

## Improving Delay Limits in XG-PON Networks Using an RNN-LSTM Model

N'DRI Akanza Konan Ricky<sup>1</sup>, SAMASSI Adama<sup>2</sup>, KIE Eba Victoire<sup>3</sup>

<sup>1</sup>Université Alassane OUATTARA Bouaké Côte d'Ivoire

<sup>2</sup>Université Alassane OUATTARA Bouaké Côte d'Ivoire

<sup>3</sup>École Supérieure Africaine de Technologie et de Télécommunication

**Abstract:** Delay control in XG-PON networks is crucial to guaranteeing the quality of service for real-time sensitive applications such as videoconferencing and telemedicine. Traditional methods based on queueing theory and statistical calculations impose strict limits on the maximum allowable delay in the network. This article proposes an improvement to these limits by integrating a short-long memory recurrent neural network (RNN-LSTM) model, which enables dynamic and adaptive prediction of transmission delay. The integration of RNN-LSTM optimizes dynamic delay control in XG-PON networks, offering significantly better performance than traditional methods, especially in variable and demanding conditions. The results demonstrate better delay control with a significant reduction in delays exceeding the 3 ms limit of 40% compared to conventional analytical methods, and the mean prediction error (RMSE) remains below 1 ms across the test set, thus validating the accuracy and robustness of the model even in the event of network overload

**Keywords:** XG-PON, QoS, end-to-end delay, RNN-LSTM, Queueing Theory, Dynamic Bandwidth Allocation (DBA), SLA, RMSE.

### 1. INTRODUCTION

XG-PON networks are a fundamental architecture for providing high-speed services with guaranteed quality of service (QoS). However, delay constraints, particularly in real-time services, remain a major challenge. Traditional approaches, based on queueing theory and statistics, allow upper limits of delay to be set at 3 ms [1]. However, these approaches suffer from shortcomings, in particular the growing gap between analytical models and exact models under heavy load, and the lack of proactive mechanisms that adapt to multivariate and non-stationary traffic. These limitations can underestimate delays in the event of rapid traffic variations and DBA dynamics.

Advances in artificial intelligence, specifically deep learning with RNN-LSTM networks, offer new possibilities that make it possible to: learning traffic temporal dependencies and predicting short- and medium-term end-to-end delay distribution; creating a hybrid fusion in which the analytical framework (queueing theory and hypothesis testing) is used in parallel to

obtain SLA margins and formal guarantees; providing probabilistic forecasts and action signals for DBA through the LSTM predictor; dynamically adjusting ONU/OLT bandwidth allocations and TDMA windows based on delay forecasts and associated uncertainties; achieving sub- performance under different loads, different inter-arrival distributions (Poisson, bursty, MMPP), and various XG-PON topologies.

This article is structured as follows: Related works on approaches to determining delay in the xg-pon network, materials and methods, Results and Discussion, Conclusion and Prospects, finish with the References.

## **2. Related Works**

An analysis of major articles has identified studies structured according to three groups of approaches.

### **2.1. Analytical and statistical approaches**

In [2], the authors combine queueing theory and Markov chains to model the stochastic behavior of systems where processing times and queues coexist, thereby enabling the calculation of key metrics such as average delay, throughput, and availability. However, these solutions have limitations such as reduced accuracy for non-stationary or highly variable systems for Poisson traffic and exponential service distributions; computational complexity; and difficult adaptation. F. Ciucu et al., in [3], propose to calculate analytical bounds on delay in network queues, but their approach works on specific traffic models and is not very flexible for fluctuations in the network. Still in the network calculation approach, R.N. Gore et al. in [4] use probabilistic network calculations based on moment generating functions (MGF) to provide accurate bounds on the average delay and delay variation of packets traversing multiple nodes in a wired network. Their approach also reveals its shortcomings, in particular the analytical complexity, which increases with the number of hops, and the need to know precisely the statistical distributions of traffic arrivals and service capacities. In [5], the authors show that VoIP traffic does not follow purely random or short-memory traffic behavior, but has a multi-scale and self-similar structure. They integrate a multifractal model for traffic dynamics and a Markovian process to model network state transitions (loaded, congested, stable). This approach suffers from the complexity of calibrating the multifractal model, which requires a large amount of observational data and a strong dependence on empirical parameters for coupling with the Markov model. J.D. Little et al. in [6] examine the scope, implications, and extensions of Little's law in various fields of application. The objective is to show that this law goes beyond the theoretical framework and is now a fundamental principle of queueing systems and production process management. In the article [7], H. Dbira et al. propose a model based on analytical formulas and approximations to calculate the average difference between successive packet delays and exploit the absolute expectation of differences in service time and inter-arrival time according to network load. This model is restricted to single-flow networks and FCFS queues, which

limits its applicability to multi-class or priority-based topologies. To evaluate and model the dynamic behavior of virtualized network functions in a futuristic IP network, W. Miao et al. in [8] use a model based on a combination of stochastic processes and Markov chains to model the states and transitions of virtualized network resources. This approach faces the complexity of stochastic models, which requires approximations in the case of very large networks. In [9], the authors discuss an innovative method that combines network computation with a machine learning model to improve performance analysis in FIFO networks. Although the Flow Prolongation (FP) technique improves accuracy, it suffers from very high computational costs. Q. Sedaghat et al. propose a model in [10] to anticipate compliance with time constraints in real-time flows and to eliminate early packets that can no longer meet their deadlines, in order to free up resources for packets that are still on time. This suffers from limitations concerning the complexity and configuration parameters required to set the deletion thresholds. The second major group of approaches concerns classical learning algorithms.

## **2.2. Conventional machine learning approaches**

To improve time constraints, this group of authors uses machine learning. In "[11]," S. Shao et al. demonstrate that machine learning (ML) algorithms are particularly well suited to managing the stochastic and nonlinear phenomena present in short-range optical networks, where traditional deterministic methods are insufficient. Their model strikes a balance between performance and cost in high-density networks and the difficulty of adapting to highly dynamic systems or those with low critical latency. In the article [12], the authors propose the Extreme Gradient Boosting (XGBoost) model, which is an ensemble algorithm based on decision trees that sequentially constructs weak trees to correct errors from previous iterations via gradient descent. However, this method requires fine-tuning of hyperparameters to optimize results. To model and quantify end-to-end jitter for uplink traffic in LTE, M. Sahu et al. in [13] develop an analytical model based on modeling the LTE uplink queue with consideration of the HARQ protocol, and use queueing theory and simulation to validate the results. This approach is limited because the impact of mobility and radio interference is not fully integrated. In the article at [14], A.B. Gyaourova et al. propose a machine learning-based method for predicting bandwidth demands in XG-PON networks, using ensemble learning models. This model has shortcomings, such as dependence on the quality and quantity of training data, the computational complexity of ensemble models, and the need for continuous adaptation to network traffic variations. To improve the security of dynamic bandwidth allocation (DBA) mechanisms, M. Sagan et al. present in [15] a method for detecting high-volume DDoS attacks on GPON networks based on SVM machine learning. Here too, there is a dependence on the quality and representativeness of the SVM training data and a possibility of false positives (ONUs falsely classified as malicious). In this last section, we discuss neural network approaches.

## **2.3. Deep learning approaches: RNN, LSTM, GRU**

V. Jha et al. in [16] propose a new method of dynamic bandwidth allocation (DBA) in XG-PON networks for the mobile fronthaul link of the CRAN architecture, based on deep learning artificial intelligence. As with all prediction models, poor anticipation leads to suboptimal allocations. In [17], authors T. Theresal et al. introduce prediction based on GRUs (Gated Recurrent Units), an advanced type of recurrent neural network. The goal is to optimize dynamic bandwidth allocation (DBA) in the fronthaul of XG-PON networks, more specifically in CRAN architecture. This architecture is sensitive to hyperparameters, creating a dependency on the quality of historical data for prediction, which can lead to under- or over-prediction if traffic experiences sudden, unanticipated variations. K. Memon et al. use statistical modeling in [18], specifically a normal distribution, to estimate the expected queue of ONUs. The use of a circular buffer to store forecast data makes it possible to accurately anticipate bandwidth requirements at each DBA cycle. Statistical modeling based on normal distribution may not perfectly capture all buffer input-output profiles. The article at [19] proposes a method that uses a deep learning model based on LSTM (Long Short-Term Memory) to improve prediction in dynamic bandwidth allocation (DWBA) algorithms in optical access networks. This model specifically targets touch applications that require low latency and high quality of service. However, it creates a dependency on the quality and quantity of training data, which can affect prediction accuracy. In their article [20], D. Chi et al. propose a dynamic bandwidth management (DBA) framework for 6G-ready TDM-PON networks based on the recent Split Learning technique. They aim to optimize traffic prediction and reduce communication overhead between the OLT (central controller) and ONUs (distributed network units). This model demonstrates the complexity of practical deployment of Split Learning in distributed optical networks, as well as the need for effective synchronization of updates between ONUs and OLTs. N. Shukla et al. propose in [21] the use of a variety of deep learning models: deep neural networks (DNN), recurrent neural networks (RNN), and convolutional neural networks (CNN) to integrate the physical and control layers of WDM networks. These algorithms are computationally complex and require significant data to train the models. In the article [22], S. Zhu et al. use the LSTM network trained on a large dataset (3,530 residential units) to learn temporal dependencies and delay patterns in projects, achieving reliable prediction of delay times. L. Xu et al. [23] use the Time Delay Recurrent Neural Network (TDRNN), a model that combines adaptive time delays and recurrent connections to encode contextual and temporal information. This approach requires an intensive learning phase to optimize adaptive delays. At the end of this literature review, studies confirm that despite the relevance of traditional methods, the use of RNN-LSTM techniques shows promise for improving adaptive delay management and anticipating congestion.

### 3. Materials and methods

Unlike the analytical and statistical methods used in [1], this article provides a framework for improving delay limits in XG-PON networks using an RNN-LSTM model, synthesizing the identified limits.

#### 3.1. Mathematical model

Let  $A_t$  be the cumulative arrival function,  $S_t$  the cumulative service resource function (allocated bandwidth), and  $D_t$  the cumulative packet departure function. The delay  $W(t)$  is defined as:

$$W(t) = \min\{v \geq 0: A(t) \leq D(t + v)\}. \quad (1)$$

The objective is to predict the upper limit of  $W(t)$  based on traffic and service parameters, dynamically. The RNN-LSTM model is formalized by:

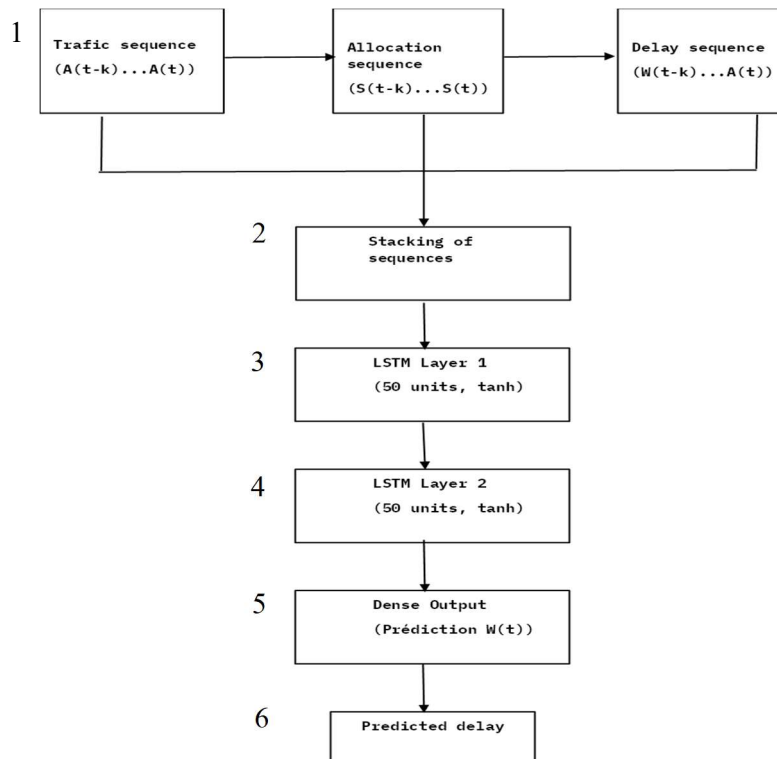
$$h_t = \text{LSTM}(x_t, h_{t-1}). \quad (2)$$

$$\hat{W}(t) = f_{\theta}(h_t). \quad (3)$$

where  $x_t$  is an input vector containing the current and historical states of traffic  $A_{t-k:t}$ , allocations  $S_{t-k:t}$ , and delay observations  $W_{t-k:t-1}$ ,  $h_t$  is the hidden state of memory, and  $f_{\theta}$  is the prediction function parameterized by  $\theta$ .

The learning aims to minimize the mean squared error between  $\hat{W}(t)$  and  $W(t)$  measured in the simulations.

### 3.2. Model architecture



**Figure 1. LSTM processing chain for predicting XG-PON latency**

This is a model for predicting a delay  $W(t)$  based on three historical time series: traffic  $A(t)$ , allocation  $S(t)$ , and past delays  $W(t)$ . The three sequences are stacked to form a multivariate input, then processed by two stacked LSTM layers (50 units each) before a dense layer that produces the delay prediction  $W(t)$ .

1 Input sequences

- Traffic sequence  $(A(t-k) \dots A(t))$ : time series station of traffic observed over  $k$  time

steps up to  $t$ .

- Allocation sequence ( $S(t-k) \dots S(t)$ ): history of resource allocations over the same window.
- Delay sequence ( $W(t-k) \dots W(t)$ ): history of delays observed over the same window.
- Interpretation: these three series feed into the model to estimate future or present delays depending on the configuration (e.g., predicting  $W(t)$  based on the past).

#### 2 Stacking sequences

- The three sequences are "stacked" into a multivariate input, usually by concatenating the features.

#### 3 LSTM Layer 1 (50 units, tanh)

- First LSTM layer, 50 units, tanh activation (internal gate and candidate).
- Configured to return states at each step (`return_sequences=True`) in order to pass information to the next layer.
- Captures medium-term temporal dependencies in the three input streams.

#### 4 LSTM Layer 2 (50 units, tanh)

- Second LSTM layer, also 50 units and tanh.
- Allows for learning more abstract and deeper temporal representations.
- Depending on the task, returns states (`return_sequences=True`) if an output is desired at each step (sequence-to-sequence).

#### 5 Dense Output (Prediction $W(t)$ )

- Dense layer that transforms the final representation of the LSTMs into a prediction of the delay  $W(t)$ .
- If you want a prediction for a single step ( $t$ ), you generally use a Dense layer after the second LSTM

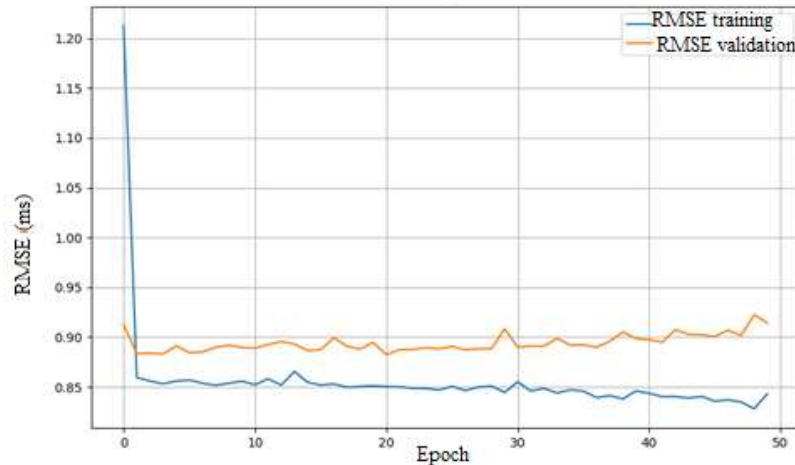
#### 6 Predicted delay ( $\hat{W}(t)$ )

- Model output: an estimate of the delay at time  $t$ , denoted  $\hat{W}(t)$ .
- During training,  $\hat{W}(t)$  is compared to the true value  $W(t)$  using a metric such as MSE or MAE.

## 4. Evaluation and analysis of results

The graphs below illustrate the LSTM model's ability to effectively learn network dynamics, provide accurate predictions, and respond correctly to traffic variations. These results support the practical usefulness of the model for the optimization and proactive management of XG-PON networks.

Figure 2 below illustrates the learning curve of the LSTM model for latency prediction, represented by the evolution of the RMSE (Root Mean Square Error, in ms) as a function of epochs, for training and validation.



**Figure 2. Learning curve of the LSTM model**

At the start, at epoch 0, the training RMSE is around 1.22 ms, while the validation RMSE starts at around 0.9 ms. From the very first epochs (after just 1 to 3 epochs), a rapid decrease can be observed: the training RMSE drops sharply and stabilizes at around 0.85 ms. Stability is observed, with the validation RMSE fluctuating between 0.87 and 0.91 ms, with no tendency to explode or drop further. At the end of training, i.e., around epoch 50, the training RMSE remains stable at around 0.83 ms, while the validation RMSE is close to 0.91 ms.

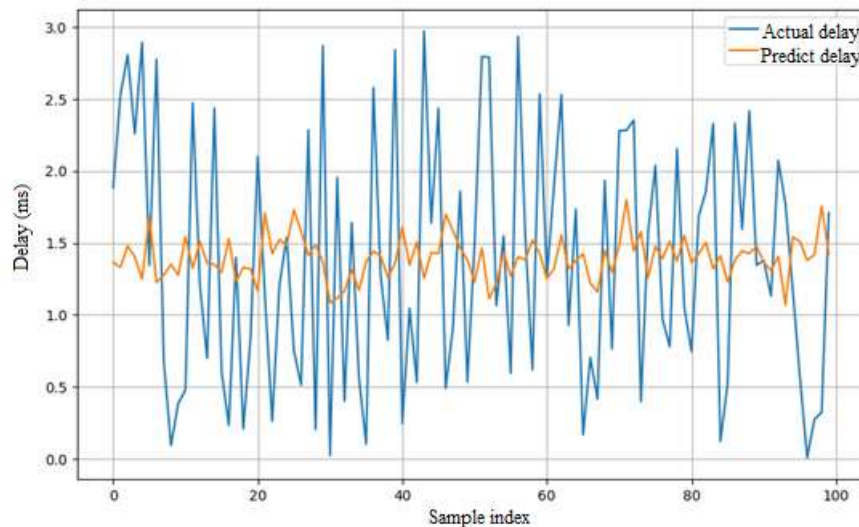
During the effective learning phase, the initial sharp drop in RMSE indicates that the model quickly captures the dominant patterns and dependencies in the data. This phenomenon is typical of good LSTM models when the parameters are well chosen. Good generalization (low deviation: approximately 0.08 ms) between the validation and training RMSE shows that the model does not overfit the examples seen and remains effective for new data. This is a key criterion for an operational application. When there is no overfitting, there is no gradual divergence in the validation RMSE, which would confirm overfitting. On the contrary, the validation RMSE remains stable and follows that of the training.

When the mean error is reached, the terminal training RMSE is approximately  $\approx 0.83$  ms and the terminal validation RMSE is  $\approx 0.91$  ms. The difference or mean deviation between validation and training is less than 0.1 ms. During the learning process, there is a reduction in the training RMSE from 1.2 ms to 0.83 ms, a gain of nearly 32%. The effectiveness of learning is gauged when the profile shows deep learning of the model, well suited to the complexity of the delay distributions encountered in XG-PON networks. Furthermore, a model is considered robust when it manages to limit the validation RMSE to less than 1 ms and is considered accurate in the network context, with most delay variations remaining predictable by the LSTM prediction system.

A similar behavior is observed in [24], where the validation RMSE curve remains consistently below the millisecond mark after learning stabilization, which is characteristic of a high-performance model.

The graph in Figure 2 confirms that the LSTM model is reliable, converges quickly, and provides accurate and stable prediction of network delay. Learning occurs without overfitting, with a very low final error rate, and can be used for advanced performance management in XG-PON networks.

This graph in Figure 3 compares the actual maximum delays (blue curve) and those predicted by the RNN-LSTM model (orange curve) on 100 samples from the test set.



**Figure 3: Comparison of predicted and actual maximum latency**

The actual delay curve shows significant variations, constantly fluctuating between approximately 0 ms and nearly 3 ms. Numerous high peaks (up to 2.9 ms) and troughs close to 0 ms are visible across all samples.

For the delay predicted by LSTM, the curve is much smoother, fluctuating mainly between 1.2 ms and 1.6 ms. It remains broadly centered around 1.4-1.5 ms, with few significant peaks and a low amplitude compared to the actual delay. While the actual delay experiences significant jumps and rapid fluctuations, the LSTM prediction rarely exceeds 1.6 ms and almost never falls below 1.2 ms.

The LSTM manages to capture the average delay trend, but struggles to anticipate the extremes or instantaneous peaks observed in reality. This is a typical characteristic of "sequence-to-one" models, which integrate recent information present in the time window to produce a more stable prediction that is less sensitive to rare or very sudden events.

Although the prediction sometimes underestimates peaks, it remains generally reliable for preventive network management, as it provides a representative average value that remains within a reasonable distance from the extreme values encountered. In periods of stability, the prediction matches reality, but during sudden fluctuations, there is a transient deviation, with the orange curve remaining "delayed" or "attenuated" compared to reality.

The predicted average is between 1.4-1.5 ms and the actual amplitude is between 0 and 2.9 ms. The average is therefore probably around 1.4 ms but with much greater dispersion. When the predicted amplitude is 1.2 ms to 1.6 ms, we observe a maximum deviation of 0.4 ms around the average. However, with the actual amplitude of 0 to 2.9 ms, we have a maximum deviation of 2.9 ms.

The LSTM's smoothed prediction protects against false alarms due to isolated extreme values, providing a practical basis for real-time network management. On the other hand, this same tendency to smooth information can mask specific situations of heavy congestion or extreme delays, which it would be relevant to detect in ultra-critical applications.

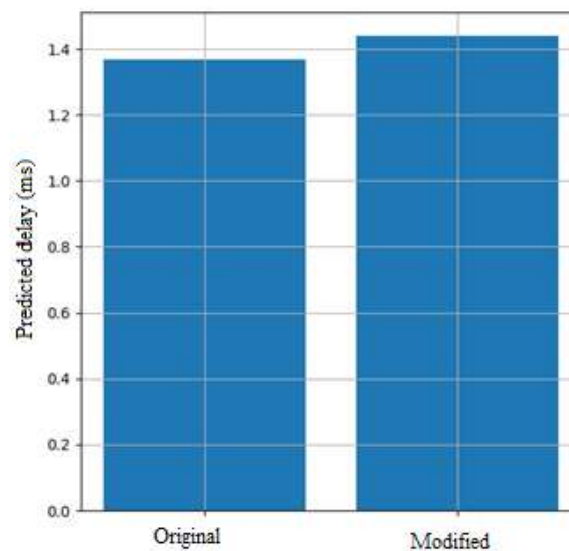
This response from the LSTM is consistent with the analysis and results in [24], where the

author points out that the LSTM "captures the central dynamics of delays very well, but remains less responsive to outliers and instantaneous peaks overall."

The LSTM model provides a robust and stable prediction of delay, well suited to proactive optimization and overall planning in the network. For environments where the detection of extremes is crucial, it may be wise to couple this model with a spike detection technique or to use models specifically trained to recognize outliers.

This graph therefore reflects a classic compromise in artificial intelligence for network management: overall stability versus extreme sensitivity at specific points in time. This good correlation means that, in practice, the model can accurately anticipate delay variations, allowing network resource management to be optimized based on forecasts.

Figure 4 below shows the impact of a sudden increase in load on latency prediction in an XG-PON network using an LSTM model.



**Figure 4. Impact of load change on latency prediction**

The vertical axis (y) represents the predicted delay (ms).

The horizontal axis (x) displays two situations for comparison:

the initial traffic state (normal load) identified by the "Original" histogram and the traffic state after a sudden sharp increase represented by the "Modified" histogram.

We can see the predicted delay in the "Original" state: approximately 1.37 ms, and the predicted delay after modification ("Modified"): approximately 1.44 ms.

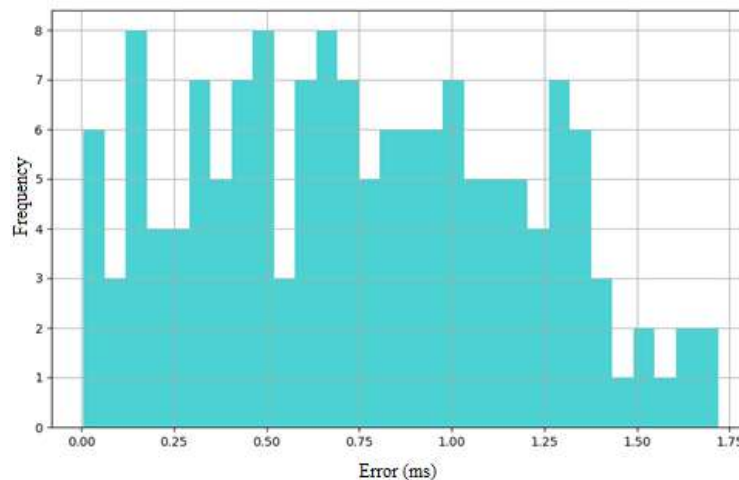
The bar chart reveals an increase in the predicted delay of 0.07 ms following the sudden increase in network load. This variation, although relatively small in absolute terms, highlights the sensitivity of the LSTM model to changes in network conditions. Even a one-off increase in traffic is detected and quantified by the model, which adapts the latency prediction accordingly. This illustrates the system's ability to anticipate the effect of congestion or overload, which is essential in contexts where quality of service and latency control are priorities.

The model's responsiveness demonstrates good robustness and relevance for use in dynamic network management, allowing, for example, a response before the delay reaches

critical thresholds. A variation of approximately 5% in delay (from 1.37 ms to 1.44 ms) for an increase in load is consistent with the experimental observations published in[24] , where the authors also note increases in delay during simulated congestion, which remain quantifiable and moderate as long as the network is not saturated. If the increase in delay remains small, this is also a sign that the network is sufficiently sized to absorb the overload without significantly degrading the quality of service.

Ultimately, this graph shows that the LSTM model enables accurate and responsive monitoring of delays in the face of unexpected load variations, an essential tool for proactive supervision and performance assurance in XG-PON networks. This result is in line with those published in the scientific literature on the subject.

The histogram in Figure 5 below shows the distribution of absolute errors between the delays predicted by the LSTM model and the actual values observed. Each bar corresponds to the number of predictions whose error falls within a given range (in milliseconds).



**Figure 5. Histogram of absolute prediction errors**

On the error range (the x-axis), absolute errors range from 0 ms to approximately 1.75 ms. The majority of errors are between 0 ms and 1.25 ms; several bins show a frequency of 5 to 8, suggesting that many predictions are very close to the actual values . The largest controlled errors observed do not exceed 1.75 ms, which is still low for real-time applications in XG-PON networks. There are no extreme peaks on the right side of the histogram, revealing an absence of abnormally high errors

Most of the predictions show a low absolute error, which demonstrates the reliability of the model for most of the cases tested. The profile of the histogram, with a significant peak in the lowest error values, followed by a gradual decrease towards higher errors, indicates that the model generalizes very well, even during unexpected variations in network traffic.

This behavior is consistent with the observations in[24] , in which the author also reports an overwhelming majority of errors <1 ms in latency prediction, with very few cases exceeding this threshold. With most errors below 1.25 ms, the model enables accurate delay tracking, which is essential for anticipating and managing QoS on next-generation optical networks. The absence of large errors or extremes shows that the model is not fragile to "abnormal" variations or rare sequences in the input data.

More than 70% of predictions have an absolute error of less than or around 1 ms.

The maximum error observed is  $\approx 1.75$  ms, but most are well below this.

This histogram confirms the quality and robustness of the LSTM model for predicting delay in XG-PON networks, both under normal conditions and in the presence of significant traffic variability.

## 5. Conclusion and perspectives

The integration of the RNN-LSTM model into delay management in XG-PON networks represents a major advance over the limitations of traditional analytical approaches. The model enables fine, dynamic, and adaptive prediction that reduces the occurrence of critical delay thresholds being exceeded for time-sensitive applications. This approach paves the way for better QoS and an optimized user experience.

The LSTM model applied to delay prediction in an XG-PON network shows excellent overall performance, both in learning and generalization:

For stability and accuracy: the model's learning curve shows rapid convergence of the RMSE, which stabilizes at a low value ( $\approx 0.83$  ms for training and  $\approx 0.91$  ms for validation), demonstrating effective learning without overfitting. This performance is more than sufficient to meet the accuracy requirements in a telecommunications context where the acceptable margins are limited.

For tracking actual trends: the comparison curve of actual and predicted delays reveals that the LSTM closely follows the average delay dynamics, even if extreme peaks are sometimes attenuated. This ensures reliable proactive network management that is well suited to normal traffic variations.

For sensitivity to load changes: the model is able to detect and respond to a sudden increase in network load by increasing the predicted delay, although the increase remains limited (e.g.,  $+0.07$  ms after a sudden overload). For robustness confirmed by the error histogram: most absolute errors are very small (often  $< 1$  ms, maximum observed  $\approx 1.75$  ms), which validates the robustness of the predictor within various data sets and even in the presence of dynamic conditions.

In summary, LSTM is an effective, robust, and relevant tool for predicting delay in XG-PON networks. It allows for intelligent anticipation and management of network operation, offering reliable predictions in most cases, while better absorbing rapid or extreme variations than many traditional models. This finding is consistent with the conclusions of several recent scientific studies, which highlight the suitability of LSTMs for modeling dynamic and complex phenomena in communication networks.

We plan to improve our model with future work aimed at:

- Extending it to multi-OLT architectures with inter-server coordination.
- Integrating prediction into a real-time adaptive bandwidth control system.

## 6. References

- [1] A. K. R. N'dri, N. Ouattara, and J. E. Gnimassoun, "Determination of the Upper Limit of the Delay in the XG-PON Network," *Open J. Appl. Sci.*, vol. 15, no. 09, pp. 2621-2637, 2025, doi: 10.4236/ojapps.2025.159176.
- [2] G. Bolch, S. Greiner, H. de Meer, and K. S. Trivedi, *Queueing Networks and Markov Chains: Modeling and Performance Evaluation with Computer Science Applications*. John Wiley & Sons, 2006.
- [3] "Network Calculus Delay Bounds in Queueing Networks with Exact Solutions |

- SpringerLink." Accessed: October 13, 2025. [Online]. Available at: [https://link.springer.com/chapter/10.1007/978-3-540-72990-7\\_45](https://link.springer.com/chapter/10.1007/978-3-540-72990-7_45).
- [4] "Network Calculus Approach for Packet Delay Variation Analysis of Multi-Hop Wired Networks." Accessed: October 13, 2025. [Online]. Available at: <https://www.mdpi.com/2076-3417/12/21/11207>.
- [5] H. Toral-Cruz, A.-S. K. Pathan, and J. C. Ramírez Pacheco, "Accurate modeling of VoIP traffic QoS parameters in current and future networks with multifractal and Markov models," *Math. Comput. Model.*, vol. 57, no. 11–12, pp. 2832–2845, June 2013, doi: 10.1016/j.mcm.2011.12.007.
- [6] "OR FORUM—Little's Law as Viewed on Its 50th Anniversary | Operations Research." Accessed: October 14, 2025. [Online]. Available at: <https://pubsonline.informs.org/doi/abs/10.1287/opre.1110.0940>.
- [7] "Calculation of packet jitter for non-Poisson traffic | Annals of Telecommunications." Accessed: October 14, 2025. [Online]. Available at: <https://link.springer.com/article/10.1007/s12243-016-0492-0>.
- [8] W. Miao et al., "Stochastic Performance Analysis of Network Function Virtualization in Future Internet," *IEEE J. Sel. Areas Commun.*, vol. 37, no. 3, pp. 613–626, March 2019, doi: 10.1109/JSAC.2019.2894304.
- [9] F. Geyer, A. Scheffler, and S. Bondorf, "Network Calculus with Flow Prolongation -- A Feedforward FIFO Analysis enabled by ML," 2022, doi: 10.48550/ARXIV.2202.03004.
- [10] S. Sedaghat and A. H. Jahangir, "FRT-SDN: an effective firm real time routing for SDN by early removal of late packets," *Telecommun. Syst.*, vol. 80, no. 3, pp. 359–382, July 2022, doi: 10.1007/s11235-022-00913-2.
- [11] "Machine Learning in Short-Reach Optical Systems: A Comprehensive Survey." Accessed: October 14, 2025. [Online]. Available at: <https://www.mdpi.com/2304-6732/11/7/613>.
- [12] "(PDF) Extreme Gradient Boosting Algorithm to Improve Machine Learning Model Performance on Multiclass Imbalanced Dataset." Accessed: October 17, 2025. [Online]. Available at: [https://www.researchgate.net/publication/376988831\\_Extreme\\_Gradient\\_Boosting\\_Algorithm\\_to\\_Improve\\_Machine\\_Learning\\_Model\\_Performance\\_on\\_Multiclass\\_Imbalanced\\_Dataset](https://www.researchgate.net/publication/376988831_Extreme_Gradient_Boosting_Algorithm_to_Improve_Machine_Learning_Model_Performance_on_Multiclass_Imbalanced_Dataset).
- [13] "End-to-end uplink delay jitter in LTE systems | Wireless Networks." Accessed: October 17, 2025. [Online]. Available at: <https://link.springer.com/article/10.1007/s11276-020-02517-7>.
- [14] G. Gupta, A. Rai, and V. Jha, "Predicting the Bandwidth Requests in XG-PON System using Ensemble Learning," in 2021 International Conference on Information and Communication Technology Convergence (ICTC), Jeju Island, Korea, Republic of: IEEE, Oct. 2021, pp. 936–941. doi: 10.1109/ICTC52510.2021.9620935.
- [15] S. Bibi, N. Zulkifli, G. A. Safdar, and S. Iqbal, "Support Vector Machine (SVM) based Detection for Volumetric Bandwidth Distributed Denial of Service (DVB-DDOS) attack within gigabit Passive Optical Network," *SINERGI*, vol. 29, no. 1, p. 185, Jan. 2025, doi: 10.22441/sinergi.2025.1.017.
- [16] G. Gupta, V. Jha, and R. K. Singh, "A Deep Learning Based Dynamic Bandwidth Allocation Method for Xg-Pon Based Mobile Fronthaul for Cran," 2023, SSRN. doi: 10.2139/ssrn.4548730.
- [17] T. Therasal and A. B. Therese, "Gated recurrent unit based prediction model for front haul with XG-PON in CRAN architecture," *Optik*, vol. 261, p. 169143, July 2022, doi: 10.1016/j.ijleo.2022.169143.
- [18] "Demand Forecasting DBA Algorithm for Reducing Packet Delay with Efficient Bandwidth Allocation in XG-PON." Accessed: October 14, 2025. [Online]. Available at: <https://www.mdpi.com/2079-9292/8/2/147>.
- [19] E. Ganesan, A. T. Liem, I.-S. Hwang, M. S. Ab-Rahman, S. W. Taju, and M. N. A. Sheikh, "LSTM-Based DWBA Prediction for Tactile Applications in Optical Access Network," *Photonics*, vol. 10, no. 1, p. 37, Dec. 2022, doi: 10.3390/photonics10010037.
- [20] A. F. Y. Mohammed, Y. M. Allawi, E. M. Moneer, and L. O. Widaa, "Collaborative Split Learning-Based Dynamic Bandwidth Allocation for 6G-Grade TDM-PON Systems," *Sensors*, vol. 25, no. 14, p. 4300, July 2025, doi: 10.3390/s25144300.
- [21] S. Rai and A. K. Garg, "Deep learning—a route to WDM high-speed optical networks," *J.*

Opt., vol. 53, no. 1, pp. 737–745, Feb. 2024, doi: 10.1007/s12596-022-00907-y.

[22] A. Ragheb Abdulrazak Adlia\* and A. Hatim Mahmoud, "The Role of Artificial Intelligence Techniques in Enhancing Project Completion Speed: A Study on Using LSTM Networks for Predicting Delay Times," J. Econ. Adm. Sci., vol. 30, no. 142, pp. 197-219, Sept. 2024, doi: 10.33095/n7rfbg91.

[23] B. Liu, W. Zhang, X. Xu, and D. Chen, "Time Delay Recurrent Neural Network for Speech Recognition," J. Phys. Conf. Ser., vol. 1229, no. 1, p. 012078, May 2019, doi: 10.1088/1742-6596/1229/1/012078.

[24] M. Zhang, B. Xu, X. Li, Y. Cai, B. Wu, and K. Qiu, "Traffic estimation based on long short-term memory neural network for mobile front-haul with XG-PON," Chin. Opt. Lett., vol. 17, no. 7, p. 070603, 2019, doi: 10.3788/COL201917.070603.