

Electric power demand forecasting in university facilities

Jorge A. Turina¹, Patricio G. Donato^{1,2}, Carlos M. Orallo^{1,2} y Marcos A. Funes^{1,2}

¹ Universidad Nacional de Mar del Plata (UNMDP)

² ICYTE (CONICET/UNMDP)

jorgeturina@gmail.com, pgdonato@conicet.gov.ar, orallo@fi.mdp.edu.ar, mfunes@fi.mdp.edu.ar

Abstract. The increasing availability of data from advanced measurement systems, largely driven by the deployment of Smart Grids, has promoted the application of machine learning techniques for electricity demand forecasting. This paper analyses the performance of different recurrent neural network architectures, specifically Long Short-Term Memory (LSTM) models trained on electricity demand time series from several buildings of the University of Mar del Plata (UNMDP). The objective is to characterize the ability of these architectures to capture the temporal dynamics of electricity consumption using different forecasting horizons and historical input data. The results highlight the challenges and limitations of predictions over extended future horizons. These findings reveal that forecasting quality depends not only on the chosen architecture but also on the representativeness and stability of the available time series.

Keywords: Power Demand, Machine Learning, Forecasting.

Pronóstico de demanda de potencia eléctrica en ambientes universitarios

Resumen. La creciente disponibilidad de datos provenientes de sistemas de medición avanzados ha impulsado, gracias al avance de las Redes Eléctricas Inteligentes, el uso de técnicas de aprendizaje automático para el pronóstico de la demanda eléctrica. Este trabajo analiza diferentes arquitecturas de redes neuronales recurrentes, específicamente del tipo Long Short-Term Memory (LSTM), aplicadas a las series temporales de demanda eléctrica de diferentes suministros de la Universidad Nacional de Mar del Plata (UNMDP). El objetivo es caracterizar la capacidad de estas arquitecturas para captar la dinámica temporal de los consumos eléctricos empleando diferentes horizontes de predicción y datos históricos. Los resultados evidencian la complejidad y limitaciones de las predicciones con horizontes futuros extensos. Estos hallazgos muestran que la calidad del pronóstico depende no solo de la arquitectura empleada, sino también de la representatividad y estabilidad de las series temporales disponibles.

Palabras clave: Demanda de Potencia, Aprendizaje Automático, Pronóstico.

1 Introducción

La transición hacia las Redes Eléctricas Inteligentes (REI, o Smart Grids) está transformando la gestión de los sistemas eléctricos gracias a la disponibilidad de grandes volúmenes de datos adquiridos mediante infraestructuras avanzadas de medición. Esto ha impulsado el desarrollo de herramientas basadas en aprendizaje automático (Machine Learning) para abordar problemas complejos relacionados con la operación y planificación de los sistemas eléctricos, incluyendo el pronóstico de la demanda. La disponibilidad de información detallada sobre el comportamiento del consumo energético permite desarrollar modelos capaces de anticipar la evolución temporal de la demanda, contribuyendo a mejorar la eficiencia operativa, la planificación energética y la integración de recursos distribuidos.

En el ámbito particular de las instituciones universitarias, la gestión eficiente de la demanda eléctrica representa un desafío significativo debido a la complejidad y variabilidad de los patrones de consumo. Los complejos universitarios comprenden un espectro heterogéneo de cargas, desde laboratorios y bibliotecas hasta residencias estudiantiles y edificios administrativos. Esta diversidad funcional se traduce en perfiles de consumo altamente variables, caracterizados por fluctuaciones tanto diarias como estacionales. Dichas variaciones se encuentran influenciadas por múltiples factores, entre los que se destacan los horarios académicos, los ciclos del calendario universitario (períodos de exámenes, recesos), eventos especiales y las condiciones climáticas. Pronosticar con precisión la demanda de potencia eléctrica en estos contextos no solo es importante para optimizar la operación de la infraestructura energética, sino también para reducir costos y minimizar el impacto ambiental.

Pronosticar una serie temporal consiste en extender sus valores históricos conocidos hacia instantes futuros, lo que metodológicamente se traduce en un problema de regresión. Si bien este enfoque es ampliamente utilizado en diversos dominios estadísticos (Shumway et al., 2006), en este trabajo se aborda específicamente el modelado de variables eléctricas.

Para el modelado de series temporales complejas, los métodos de aprendizaje automático permiten capturar relaciones no lineales donde los enfoques tradicionales fallan. Aunque el estado del arte reciente destaca arquitecturas como Transformers temporales o N-BEATS, su elevada demanda computacional compromete la viabilidad operativa en la optimización energética local, la cual exige actualizaciones rápidas ante cambios de consumo (Casals Cunill et al., 2024). Por ello, redes recurrentes tales como Long Short-Term Memory (LSTM) (Hochreiter et al., 1997) representan una alternativa eficiente por su capacidad para procesar dependencias secuenciales de corto y largo plazo.

En trabajos previos realizados por los autores se ensayaron diferentes propuestas para pronosticar demanda a futuro basadas tanto en redes tipo perceptrón multicapa como celdas LSTM (Donato et al., 2023; Donato et al., 2022). Estos estudios permitieron verificar la viabilidad del uso de técnicas de aprendizaje profundo para modelar el comportamiento de la demanda eléctrica en un suministro de la UNMDP (Facultad de Ingeniería). Paralelamente, se evaluaron otros modelos de pronóstico como K-Nearest Neighbors, Ridge y Lasso, Random Forest y Stacking, y se evaluaron las

prestaciones de cada uno de ellos frente al suministro mencionado (Maza Díaz et al., 2024). Los resultados reflejaron que los modelos Random Forest y Stacking presentaban las mejores métricas de error. Posteriormente, se aplicó el modelo Random Forest a distintos suministros pertenecientes a la universidad y se observó que, dependiendo de las características poblacionales, de infraestructura y de actividades, las magnitudes de los errores podían cambiar sustancialmente (Maza Díaz et al., 2025).

El presente trabajo profundiza sobre esta línea de investigación mediante un análisis comparativo de diversas topologías aplicadas a múltiples suministros de la UNMDP. El objetivo es identificar las configuraciones que optimicen el compromiso entre precisión predictiva, costo computacional y capacidad de generalización.

2 Materiales y métodos

En esta sección se describen tanto las bases de datos empleadas en este estudio como las arquitecturas evaluadas y el proceso para llegar a los resultados obtenidos.

2.1 Preprocesamiento de la base de datos

En este trabajo se utilizaron las series temporales de mediciones de potencia activa correspondientes a los suministros de la Facultad de Ingeniería (FI), la Facultad de Ciencias Exactas (FCE), la Facultad de Ciencias Económicas y Sociales (FCEyS), y el Instituto de Investigaciones en Ciencia y Tecnología de Materiales (INTEMA) de la UNMDP. Dichas series comprenden un período de cuatro años y están muestreadas con un período de 15 minutos. Inicialmente, el estudio se centró en los datos de FI para luego extenderse al análisis de los demás edificios.

La primera parte del preprocesamiento consiste en la limpieza y promediado de las bases de datos. Se eliminaron valores atípicos (potencia nula) derivados de fallos en la adquisición y, posteriormente, se promediaron las muestras originales cada 15 minutos para obtener una resolución horaria, reduciendo variaciones de alta frecuencia. En la Tabla 1 se resumen algunos de los valores estadísticos que caracterizan a los diferentes conjuntos de datos que se emplearán para el entrenamiento y evaluación de los algoritmos. Los valores del resto de los suministros se consignan para tener dimensión de cómo varían en diferentes ambientes universitarios.

Tabla 1. Estadística del conjunto de datos de potencia activa.

	FI	FCE	FCEyS	INTEMA
Promedio [kW]	26.62	87.26	53.89	73.90
Desviación estándar [kW]	15.52	41.86	22.33	41.16
Valor mínimo [kW]	2.86	17.46	26.06	12.91
25 % (Primer cuartil) [kW]	16.86	61.14	37.54	48.11
50 % (Mediana) [kW]	19.81	80.00	43.82	59.68
75 % (Tercer cuartil) [kW]	32.74	99.69	67.29	82.32
Valor máximo [kW]	102.45	336.30	178.48	305.99

La segunda parte del preprocesamiento consiste en hallar y utilizar las características correctas a partir de los datos en crudo. Esto se conoce como feature engineering o ingeniería de características y es un paso fundamental para reducir la complejidad del modelado y obtener resultados de mayor calidad. Las características extraídas fueron las siguientes:

- **Codificación temporal:** La naturaleza periódica del consumo eléctrico se representó mediante una codificación temporal cíclica, utilizando funciones seno y coseno para la hora del día (período de 24 h) y el día de la semana (período de 7 días). Esta estrategia permite que el modelo capture la continuidad temporal entre el final y el inicio de cada ciclo, de modo que instantes como las 23:00 h y las 00:00 h sean interpretados como cercanos en el dominio temporal. Para complementar esta representación, se incorporaron variables dummy (one-hot encoding) para cada hora y día. En este esquema, cada instante temporal se representa mediante un conjunto de variables binarias, donde únicamente la variable asociada a la hora o día considerado toma el valor 1, mientras que las restantes permanecen en 0. Esta combinación busca potenciar la capacidad del modelo. Mientras la codificación cíclica aporta información sobre la continuidad del tiempo, las variables binarias permiten identificar y aprender patrones específicos o picos de consumo asociados a momentos puntuales (como el inicio de una jornada académica), activando con valor 1 únicamente el intervalo temporal correspondiente.

- **Identificación de días hábiles:** Se implementó una variable booleana para distinguir días laborables. El umbral se fijó en 30 kW, tras un análisis empírico de los consumos históricos en fines de semana, asegurando una discriminación robusta frente a la variación de la carga base.

- **Variables de retardo (lags):** Se incorporaron variables lag (mediciones que se arrastran de valores pasados) de las últimas tres horas y de la muestra número 24 (mismo horario, día anterior). Para cada una, se calculó la diferencia incremental respecto al valor actual, permitiendo al modelo capturar tanto el estado reciente como la dinámica de variación temporal (delta).

- **División de datos:** El conjunto de datos se dividió en entrenamiento, validación y test con el objetivo de garantizar la capacidad de generalización del modelo y mitigar el sobreajuste estacional. Se reservó un año calendario completo para la validación, permitiendo evaluar el desempeño frente a la totalidad de las variaciones climáticas y los ciclos de actividad académica anual (ver figura 1). Se tomó la decisión estratégica de excluir el periodo de vacaciones de verano del conjunto de entrenamiento. Se observó que, durante este intervalo, la demanda eléctrica presenta una variabilidad mínima, lo que induce al modelo a converger hacia lo que se denomina como una predicción por persistencia, donde se repite el último valor de la serie (un tipo de error común en series temporales con baja variabilidad). Al eliminar este sesgo de baja actividad, se evita que el modelo aprenda patrones triviales, ya que estos intervalos pueden ser modelados adecuadamente mediante un comportamiento persistente.

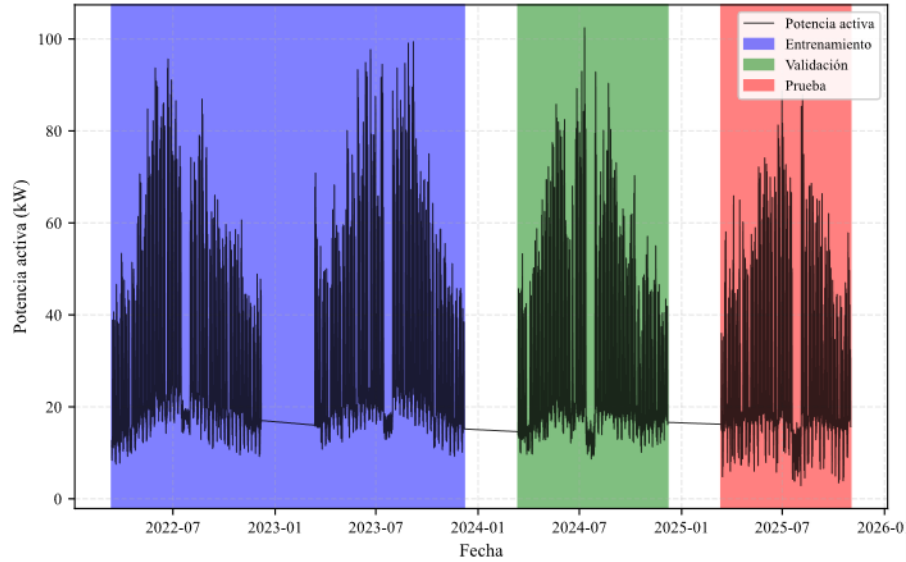


Fig. 1. Potencia activa en función de muestras ordenadas en el tiempo (FI).

2.2 Dimensión de la ventana de datos de entrada

El pronóstico de demanda de potencia se realizó a partir de ventanas de datos de entrada con diferentes horizontes temporales de predicción (ver figura 2). Gracias al promediado horario realizado en el preprocesamiento, se definieron longitudes de 6, 12, 24 y 48 muestras, equivalentes a la misma cantidad de horas de historia reciente.

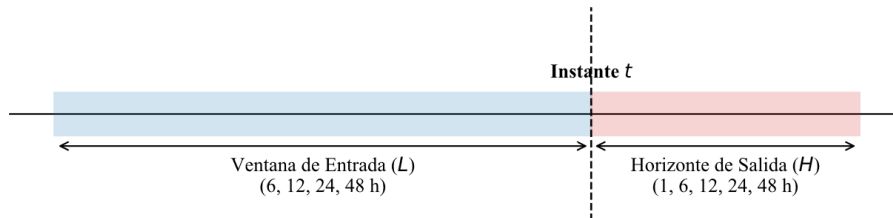


Fig. 2. Configuración temporal del modelo, donde L y H representan la historia de entrada y el horizonte de pronóstico, respectivamente.

Con el fin de enriquecer la información contextual de cada ventana, se incorporaron los siguientes atributos adicionales:

- **Flags de Integridad y Continuidad:** Se añadió una variable binaria (flag) para indicar si la ventana de datos está completa (sin saltos temporales por fallos en el sistema de medición o el suministro eléctrico). Asimismo, se incluyó un segundo

indicador para señalar si la ventana abarca más de un ciclo diario, para ayudar al modelo a discernir cambios de día.

- **Análisis en el Dominio de la Frecuencia:** Para capturar periodicidades ocultas y el comportamiento espectral de la demanda, se calcularon las tres componentes principales en magnitud de la Transformada Discreta de Fourier (DFT) sobre cada ventana. Se optó por utilizar las primeras tres componentes de la DFT tras verificar experimentalmente que mejoraban la precisión del pronóstico.

En la figura 3 se muestra en forma esquemática cómo es el vector de datos con el que se alimenta el algoritmo.

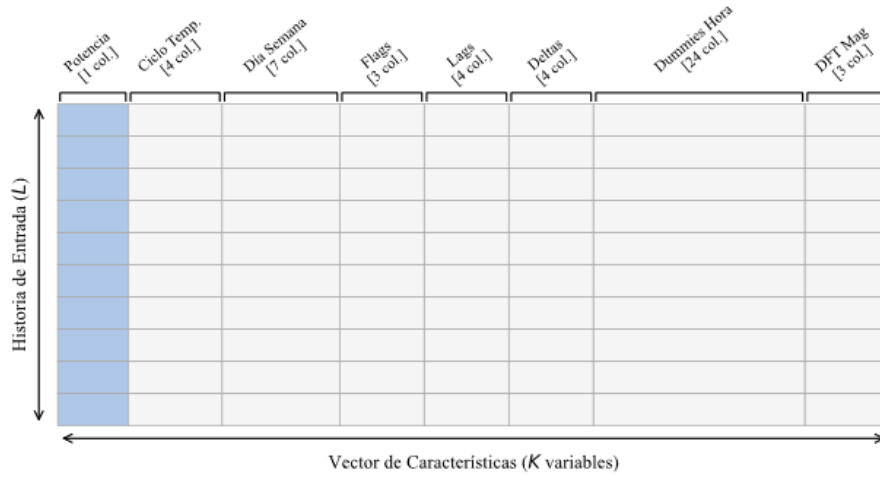


Fig. 3. Composición de los datos de entrada al modelo. Cada fila representa un paso temporal dentro de la ventana L , mientras que las columnas desglosan las $K=50$ variables características (features) que conforman el contexto informativo para el pronóstico.

2.3 Elección de las métricas del error

La selección de las variables características y las arquitecturas en este trabajo fueron el resultado de un proceso iterativo de pruebas, analizando el impacto de cada variable sobre una configuración base. Para comparar a los diferentes modelos se emplearon las siguientes métricas:

- **Error Porcentual Absoluto Medio (MAPE):** valor que representa el promedio de las diferencias absolutas relativas entre los valores reales y las predicciones.

$$MAPE_h = \frac{1}{n} \sum_{i=1}^n \left| \frac{\hat{y}_{i,h} - y_{i,h}}{y_{i,h}} \right| \times 100$$

- **Desviación Estándar del Error Relativo (DEER):** valor que permite determinar la dispersión del pronóstico en las diferentes condiciones de ensayo.

$$DEER_h = \sqrt{\frac{1}{n} \sum_{i=1}^n \left(\left| \frac{\hat{y}_{i,h} - y_{i,h}}{y_{i,h}} \right| - \frac{MAPE_h}{100} \right)^2} \times 100$$

Donde i representa cada muestra de cada paso horario (h) dentro del horizonte de predicción. En este contexto, el $MAPE_h$ interviene como la media aritmética de la distribución de los errores absolutos relativos. El uso del operador de valor absoluto en ambas métricas evita que los errores por sobreestimación (positivos) y subestimación (negativos) se cancelen algebraicamente entre sí, permitiendo cuantificar la magnitud real del impacto de las desviaciones. Además, dado que estas métricas se calcularon para cada horizonte de predicción (h), en los escenarios multi-step (como 6, 12, 24 y 48 h) se procedió a promediar los resultados de los distintos horizontes a fin de obtener un valor puntual comparable con el resto.

Estas métricas permiten cuantificar la dispersión y errores de predicción de forma normalizada, eliminando la dependencia de la escala absoluta (kW). El criterio empleado para ponderar los modelos fue minimizar la DEER, y en segunda instancia reducir el valor del MAPE. Este enfoque garantiza que el modelo elegido no solo sea exacto en promedio, sino que sea el más consistente y estable en sus predicciones.

2.4 Arquitecturas y algoritmos empleados

Para el desarrollo del modelo de predicción, se evaluaron diferentes configuraciones de redes LSTM, con el objetivo de evitar dos escenarios críticos:

- **Underfitting:** Se buscó que el modelo supere el desempeño de las predicciones persistentes (donde el valor futuro se asume igual al último valor conocido), demostrando una verdadera capacidad de aprendizaje de la serie temporal.
- **Overfitting:** Para prevenir la sobreadaptación a los datos de entrenamiento, se implementó la técnica de Early Stopping. Esta función monitorea la pérdida en el conjunto de validación y detiene el entrenamiento de manera automática cuando el desempeño deja de mejorar.

Con este propósito, se definieron las siguientes variantes arquitectónicas, las cuales representan distintos niveles de complejidad y capacidad de modelado:

- **Configuración A (Base):** Una única capa LSTM con 4 unidades en su estado oculto (hidden state).
- **Configuración B (Dos capas):** Una arquitectura con una primera capa LSTM de 32 unidades, cuya salida se conecta a una segunda capa LSTM de 16 unidades.
- **Configuración C (Tres capas simétricas):** Una estructura de tres niveles con capas LSTM de 32, 16 y 32 unidades consecutivamente.
- **Configuración D (Alta capacidad):** Una única capa LSTM robusta con 256 unidades de hidden state.
- **Configuración E (Dos capas de alta capacidad):** Una arquitectura profunda compuesta por una primera capa LSTM de 256 unidades seguida de una de 128 unidades.

En todos los casos, la etapa final de la red consiste en una capa densa (fully connected). La dimensión de salida de esta capa se ajustó de manera dinámica para coincidir con el horizonte de predicción (cantidad de muestras a predecir) requerido en cada variación planteada.

Para el entrenamiento de estas arquitecturas y con el fin de garantizar la reproducibilidad de los experimentos, se utilizó la librería TensorFlow/Keras. Con este propósito, se estableció un control de aleatoriedad global fijando una semilla determinista (seed=47) tanto en el entorno de TensorFlow como en las operaciones de la librería NumPy. El proceso de optimización se llevó a cabo mediante el algoritmo Adam (Kingma et al., 2015), configurado con su tasa de aprendizaje (learning rate) inicial base de 0.001, empleando el Error Cuadrático Medio (MSE) como función de pérdida (loss function) tanto para las etapas de entrenamiento como de validación y prueba. El proceso se definió para un máximo de 300 épocas de entrenamiento con un tamaño de lote (batch size) de 32. Para la estrategia de Early Stopping se utilizó una paciencia (patience) de 10 épocas, además se configuró para restaurar automáticamente los mejores pesos obtenidos por la red (restore_best_weights). Cabe destacar que las arquitecturas se entrenaron sin la inclusión de capas o tasas de dropout.

3 Resultados

Como se mencionó anteriormente se plantearon variaciones de arquitecturas, horizontes de predicción y cantidad de muestras a la entrada. Los resultados se muestran en las figuras 4 y 5. Cada figura muestra los resultados obtenidos con cada arquitectura, empleando todas las variantes de horizontes y muestras de entrada. El eje de las abscisas representa el MAPE, mientras que el eje de las ordenadas representa el DEER, por lo cual el mejor pronóstico es el que se obtiene con las arquitecturas que se mapean en la esquina inferior izquierda, mientras que los peores desempeños corresponden a la esquina superior derecha. En cada gráfico se visualizan los resultados obtenidos con cada una de las ventanas de datos de entrada y horizontes de predicción, los cuales se representan como un círculo cuyo tamaño es proporcional al horizonte de predicción, mientras que su color identifica la ventana de datos de entrada. La mejor combinación de horizonte/ventana se ha señalado mediante un símbolo de una estrella.

En todos los casos se observa que los mejores resultados se obtienen para los horizontes de predicción de una hora, con ventanas de entrada de 12 o 24 horas, y que al aumentar el horizonte de predicción también aumentan las métricas de error. Este comportamiento es lógico, ya que si el algoritmo comete un determinado error al pronosticar una hora hacia adelante, claramente cometerá un error mayor al pronosticar 6 o más horas en avance, debido a la propagación acumulativa de los errores de predicción. En las arquitecturas B y E se aprecia que la diferencia en el MAPE entre los horizontes de predicción de 6 y 12 horas es muy similar, fluctuando en el rango de 12 % a 15 %, pero el rango de variación en el DEER es bastante más amplio, entre 15 % y 23 %.

En la Figura 6 se comparan los mejores resultados obtenidos por cada arquitectura. Nuevamente se advierte que los horizontes de 6 y 12 horas presentan un comporta-

miento agrupado. Además, con excepción de la arquitectura más simple (A), las arquitecturas que resultan óptimas para cada horizonte muestran desempeños muy similares. A modo referencia, los índices MAPE y DEER correspondientes a una predicción persistente de horizonte igual a 1 hora son 9.79 % y 10.94 % respectivamente.

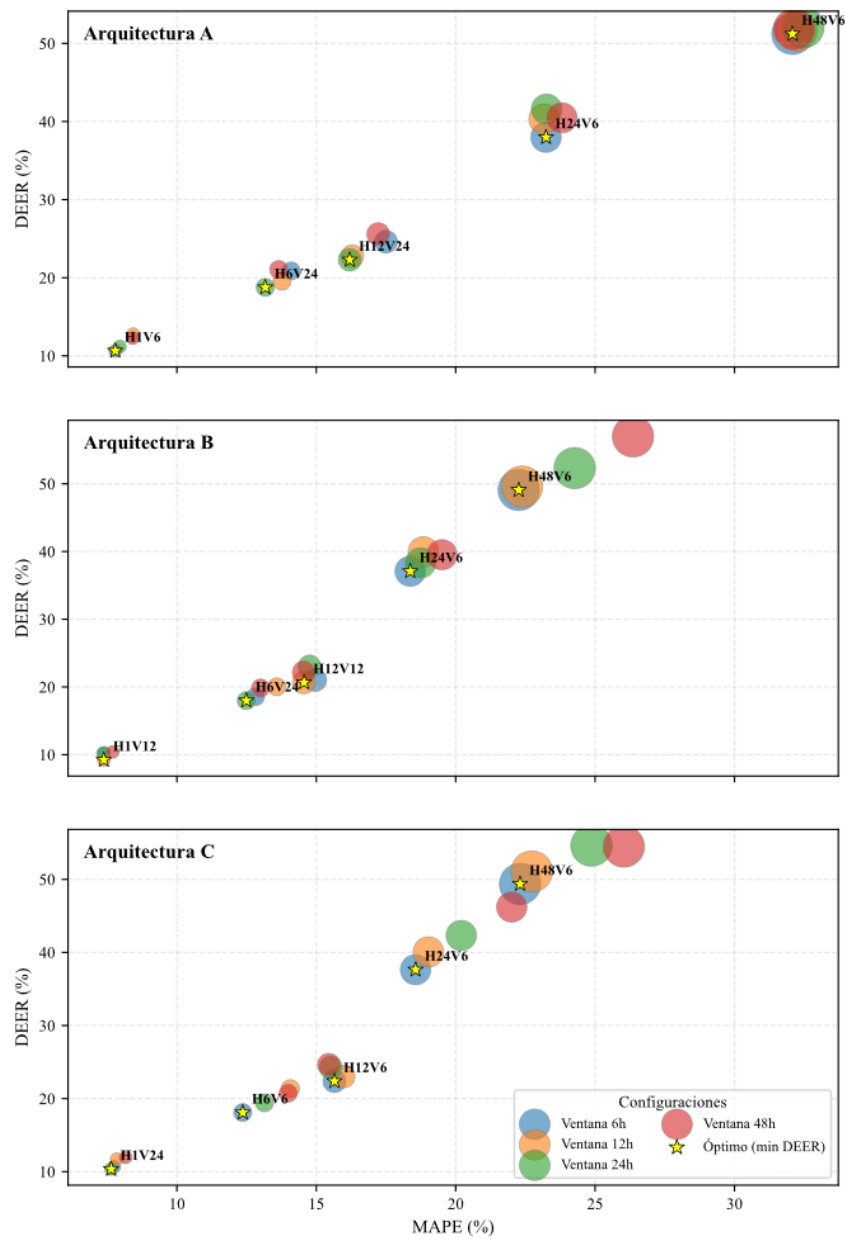


Fig. 4. Resultados de las arquitecturas A, B y C.

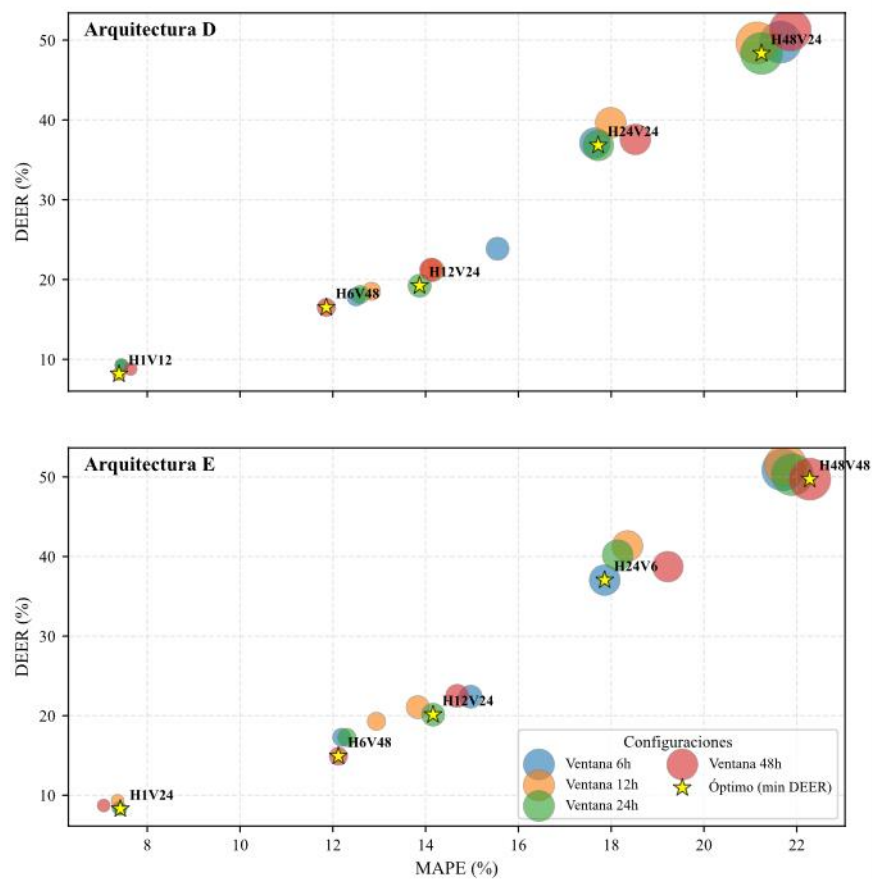


Fig. 5. Resultados de las arquitecturas D y E.

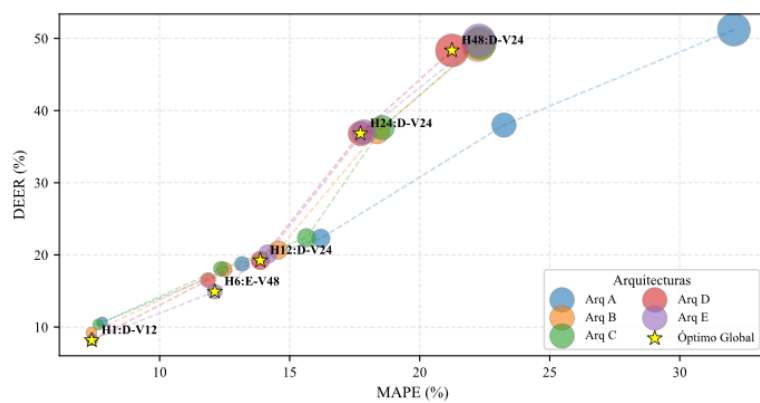


Fig. 6. Comparativa de los mejores resultados de cada arquitectura para el mismo horizonte de predicción.

Las condiciones experimentales de este estudio se llevaron a cabo de forma local utilizando exclusivamente el procesador central (CPU) de una Notebook equipada con un procesador AMD Ryzen 7 5700U (8 núcleos, 16 hilos), 12 GB de memoria RAM (de los cuales 2 GB están reservados para los gráficos integrados AMD Radeon Graphics, dejando 10 GB disponibles para el sistema), bajo el sistema operativo Windows 11 de 64 bits. No se empleó aceleración por hardware basada en unidades de procesamiento gráfico (GPU) dedicada para el proceso de entrenamiento. Bajo este entorno el modelo E requirió 37'04'' de entrenamiento en promedio para todos los casos, lo que representa más de diez veces el tiempo del modelo A cuyo tiempo promedio fue de 3'18''. Como se observa en la Figura 6, esta diferencia de capacidad no resulta significativa en horizontes cortos (hasta 6 horas), donde los modelos simples logran un desempeño competitivo. Específicamente, los tiempos promedio registrados para las arquitecturas A, B, C, D y E fueron de 3'18'', 4'58'', 7'01'', 17'45'' y 37'04'' respectivamente. En la Tabla 2 se presentan las mejores configuraciones (identificadas con una estrella en la figura 6) para cada horizonte de predicción. Cabe recordar que tanto el DEER como el MAPE se calcularon según lo expuesto en la Sección 2.3, promediando los valores de cada paso horario en los escenarios multi-step para obtener un valor puntual comparable entre los modelos de mismo horizonte. En la misma tabla se incorporan las métricas correspondientes a la persistencia según horizonte (DEER p. y MAPE p.) a modo de referencia.

Tabla 2. Resumen de métricas para las configuraciones óptimas.

Horizonte / Ventana (h)	DEER p. (%)	MAPE p. (%)	Tiempo train (minutos)	DEER (%)	MAPE (%)	Loss train (mejor época)	Loss val (mejor época)
1 / 12	10.94	9.79	11.986	8.17	7.39	0.0164	0.0186
6 / 48	29.60	27.17	69.567	14.91	12.12	0.0413	0.0569
12 / 24	43.98	44.08	17.625	19.23	13.87	0.0563	0.0821
24 / 24	57.30	48.41	18.427	36.82	17.72	0.1034	0.1292
48 / 24	67.04	53.95	16.838	48.34	21.24	0.1529	0.1756

4 Extensión del estudio a otros perfiles de demanda

Con el objetivo de demostrar cómo la calidad del pronóstico depende no solo de la arquitectura empleada, sino también de la representatividad y estabilidad de las series temporales disponibles (Porteiro et al., 2022), se evaluaron las mejores arquitecturas para cada horizonte de predicción usando como entrada las bases de datos de los otros edificios universitarios mencionados en la Tabla 1.

Bajo las mismas condiciones de configuración, división de datos, preprocesamiento y estructura de ventanas, la evaluación de las mejores arquitecturas sobre los distintos conjuntos de datos revela reducciones promedio del 18 % en el MAPE y del 32 % en el DEER, según se ilustra en la Figura 7. Las mejoras resultaron bastante notorias

en horizontes cortos donde, por ejemplo, en el horizonte de 1 hora se verifican mejoras de orden del 35 % sobre el MAPE y el DEER promediando los diferentes conjuntos de datos, en comparación con los resultados obtenidos para FI. Es destacable cómo los diferentes conjuntos de datos presentan un comportamiento bastante lineal al aumentar el horizonte de predicción, salvo el de FI.

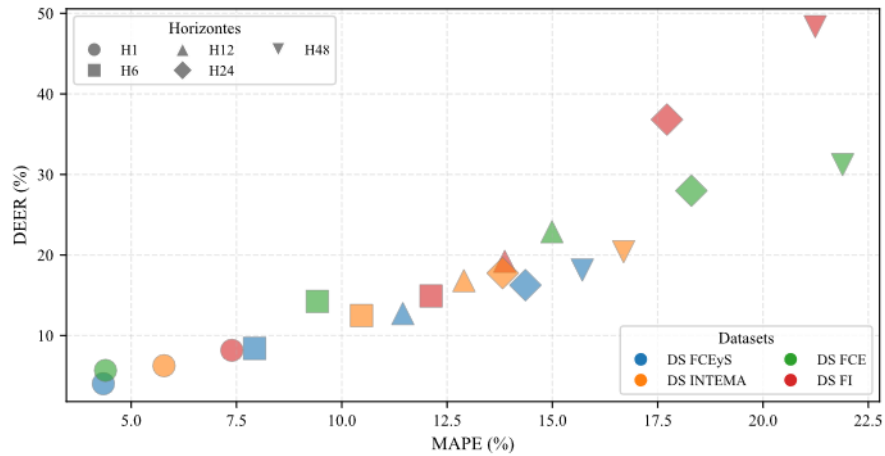
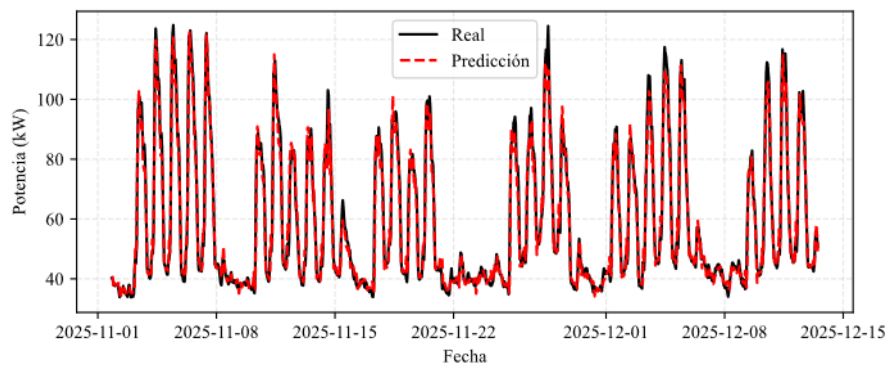


Fig. 7. Comparativa de la mejor arquitectura para cada horizonte de predicción planteado en el trabajo, aplicado a los otros conjuntos de datos mencionados.

A modo ilustrativo se muestra en la figura 8 una comparativa de la mejor arquitectura con el horizonte de 1 hora en un rango al azar de 1000 datos que comprende de 2025-11-01 20:00:00 a 2025-12-13 12:00:00 para el dataset de FCEyS y FI respectivamente. Se puede observar como la arquitectura predice correctamente la tendencia de la demanda con una pequeña dispersión en los valores estimados. En ambas dependencias, entre las cuales existen diferencias en los patrones de demanda, la arquitectura seleccionada permite una correcta predicción. Se destaca en particular la predicción en FI, en la cual durante fines de semana y/o feriados la inyección de energía renovable cambia la forma del perfil.



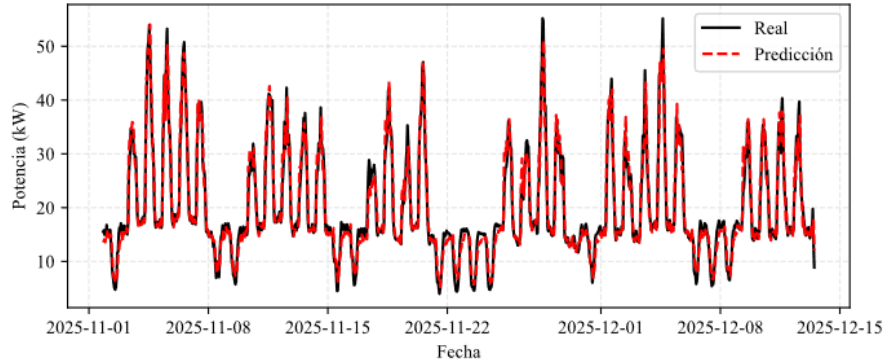


Fig. 8. Demanda real vs predicción de la mejor arquitectura el horizonte de 1 hora. Superior FCEyS, inferior FI.

5 Conclusiones

En este trabajo se ha presentado un estudio sobre algoritmos de pronóstico de demanda de potencia eléctrica en edificios universitarios. Se han ensayado arquitecturas con diferentes grados de complejidad, empleando diferentes cantidades de datos de entrada y horizontes de predicción temporales. Los resultados revelan que el error aumenta al extender el horizonte de predicción, y que ampliar la ventana de datos de entrada no siempre garantiza una mejora proporcional. Respecto de las arquitecturas, si bien las de mayor capacidad (D y E) logran un mejor desempeño general, implican un costo computacional superior.

Para mejorar el desempeño de los pronósticos, en trabajos futuros se podrían incluir variables exógenas en el modelo, tales como la temperatura, presión y humedad en el momento de la medición. Adicionalmente, el criterio de umbralización empleado para discriminar días hábiles y no hábiles podría ser profundizado, adaptándolo a los cambios en el comportamiento del consumo. Además, considerando que solo se entrenó con dos años de mediciones, el desempeño podría mejorarse en trabajos futuros utilizando más valores cuando estén disponibles, teniendo en cuenta de emplear un año para validar el modelo y no solamente en una porción del mismo.

Agradecimientos

Este trabajo fue soportado por la Universidad Nacional de Mar del Plata (UNMDP) y el Consejo Nacional de Investigaciones Científicas y Tecnológicas (CONICET). Los autores agradecen al Programa Iberoamericano de Ciencia y Tecnología para el Desarrollo (CYTED) por el apoyo brindado a través de las Redes Temáticas RIPCEL 726RT0203 y RIIIAinSG 726RT0204.

Referencias

- Casals Cunill, E. C., López Hierrezuelo, S., García, D. J. y Lago Solano, R. D. (2024). Diseño de Red Neuronal Artificial para la Predicción de la Demanda Eléctrica. *Ingeniería Energética*, 45(3), e2419. <http://scielo.sld.cu/pdf/rie/v45n3/1815-5901-rie-45-03-e2419-04.pdf>
- Donato, P.G., Orallo, C.M., Funes, M.A., Echeverría, N.I. (2022). Pronóstico de variables eléctricas en el marco del proyecto de ciudades inteligentes en Mar del Plata. *IEEE Argencon 2022*, San Juan, Argentina.
- Donato, P.G., Hernández, Á., Funes, M., Nieto, R., Orallo, C., de Diego, L. (2023). Pronóstico de consumo de energía en ámbitos institucionales y domésticos. *28º Congreso Argentino de Control Automático (AADECA 2023)*, Buenos Aires, Argentina.
- Hochreiter, S., Schmidhuber, J. (1997). Long Short-Term Memory. *Neural Computation*, 9(8), 1735-1780.
- Kingma, D. P., Ba, J. (2015). Adam: A Method for Stochastic Optimization. *Proceedings of the 3rd International Conference on Learning Representations (ICLR)*. San Diego, CA.
- Maza Díaz, Y., Donato, P.G., Funes, M.A., Orallo, C.M. (2024). Evaluating machine learning algorithms for power demand forecasting. *XV Latin-American Congress on Electricity Generation and Transmission (CLAGTEE 2024)*, Mar del Plata, Argentina.
- Maza Díaz, Y., Donato, P.G., Funes, M.A., Orallo, C.M. (2025). Evaluación del algoritmo Random Forest para predicción de demanda en suministros de energía de la UNMDP. *VI Jornadas de Ingeniería Aplicada*. N°3 Ed. 1. pp 99-104. ISBN 978-631-6662-42-2
- Porteiro, R., Hernández-Callejo, L., Nesmachnow, S. (2022). Electricity demand forecasting in industrial and residential facilities using ensemble machine learning. *Revista Facultad de Ingeniería Universidad de Antioquia*, (102), 9-25. <https://doi.org/10.17533/udea.redin.20200584>
- Shumway, R.H., Stoffer, D.S. (2006). *Time Series Analysis and Its Applications*. 2nd edn. Springer, New York.