

Probabilistic short-term load forecasting models based on a parametric approach

L.A. Fernández-Jiménez, S. Terreros-Olarte, C. Capellán-Villacian, A. Falces, E. García-Garrido, P.M. Lara-Santillán, M. Mendoza-Villena, E. Zorzano-Alba, P.J. Zorzano-Santamaría

Department of Electrical Engineering
E.T.S.I.I., University of La Rioja
C/ San José de Calasanz 31, 26004 Logroño (Spain)
Phone number: +0034 941 299473

Abstract. The forecasting of electric load plays an essential role in the effective management of electric power systems. Specifically, short-term models, which predict hourly load over the 24 hours of the subsequent day, hold significant value for applications within the realm of electricity markets. In this context, research efforts have predominantly concentrated on the development of point load forecasting models. These models provide solely the estimated value of hourly load, omitting any information regarding its associated uncertainty. Probabilistic forecasting models aim to address this limitation by offering comprehensive information on forecasted values, including their associated uncertainty, thereby enabling their more effective utilization in risky decision-making environments. This paper presents a parametric probabilistic model designed for hourly load forecasting. The model is refined through a multi-objective genetic algorithm optimization process that identifies explanatory variables from a specified set. The selected variables are combined linearly to predict the parameters of a probability distribution function for the hourly load. The process also selects the type of distribution from among those characterized by two parameters. The model is applied to data from a real distribution substation, yielding superior forecasting evaluation indexes compared to those achieved by two benchmark probabilistic load forecasting models.

Key words. Short-term load forecasting, probabilistic forecasting, multi-objective optimization, model calibration, uncertainty

1. Introduction

Short-term load forecasting models are essential tools in the management and operation of electric power systems, primarily due to their ability to predict future energy demand accurately. The main objective of these models is to forecast future electricity demand over a short term, spanning from the next few hours up to the forthcoming week. This forecasting is indispensable for various operational and strategic applications, including generation dispatch planning, grid operation, scheduled maintenance, balance supply and demand, minimize production costs and the effective management of power purchases and sales within electricity markets [1,2].

Most of the short-term load forecasting models presented in the international literature are categorized as deterministic or point models. Such models solely provide the expected load value for the future instant. It is evident that every forecast is subject to an error, a factor that deterministic models fail to assess. On the contrary, in recent years, a new category of load forecasting models has emerged. These models are designed to furnish information for assessing the uncertainty associated with the forecasted value. Specifically, they are known as probabilistic load forecasting models. Probabilistic models can provide forecasting intervals, which define the range within which the real value is expected to lie at a specified confidence level, quantiles, or a probability density functions [3]. In decision-making contexts characterized by uncertainty, this type of information offers significantly more value than that derived from deterministic forecasting models [4]. Probabilistic models that yield the probability density function of the variable to be predicted as a result of the prediction process offer the most comprehensive information.

In the development of probabilistic forecasting models, two distinct approaches can be adopted: parametric and non-parametric. The parametric approach presupposes the distribution type that the variable to be predicted will follow, whereas the non-parametric approach does not rely on such an assumption. The parametric approach primarily utilizes a technique known as “dressed model”, which involves superimposing a probability density function over the value provided by a deterministic forecasting model. In this technique, the deterministic model’s prediction serves as the mean of the distribution. In general, models based on the non-parametric approach often provide better forecasting results than those based on parametric models. However, we believe that the parametric approach, improved with the application of optimization techniques, can provide results comparable to those obtained with non-parametric approach models in the probabilistic short-term load forecast.

In the early research works focused on analysing probabilistic forecasting models within the electric power

sector, researchers identified the three primary properties that must characterize these models: Reliability, sharpness and resolution [5]. The reliability property is also known as calibration. Reliability measures the alignment between the observed and predicted distributions of the forecast variable, indicating the accuracy of the forecast. Sharpness, on the other hand, assesses the concentration of the forecasted distribution; a more concentrated (narrower) distribution signifies greater sharpness, reflecting reduced forecast uncertainty. Resolution describes the model's ability to produce distinct forecasts for different conditions while ensuring each forecast remains conditionally reliable.

In the majority of studies focused on probabilistic forecasting models for the electricity sector (load forecasting, wind or photovoltaic power production, and electricity market price forecasting), there is a notable lack of emphasis on reliability. This is despite the widely acknowledged principle that these models should aim to optimize sharpness while being subject to proper calibration [6].

In this paper, we present the results obtained in the development and evaluation of probabilistic short-term load forecasting models for a distribution substation. Our methodology adopts a parametric approach, presupposing the probability density function (PDF) that the predicted variable, hourly load, will follow. The prediction horizon of the models includes the 24 hours of the day following to the one on which the predictions are carried out, enabling the use of the forecasts in day-ahead electricity markets. The models utilize data comprising forecasts of principal weather variables related to the substation's location, hourly load data over the previous days and a set of dummy variables related to the type of day for which the forecast is carried out. A multi-objective genetic algorithm guides an optimization process with a particular focus on the reliability of models, aiming to identify the optimal model that balances accuracy with calibration effectively. This process facilitates the selection of the most suitable probability density function (PDF) and the determination of parameter values for these PDFs, drawing from a pool of available explanatory variables. The forecasting results obtained with the optimized models are compared to those obtained with two benchmark probabilistic forecasting models, proving the superiority forecasting performance of the proposed model.

2. Selected parametric models

The family of models selected for this task are the Generalized Additive Models for Location, Scale, and Shape (GAMLSS) [7]. GAMLSS represent a sophisticated class of statistical models. They provide exceptional flexibility by enabling the modelling of data across multiple distribution parameters (mean, variance, shape, among others) through smooth additive functions. These functions effectively capture non-linear relationships between variables. Unlike traditional models, GAMLSS facilitate the fitting of a diverse range of distributions to the response variable. This capability is particularly beneficial for analysing data characterized by complex behaviours, such as heteroscedasticity (varying variance) or asymmetries. In

a GAMLSS model with two parameters, the first parameter, denoted as μ , represents the conditional mean of the response variable, which is contingent upon the set of predictors incorporated within the model. The second parameter, σ , signifies the dispersion or scale of the distribution of the response variable, a factor directly associated with its variance. While more sophisticated GAMLSS models can incorporate distribution functions with additional parameters to account for skewness and kurtosis, thus offering more sophisticated distributive characteristics, these elements have not been considered within the scope of this current study.

In this study, we specifically focused on selecting a two-parameter probability distribution function, characterized by parameters μ and σ , through the exclusive use of linear functions for both parameters. This approach entails that μ and σ of the distributions are modelled as linear functions of the selected explanatory variables. Our analysis did not incorporate the use of smoothers for these variables. The restriction to linear models was deliberate, aiming to elucidate the direct linear relationships within the data, thereby simplifying the interpretation and application of the results.

Thus, the GAMLSS models with two parameters used in this study can be defined by equations (1) to (3),

$$y \sim \mathbf{D}(\mu, \sigma) \quad (1)$$

$$g_1(\mu) = \beta_{01} + \beta_{11}x_1 + \beta_{21}x_2 + \dots + \beta_{m1}x_m \quad (2)$$

$$g_2(\sigma) = \beta_{02} + \beta_{12}x_1 + \beta_{22}x_2 + \dots + \beta_{n2}x_n \quad (3)$$

where y represents the dependent variable (hourly load), $\mathbf{D}$ represents a family of distributions parametrized by μ and σ , which, in turn, represent the parameters of location and scale, β represent the linear coefficients, x_i the explanatory variable i , m and n the number of explanatory variables selected for each parameter, and g_1 and g_2 represent the link functions. These latter functions can be identity, logarithmic or logit functions, depending on the modelled distribution.

3. Evaluation indexes

The main index for evaluating a probabilistic forecast is known as the Continuous Ranked Probability Score (CRPS). The CRPS serves as a comprehensive metric that simultaneously assesses both reliability and resolution in probabilistic forecasting models [8]. The CRPS, which mathematical calculation is expressed in equation (4), quantifies the discrepancy between the cumulative distribution function (CF) of forecasted probabilities and the CF of the observed value by integrating across the entire range of possible outcomes for the forecast variable. In situations where the variable to be forecasted is represented by empirical distributions, as in the case of forecasts derived from non-parametric approaches that provide a set of quantiles, the CRPS can be computed using equation (5), where y represents the real values, q_i the value of the quantile i and M the number of quantiles per forecast [9]. Equation (4), applicable to forecasts presented as CFs (or PDFs), and equation (5), relevant for forecasts articulated through a set of quantiles, both yield the instantaneous

value of the CRPS. From the instantaneous values recorded during the evaluation period, the average CRPS is computed. This average CRPS serves as an indicator of the model's probabilistic forecasting performance and is used to compare models. A lower average CRPS indicates a more accurate probabilistic forecast.

$$\text{CRPS}(CF, y) = \int_{\mathbb{R}} [CF(x) - \mathbb{1}(x \geq y)]^2 dx \quad (4)$$

$$\text{CRPS} = \frac{1}{M} \sum_{i=1}^M |q_i - y| - \frac{1}{2M^2} \sum_{i,j=1}^M |q_i - q_j| \quad (5)$$

Although the CRPS allows a probabilistic prediction to be correctly evaluated, another indicator is needed to assess the reliability or calibration of the forecast. Reliability refers to the statistical consistency between forecasted outcomes and actual observations. The reliability of a forecasting model can be evaluated by means of the Probability Integral Transform (PIT) histogram. The PIT value is derived by computing $CF(x)$ at a specific verification point x , which corresponds to an observed value of the dependent variable. The obtained PIT value varies from 0 to 1, representing the quantiles of the distribution. By comparing the PIT values with their corresponding verification points, one can assess the level of calibration or statistical consistency between the forecasted probabilities and the observed outcomes. In an ideal scenario, perfect calibration is attained when the PIT distribution is uniform, culminating in a flat histogram when utilizing evenly spaced intervals, also known as bins. In our analysis, we employ the Reliability Index (RI) [10], which is calculated based on the frequency with which the actual value of the variable to be predicted falls within one of the intervals represented in the PIT histogram. The RI is determined using equation (6), where K denotes the number of evenly spaced intervals, or bins, utilized in the histogram, and κ_k represents the frequency of instances where the value of the outcome variable falls within the range specified by bin k . A lower value of the RI suggests a model that is better calibrated.

$$\text{RI} = \sum_{k=1}^K \left| \kappa_k - \frac{1}{K} \right| \quad (6)$$

In addition, to evaluate the deterministic (point) forecasts provided by the models, we will use the root mean square error (RMSE), mean absolute error (MAE) and mean absolute percentage error (MAPE). These indexes are defined by equations (7) to (9), where, $\hat{P}(t)$ represents the load demand point forecast for hour t , $P(t)$ the actual value and N indicates the total number of hours in the evaluation period. The values for the load demand point forecast corresponded to the mean of the distribution corresponding to each hour for the case of GAMLSS models or to the quantile 0.5 value for the case of non-parametric models.

$$\text{RMSE} = \sqrt{\frac{1}{N} \sum (\hat{P}(t) - P(t))^2} \quad (7)$$

$$\text{MAE} = \frac{1}{N} \sum |\hat{P}(t) - P(t)| \quad (8)$$

$$\text{MAPE} = 100 \frac{1}{N} \sum \left| \frac{\hat{P}(t) - P(t)}{P(t)} \right| \quad (9)$$

4. Case Study

The objective is to obtain a probabilistic forecasting model of the hourly load in a power substation 66/13.2 kV feeding a small town of about 5000 inhabitants located in the north of Spain. The consumption fed from this substation is mainly industrial, although the number of residential customers is near 3000. The forecast horizon includes the 24 hours of the next day, day d , with the forecast being carried out in the first hours of day $d-1$. The dataset, comprising hourly consumption records spanning 30 consecutive months, serves as the base for the forecasting models' development. The hourly load dataset was expanded with the values from a set of potential explanatory variables.

A set of 55 explanatory variables was employed in the analysis, encompassing a diverse range of factors as shown in Table I. These include the forecasted values of seven weather variables (V1 to V7), a logical indicator reflecting the official time for the forecast horizon, distinguishing between summer and winter time (V8), and the hourly load values at the same hour on the preceding six days (V9 to V14). Additionally, the model incorporates dummy variables to represent the hour of the day (V15-V37), the day of the week (V38-V43), the month of the year (V44 to V54), and a national, regional, or local holiday (V55).

Table I. – Available explanatory variables

VARIABLE	DESCRIPTION
V1	Temperature (K)
V2	Global horizontal irradiance (W/m ²)
V3	Wind speed (m/s)
V4	Relative humidity (per unit)
V5	Total cloud cover (per unit)
V6	Rainfall (kg/m ²)
V7	Snow (kg/m ²)
V8	European summer time (logical)
V9–V14	Load (kW) lag i hourly load lagged “ i ” hours, i
V15–V37	Dummy variables for the hour of the day
V38–V43	Dummy variables for day of the week
V44–V54	Dummy variables for the month of the year
V55	Dummy variable for holiday

The forecasts for weather variables (V1 to V7) were sourced from Meteogalicia, the meteorological service of Galicia. Utilizing a numerical weather prediction model, Meteogalicia issues daily forecasts in the early hours for these variables for all hours in the following three days. These forecasts are distributed across points of an analysis grid, which spans the entire Iberian Peninsula with grid points approximately 12 km apart. The values of variables V1 to V7 were determined by calculating the weighted average of the forecasts for the four points on the analysis grid nearest to the urban centre of the considered town. The weighting factor was based on the distance between the

urban centre and each point on the analysis grid. Variables V9 to V14 represent the hourly load recorded at the corresponding hour for which the forecast is generated, covering the period from days d-2 to d-7. It is important to note that data for the same time on day d-1 are not available when the forecasts are carried out, as these forecasts are conducted in the early hours of day d-1.

The complete dataset was partitioned into a training dataset and a testing dataset. The training dataset encompassed the data of the initial 24 months, while the testing dataset included the data of the subsequent 6-month period.

The aim of the research presented in this paper was to develop a probabilistic model able of predicting hourly load using a parametric approach. Eight two-parameter distributions were considered for model selection: Gamma, Inverse Gamma, Logistic, Log-Normal, Normal, Gumbel, Reverse Gumbel, and Weibull. All these distributions are characterized by two parameters, the location parameter, μ and the scale parameter, σ . Clearly, the objective is to identify the model that yields the most accurate probabilistic forecasts. This involves selecting from among eight possible distributions and determining the optimal linear combinations of the explanatory variables for both parameters.

To determine the most accurate probabilistic forecasting model for the hourly load demand across each of the 24 hours of the subsequent day, a selection process was implemented, controlled by a multi-objective genetic algorithm aimed at optimizing two objectives: the minimization of the average CRPS and the minimization of the RI in the evaluation period. Model selection is exclusively based on data from the training set, employing a 5-fold cross-validation procedure. This approach involves dividing the training dataset into five distinct subsets. In an iterative process, each subset is alternatively used as the evaluation period while the GAMLSS models are fitted using the data from the remaining four subsets. This process yields five values for the average CRPS and for the RI. The fitness function's output values are the means of these five obtained values for CRPS and RI, respectively.

In the genetic algorithm, each individual or potential solution was represented by a 110-element vector of real numbers in the range 0 to 1, with each element signifying the contribution of the corresponding variable to the model. Specifically, these elements indicated the exclusion (value under 0.5) or selection (value equal to or greater than 0.5) of each variable for constructing the linear combination of explanatory variables involved in the value of the μ parameter (the first 55 elements) and in the value of the σ parameter (the last 55 elements).

The multi-objective genetic algorithm used was the NSGA-II [11], with a total of 200 generations and 100 individuals per generation. The optimization process was conducted for each of the eight distributions, resulting in eight sets of non-dominated solutions after the iterative processes. These solutions represent linear combinations of the explanatory variables for the two parameters that characterize each distribution.

The optimization process for each distribution yielded a set of 100 non-dominated solutions. The eight sets were brought together to form a single consolidated set. From this consolidated set, the global non-dominated solutions were then identified. This process resulted in a reduced set of 131 non-dominated solutions, exclusively corresponding to three of the eight distributions: Gamma, Inverse Gamma, and Log-Normal. The selection of the model from the 131 candidates identified as non-dominated global solutions was based on the normalization of the CRPS and RI between their respective maximum and minimum values in the consolidated set of non-dominated solutions. The model that exhibited the lowest aggregate value of these two normalized indexes was chosen, indicating an equitable weighting between the two indexes. This optimal (proposed) model was a Gamma distribution-based model which used reduced sets of the explanatory variables for the calculation of the two parameters. Specifically, it used 51 of the explanatory variables to obtain the value of the μ parameter and only 30 to obtain the value of the σ parameter.

Based on the t-statistic, the assessment of the significance of the explanatory variables reveals that the load from the same hour seven days prior and the dummy variable for holidays are the two most crucial variables in calculating the μ parameter. Conversely, for the sigma parameter calculation, the most significant variables are the holiday indicator and relative humidity. Temperature and dummy variables denoting the hours of the day also significantly influenced both parameters. Notably, the dummy variable for Sundays exhibited a peculiar pattern; it was the fourth most significant variable for the μ calculation but was the ninth for σ .

After selecting the model characteristics, specifically, the distribution and the linear combinations of explanatory variables for determining the values of the two parameters, based on the outcomes of the 5-fold cross-validation procedure, the model was fitted using the entirety of the data from the training set. Subsequently, this final model was applied to the testing dataset, which had remained unused until this stage. Figure 1 shows the PDFs associated with probabilistic forecast of the hourly load for two distinct hours within the testing period. These two functions illustrate the different forecasting uncertainty level for each hour. Specifically, the forecast corresponding to the PDF represented in blue exhibits a higher degree of uncertainty compared to its counterpart in red. This heightened uncertainty correlates with a diminished sharpness in the blue PDF's curve. Additionally, the figure includes circular markers on both PDFs, denoting the actual hourly demand values for the respective hours.

To evaluate the effectiveness of the forecasts generated by the proposed model, two benchmark models were developed for comparative analysis. The first benchmark model, REF1, shares similar characteristics with the proposed model in the sense that it uses the same technique, that is, it is a GAMLSS model characterized by two parameters, μ and σ , but the optimization process is restricted to the selection of the distribution function. Specifically, REF1 uses linear combinations of all the

available explanatory variables for the formulation of both parameters. The second benchmark model, REF2, utilizes a non-parametric approach, using the quantile regression technique and incorporating all available explanatory variables. Quantile regression involves the estimation of parameters for functions that correlate the quantiles of a dependent variable with its explanatory variables. When these functions assume a linear form, the resulting model is referred to as a Linear Quantile Regression model [12].

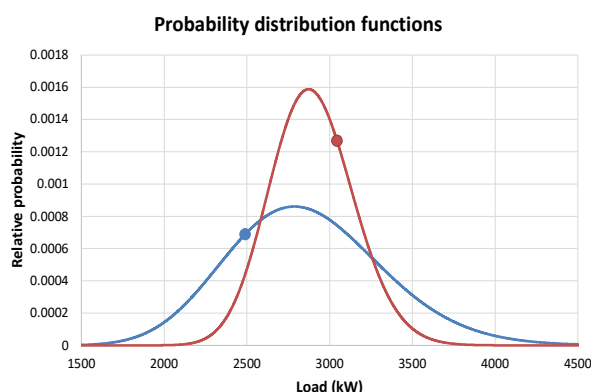


Fig. 1. PDFs corresponding to the probabilistic forecasts of two hours of the testing dataset.

The development of the REF1 model involved the selection of one of the eight possible distributions. The selection was carried out with a 5-fold cross-validation procedure using the training dataset. The selected model had a gamma distribution function. So, REF1 corresponded to a GAMLSS model with gamma distribution and used all the available explanatory variables in the formulation of the values of the two parameters. The REF2 model was also developed with the data of the training dataset using all available explanatory variables. This model facilitated the derivation of linear combinations of these variables to estimate the values for 99 quantiles of the dependent variable, ranging from 0.01 to 0.99 in 0.01 increments. Both benchmark models were subsequently applied in the probabilistic forecast in the test period.

Table II. – Forecasting results with the testing dataset

	OPTIMIZED MODEL	REF1 MODEL	REF2 MODEL
CRPS (kW)	161.17	161.88	164.44
RI	0.06261	0.06409	0.1023
RMSE (kW)	299.47	307.78	313.59
MAE (kW)	226.35	229.21	231.46
MAPE (%)	6.3219	6.3399	6.4099
Explanatory variables	51(μ) - 30 (σ)	55(μ) - 55 (σ)	55

Table II presents the forecasting results obtained with the optimized GAMLSS model and the two benchmark models in the probabilistic forecast for the testing dataset. The average CRPS value for the three models was calculated using 99 quantiles (from 0.01 to 0.99) and the RI value was calculated with 19 quantiles, from 0.05 to 0.95 in 0.05 increments, which represents a total of 20 intervals or bins. For the deterministic prediction (punctual value of hourly load), the mean of the gamma distribution was used for the first two models, and the 0.5 quantile value for the REF2

model. The optimized GAMLSS model outperforms the two benchmark models across all five evaluation indexes for the test period. Notably, the most significant enhancements were observed in the RI and RMSE indexes when compared to the REF2 model. While the distinction between the optimized GAMLSS model and the REF1 model is less marked, the optimized model demonstrates a better performance across all indexes.

Moreover, the two GAMLSS (the optimized one and REF1) can be considered calibrated, in contrast to the linear quantile regression-based model, REF2, which is not calibrated. The calibration of a model is determined by its ability to produce a flat PIT histogram, an evaluation performed using the RI index. It is important to note that, due to the inherent randomness and the finite length of the data series, a perfectly calibrated model (statistical consistency between the distributions of observed and predicted values) will exhibit an RI value exceeding 0. This phenomenon is well-documented in the international literature, which has established a critical RI index value [13]. This value delineates whether a model can be considered calibrated at a specified significance level. For instance, in the case of the testing dataset containing 4344 records, the critical RI index value at a significance level of 0.05 is identified as 0.06797. Accordingly, this criterion confirms the calibration of the two GAMLSS models.

Figure 2 presents the PIT histogram for predictions during the testing period. The histogram's intervals are determined by quantiles, ranging from 0.05 to 0.95, creating a total of 20 intervals of equal width. Ideally, in a completely flat histogram, the probability of the dependent variable's true value falling within each interval would consistently be 0.05. However, as illustrated in the figure, the observed relative frequencies of these intervals exhibit slight variations around this target value. Despite these deviations, based on the previously outlined critical value for the RI index, the model can be considered to be calibrated.

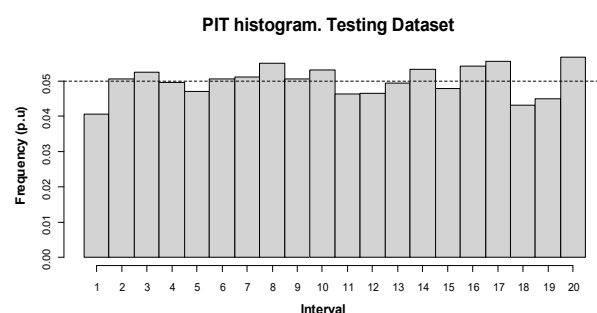


Fig. 2. PIT histogram for the forecasts with data from the testing dataset.

Figure 3 illustrates the probabilistic forecasts of hourly load for a winter week within the testing period. The figure shows the observed (actual) hourly load values in black, alongside five quantiles depicted in varying shades of grey. Notably, the median quantile (Q0.5) closely mirrors the actual load values, with this accuracy most pronounced on the five working days. The occurrence of actual values falling beneath the expected value of the 0.05 quantile is observed, aligning with the statistical anticipation that this

would happen in 5% of cases. Furthermore, the graph reveals an increase in predictive uncertainty during the weekend. This is indicated by the broader interval between

the 0.05 and 0.95 quantile values, suggesting a greater range of possible load values compared to the weekdays.

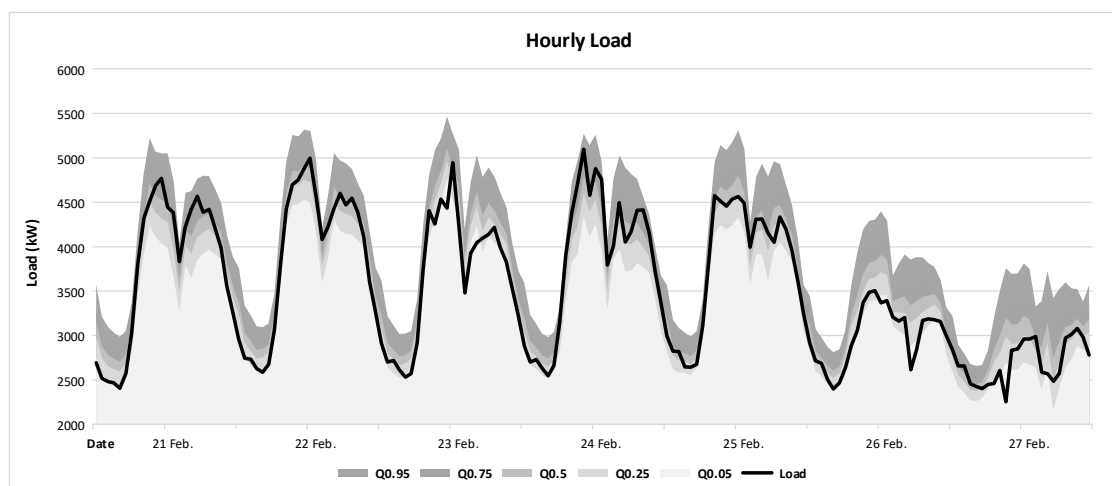


Fig. 3. Probabilistic hourly load forecasts for a week in winter.

5. Conclusions

This study details the development of a short-term probabilistic load forecasting model, focusing on a GAMLSS framework. Our methodology optimizes the selection of explanatory variables for modelling the distribution parameters, enhancing the precision of GAMLSS type models. Applied to a case study of electricity consumption at a 66/13.2 kV substation feeding a small town of 4500 inhabitants, our approach employs a dual-objective optimization strategy. This strategy facilitates the identification and linear combination of explanatory variables that most effectively model the distribution's location and scale parameters. Compared to two benchmark models, the optimized model demonstrates superior forecasting accuracy, as evidenced by a comparative analysis.

Our research efforts are presently focused on applying the proposed methodology to the development of probabilistic short-term net load forecasting models. These models are designed to predict the electricity consumption for customers who utilize photovoltaic generation systems installed behind the meters. These models will become critically important for marketers and distribution system operators. Their importance is expected to grow in parallel with the increasing penetration of photovoltaic power generation in distribution grids.

Acknowledgement

Authors are very grateful to the Spanish Government and the European Union for their support through grant TED2021-129722B-C33 funded by MICIU/AEI/10.13039/501100011033 and by the “European Union NextGeneration EU/PRTR”.

References

- [1] S. Akhtar, S. Shahzad, A. Zaheer, H.S. Ullah, H. Kilic, R. Gono, M. Jasiński and Z. Leonowicz, “Short-Term Load

Forecasting Models: A Review of Challenges, Progress, and the Road Ahead”, *Energies* (2023). Vol. 16(10), 4060.

- [2] M. López, S. Valero, C. Senabre, J. Aparicio and A. Gabaldon, “Application of SOM neural networks to short-term load forecasting: The Spanish electricity market case study”, *Electric Power Systems Research* (2012). Vol. 91, pp. 18–27.
- [3] T. Hong and S. Fan, “Probabilistic electric load forecasting: A tutorial review”, *International Journal of Forecasting* (2016). Vol. 32(3), pp. 914–38.
- [4] A. Bracale, P. Caramia, P. de Falco and T. Hong, “Multivariate Quantile Regression for Short-Term Probabilistic Load Forecasting”, *IEEE Trans. Power Syst.* (2020). Vol. 35(1), pp. 628–638.
- [5] P. Pinson, H.A. Nielsen, J.K. Møller, H. Madsen and G.N. Kariniotakis, “Non-parametric probabilistic forecasts of wind power: required properties and evaluation,” *Wind Energy* (2007). Vol. 10(6), pp. 497–516.
- [6] T. Gneiting, S. Lerch and B. Schulz, “Probabilistic solar forecasting: Benchmarks, post-processing, verification”, *Solar Energy* (2023). Vol. 252, pp. 72–80.
- [7] M.D. Stasinopoulos, R.A. Rigby, G.Z. Heller, V. Voudouris and F. De Bastiani, *Flexible Regression and Smoothing. Using GAMLSS in R*. CRC Press. Boca Raton (2017).
- [8] P. Lauret, M. David, and P. Pinson, “Verification of solar irradiance probabilistic forecasts”, *Solar Energy* (2019). Vol. 194, pp. 254–271.
- [9] M. Zamo and P. Naveau, “Estimation of the Continuous Ranked Probability Score with Limited Information and Applications to Ensemble Weather Forecasts,” *Mathematical Geosciences* (2018). Vol. 50(2), pp. 209–234.
- [10] M. Gebetsberger, J.W. Messner, G.J. Mayr and A. Zeileis, “Estimation methods for nonhomogeneous regression models: Minimum continuous ranked probability score versus maximum likelihood”. *Monthly Weather Review* (2018). Vol. 146(12), pp. 4323–4338.
- [11] K. Deb, A. Pratap, S. Agarwal, and T. Meyarivan, “A fast and elitist multiobjective genetic algorithm: NSGA-II”, *IEEE Transactions on Evolutionary Computation* (2002). Vol. 6(2), pp. 182–197.
- [12] R. Koenker and G. Bassett, “Regression quantiles”, *Econometrica* (1978). Vol. 46, pp. 33–50.
- [13] D.S. Wilks, “Indices of rank histogram flatness and their sampling properties”, *Monthly Weather Review* (2019). Vol. 147(2), pp. 763–769.