs Chirp Spread Spectrum (CSS) modulation to provide robust communication under low signal conditions [Kufakunesu *et al.*, 2024]. Owing to its long communication range and low energy consumption, LoRaWAN has become widely adopted in applications such as environmental monitoring, smart agriculture, and utility

metering. These characteristics make efficient resource management essential to maintain scalability and network performance.

LoRaWAN networks follow a star-of-stars architecture in which End Devices (EDs) communicate with one or more Gateways (GWs) that forward packets to a centralized Network Server (NS) through an IP-based backhaul [Jouhari *et al.*, 2023]. Since multiple GWs may receive the same uplink (UL), the NS performs packet deduplication and manages functions such as ADR. Communication is primarily UL-driven, with downlink (DL) opportunities occurring after each UL [LoRa Alliance, 2020]. Consequently, network performance depends strongly on the transmission parameters, which directly affect data rate, coverage, reliability, energy consumption, and channel occupancy.

Among these parameters, the Spreading Factor, Bandwidth, Coding Rate, and Transmission Power have the greatest influence on link performance [Benkahla *et al.*, 2021]. The SF, typically ranging from SF7 to SF12, determines the number of chips per symbol, where each symbol carries SF bits of information. Higher SF values improve receiver sensitivity and communication range but reduce the achievable data rate. BW controls the transmission speed, while CR introduces Forward Error Correction to improve robustness at the cost of additional packet overhead [Soy, 2023].

The combined effect of SF, BW, and CR determines the nominal physical bit rate (R_b) [Kufakunesu *et al.*, 2024], given by

$$R_b = SF \cdot \frac{BW}{2^{SF}} \cdot CR \quad (1)$$

where the data rate increases with BW and CR and decreases as SF grows. TP provides an additional degree of control over link reliability and energy expenditure, allowing devices to balance connectivity and battery consumption [Farhad *et al.*, 2020].

These parameters also determine the packet Time-on-Air (ToA), defined as the total time a packet occupies the wireless channel [Acosta-Garcia *et al.*, 2025]. ToA directly influences energy consumption and compliance with duty-cycle restrictions. The symbol duration (T_{sym}) is calculated as

$$T_{sym} = \frac{2^{SF}}{BW} \quad (2)$$

and the total ToA is obtained from the sum of the preamble and payload transmission times [Teymuri *et al.*, 2023]:

$$\begin{aligned} ToA &= T_{preamble} + T_{payload} \\ &= (N_{preamble} + 4.25) \cdot T_{sym} + N_{payload} \cdot T_{sym} \end{aligned} \quad (3)$$

As a result, increasing SF extends communication range but substantially increases channel occupancy, whereas larger bandwidths reduce transmission time. Consequently, selecting appropriate transmission parameters requires balancing coverage, reliability, throughput, and energy efficiency.

Table 1. Summary of Related Works on Link Quality Estimation and Optimization

Reference	Methodology	Metrics Used
[Liu et al., 2020]	Exponential Weighted Kalman Filtering	SNR, LQI, PRR
[Isyaku et al., 2021]	Multi-constraint Optimization	Link Latency, Delivery Ratio, Criticality
[Liu et al., 2021]	Weighted Extreme Learning Machine with PSO	LQI, RSSI, SNR, Grade
[Pham et al., 2021]	Support Vector Machine	LQI, RSSI, SNR, PRR
[Wang et al., 2022]	Genetic Algorithm-optimized Backpropagation Neural Network	RSSI, SNR, PRR
[Iqbal et al., 2023]	EWMA + Beta Probability Density Function	Link Reliability, Delay
[Tesfaw et al., 2023]	Gated Recurrent Units	Channel Data, RIS Phase Shift
[Shahid et al., 2024]	Fuzzy Logic	Link Quality, Residual Energy
[Dev et al., 2025]	Composite Cost Function	Link Stability, Residual Energy, Drain Rate
[Moreira et al., 2025]	Binary Search Heuristic	SNR, RSSI, PDR
This work	Heuristic LQE-based adaptive scheme	SNR, SF, TP

3.2 Adaptive Data Rate

To ensure scalability and efficiency in large-scale deployments, LoRaWAN employs the ADR mechanism, a dynamic procedure for optimizing transmission parameters. Although manual configuration of Spreading Factors and Transmission Power can guarantee connectivity, it typically results in suboptimal resource usage, as devices may transmit with unnecessarily high power or low data rates to accommodate worst-case channel conditions. ADR mitigates this inefficiency by automatically adjusting these parameters according to the actual link budget required [Kufakunesu et al., 2024]. By lowering the SF and reducing TP for devices experiencing favorable link conditions, the network minimizes ToA and significantly extends the battery lifetime of EDs.

As illustrated in the flowchart in Figure 1, the control logic for this optimization resides primarily within the NS. The server continuously monitors the SNR of uplink packets received from each ED. By averaging these values over a history of frames defined by the packet history size M , which is set to 20 in the standard ADR algorithm [de Jesus et al., 2021], the server estimates the link margin available beyond the demodulation floor. If sufficient margin is available, the server determines an appropriate combination of a lower SF and reduced TP while maintaining reliable connectivity. These parameters are then encapsulated in a MAC command and transmitted to the device through a downlink message [Anwar et al., 2021]. This centralized approach keeps the complexity of channel estimation at the server side, allowing EDs to remain lightweight and energy-efficient.

On the ED side, the protocol includes a failsafe mechanism known as ADR Backoff to handle sudden changes in link quality or the loss of downlink connectivity. If a device operating under ADR control transmits a sequence of uplink frames without receiving any downlink acknowledgment, it infers that the current data rate is too high or the transmit power is too low to reach the gateway. To recover connectivity, the device autonomously triggers the backoff algorithm. This process involves stepwise increments of the SF—effectively lowering the data rate to gain receiver sensitivity—and resetting the TP to its maximum level [Louati et al., 2024]. This autonomous reaction allows the node to re-establish communication with the network without human intervention, even after significant environmental changes.

It is important to note that the effectiveness of the standard

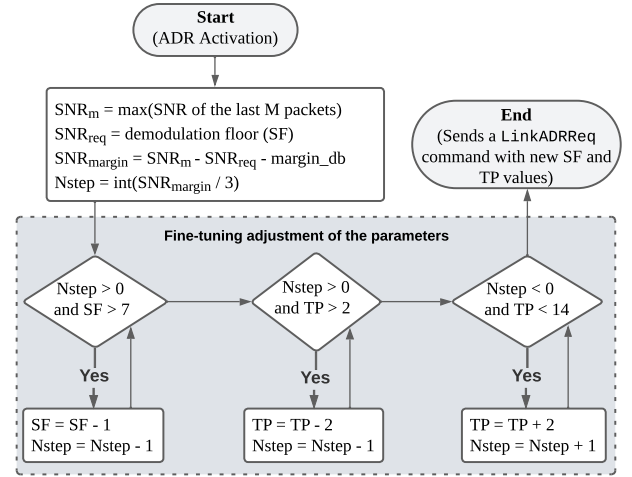


Figure 1. Flowchart demonstrating the Standard ADR logic.

ADR depends heavily on the stationarity of the ED. The algorithm assumes that the radio channel remains relatively stable over the convergence period. Consequently, while highly effective for static sensor nodes, the mechanism struggles with mobile devices where the radio environment changes faster than the algorithm can adapt. In such mobility scenarios, the latency inherent in the response loop often leads to link instability [Farhad et al., 2020]. Therefore, addressing these dynamic conditions requires the implementation of alternative management strategies that go beyond the limitations of the default protocol, offering more robust adaptation mechanisms suited for non-stationary nodes.

4 Proposed Scheme

This section details the architectural components of the proposed solution, focusing on the integration of link assessment techniques within the network operation. The discussion begins with the conceptual foundations of Link Quality Estimation and advances to the specific implementation of the LoRa-LQE algorithm.

4.1 Link Quality Estimation

LQE is a fundamental mechanism for inferring the operational state of wireless links by statistically modeling physical metrics—such as RSSI, SNR, and LQI. By applying

filtering and mapping techniques, LQE estimates the packet delivery probability efficiently, ensuring robustness against channel fluctuations [Liu *et al.*, 2020]. The accuracy of these estimators relies fundamentally on capturing both large-scale phenomena, such as path loss and shadowing, and fast variations caused by multipath fading and interference, which are characteristic of dynamic and heterogeneous environments [Dev *et al.*, 2025].

LQE techniques are critical in networks subject to high temporal and spatial variability. In WSNs, precise link assessment drives the selection of stable routes, reduces retransmissions, and optimizes energy consumption, particularly in urban, industrial, or environmental monitoring scenarios where propagation conditions are highly irregular. In advanced architectures, such as RISs-assisted communications or links mediated by UAVs, LQE is essential due to node mobility and the requirement for frequent adjustments in phase, direction, or trajectory [Tesfaw *et al.*, 2023]. In these systems, LQE functions as a fundamental enabler for adaptive decisions that enhance link reliability and spectral efficiency.

In the IoT context, device energy constraints and the need for long-range, low-power operation intensify the demand for precise estimation. Technologies like LoRaWAN depend on reliable channel state assessments to tune critical transmission parameters—including SF and TP—which are vital for the effective operation of ADR mechanisms. Consequently, data-driven approaches and machine learning techniques have become central to ensuring scalability, predictability, and robustness in dense or massive IoT deployments [Wang *et al.*, 2022].

Structuring LQE to support the adaptive tuning of physical parameters enables more efficient resource management strategies. The following subsection introduces the LoRa-LQE scheme, which exploits this potential by integrating LQE-derived intelligence directly into the LoRaWAN link configuration and optimization process.

4.2 LoRa-LQE

LoRaWAN networks exhibit significant susceptibility to environmental variations as they operate under low data rates, high receiver sensitivity, and severe energy constraints [Bonilla *et al.*, 2023]. Consequently, LQE techniques are indispensable for complementing resource management mechanisms by enabling continuous assessment of propagation conditions and the effective degree of link degradation or stability. Unlike generic literature approaches, applying LQE in LoRaWAN demands specific consideration of IoT characteristics, such as sparse sampling, high spatial variability, and the direct impact of SF and transmission power on energy consumption and network scalability [Moreira *et al.*, 2025].

To address these challenges, this study proposes LoRa-LQE, a mechanism based on a heuristic model expressed through analytical functions that utilize indicators derived from SNR time series to dynamically adjust LoRaWAN transmission parameters in varying environments. The scheme incorporates both the historical SNR values from recently received packets and their exponential moving

average [Iqbal *et al.*, 2023], enabling the simultaneous capture of instantaneous fluctuations and channel stability trends.

Furthermore, the design integrates normalized metrics [Liu *et al.*, 2021], logistic penalties [Zhang and Qiu, 2023], and adaptive boosts [Al-Qassas and Qasaimeh, 2024], while employing cooldown windows to prevent excessive adjustment oscillations [Quattrocchi *et al.*, 2024]. Collectively, these elements endow LoRa-LQE with a decision-making process capable of balancing reliability, coverage, and energy efficiency.

The proposed mechanism is outlined in Algorithm 1. In Line 1, the *SNRlist* variable stores the historical SNR values obtained from the last M packets received by the device *nodeID*. In Line 2, the *EMA* variable stores the Exponential Moving Average computed based on the stored SNR samples. Lines 3–5 determine whether the node is in a cooldown period. If the cooldown counter is positive, the system decrements it and immediately returns the current SF and TP values, thereby preventing premature adjustments.

The function presented in Line 6 computes the LQE using the SNR window, the updated *EMA*, and the current SF and TP settings. Lines 7–14 apply the LQE rules to adjust transmission parameters by comparing the estimated quality against thresholds that define distinct operational zones. Specifically, the *LqeThsDown* threshold establishes an excellence level above which robustness is relaxed to prioritize energy efficiency. Conversely, the *LqeThsUp* threshold determines the minimum critical level below which robustness is increased to restore reliability. Intermediate values maintain the current configuration, creating a stability zone that mitigates unnecessary oscillations.

Lines 15–16 check for parameter modifications. If changes occur, the cooldown counter resets to the *cooldownSteps* value, ensuring that further adaptations occur only after a minimum number of uplink transmissions. Finally, in Line 17, the algorithm returns the updated SF and TP values, which are subsequently transmitted to the device via the *LinkADRRReq* MAC command, similar to the standard ADR mechanism.

Algorithm 2 details the *ComputeLQE* function of the main algorithm, focusing on the core application of the LQE technique. In Line 1, the variable SNR_{req} stores the minimum SNR threshold required for demodulation at the current spreading factor, obtained via the *DemodulationFloor* function. In Line 2, the variable SNR_{marg} captures the available channel margin, defined as the difference between the current *EMA* value and the demodulation threshold. This metric indicates the extent to which the link operates above or below the minimum requirement for reliable reception. In Line 3, the variables $Marg_{norm}$ and EMA_{norm} store the normalized versions of the SNR_{marg} and the *EMA*, respectively, ensuring that heterogeneous attributes remain comparable within the final LQE composition. In Line 4, the *SD* variable holds the standard deviation of the SNR history, quantifying the degree of channel instability. In Line 5, SD_{pen} represents the penalty factor associated with the standard deviation, calculated by the *GetSdPenalty* function. This factor reduces the LQE value whenever the channel exhibits significant instability.

In Line 6, the SF_{norm} variable stores the normalized

Algorithm 1: LoRa-LQE

```

Input:  $SF, TP, M, nodeID, cooldownSteps$ 
1  $SNRlist \leftarrow GetSNRHistory(M)$ 
2  $EMA \leftarrow ComputeEMA(SNRlist)$ 
   // If in cooldown, skip adaptation but return
   current values
3 if  $cooldownCnt[nodeID] > 0$  then
4    $cooldownCnt[nodeID] --$ 
5   return  $SF, TP$ 
6  $LQE \leftarrow ComputeLQE(SNRlist, EMA, SF, TP)$ 
7  $SF' \leftarrow SF$ 
8  $TP' \leftarrow TP$ 
9 if  $LQE > LqeThsDown$  then
10   $SF' \leftarrow SF' - 1$ 
11   $TP' \leftarrow TP' - 2$ 
12 else if  $LQE < LqeThsUp$  then
13   $SF' \leftarrow SF' + 1$ 
14   $TP' \leftarrow TP' + 2$ 
   // Apply cooldown after change
15 if  $(SF' \neq SF) \text{ or } (TP' \neq TP)$  then
16   $cooldownCnt[nodeID] \leftarrow cooldownSteps$ 
17 return  $SF', TP'$ 

```

version of the current SF. This facilitates the appropriate calculation of penalties, acknowledging that the impact of the SF is non-linear across its operational range. In Line 7, SF_{pen} contains the penalty term associated with the spreading factor, determined by the *GetSfPenalty* function. Since SF_{pen} grows with SF and is combined with a negative weight $w_3 = -0.3$ in Line 9, its net effect is to reduce the LQE score as the spreading factor increases, reflecting the associated costs of extended channel occupancy time and higher energy consumption.

Line 8 defines the $Boost_{ee}$ variable, which adds a supplementary gain to the LQE to prioritize configurations that favor energy efficiency. In Line 9, all normalized components, penalties, and the boost factor are linearly combined to derive the final LQE value. This composition employs weights w_1, w_2, w_3, w_4 , and w_5 to adjust the relative influence of the primary metrics, whereas the previously calculated increments are added directly as they represent additional characteristics of the link state. Finally, in Line 10, the algorithm returns the computed LQE value, which serves as the input for the adaptation mechanism to determine adjustments in transmission parameters.

Algorithm 2: ComputeLQE function

```

Input:  $SNRlist, EMA, SF, TP$ 
1  $SNR_{req} \leftarrow DemodulationFloor(SF)$ 
2  $SNR_{marg} \leftarrow EMA - SNR_{req}$ 
3  $Marg_{norm}, EMA_{norm} \leftarrow$ 
    $Normalize(SNR_{marg}, EMA)$ 
4  $SD \leftarrow GetStandardDeviation(SNRlist)$ 
5  $SD_{pen} \leftarrow GetSdPenalty(SD)$ 
6  $SF_{norm} \leftarrow (SF - SF_{min}) / (SF_{max} - SF_{min})$ 
7  $SF_{pen} \leftarrow GetSfPenalty(SF_{norm})$ 
8  $Boost_{ee} \leftarrow GetEnergyEffBoost(SF, TP)$ 
9  $LQE \leftarrow w_1 \cdot EMA_{norm} + w_2 \cdot Marg_{norm} + w_3 \cdot$ 
    $SF_{pen} + w_4 \cdot SD_{pen} + w_5 \cdot Boost_{ee}$ 
10 return  $LQE$ 

```

4.2.1 Exponential Moving Average

In Algorithm 1, the *ComputeEMA* function maintains a running channel-quality estimate by combining the most recent SNR measurement with the previously accumulated average, without storing the full packet history. The recurrence relation is defined by:

$$EMA_t = \alpha \cdot SNR_t + (1 - \alpha) \cdot EMA_{t-1}, \quad (4)$$

where SNR_t is the SNR of the most recently received packet, EMA_{t-1} is the previous estimate, and $\alpha \in (0, 1)$ is the smoothing factor. The influence of each past observation decays geometrically, with older measurements contributing progressively less to the current estimate [Farhad *et al.*, 2020], which enables a lightweight yet responsive tracking of channel variations without the memory overhead associated with fixed-size packet history windows.

The parameter α controls a trade-off between noise rejection and tracking speed. A low value such as $\alpha = 0.1$ suppresses short-term fluctuations but reacts slowly to channel degradation, risking persistent overestimation of link quality in mobile scenarios. A high value such as $\alpha = 0.7$ follows each measurement more closely but introduces oscillations under stable conditions, potentially triggering unnecessary adaptation commands [Liu *et al.*, 2020].

Figure 2 illustrates these behaviors using a synthetic SNR trace structured according to the speed classes presented in Section 5.3. The horizontal axis represents the packet index, where each unit corresponds to one uplink transmission at a 600 s interval.

Phase 1 spans packet indices from 0 to 29 and corresponds to a node speed of approximately 4 m/s. This phase represents a stable low-speed regime in which the node moves slowly enough that the SNR remains roughly constant near 8 dB across consecutive uplinks.

Phase 2 spans packet indices from 30 to 49 and corresponds to node speeds between approximately 12 m/s and 16 m/s. This phase represents the high-speed degradation scenario at the upper end of the evaluation range. At these speeds, the node experiences substantially different propagation conditions between successive transmissions, exposing each uplink to distinct path-loss and shadowing conditions. As a result, the figure exhibits an abrupt SNR reduction of approximately 7 dB.

Phase 3 spans packet indices from 50 to 79 and corresponds to node speeds between approximately 4 m/s and 8 m/s. This phase simulates a partial recovery as the node decelerates and moves into a less obstructed area. The recovery is intentionally incomplete, reflecting realistic deployments in which distance and obstacles prevent full restoration of link quality within a short observation window.

The adopted value $\alpha = 0.3$ provides a middle ground well suited to this context: it assigns approximately 66% of the weight to the three most recent packets and reduces any individual past sample below 5% influence after nine packets—a horizon consistent with the rapid channel variations observed in the evaluated mobility scenarios.

As visible in Figure 2, $\alpha = 0.3$ converges toward the true channel level within a few packets at the onset of Phase 2, while remaining smoother than $\alpha = 0.7$

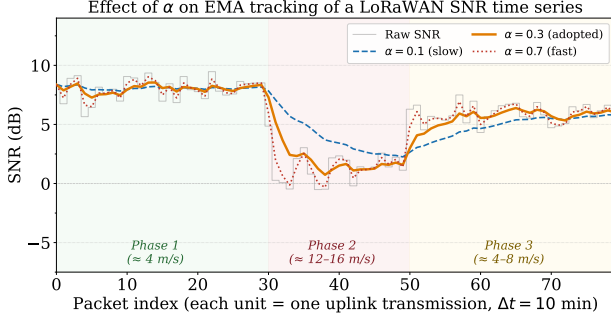


Figure 2. Effect of α on EMA tracking of a synthetic SNR series.

during the oscillatory recovery of Phase 3—a behavior consistent with LoRaWAN-specific recommendations for mobile deployments [Iqbal et al., 2023].

4.2.2 LQE Component Functions

In addition to the Exponential Moving Average, several auxiliary mechanisms contribute to the construction of the LoRa-LQE metric. Accordingly, this section details the metrics and processing steps involved in the LQE calculation, clarifying how each component contributes to the decision-making process of Algorithms 1 and 2.

Normalization. To ensure numerical stability and comparable scaling between features, both the SNR margin and the exponential moving average of the SNR are normalized. The normalized quantities are defined by:

$$Marg_{\text{norm}} = \frac{SNR_{\text{margin}}}{6.0} \quad (5)$$

$$EMA_{\text{norm}} = \frac{EMA}{12.0}. \quad (6)$$

In Equation 5, the factor 6.0 reflects the characteristic SNR separation associated with the capture effect in LoRa receivers [Kufakunesu et al., 2024; Benkahla et al., 2021]. These studies consistently indicate that successful signal dominance in concurrent transmissions emerges when the received power difference approximates 6 dB, a value also aligned with typical receiver noise figures.

The factor 12.0 in Equation 6 represents the effective SNR variation observed across the practical range of spreading factors, from SF7 to SF12, which yields a gain close to 12.5 dB in practical measurements [Azevedo and Mendonça, 2024]. The resulting normalization aligns both metrics within comparable numerical ranges, producing stable features consistent with contemporary literature.

Penalties. The LQE calculation incorporates specific penalties to capture adverse effects associated with unfavorable parameters or channel instability. These penalties reduce the final LQE value whenever the link exhibits suboptimal conditions, acting as correction mechanisms capable of preventing overly optimistic transmission parameter adjustments.

As presented in Line 5 of Algorithm 2 via the *GetSdPenalty* function, the first penalty applies to the standard deviation proportionally to SNR instability, expressed as:

$$SD_{\text{pen}} = \frac{SD}{Norm_{\text{factor}}} \quad (7)$$

thus maintaining high sensitivity to link fluctuations. This study adopts $Norm_{\text{factor}} = 6$ because real-world LoRaWAN scenarios typically exhibit SNR standard deviations below 6 dB, even under mobility conditions [Harinda et al., 2022; Jeftenić et al., 2020]. Consequently, dividing by this value normalizes the penalty approximately to the $[0, 1]$ interval, preventing the standard deviation term from unduly dominating the LQE calculation while maintaining its contribution order of magnitude comparable to other normalized model components.

The second penalty addresses high spreading factors, represented by the *GetSfPenalty* function in Line 7 of Algorithm 2. It aims to discourage SF increases when channel conditions do not justify such choices, avoiding unnecessary energy consumption and capacity reduction. This penalty utilizes a logistic model, where the sigmoidal shape allows for a smooth, continuous, and controlled reduction of the indicator value without introducing abrupt transitions [Zhang and Qiu, 2023]. The corresponding formulation is expressed as:

$$SF_{\text{pen}} = \frac{1}{1 + e^{-k(SF_{\text{norm}} - c)}}, \quad (8)$$

where SF_{norm} denotes the normalized spreading factor. The parameters $k = 8.0$ and $c = 0.55$ control the slope and transition point of the sigmoid function, respectively. The adopted values produce a relatively sharp penalty increase for higher SF levels, with the transition region located near the upper-middle portion of the normalized range. Consequently, low and intermediate SF values receive only a limited penalty, whereas higher SF levels are increasingly penalized due to their greater impact on channel occupancy and energy consumption.

Energy Efficiency Boost. The energy efficiency incentive, represented by the *GetEnergyEffBoost* function in Line 8 of Algorithm 2, integrates a bonus mechanism combining two main terms: one associated with the spreading factor and another related to the transmission power level, as presented below:

$$B_{\text{SF}} = \max(0.075, 0.15 - 0.025 \cdot (SF - 7)) \quad (9)$$

$$B_{\text{TP}} = 0.2 \cdot \min\left(4, \max\left(0, \left\lceil \frac{14 - \text{TP}}{2} \right\rceil\right)\right) \quad (10)$$

$$Boost_{\text{ee}} = B_{\text{SF}} + B_{\text{TP}} \quad (11)$$

where the term B_{SF} gradually reduces the incentive as the spreading factor increases, and is lower-bounded to prevent excessive penalization in high-robustness regimes. Conversely, the term B_{TP} follows a 2 dBm granularity, as in Algorithm 1, assigning progressively larger increments for lower transmission power levels, saturating at levels below 6 dBm. The total increment $Boost_{\text{ee}}$ results from the

sum of these two effects, promoting more energy-efficient configurations whenever adequate link quality conditions are maintained.

Notably, the spreading factor is sufficiently significant in this mechanism that its effect is captured by two distinct components: one acting as a penalty, denoted by $SfPenalty$, and another serving as an energy efficiency incentive, represented by B_{SF} . While $SfPenalty$ acts as a reducing factor that grows with the use of high SF values—reflecting the negative impact of longer transmission times on link reliability—the term B_{SF} introduces a controlled incentive for energy-saving configurations, assigning greater benefit to lower SF values. This differentiation allows the model to preserve the robustness properties inherent to SF increases while simultaneously promoting efficient device energy resource usage.

Computational complexity. In the proposed LoRa-LQE scheme, the computational complexity per uplink evaluation is dominated by the SNR history window of size M and by the G gateways whose reception reports are used to reconstruct it. In each decision cycle, the algorithm processes up to M recent packets and aggregates information from all G gateways, yielding a total cost of $\mathcal{O}(MG)$. Additional operations—such as the exponential moving average update, normalization, and threshold-based decision logic—run in constant time.

In single-gateway scenarios, the complexity reduces to $\mathcal{O}(M)$, as each packet provides only one reception record. Since both M and G are bounded in practical LoRaWAN deployments, the effective runtime per decision is effectively constant, while the asymptotic complexity remains $\mathcal{O}(MG)$ [Cormen et al., 2022]. All processing occurs on the NS side, imposing no additional burden on the EDs.

4.3 Analytical Design of LoRa-LQE

The LoRa-LQE mechanism is designed to simultaneously improve two fundamentally competing objectives: communication reliability—captured by the PDR—and energy efficiency. These objectives are in tension because configurations that maximize PDR, e.g., higher SF or TP, tend to increase energy consumption, whereas configurations that minimize energy expenditure often sacrifice reliability. This inherent conflict constitutes a classical multi-objective optimization problem [Bisio et al., 2026].

Decision Variables and Feasible Space

Let $\mathbf{x} = (SF, TP)$ denote the decision vector [Deb, 2009], where $SF \in \{7, 8, 9, 10, 11, 12\}$ is the Spreading Factor and $TP \in \{2, 4, 6, 8, 10, 12, 14\}$ dBm is the Transmission Power. The feasible set is therefore:

$$\mathcal{X} = \{7, \dots, 12\} \times \{2, 4, 6, 8, 10, 12, 14\}, \quad (12)$$

yielding $|\mathcal{X}| = 42$ candidate configurations at each adaptation step.

Objective Functions

The first objective maximizes the estimated PDR, defined as the fraction of successfully received packets over the total number of transmitted packets:

$$\text{PDR} = \frac{N_{\text{received}}}{N_{\text{sent}}}, \quad (13)$$

where N_{received} is the number of packets successfully decoded at the gateway and N_{sent} is the total number of uplink transmissions [Anwar et al., 2021]. Since computing PDR directly requires accumulating reception outcomes over time, the LQE model provides a real-time estimate $\widehat{\text{PDR}}(\mathbf{x})$ derived from the instantaneous channel state, so that the first objective becomes:

$$f_1(\mathbf{x}) = \widehat{\text{PDR}}(\mathbf{x}). \quad (14)$$

As established in the literature [Farhad et al., 2020], packet reception probability is a monotonically increasing function of the effective link margin $\Delta_{\text{SNR}} = \text{EMA} - \delta(SF)$, where $\delta(SF)$ is the demodulation floor at a given spreading factor.

The second objective maximizes energy efficiency:

$$f_2(\mathbf{x}) = \frac{N_{\text{bits}}}{E(\mathbf{x})}, \quad (15)$$

where $E(\mathbf{x})$ is the energy consumed per transmission cycle and N_{bits} is the payload size in bits, held constant across configurations [Loubany et al., 2023]. The consumed energy is proportional to the transmit power and the ToA:

$$E(\mathbf{x}) \propto P_{\text{TX}}(TP) \cdot \text{ToA}(SF), \quad (16)$$

where $P_{\text{TX}}(TP) = 10^{TP/10}$ mW and $\text{ToA}(SF)$ grows exponentially with SF as formalized in Section 3.1. Combining Equations 15 and 16 makes explicit that both higher SF and higher TP degrade energy efficiency, creating the trade-off the mechanism must navigate.

Multi-Objective Problem Statement

The joint optimization problem can be formally stated as:

$$\begin{aligned} &\underset{\mathbf{x} \in \mathcal{X}}{\text{maximize}} && \mathbf{F}(\mathbf{x}) = [f_1(\mathbf{x}), f_2(\mathbf{x})]^\top \\ &\text{subject to} && SF \in \{7, \dots, 12\}, \\ & && TP \in \{2, 4, 6, 8, 10, 12, 14\} \text{ dBm}, \\ & && \text{EMA} \geq \delta(SF) + \epsilon, \end{aligned} \quad (17)$$

where $\epsilon \geq 0$ is a reliability margin that prevents the selection of configurations whose demodulation sensitivity is insufficient for the current channel conditions. Because f_1 and f_2 are conflicting objectives, no single solution simultaneously maximizes both; instead, the solution set forms a Pareto front—a set of non-dominated configurations in which any improvement in one objective necessarily causes degradation in the other [Deb, 2009].

Weighted-Sum Scalarization and the LQE Score

Exact computation of the Pareto front in real time is computationally impractical for resource-constrained IoT server-side implementations. A common and theoretically grounded strategy is weighted-sum scalarization, in which both objectives are aggregated into a single scalar criterion [Zhu *et al.*, 2022]:

$$\max_{\mathbf{x} \in \mathcal{X}} \lambda_1 f_1(\mathbf{x}) + \lambda_2 f_2(\mathbf{x}), \quad \lambda_1, \lambda_2 \geq 0, \quad \lambda_1 + \lambda_2 = 1. \quad (18)$$

This formulation provides the design rationale for the composite LQE score defined in Algorithm 2. The LQE expression

$$LQE = w_1 \text{EMA}_{\text{norm}} + w_2 \text{Marg}_{\text{norm}} + w_3 SF_{\text{pen}} + w_4 SD_{\text{pen}} + w_5 \text{Boost}_{\text{ee}} \quad (19)$$

embodies the same scalarization structure, with each term encoding a contribution to either f_1 or f_2 . The terms EMA_{norm} and $\text{Marg}_{\text{norm}}$, weighted by $w_1 = 0.1$ and $w_2 = 0.7$, approximate the reliability objective f_1 by capturing the magnitude and margin of the received SNR relative to the demodulation floor. The terms SF_{pen} and SD_{pen} , weighted by $w_3 = -0.3$ and $w_4 = -0.35$, approximate the energy-related objective f_2 by penalizing high spreading factors and channel variability, both of which increase ToA and, consequently, energy consumption. Finally, Boost_{ee} introduces an explicit incentive for energy-efficient configurations: its value is computed conditionally as a decreasing function of both SF and TP , assigning a larger contribution to the LQE score whenever the current configuration operates at lower spreading factor or transmission power. The thresholds $\theta_d = \text{LqeThsDown}$ and $\theta_u = \text{LqeThsUp}$ then implement a practical decision rule that selects the most energy-efficient configuration whose LQE exceeds θ_d , or the most reliable configuration when LQE falls below θ_u , approximating the trade-off boundary defined by the Pareto front.

It is important to highlight that evaluating f_1 and f_2 directly at runtime is precluded by the lack of instantaneous real-time PDR and energy metrics at the network server. To maintain multi-objective rigor under this constraint, the composite LQE score acts as a surrogate scalarization function [Deb, 2009] that maps abstract goals to physical proxies derived from the uplink SNR history. The weights and penalties in Equation 19 guide the optimization trajectory along the conceptual Pareto front, translating reliability and energy expenditures into deterministic decisions. Thus, Equation 17 provides the formal mathematical foundation for the LQE score structure under strict real-time constraints. Post-hoc evaluations in Section 6 empirically confirm that maximizing this surrogate function successfully drives the mobile network toward non-dominated configurations.

This scalarization approach is well-suited to the LoRa-LQE context: it operates with $\mathcal{O}(MG)$ complexity (Section 4), requires no offline training or population maintenance as in evolutionary approaches, and produces deterministic decisions compatible with the event-driven architecture of

LoRaWAN network servers. The following subsection details each component function constituting the LQE score.

5 Methodology

This section outlines the methodology employed to assess the performance and scalability of the LoRa-LQE mechanism under different mobility conditions in LoRaWAN networks. The methodological flow comprises the evaluation strategy, details of the simulation-based experimental environment, the speed classes adopted, and the metrics applied to evaluate the effectiveness of the proposed scheme.

5.1 Evaluation Strategy

The evaluation of LoRa-LQE was performed through simulations to overcome the practical constraints of deploying large-scale, mobile LoRaWAN systems in real-world settings. This method allows controlled experimentation with various mobility behaviors, device densities, and transmission conditions while preserving reproducibility and cost-effectiveness.

The methodological emphasis lies in contrasting the behavior of LoRa-LQE with several recent ADR mechanisms designed specifically for mobile scenarios, including Mobile ADR [Wang *et al.*, 2024], PF-ADR [Sarmiento Neto *et al.*, 2025], and P-ADR [Sarmiento Neto *et al.*, 2024], all of which employ distinct strategies to adjust transmission parameters under varying mobility conditions. These schemes have demonstrated effectiveness in dynamic environments and thus provide a diverse set of reference points for comparison. Additionally, the Standard ADR mechanism is included as a baseline to represent the conventional, mobility-agnostic approach widely used in LoRaWAN deployments. By varying key parameters such as node density and speed, the assessment aims to capture the adaptability, scalability, and efficiency of LoRa-LQE across heterogeneous IoT scenarios.

Each simulation was executed 10 times with distinct random seeds to capture variability and ensure robust outcomes. This repetition improves the accuracy of performance estimates and supports the analysis of result convergence. To ensure statistical soundness, 95% confidence intervals were computed, quantifying the standard error of the mean and reinforcing the reliability of the findings.

5.2 Experimental Setup

In the simulations, all scenarios were implemented in C++ and executed using NS-3¹, integrating the LoRaWAN module² [Farhad *et al.*, 2020; Kufakunesu *et al.*, 2024]. The simulation environment consists of a single gateway and a network server, with 200 to 1,000 LoRa end devices distributed across a circular region with a radius of 5000 m. Scenarios with two gateways were additionally considered for comparison. The end devices transmit 30-byte unconfirmed uplink payloads at an application transmission interval of 600 s, equivalent to 144 packets per day, following a

¹<https://www.nsnam.org/releases/ns-3-48/>

²<https://github.com/signetlabdei/lorawan/releases/tag/v0.3.6>

configuration similar to that adopted in [Sarmiento Neto *et al.*, 2025].

For the performance analysis of the considered schemes, the packet history size M , as introduced in Section 3.2, is set to $M = 2$, considering the highly dynamic nature of the studied environment, which requires more timely adaptations. A reduced window size prioritizes recent observations and enables mobile mechanisms to react faster to abrupt channel variations, improving short-term link estimation and reducing the time required to adjust transmission parameters under mobility and rapidly varying interference conditions [Sarmiento Neto *et al.*, 2024]. A summary of the simulation parameters used in the evaluation of the LoRa-LQE scheme is provided in Table 2.

Table 2. Parametrization adopted.

Parameter	Value
Number of EDs	200, 400, 600, 800, 1 000
Number of GWs	1, 2
Simulation time	24 hours
Area radius	5000 m
App transmission interval	600 s
Packet size	30 bytes
Packet history size (M)	2
Carrier frequency	868 MHz (EU-868)
Bandwidth	125 kHz

The behavior of the EDs was represented through a two-dimensional Random-Walk mobility model [Nouar *et al.*, 2024], with random speed and direction and step-based updates applied every 1000 m, resulting in piecewise linear trajectories. To capture more realistic attenuation effects, the wireless channel followed a log-distance path loss formulation. Moreover, shadowing effects were simulated by introducing stochastic signal fluctuations caused by environmental obstructions and reflective surfaces [Anwar *et al.*, 2021]. The combined path loss and shadowing model is expressed as:

$$L(d) = L(d_0) + 10\gamma \log_{10} \left(\frac{d}{d_0} \right) + X_\sigma \quad (20)$$

where $L(d)$ denotes the path loss in decibels at distance d from the transmitter, $L(d_0)$ is the reference loss at distance d_0 , γ is the path loss exponent, and X_σ represents the shadowing component introduced by the correlated shadowing model [Acosta-Garcia *et al.*, 2025]. The propagation model parameters, selected to represent an urban environment [Al-Gumaei *et al.*, 2022; Farhad *et al.*, 2020], are summarized in Table 3.

Table 3. Propagation and Shadowing Model Parametrization.

Parameter	Value
Reference distance (d_0)	1 m
Path loss exponent (γ)	3.76
Shadowing model	Correlated Shadowing
Shadowing standard deviation (σ)	4.0 dB
Correlation distance (d_{cor})	110 m

In each simulation scenario, all schemes were initially

configured with $SF12$ and a transmission power of $TP = 14$ dBm for every ED. Beyond the rationale for selecting $SF12$ — whose lower data rate and longer airtime yield the highest receiver sensitivity and the broadest communication range [Benkahla *et al.*, 2021] — the choice of $TP = 14$ dBm follows a similar conservative principle. Starting from the maximum permissible transmit power ensures that even end devices under severely degraded channel conditions can reliably reach the gateway during the early stages of operation. Such an initialization strategy guarantees that devices with poor channel conditions are able to establish a dependable connection from the outset, minimizing early packet losses and providing a robust reference point from which transmission parameters can be gradually optimized.

5.2.1 LoRa-LQE Parametrization

The parametrization adopted for the LoRa-LQE evaluation was selected to emphasize energy-efficient behavior while preserving stability under realistic LoRaWAN channel variations, as presented in Table 4. The rationale for the EMA smoothing factor $\alpha = 0.3$ is provided in Section 4.2.1. The decision thresholds were set to $LqeThsDown = 1.5$ and $LqeThsUp = -2.5$, creating an asymmetric decision window that favors conservative downlink adjustments and avoids premature increases in spreading factor or transmission power. A cooldown period of 3 uplinks prevents oscillatory behavior, reducing redundant MAC commands and unnecessary energy expenditure on the end devices.

The weighting coefficients used in the LQE composition were calibrated to prioritize link margin as the dominant indicator of safe energy-efficient operation. The weights $w_1 = 0.1$ and $w_2 = 0.7$ reflect the minor contribution of the EMA relative to the normalized SNR margin, while negative weights $w_3 = -0.3$ and $w_4 = -0.35$ penalize high spreading factors and unstable channels, respectively. The energy-efficiency reinforcement term with $w_5 = 1.0$ ensures that favorable operating conditions are exploited to reduce SF and TP whenever possible. Together, these parameters produce a decision model that emphasizes conservative adaptation under uncertainty and aggressive optimization when channel conditions are reliably favorable.

Table 4. LoRa-LQE parametrization adopted.

Parameter	Value
EMA smoothing factor (α)	0.3
$LqeThsDown$	1.5
$LqeThsUp$	-2.5
Cooldown steps	3
EMA weight (w_1)	0.1
Margin weight (w_2)	0.7
SF penalty weight (w_3)	-0.3
SD penalty weight (w_4)	-0.35
Energy-efficiency boost weight (w_5)	1.0

5.3 Speed Classes

Dynamic environments are characterized by heterogeneous mobility patterns, encompassing diverse movement speeds

that impose distinct constraints on wireless connectivity. Evaluating network reliability across this operational spectrum facilitates the validation of communication protocols in realistic deployment scenarios. To assess LoRa-LQE performance under varying mobility levels, the scheme was applied progressively to higher speed classes, following an approach similar to [Teymuri *et al.*, 2023]. These speed classes are modeled to represent a broad range of application contexts involving various moving objects:

- *Class 1: 0.1 to 4 m/s*: this interval represents very low mobility, typically associated with pedestrian motion in both urban and rural environments, as well as movement assisted by devices such as powered wheelchairs or slow ambulation in parks and recreational zones [Arumukhom Revi *et al.*, 2021];
- *Class 2: 4 to 8 m/s*: corresponds to speeds commonly observed in urban cycling, electric scooter usage on public pathways, and small delivery vehicles traveling at reduced velocities in residential or commercial districts [Ren *et al.*, 2022];
- *Class 3: 8 to 12 m/s*: relates to speeds attained by lightweight motorized transport, including compact electric cars operating on short trips within cities, drones in low-altitude missions, as well as motorcycles circulating under moderate traffic conditions [Li *et al.*, 2024];
- *Class 4: 12 to 16 m/s*: covers velocities typically associated with vehicles navigating urban and suburban highways under average traffic loads [Ferrari *et al.*, 2020];
- *Class 5: 16 to 20 m/s*: represents mobility levels compatible with vehicles traveling on highways or fast-moving routes connecting urban and rural regions [Di Renzone *et al.*, 2024].

Evaluating LoRa-LQE across different speed classes is essential to understand how it adapts to fast-changing channel conditions. Since each speed range brings distinct propagation dynamics, this analysis more fully assesses the scheme's robustness and its ability to maintain reliable performance in realistic mobile LoRaWAN deployments.

5.4 Performance Metrics

To analyze the obtained results, a set of metrics commonly employed in simulations of wireless communication protocols—particularly in LoRaWAN scenarios—was considered. These metrics encompass packet delivery ratio, energy efficiency, and packet loss ratio. Each one offers meaningful insights into distinct facets of network behavior, including communication reliability, energy consumption, and the influence of interference and varying channel conditions.

The PDR corresponds to the proportion between the number of packets correctly received at the gateway and the total number of packets transmitted by the device, as presented in Equation 13. This metric is fundamental for assessing communication robustness, as it reflects the effectiveness of data transmission within the network [Anwar *et al.*, 2021].

Energy efficiency is another relevant performance indicator, expressing the amount of successfully delivered data, in bits, relative to the energy consumed, measured in joules (J), as denoted in Equation 15. This metric reflects how effectively the device utilizes its energy resources to achieve successful communication [Al-Gumaei *et al.*, 2022].

Finally, the Packet Loss Ratio (PLR) is a metric used to evaluate network reliability by indicating the percentage of transmitted packets that fail to reach the gateway. Multiple factors contribute to this metric [Anwar *et al.*, 2021]. PLR-T occurs when the gateway is engaged in downlink activities, such as transmitting acknowledgment (ACKs) packets, which prevents simultaneous reception of uplink messages. PLR-S takes place when packets are lost because their signal strength drops below the gateway's sensitivity limit, producing signals too weak to be correctly decoded. PLR-R denotes packet loss that happens when all reception paths at the gateway are occupied by other ongoing transmissions. Lastly, PLR-I results from interference when packets overlap on identical or adjacent channels, a situation that may arise even with different spreading factors. Such collisions can cause packet loss, especially in scenarios with elevated uplink traffic or network congestion.

Examining these metrics provides a broad view of network behavior with respect to reliability, energy usage, and operational efficiency under different conditions. These observations make it possible to thoroughly assess how effectively the proposed solution enhances LoRaWAN communication.

6 Results and Analysis

This section presents and analyzes the simulation results obtained following the methodology described in Section 5. The evaluation investigates the performance of LoRa-LQE under different network and mobility conditions. Initially, the impact of node mobility, application transmission interval, network density, and gateway deployment is analyzed and compared with state-of-the-art ADR schemes. Subsequently, sensitivity analyses are conducted to assess the influence of payload size, packet history size, and mobility models on the performance of the proposed mechanism.

Based on the speed classes defined in Section 5.3 and considering an intermediate network density of 500 EDs, Figures 3a and 3b evaluate the impact of node mobility on packet delivery ratio and energy efficiency. Across all speed classes, LoRa-LQE consistently achieves the highest performance, maintaining an average PDR of 0.769. Compared with Mobile ADR, PF-ADR, P-ADR, and standard ADR, the proposed scheme improves PDR by 11.8%, 31.3%, 57.6%, and 122.4%, respectively. The performance gap becomes more pronounced at higher speed classes, where the standard ADR fails to adapt effectively to dynamic environments.

A similar trend is observed for energy efficiency in Figure 3b. LoRa-LQE achieves an average efficiency of 3469 bits/J, corresponding to gains of 26.4%, 82.7%, 109.1%, and 103.9% over Mobile ADR, PF-ADR, P-ADR, and standard ADR, respectively. These results indicate that the

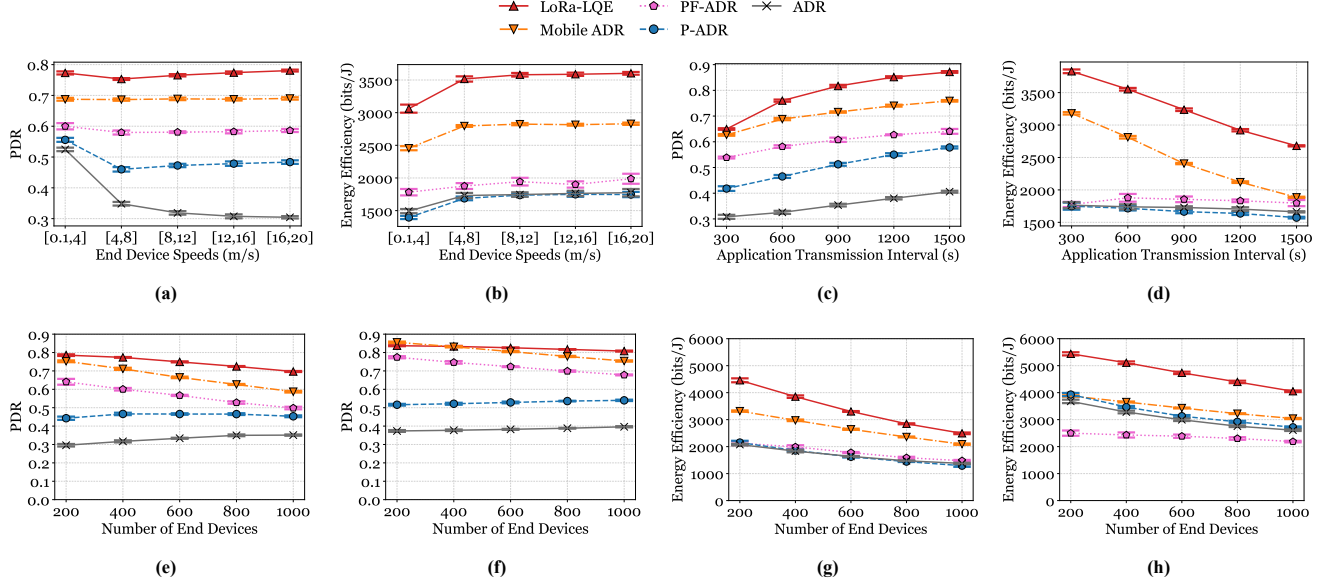


Figure 3. Comparison of ADR schemes under varying network conditions: (a) PDR versus speed classes, (b) energy efficiency versus speed classes, (c) PDR versus transmission interval, (d) energy efficiency versus transmission interval, (e) PDR versus number of EDs (1 GW), (f) PDR versus number of EDs (2 GWs), (g) energy efficiency versus number of EDs (1 GW), and (h) energy efficiency versus number of EDs (2 GWs).

proposed link-quality estimation mechanism enables more accurate transmission parameter selection under mobility, simultaneously improving PDR and reducing the energy cost per successfully delivered bit.

Figures 3c and 3d evaluate the impact of the application transmission interval on PDR and energy efficiency for a network comprising 500 EDs and an intermediate mobility range of 4–16 m/s, corresponding to the three middle speed classes defined in Section 5.3. As the transmission interval increases from 300 s to 1500 s, all ADR schemes benefit from the reduced channel occupancy and lower collision probability. Nevertheless, LoRa-LQE consistently achieves the best overall performance, reaching an average PDR of 0.783 and an average energy efficiency of 3299 bits/J. Compared with Mobile ADR, PF-ADR, P-ADR, and standard ADR, the proposed scheme improves the average PDR by 11.6%, 31.5%, 56.7%, and 122.8%, respectively. Similarly, average energy efficiency gains of 32.2%, 77.5%, 93.7%, and 88.2% are observed over the same schemes. These results indicate that LoRa-LQE is able to exploit the reduced network contention more effectively, leveraging its link-quality estimation mechanism to select transmission parameters that simultaneously maximize reliability and energy efficiency.

Subsequently, keeping the same mobility range of 4–16 m/s, the number of EDs was varied to assess network scalability. The resulting PDR performance is shown in Figures 3e and 3f for one- and two-gateway deployments, respectively. As network density increases, all ADR schemes experience a gradual reduction in PDR due to the higher level of channel contention and packet collisions. Nevertheless, LoRa-LQE consistently maintains the highest reliability. In the 1-GW scenario, it achieves an average PDR of 0.746, outperforming Mobile ADR, PF-ADR, P-ADR, and standard ADR by 12.1%, 32.1%, 62.8%, and 128.0%, respectively. When a second gateway is deployed, the average PDR increases to 0.824, highlighting the benefits of additional reception diversity. Under this configuration, LoRa-LQE maintains its

superiority, providing gains of 14.0%, 56.1%, and 115.1% over PF-ADR, P-ADR, and standard ADR, respectively, while remaining slightly ahead of Mobile ADR by 2.5%. These results demonstrate that LoRa-LQE preserves its reliability advantages under increasing network density and can effectively leverage additional gateway infrastructure.

A similar trend can be observed for energy efficiency in Figures 3g and 3h. As network density increases, energy efficiency decreases for all schemes due to the larger number of retransmissions and unsuccessful transmissions caused by congestion. Despite this degradation, LoRa-LQE consistently achieves the highest efficiency. In the 1-GW scenario, it reaches an average energy efficiency of 3390 bits/J, providing gains of 25.8%, 87.3%, 101.7%, and 100.1% over Mobile ADR, PF-ADR, P-ADR, and standard ADR, respectively. The addition of a second gateway further increases the average efficiency to 4746 bits/J, representing an improvement of approximately 40.0% relative to the 1-GW deployment. Under the 2-GW configuration, LoRa-LQE outperforms Mobile ADR, PF-ADR, P-ADR, and standard ADR by 37.6%, 100.5%, 47.4%, and 55.3%, respectively. The larger gains observed in the dual-gateway scenario suggest that LoRa-LQE is able to exploit the additional reception diversity more effectively, translating the improved link conditions into higher energy efficiency while maintaining reliable communication.

Keeping the mobility range fixed at 4–16 m/s and considering an intermediate network density of 500 EDs, Figure 4 evaluates the impact of packet payload size on the performance of the ADR schemes. As expected, increasing the payload size generally reduces reliability due to longer channel occupancy and a higher probability of collisions. Nevertheless, LoRa-LQE consistently achieves the highest PDR across all payload sizes, reaching an average value of 0.7425 and outperforming Mobile ADR, PF-ADR, P-ADR, and standard ADR by 9.2%, 33.7%, 68.5%, and 139.4%, respectively. A similar trend is observed for energy

efficiency, where LoRa-LQE achieves the best performance and increasingly larger gains as the payload size grows, reaching 6833 bits/J at 50 bytes. On average, it improves energy efficiency by 46.6%, 111.9%, 148.9%, and 146.8% compared with Mobile ADR, PF-ADR, P-ADR, and standard ADR, respectively. Although the average PDR decreases more noticeably at 50 bytes, this reduction is accompanied by a substantial increase in energy efficiency, suggesting that the transmission parameter selections made by LoRa-LQE become increasingly advantageous as the energy cost of larger packets grows. Overall, these results show that LoRa-LQE remains effective under higher channel occupancy, preserving reliability while improving energy utilization.

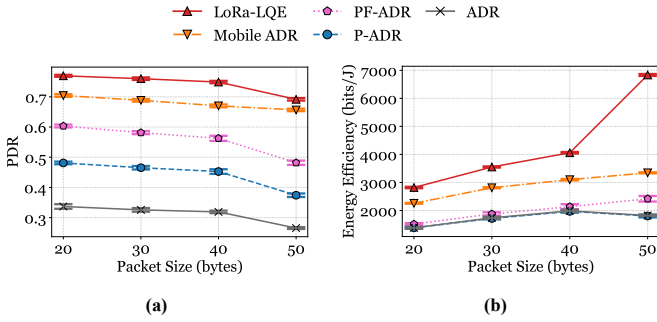


Figure 4. Impact of packet payload size on the performance of ADR schemes: (a) PDR and (b) energy efficiency.

To analyze how the packet history size influences the performance of the proposed scheme, a scenario with a network density of 500 EDs is considered. Figure 5a shows that increasing M yields only modest PDR improvements, with $M = 8$ achieving the highest average PDR and outperforming $M = 2$ by approximately 1.4%. However, energy efficiency is much more sensitive to this parameter, as shown in Figure 5b. Despite their similar PDR values, $M = 2$ improves energy efficiency by 31.3%, 61.2%, and 94.1% compared to $M = 8$, $M = 14$, and $M = 20$, respectively. These results indicate that larger history sizes provide limited reliability gains while significantly reducing energy efficiency, making $M = 2$ the most balanced configuration for the evaluated scenario.

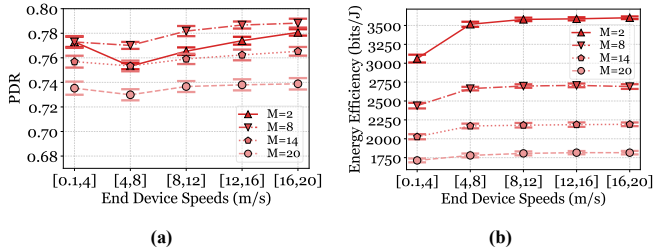


Figure 5. Impact of different packet history sizes M on LoRa-LQE performance across end-device speed classes: (a) PDR and (b) energy efficiency.

To further assess the robustness of LoRa-LQE under different mobility dynamics, its performance was evaluated using the Random Walk (RW), Steady-State Random Waypoint (SSRWP), and Gauss-Markov (GM) mobility models [Nouar et al., 2024] for a network density of 500 EDs. While RW was adopted as the reference mobility model in the previous experiments, SS-RWP and GM provide more

realistic representations of node movement [Sarmiento Neto et al., 2025]. As shown in Figure 6a, SS-RWP consistently achieves the highest PDR across most speed classes, reaching an average value of 0.813, which corresponds to gains of 5.8% and 3.3% over RW and GM, respectively. The performance advantage is particularly evident at low and moderate speeds, although GM gradually approaches SS-RWP as mobility increases. A similar trend is observed in Figure 6b, where SS-RWP attains the highest average energy efficiency, reaching 4601 bits/J and outperforming RW and GM by 33.0% and 30.8%, respectively. Overall, these results suggest that LoRa-LQE tends to perform better under more realistic mobility conditions, where movement exhibits greater spatial and temporal correlation, enabling more accurate link-quality estimation and improved network performance.

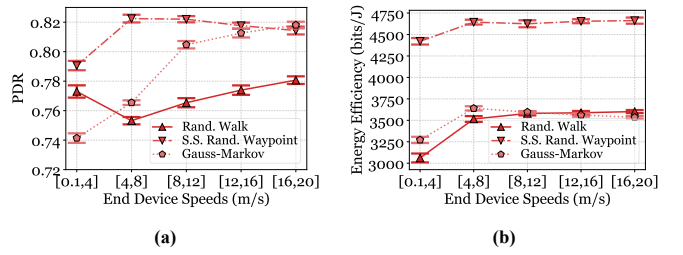


Figure 6. Impact of mobility models on LoRa-LQE performance across end-device speed classes: (a) PDR and (b) energy efficiency.

To better interpret these results, consider the analysis scenario with 500 EDs in the central speed class of 8 to 12 m/s, where the largest average energy efficiency gap between LoRa-LQE and Mobile ADR occurs. As shown in Figure 7a, LoRa-LQE gradually seeks to reduce the highest spreading factor indices, SF11 and SF12, by distributing transmissions across the intermediate spreading factor levels, SF7, SF8, and SF9. Consequently, the scheme favors moderate energy consumption while attempting to optimize packet delivery ratio, resulting in superior average energy efficiency. In contrast, Figure 7b shows that Mobile ADR seeks to reduce the maximum spreading factor level, converging almost exclusively to the intermediate level, SF9. This behavior leads to higher packet collision rates since most packets arrive at the gateway with the same SF, except for those favored by the capture effect [Kufakunesu et al., 2024].

To demonstrate the impact of this spreading factor distribution, it is useful to analyze the distinct packet loss factors. Observing LoRa-LQE in Figure 8a, its policy of balancing intermediate SF levels results in a relatively low packet loss proportion. However, the average PLR-S tends to be slightly higher for the most distant end devices, as the mechanism avoids the highest SF levels, thereby increasing the probability of packet loss due to lower reception sensitivity. On the other hand, as shown in Figure 8b, the losses in PLR-S and PLR-I associated with Mobile ADR are significantly higher, as this mechanism utilizes the same intermediate SF level for all devices.

In summary, simulation results demonstrate that LoRa-LQE effectively balances reliability and energy efficiency in dynamic LoRaWAN scenarios, consistently outperforming state-of-the-art schemes such as Mobile ADR and PF-ADR. The proposed mechanism leverages

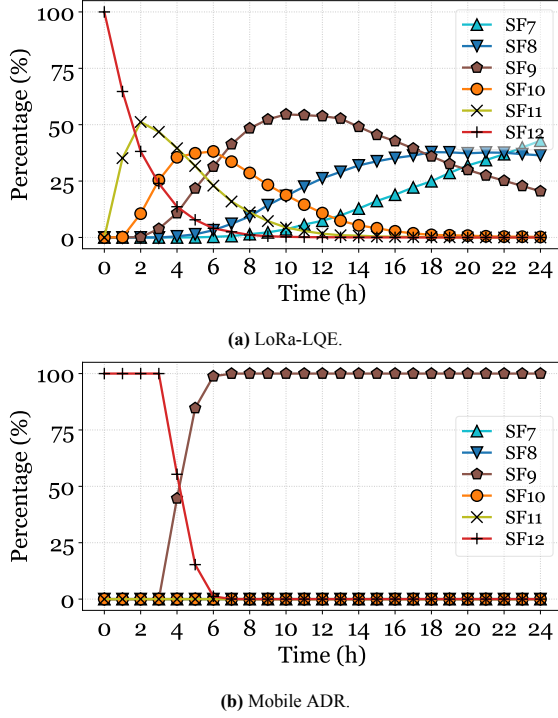


Figure 7. Temporal evolution and final distribution of the average SF for 500 end devices, considering speeds in the range [8, 12] m/s, across the evaluated schemes.

filtered SNR measurements, channel stability, and packet history to dynamically optimize transmission parameters, mitigating mobility and network congestion effects. However, a challenge lies in the reliance on fixed weighting factors and heuristic constants for LQE calculation and decision thresholds. While robust in tested scenarios, these static parameters may require manual fine-tuning to achieve optimal performance across heterogeneous environments with different interference patterns or propagation characteristics.

7 Conclusion

This work presented LoRa-LQE, a Link Quality Estimation mechanism designed to overcome inefficiencies of the standard ADR in mobile LoRaWAN scenarios. The study showed that the latency of ADR's adaptation dynamics leads to suboptimal transmission parameter choices, causing packet loss and unnecessary energy consumption. To address this, the proposed solution uses a heuristic that combines SNR history with stability metrics to estimate link quality more accurately. By prioritizing stable connection windows and reducing the aggressive backoff behavior of the default protocol, LoRa-LQE improves connectivity reliability for mobile end devices.

The evaluation, conducted through extensive simulations, confirmed the superior performance of the proposed solution over both the standard ADR and other mobility-aware schemes. Across scenarios from low to high node density, the algorithm consistently maintained robust communication. Results showed notable reliability gains, with PDR improvements of approximately 12.1% and 139.4% compared to Mobile ADR and standard ADR, respectively.

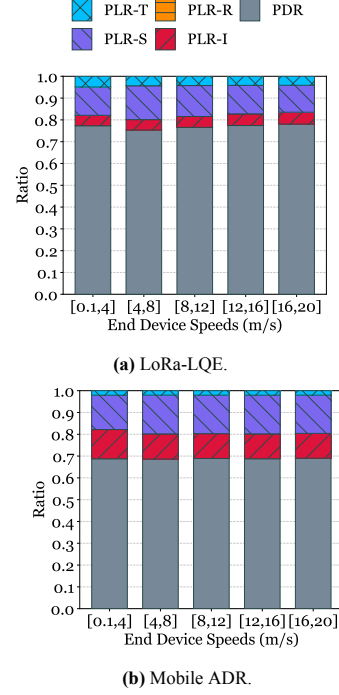


Figure 8. Distribution of packet loss origins for 500 end devices, considering speeds in the range [8, 12] m/s, across the evaluated schemes.

The proposed scheme also improved energy efficiency by approximately 46.6% and 146.8% relative to these approaches. These results demonstrate the benefits of incorporating short-term channel stability indicators into parameter allocation decisions.

Despite these promising results, the current iteration of the algorithm relies on fixed weighting factors and predefined constants for link-quality estimation. Although effective in the simulated environments, these static parameters may not adapt optimally to highly heterogeneous scenarios with unpredictable interference patterns. Future work will therefore focus on dynamic parameter tuning based on environmental feedback, as well as on the implementation and validation of LoRa-LQE in a real-world testbed using commercial off-the-shelf hardware under physical channel conditions.

Declarations

Funding

No funding was received for this study.

Authors' Contributions

G.A.S.N. conceptualized the study and wrote the main manuscript text; G.A.S.N. and T.A.R.S. formalized and validated the proposed model; G.A.S.N., F.J.V.S., and P.F.F.A. developed and conducted the simulations; G.A.S.N., F.J.V.S., and L.H.O.M. prepared and reviewed all figures and tables; J.V.R.J. contributed to the final version of the manuscript. All authors reviewed and approved the manuscript.

Competing interests

The authors declare no competing interests.

Availability of data and materials

No datasets were used in this study. The research was conducted entirely through a simulated network environment.

References

- Acosta-Garcia, L., Aznar-Poveda, J., Garcia-Sanchez, A. J., Hollenstein, J., Garcia-Haro, J., and Fahringer, T. (2025). ADRL: A reconfigurable energy-efficient transmission policy for mobile LoRa devices based on reinforcement learning. *Internet of Things*, 31:101577. DOI: 10.1016/j.iot.2025.101577.
- Al-Gumaei, Y. A., Aslam, N., Aljaidi, M., Al-Saman, A., Alsarhan, A., and Ashyap, A. Y. (2022). A Novel Approach to Improve the Adaptive-Data-Rate Scheme for IoT LoRaWAN. *Electronics*, 11(21). DOI: 10.3390/electronics11213521.
- Al-Qassas, R. and Qasaimeh, M. (2024). An empirical evaluation of link quality utilization in ETX routing for VANETs. *PeerJ Computer Science*, 10:e2259. DOI: 10.7717/peerj-cs.2259.
- Almuhaya, M. A. M., Jabbar, W. A., Sulaiman, N., and Abdulmalek, S. (2022). A Survey on LoRaWAN Technology: Recent Trends, Opportunities, Simulation Tools and Future Directions. *Electronics*, 11(1). DOI: 10.3390/electronics11010164.
- Anwar, K., Rahman, T., Zeb, A., Khan, I., Zareei, M., and Vargas-Rosales, C. (2021). RM-ADR: Resource Management Adaptive Data Rate for Mobile Application in LoRaWAN. *Sensors*, 21(23). DOI: 10.3390/s21237980.
- Arumukhom Revi, D., De Rossi, S. M. M., Walsh, C. J., and Awad, L. N. (2021). Estimation of Walking Speed and Its Spatiotemporal Determinants Using a Single Inertial Sensor Worn on the Thigh: from Healthy to Hemiparetic Walking. *Sensors*, 21(21). DOI: 10.3390/s21216976.
- Azevedo, J. A. and Mendonça, F. (2024). A Critical Review of the Propagation Models Employed in LoRa Systems. *Sensors*, 24(12). DOI: 10.3390/s24123877.
- Benkahla, N., Tounsi, H., Song, Y.-Q., and Frikha, M. (2021). Review and experimental evaluation of ADR enhancements for LoRaWAN networks. *Telecommunication Systems*, 77(1):1–22. DOI: 10.1007/s11235-020-00738-x.
- Bisio, I., Garibotto, C., Lavagetto, F., Sciarrone, A., and Zerbino, M. (2026). Distributed Multiobjective Optimization for Edge Computing in Resource-Constrained Social IoT Networks. *IEEE Internet of Things Journal*, 13(8):15549–15558. DOI: 10.1109/JIOT.2026.3657050.
- Bonilla, V., Campoverde, B., and Yoo, S. G. (2023). A Systematic Literature Review of LoRaWAN: Sensors and Applications. *Sensors*, 23(20). DOI: 10.3390/s23208440.
- Cormen, T. H., Leiserson, C. E., Rivest, R. L., and Stein, C. (2022). *Introduction to Algorithms*. MIT Press, Cambridge, MA, 4 edition. DOI: 10.1017/cbo9781107279667.013.
- de Jesus, G. G. M., Souza, R. D., Montez, C., and Hoeller, A. (2021). LoRaWAN Adaptive Data Rate With Flexible Link Margin. *IEEE Internet of Things Journal*, 8(7):6053–6061. DOI: 10.1109/JIOT.2020.3033797.
- Deb, K. (2009). *Multi-Objective Optimization Using Evolutionary Algorithms*. John Wiley & Sons, Wiley paperback edition. Book.
- Dev, A., Khan, K. N., Patil, N., Sa, D. S., Arora, N., and Singh, A. (2025). Energy-efficient routing algorithms for mobile ad-hoc networks. *Journal of Wireless Mobile Networks, Ubiquitous Computing, and Dependable Applications*, 16(2):332–345. DOI: 10.58346/jowua.2025.i2.021.
- Di Renzone, G., Parrino, S., Peruzzi, G., Pozzebon, A., and Vangelista, L. (2024). LoRaWAN for Vehicular Networking: Field Tests for Vehicle-to-Roadside Communication. *Sensors*, 24(6). DOI: 10.3390/s24061801.
- Farhad, A., Kim, D.-H., Subedi, S., and Pyun, J.-Y. (2020). Enhanced LoRaWAN Adaptive Data Rate for Mobile Internet of Things Devices. *Sensors*, 20(22). DOI: 10.3390/s20226466.
- Ferrari, P., Sisinni, E., Carvalho, D. F., Depari, A., Signoretti, G., Silva, M., Silva, I., and Silva, D. (2020). On the use of LoRaWAN for the Internet of Intelligent Vehicles in Smart City scenarios. In *2020 IEEE Sensors Applications Symposium (SAS)*, pages 1–6. DOI: 10.1109/SAS48726.2020.9220069.
- Harinda, E., Wixted, A. J., Qureshi, A.-U.-H., Larijani, H., and Gibson, R. M. (2022). Performance of a Live Multi-Gateway LoRaWAN and Interference Measurement across Indoor and Outdoor Localities. *Computers*, 11(2). DOI: 10.3390/computers11020025.
- Iqbal, S., Ahmed, A., Siraj, M., Tamimi, M. A., Bhangwar, A. R., and Kumar, P. (2023). A Multi-Hop QoS-Aware and Predicting Link Quality Estimation (PLQE) Routing Protocol for Reliable WBSN. *IEEE Access*, 11:35993–36003. DOI: 10.1109/ACCESS.2023.3266067.
- Isyaku, B., Bakar, K. A., Mohd Zahid, M. S., Alkhamash, E. H., Saeed, F., and Ghaleb, F. A. (2021). Route Path Selection Optimization Scheme Based Link Quality Estimation and Critical Switch Awareness for Software Defined Networks. *Applied Sciences*, 11(19). DOI: 10.3390/app11199100.
- Jeftenić, N., Simić, M., and Stamenković, Z. (2020). Impact of Environmental Parameters on SNR and RSS in LoRaWAN. In *2020 International Conference on Electrical, Communication, and Computer Engineering (ICECCE)*, pages 1–6. DOI: 10.1109/ICECCE49384.2020.9179250.
- Jouhari, M., Saeed, N., Alouini, M.-S., and Amhoud, E. M. (2023). A Survey on Scalable LoRaWAN for Massive IoT: Recent Advances, Potentials, and Challenges. *IEEE Communications Surveys Tutorials*, 25(3):1841–1876. DOI: 10.1109/COMST.2023.3274934.
- Kufakunesu, R., Hancke, G. P., and Abu-Mahfouz, A. M. (2024). Collision Avoidance Adaptive Data Rate Algorithm for LoRaWAN. *Future Internet*, 16(10). DOI: 10.3390/fi16100380.
- Li, D., Wang, Z., Lai, Y., and Shen, H. (2024). Dataset Augmentation and Fractional Frequency Offset Compensation Based Radio Frequency Fingerprint Identification in Drone Communications. *Drones*, 8(10).

- DOI: 10.3390/drones8100569.
- Liu, L., Lv, H., Xu, J., and Shu, J. (2021). A Link Quality Estimation Method Based on Improved Weighted Extreme Learning Machine. *IEEE Access*, 9:11378–11392. DOI: 10.1109/ACCESS.2021.3051169.
- Liu, W., Xia, Y., Luo, R., and Hu, S. (2020). Lightweight, Fluctuation Insensitive Multi-Parameter Fusion Link Quality Estimation for Wireless Sensor Networks. *IEEE Access*, 8:28496–28511. DOI: 10.1109/ACCESS.2020.2972326.
- LoRa Alliance (2020). TS001-1.0.4 LoRaWAN® L2 1.0.4 Specification. Available at: <https://resources.lora-alliance.org/technical-specifications/ts001-1-0-4-lorawan-l2-1-0-4-specification>. Technical Specification.
- Louati, R., Thaalbi, M., and Tabbane, N. (2024). An Adaptive Data Rate scheme for multi-hop forwarding in LoRaWAN networks. In *2024 International Wireless Communications and Mobile Computing (IWCMC)*, pages 1791–1796. DOI: 10.1109/IWCMC61514.2024.10592512.
- Loubany, A., Lahoud, S., Samhat, A. E., and El Helou, M. (2023). Improving Energy Efficiency in LoRaWAN Networks with Multiple Gateways. *Sensors*, 23(11). DOI: 10.3390/s23115315.
- Moreira, D. S., Santos, G., Yanai, A. E., Souza, P. B. d., Melo, P. V. F. d., and Mota, E. (2025). LoRaBB: an Algorithm for Parameter Selection in LoRa-Based Communication for the Amazon Rainforest. *Sensors*, 25(4). DOI: 10.3390/s25041200.
- Nouar, A., Abbes, M. T., Boumerdassi, S., and Chaib, M. (2024). Impact of Mobility Model on LoRaWAN Performance. *Journal of Communications*, 19(1). DOI: 10.12720/jcm.19.1.7-18.
- Pham, V. D., Hao Do, P., Le, D. T., and Kirichek, R. (2021). LoRa Link Quality Estimation Based on Support Vector Machine. In *Distributed Computer and Communication Networks: Control, Computation, Communications*, pages 92–102, Cham. Springer International Publishing. DOI: 10.1007/978-3-030-92507-9.
- Quattrocchi, G., Incerto, E., Pincirolì, R., Trubiani, C., and Baresi, L. (2024). Autoscaling Solutions for Cloud Applications Under Dynamic Workloads. *IEEE Transactions on Services Computing*, 17(3):804–820. DOI: 10.1109/TSC.2024.3354062.
- Ren, Y., Liu, L., Li, C., Cao, Z., and Chen, S. (2022). Is LoRaWAN Really Wide? Fine-grained LoRa Link-level Measurement in an Urban Environment. In *2022 IEEE 30th International Conference on Network Protocols (ICNP)*, pages 1–12. DOI: 10.1109/ICNP55882.2022.9940375.
- Sarmiento Neto, G. A., da Silva, T. A. R., Abreu, P. F. F., Veloso, A. F. D. S., Mendes, L. H. D. O., and dos Reis Junior, J. V. (2024). A Percentile Based ADR for Mobile LoRaWAN Applications. In *XLII Brazilian Symposium on Computer Networks and Distributed Systems (SBRC)*, pages 43–56, Porto Alegre, RS, Brasil. SBC. DOI: 10.5753/sbrc.2024.1248.
- Sarmiento Neto, G. A., da Silva, T. A. R., Veloso, A. F. d. S., de Abreu, P. F., Mendes, L. H. d. O., Rabelo, R. A. L., and dos Reis Jr, J. V. (2025). An Alternative Particle Filter-Driven ADR for Mobile Devices in LoRaWAN Networks. *Journal of Internet Services and Applications*, 16(1):194–208. DOI: 10.5753/jisa.2025.5043.
- Shahid, M., Tariq, M., Iqbal, Z., Albarakati, H. M., Fatima, N., Khan, M. A., and Shabaz, M. (2024). Link-Quality-Based Energy-Efficient Routing Protocol for WSN in IoT. *IEEE Transactions on Consumer Electronics*, 70(1):4645–4653. DOI: 10.1109/TCE.2024.3356195.
- Soy, H. (2023). An adaptive spreading factor allocation scheme for mobile LoRa networks: Blind ADR with distributed TDMA scheduling. *Simulation Modelling Practice and Theory*, 125:102755. DOI: 10.1016/j.simpat.2023.102755.
- Tesfaw, B. A., Juang, R.-T., Tai, L.-C., Lin, H.-P., Tarekegn, G. B., and Nathanael, K. W. (2023). Deep Learning-Based Link Quality Estimation for RIS-Assisted UAV-Enabled Wireless Communications System. *Sensors*, 23(19). DOI: 10.3390/s23198041.
- Teymuri, B., Serati, R., Anagnostopoulos, N. A., and Rasti, M. (2023). LP-MAB: Improving the Energy Efficiency of LoRaWAN Using a Reinforcement-Learning-Based Adaptive Configuration Algorithm. *Sensors*, 23(4). DOI: 10.3390/s23042363.
- Wang, H., Jia, J., Liu, M., and Pan, R. (2022). Research on Data-Driven LoRa Link Quality Estimation. In *2022 4th International Conference on Natural Language Processing (ICNLP)*, pages 580–586. DOI: 10.1109/ICNLP55136.2022.00106.
- Wang, H., Zhang, X., Liao, J., Zhang, Y., and Li, H. (2024). An improved adaptive data rate algorithm of LoRaWAN for agricultural mobile sensor nodes. *Computers and Electronics in Agriculture*, 219:108773. DOI: 10.1016/j.compag.2024.108773.
- Zhang, Y. and Qiu, H. (2023). Delay-Aware and Link-Quality-Aware Geographical Routing Protocol for UANET via Dueling Deep Q-Network. *Sensors*, 23(6). DOI: 10.3390/s23063024.
- Zhu, S., Xu, L., Goodman, E. D., and Lu, Z. (2022). A New Many-Objective Evolutionary Algorithm Based on Generalized Pareto Dominance. *IEEE Transactions on Cybernetics*, 52(8):7776–7790. DOI: 10.1109/TCYB.2021.3051078.