

Estudio para la proyección de llamadas recibidas en un centro de experiencia telefónica de una entidad financiera

Autor: Juan Pablo Ruiz Cotrino

Tutor: Deniz Andrea Sanchez Segura

Universidad Santo Tomas

01/05/2023

Contents

1. Introduccion	4
2. Objetivos de la investigación	5
2.1 Objetivo General	5
2.2 Objetivos Específicos	5
3. Marco Teorico	6
3.1 Antecedentes	6
3.2 Bases teoricas	7
3.2.1 Series temporales	7
3.2.1.1 Modelo AR	7
3.2.1.2 Modelo MA	8
3.2.1.3 Modelos ARIMA	9
3.2.2 Random Forest	9
3.2.3 Redes Neuronales	11
3.2.3.1 Distribucion de la red neuronal	11
3.2.4 Criterios de seleccion modelo	13
3.2.4.1 Criterio de informacion de Akaike (AIC)	14
3.2.4.2 Error Cuadratico Medio	14
3.2.4.3 Coeficiente de Determinacion	14
4. Resultados	16
4.1 Fuentes de información y Datos	16

4.2 Desarrollo	19
4.2.1 Modelo Arima	19
4.2.1.1 Transformacion de la serie	19
4.2.1.2 Seleccion del modelo	23
4.2.1.3 Supuestos del modelo	26
4.2.1.4 Pronosticos del modelo	30
4.2.2 Random Forest	32
4.2.3 Redes Neuronales	39
4.3 Comparacion de los modelos aplicados	43
5. Conclusiones	47
6. Referencias	47

1. Introduccion

En la actualidad, muchas empresas en todo el mundo utilizan diversos medios de comunicación para responder a solicitudes de los clientes y contactar a clientes potenciales para ofrecer productos. Entre estos canales, se encuentra el contacto vía telefónica, ya que es ampliamente utilizado por la mayoría de la población, lo que permite a las empresas ampliar su alcance y los beneficios que pueden ofrecer. Para llevar a cabo este contacto, existen varios call centers que actúan como enlace entre la empresa y el cliente final. Sin embargo, uno de los mayores desafíos en la operación de estos centros es la determinación de los flujos de llamadas durante un período determinado, ya que es difícil saber cuándo aumentarán o disminuirán significativamente. Además, no siempre se conoce la cantidad de agentes que están trabajando, lo que dificulta la determinación de la eficacia de la respuesta. Como resultado, los clientes pueden verse afectados o insatisfechos debido a los tiempos de espera prolongados y la falta de disponibilidad para atender sus llamadas. Por otro lado, la entidad también tiene ciertas dudas sobre cómo distribuir los agentes, como cuántos agentes se necesitan en un día para garantizar el más alto nivel de atención, en qué momentos del día disminuye el flujo de llamadas para programar los descansos de los agentes y determinar la duración promedio de cada llamada.

2. Objetivos de la investigación

2.1 Objetivo General

Determinar los posibles patrones de comportamiento del flujo de llamadas del centro de contacto de una entidad en diferentes periodos de tiempo, con el fin de establecer la forma óptima de realizar la operación diaria y mensual.

2.2 Objetivos Específicos

- Identificar las posibles causas o factores que puedan modificar el flujo de llamadas.
- Reconocer los modelos que se pueden emplear para predecir la entrada de llamadas en el futuro.
- Caracterizar los datos previos con los que se buscaran identificar ciertos patrones que puedan orientarnos hacia la selección del mejor modelo.
- Comparar los resultados arrojados por cada modelo y evaluar el ajuste de cada uno de ellos en relación con los datos empleados.

3. Marco Teorico

3.1 Antecedentes

Durante el planteamiento de esta investigación se busca desarrollar los posibles métodos o técnicas que nos expliquen un poco mejor el comportamiento de los datos para así realizar su respectivo pronóstico, así mismo se pretende caracterizar cada método para reconocer en qué condiciones el modelo puede generar una mejor respuesta, y así determinar cuál es la opción más viable en base a los datos obtenidos. Ahora bien, en la literatura se encuentran múltiples estudios para el pronóstico de un call center a través de diferentes modelos, en primera instancia se tomó la investigación realizada por (Gaviria, 2020) la cual se basaba en el análisis y recontacto a través de modelos como Gradient Boosting, árboles de decisión, regresión, Random Forest y redes neuronales. En este estudio el criterio de selección se realizó a través del AUC es decir el área bajo la curva para lo cual se seleccionó el Gradient Boosting, ya que presento una mejor respuesta respecto a las demás modelos evaluadas. Cabe aclarar que también se desarrolló un rebalanceo con lo que se encontró un mayor aumento del AUC en los modelos y que para este caso en especial se asociaron variables que pueden explicar aún mejor el modelo. Por otro lado, se encontramos un estudio para la proyección de llamadas a través de series temporales realizado por (Hernández, 2019), específicamente un modelo ARIMA con lo cual se abre otra posibilidad de estudio dado que no se evalúan dicha información en el estudio anterior. Además, en esta ocasión no se tomarán en cuenta las asociaciones entre las variables sino únicamente el análisis de la información del pasado. Esto nos permitirá identificar que enfoque es más conveniente al realizar las proyecciones, se plantearon alrededor de nueve modelos a los cuales se les realizó la evaluación de los supuestos de normalidad y auto correlación. En este caso, el criterio de selección al momento de realizar la comparación de cada uno de los modelos se elaboró en base al criterio de información y para el caso de los pronósticos se tuvo en cuenta el error medio para determinar qué tan acertadas fueron las proyecciones. Por último, pero no menos importante y siguiendo con el enfoque hacia las series temporales, nos encontramos con una investigación realizada en la cual encontramos diferentes modelos como un modelo SARIMA, Random Forest, Gradient Boosting y ARIMA. Con un panorama un poco más amplio nos encontramos con que (Benavides, 2022), expone en la investigación que los mejores ajustes se presentan en los modelos ARIMA y Gradient Boosting. Ahora bien la estructura del presente documento se desarrollara de la siguiente manera, en primera instancia se realizara la respectiva caracterización de cada uno de los métodos determinando las variables que se emplearan, así

mismo se realizaran diversas combinaciones entre el método para poder tener diversos puntos de vista al abordar la serie, posteriormente una comparación de los métodos con medidas como la varianza, el error cuadrático medio y el error cuadrático medio porcentual para determinar la conveniencia del modelo y por último la respectiva conclusión para futuros análisis o posibles modificaciones que mejoren los resultados deseados.

3.2 Bases teoricas

Cuando se habla de centros de experiencia o atención telefónica, el principal enfoque al momento de determinar qué tan exitosa será la operación se basa en la capacidad con la que se cuenta para poder atender y dar respuesta a cada uno de los clientes con los que se tiene contacto. Entre los parámetros de medida que encontramos para medir el desempeño encontramos el nivel de atención, el nivel de servicio, llamadas contestadas, entrantes, etc. Ahora bien, una vez definida esta ruta es importante tener una aproximación que nos permita dimensionar las llamadas que ingresan al centro de experiencia telefónica para poder determinar el personal que se requiere y minimizar los tiempos de espera por parte de los clientes.

En la actualidad, encontramos múltiples técnicas para la proyección o pronóstico de las llamadas, por lo cual a continuación se explicaran de manera detallada el funcionamiento de algunas herramientas con las que se pretende realizar el estudio y que pueden generar mejores resultados en base a la literatura encontrada.

3.2.1 Series temporales

Cuando hablamos de pronósticos, podemos asociar diversos factores con la finalidad de explicar cómo se van desarrollando los comportamientos a través del tiempo. Es aquí cuando surgen las series temporales, que básicamente consisten en un conjunto de observaciones tomadas de manera uniforme a través de un determinado periodo de tiempo. En estas observaciones se pueden a través de diferentes modelos lineales, a continuación, se presentan algunos:

3.2.1.1 Modelo AR Un modelo autorregresivo es aquel que permite predecir o pronosticar las futuras observaciones en base a los datos recopilados en el pasado, para ello se realiza un desglose de las observaciones pasadas con su respectivo error aleatorio o ruido blanco, el orden del modelo dependerá directamente de la cantidad de retrasos que se tengan en cuenta para realizar el pronóstico y se expresa de la siguiente manera

$$X_t = \phi_0 + \phi_1 X_{t-1} + \phi_2 X_{t-2} + \dots + \phi_p X_{t-p} + a_t$$

$$a_t \sim RB(0, \sigma^2)$$

Donde:

- X_t : es el valor de la serie de tiempo en el momento t.
- ϕ_0 : es una constante.
- $\phi_1, \dots, \phi_p$: son los coeficientes de autorregresión que indican cómo la serie de tiempo se correlaciona consigo misma a lo largo del tiempo. El número de términos autorregresivos se denota por “p”.
- a_t : es el término de error en el momento

3.2.1.2 Modelo MA Un modelo de media móvil se genera a través de una serie estacionaria la cual se encuentra fuertemente influenciada por los choques aleatorios de las observaciones previamente tomadas, es decir que hay una fuerte dependencia en los errores previos de la serie. Generalmente viene dada por la siguiente expresión:

$$X_t = \theta_0 + a_t - \theta_1 a_{t-1} - \theta_2 a_{t-2} - \dots - \theta_q a_{t-q}$$

$$a_t \sim RB(0, \sigma^2)$$

Donde:

- $\theta_1, \dots, \theta_q$: son los coeficientes de media móvil que indican cómo los errores se correlacionan con la serie de tiempo en los momentos anteriores. El número de términos de media móvil se denota por “q”.
- $a_{t-1}, \dots, a_{t-q}$: son los errores en los momentos anteriores.
- a_t : es el término de error en el momento

- θ_0 : es una constante.
- X_t : es el valor de la serie de tiempo en el momento t.

3.2.1.3 Modelos ARIMA Un modelo de series temporales ARIMA (Autoregressive Integrated Moving Average) se basa en 3 componentes los cuales se encargan de explicar la información de diferentes formas, entre los componentes encontramos el autorregresivo (AR) y el componente de media móvil (MA) como se ha mencionado anteriormente. Únicamente se adhiere un nuevo orden el cual corresponde al componente de diferenciación integrada (I) con el cual se forma en conjunto las siglas (p,d,q) donde cada uno representa el orden del polinomio. En este caso al unir los dos modelos tanto AR como MA, el componente de integración busca garantizar la estacionariedad de la serie, ya que así es posible asegurar que la media y la varianza sean más estables a través del tiempo. También se debe suponer que el ruido blanco es completamente independiente por lo que no existirá correlación entre los errores aleatorios. A continuación, se presenta la notación completa del modelo:

$$X'_t = \phi X'_{t-1} + \phi X'_{t-2} + \phi X'_{t-p} + a_t - \theta_1 a_{t-1} - \theta_2 a_{t-2} - \dots - \theta_q a_{t-q}$$

$$a_t \sim RB(0, \sigma^2)$$

3.2.2 Random Forest

La técnica de Random Forest es un algoritmo de aprendizaje automático que generalmente se emplea para realizar cualquier tipo de clasificación o regresión mediante la combinación de varios árboles de decisión para llegar a un pronóstico final. Ahora bien, cuando hablamos de un árbol de decisión, nos referimos a un conjunto de decisiones tomadas bajo ciertos parámetros para llegar a una respuesta como se puede apreciar en la figura 1. Esta toma este nombre debido a la forma que se crea después de realizar cada pregunta, ya que a medida que avanza se crean ramas que representan las posibles respuestas a las preguntas. Las hojas del árbol representan las etiquetas finales de la clasificación o los valores finales de la regresión o clasificación.

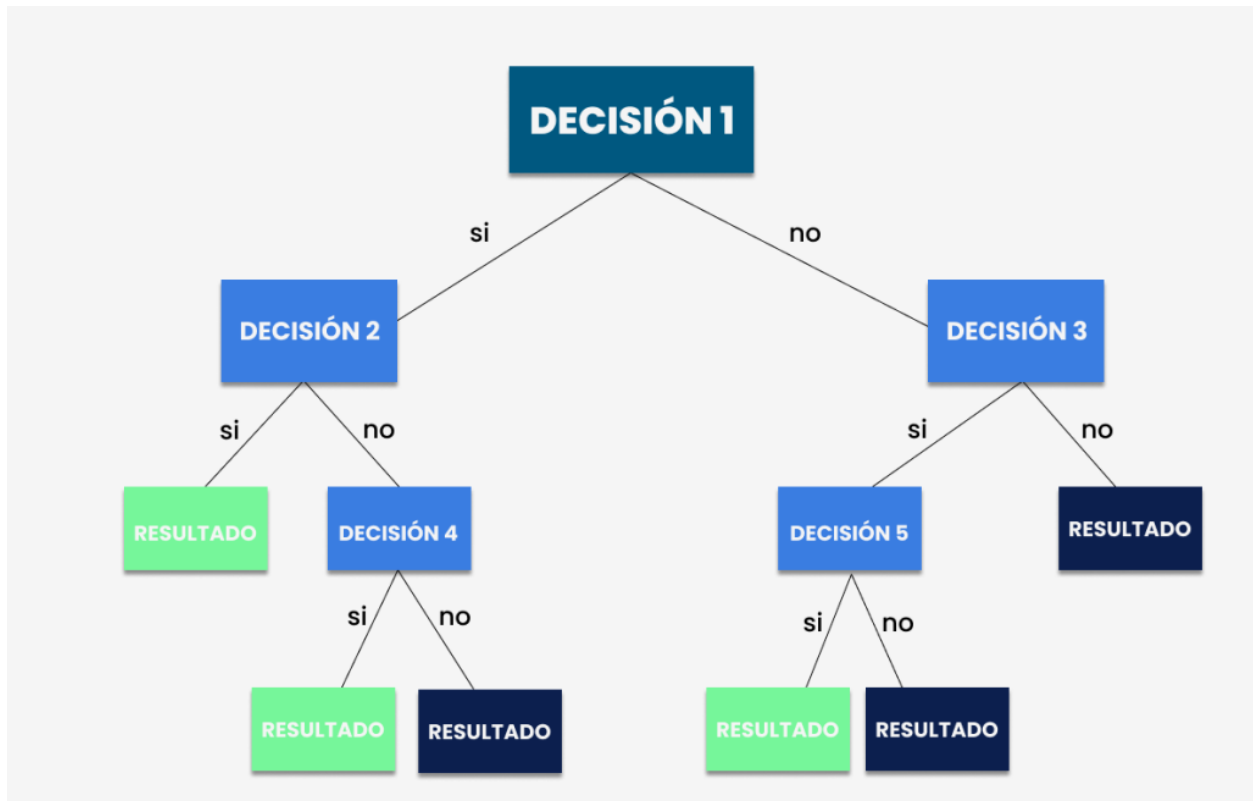


Figure 1: Estructura Arbol de Decision

Cuando se habla de un bosque aleatorio o random forest se agrupan multiples arboles de decision y para la prediccion final en el caso de ser un modelo de clasificacion se obtiene mediante la frecuencia en cada uno de los arboles empleados, mientras que para la regresion corresponde al promedio de los casos evaluados. Cabe aclarar que en cada arbol se genera de manera aleatoria una muestra con reemplazo con la finalidad de asegurar que cada arbol tenga diferentes datos de entrenamiento. En la figura 2 se puede apreciar como se realiza el proceso del bosque aleatorio para llegar a realizar el pronostico

Las ventajas de Random Forest incluyen:

- Puede manejar grandes conjuntos de datos con alta dimensionalidad.
- Es robusto ante valores atípicos y ruido en los datos.
- Proporciona una medida de importancia de las características para la selección de características.
- Es relativamente rápido en la predicción.

Algunas de las desventajas de Random Forest son:

- Puede ser difícil de interpretar y explicar debido a la complejidad de los árboles individuales.
- El entrenamiento de múltiples árboles puede requerir una mayor capacidad de cómputo y tiempo de procesamiento.

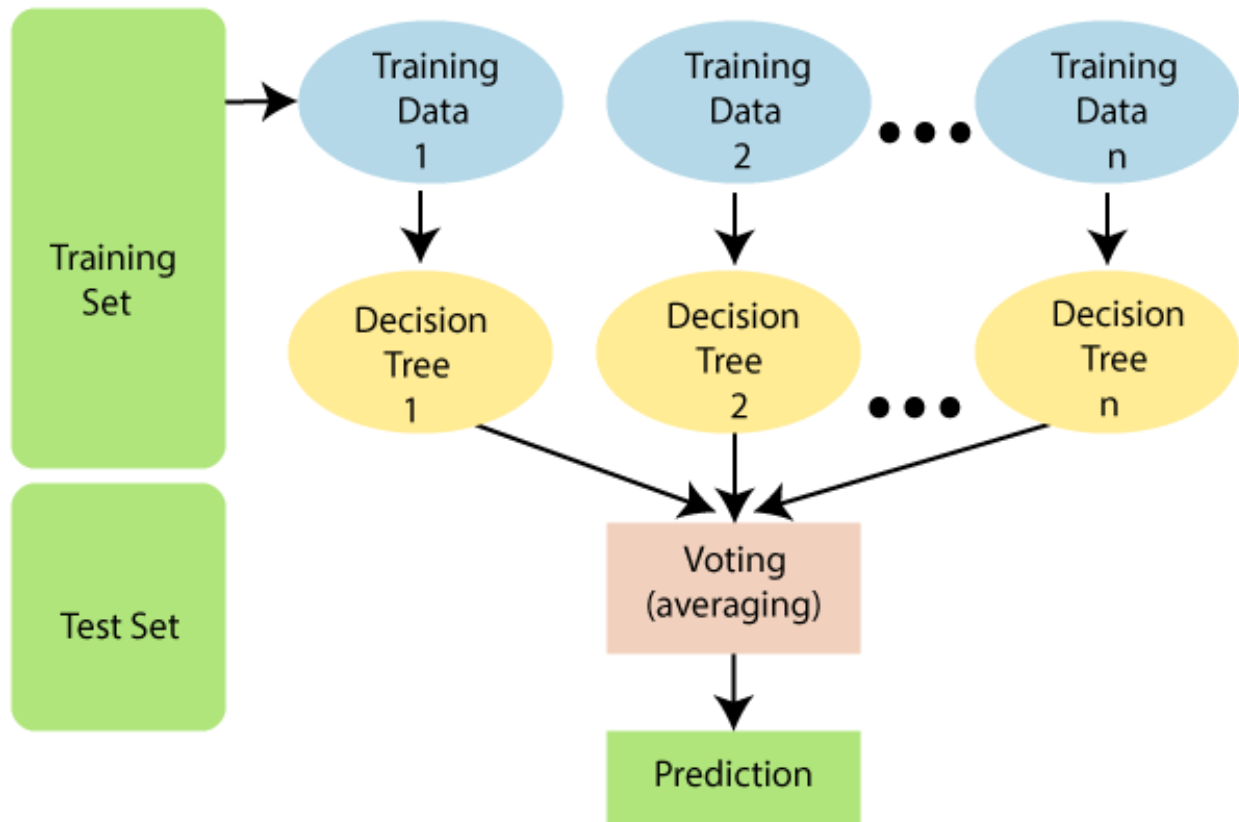


Figure 2: Estructura Random Forest

3.2.3 Redes Neuronales

Las redes neuronales se refieren a un conjunto de conexiones de entrada y salida que se establecen a través de aprendizaje computacional, con el objetivo de simular el comportamiento del cerebro humano. A través de este proceso, las redes neuronales tienen la capacidad de aprender y reconocer patrones específicos con el fin de generalizar la información proporcionada. Esto les permite realizar predicciones o clasificaciones en base a datos aprendidos durante el proceso.

3.2.3.1 Distribución de la red neuronal

- Entrada: Consiste en la recopilación de todos los datos con los cuales se va a entrenar el modelo, por lo que dicha entrada de información global es fraccionada en diferentes combinaciones de entradas simples. Posteriormente se le realiza su respectiva asignación de pesos con lo cual cada entrada puede generar una mayor o menor influencia durante el proceso de clasificación o regresión a esto se le denomina función de entrada. Entre las más comunes encontramos las funciones de la sumatoria, productoria y máximo, cabe aclarar que cada uno se ajusta a los pesos sinápticos y se expresan así:

Sumatoria de entradas

$$\sum_{j=1}^n (n_{ij} w_{ij})$$

Productoria de entradas

$$\prod_{j=1}^n (n_{ij} w_{ij})$$

Maximo de entradas

$$\text{Max}(n_{ij} w_{ij})$$

Donde n_{ij} son las entradas de la red neuronal y w_{ij} son los pesos sinópticos

- Función de activación: El objetivo de la función de activación una vez se ha recibido la información en la red es aplicar una transformación no lineal al modelo para garantizar que no se limite y sea capaz de capturar relaciones mucho más complejas. Entre las funciones destacadas tenemos:

Función Lineal a trozos

$$f(x) = \begin{cases} -1 & x \leq -1/a \\ a * x & -1/a < x < 1/a \\ 1 & x \geq 1/a \end{cases}$$

Función Sigmoidal

$$f(x) = \frac{1}{1 + e^{-x}}$$

Funcion tangente hiperbolica

$$f(x) = \frac{e^{gx} - e^{-gx}}{e^{gx} + e^{-gx}}$$

Funcion Sinusoidal

$$f(x) = A \sin(\omega x + \varphi)$$

- Salida: En este apartado se compilan todas las respuestas en base al modelo y función de activación aplicada, generalmente las variables de respuesta en estos modelos de redes neuronales se encuentran entre $[0,1]$ o $[-1,1]$

En la figura 3 se puede ver una representación de la estructura de una red neuronal teniendo en cuenta las entradas, capas o función de activación y sus respectivas salidas que en este caso serían los pronósticos de la serie.

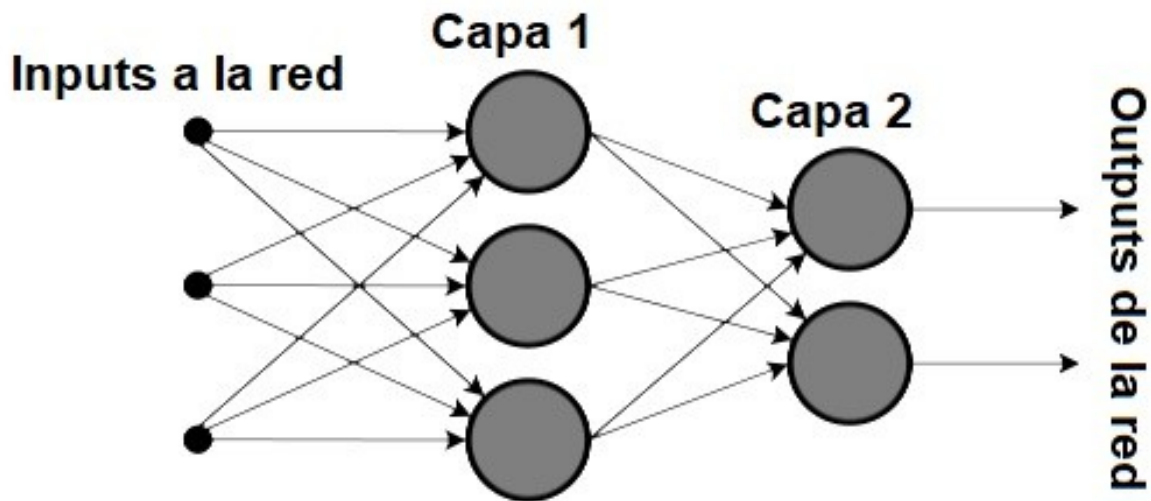


Figure 3: Estructura Red Neuronal

3.2.4 Criterios de seleccion modelo

Durante el diseño de un modelo es necesario generar ciertas medidas de calidad con las que se pueda comparar los modelos entre sí, de esta forma se puede garantizar la selección del modelo más adecuado. Entre las medidas a evaluar a partir de los datos de prueba contra las predicciones encontramos los siguientes

3.2.4.1 Criterio de informacion de Akaike (AIC) Es una medida estadística la cual evalúa la verosimilitud y complejidad del modelo en base a la cantidad de parámetros, destaca ya que busca seleccionar el modelo más simple pero que explique mejor la variable de interés por lo que se busca evitar sobreajustes. La ecuación es la siguiente

$$AIC = 2K - 2\ln(L)$$

Donde:

- K corresponde a la cantidad de parametros del modelo

- L la funcion de verosimilitud

3.2.4.2 Error Cuadratico Medio El error cuadrático medio es una medida empleada para determinar la precisión del modelo el cual consiste en la media de los errores al cuadrado entre las observaciones y los pronósticos generados, es decir que se busca medir la distancia o que tan alejados se encuentran las predicciones al valor original. Se expresa así:

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \tilde{y}_i)^2$$

Donde:

- $\tilde{y}_i$ Valores pronosticados

- y_i valores originales

- n Cantidad de observaciones

3.2.4.3 Coeficiente de Determinacion Es una medida que mide el ajuste del modelo a determinado grupo de datos, es decir que tanto cambia la variable respuesta o dependiente, este coeficiente se mide de 0 a 1. Cuando se encuentra cerca de 0 indica que el modelo no explica bien la variable respuesta, mientras que si se acerca al 1 el modelo captura la variabilidad de los datos. La ecuación del coeficiente se expresa así:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \tilde{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

Donde:

$\bar{y}$ Promedio de las observaciones

y_i valores originales

$\tilde{y}_i^2$ Valores pronosticados

4. Resultados

4.1 Fuentes de información y Datos

Los datos empleados para el estudio corresponden al tráfico de llamadas diario por cierta entidad financiera durante el 1 de agosto de 2016 a el 17 de marzo de 2017. En cuanto a las variables tenemos la fecha, llamadas de ingreso, contestadas, entrantes, abandonadas, transferidas y contestadas en menos de 20 segundos. Así mismo, se cuenta con el nivel de atención que corresponde a la relación entre las llamadas contestadas y las llamadas entrantes, para el nivel de servicio encontramos la relación entre las llamadas contestadas en menos de 20 segundos y las llamadas entrantes y por último el nivel de abandono el cual es equivalente a la relación entre las llamadas abandonas y las llamadas entrantes. Cabe aclarar que estos niveles se encuentran correlacionados entre ellos mismos. Estos fueron los datos obtenidos:

Table 1: Base de registros de llamadas

Fechas	Ingresos	Entrante	Contestada	Abandonada	Contestadas en 20 seg	Transferidas	Nivel Atencion	Nivel Servicio	Nivel Abandono	dia
2016-08-01	2883	3049	2780	269	1859	264	0.912	0.610	0.088	lunes
2016-08-02	2676	2863	2720	143	1939	262	0.950	0.677	0.050	martes
2016-08-03	2526	2795	2666	129	2011	248	0.954	0.720	0.046	miércoles
2016-08-04	2330	2586	2461	125	1888	207	0.952	0.730	0.048	jueves
2016-08-05	2243	2434	2314	120	1652	194	0.951	0.679	0.049	viernes
2016-08-06	706	754	722	32	599	47	0.958	0.794	0.042	sábado
2016-08-07	124	115	98	17	63	0	0.852	0.548	0.148	domingo
2016-08-08	3164	3388	3083	305	1883	335	0.910	0.556	0.090	lunes
2016-08-09	2739	2949	2780	169	2034	275	0.943	0.690	0.057	martes
2016-08-10	2822	3017	2780	237	1769	256	0.921	0.586	0.079	miércoles

Posteriormente se realizó un gráfico temporal con todas las variables para determinar posibles tendencias que se muestran en la Figura 4, en esta oportunidad podemos observar que todos los datos presentan un comportamiento muy similar durante todo el año por lo cual a primera vista no se podría señalar algún comportamiento inusual excepto por ciertos datos atípicos que se pueden encontrar en los primeros días del mes de enero los cuales podrían estar asociados a las festividades que se llevan a cabo en la región. Así mismo, podemos caracterizar que las llamadas que, en su mayoría, se encuentran por debajo de las demás variables corresponden a las llamadas que fueron contestadas en menos de 20 segundos.

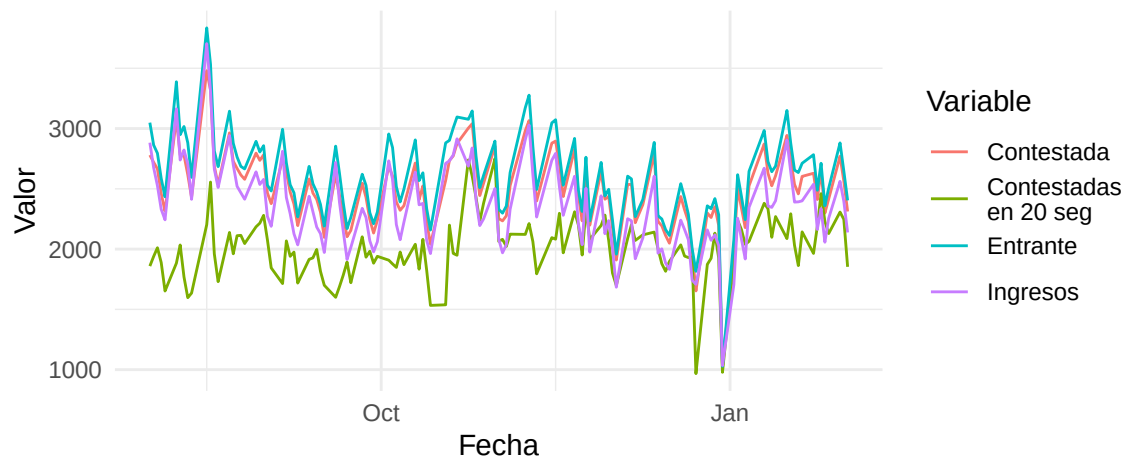


Figure 4: Grafico historico de los registros de llamadas por categoria

Por otro lado, también se observó el comportamiento de cada día de la semana históricamente como se aprecia en la figura 4 y 5, esto con el fin de caracterizar ciertas tendencias o puntos en los que podamos identificar alta variabilidad o que puedan de alguna manera sesgar la informacion al momento de proponer un modelo.

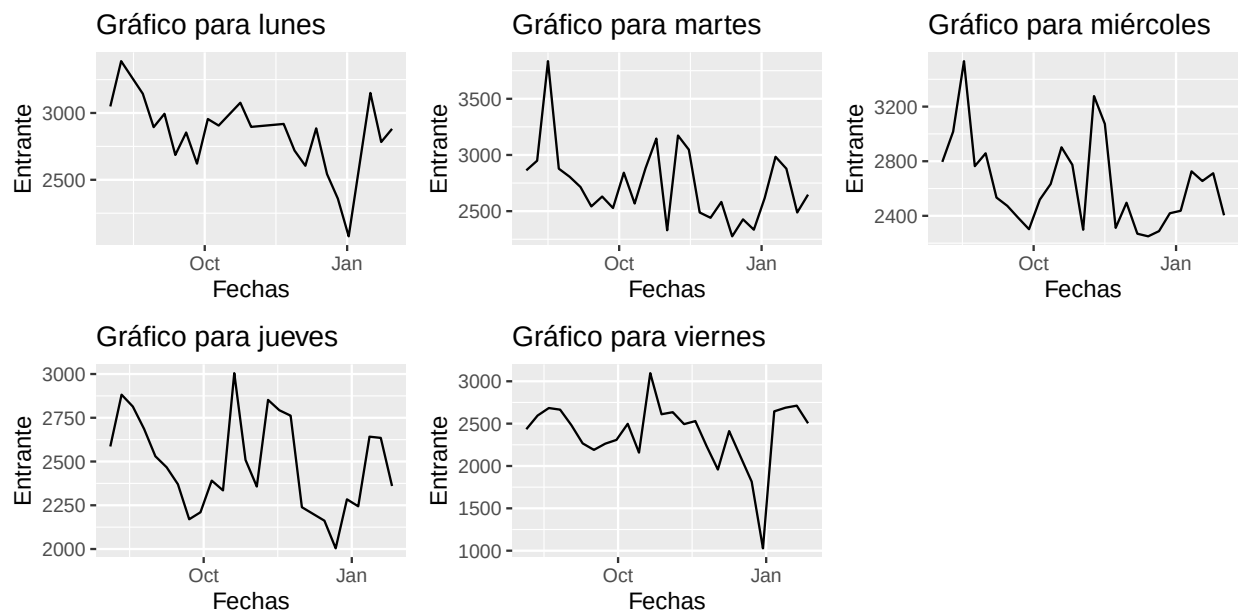


Figure 5: Grafico historico de los registros de llamadas por día de la semana

Encontramos que al realizar un desglose de los días de la semana contra las llamadas entrantes se puede ver una alta variabilidad en los datos y cierta tendencia a la baja del periodo de agosto a enero. Cabe resaltar también que para los días miércoles y jueves hay cierta caída para el mes de octubre, mientras que los días lunes y viernes conservan tendencias relativamente similares durante todo el periodo evaluado. Así mismo, se encuentra que para los días viernes disminuye la cantidad de llamadas alcanzando un máximo de 3000 aproximadamente durante el periodo evaluado y un mínimo de 1000. Se puede destacar que los días con mayor cantidad de llamadas corresponden a los martes y miércoles con mínimos sobre 2300 y máximos de 3500 a 3800.

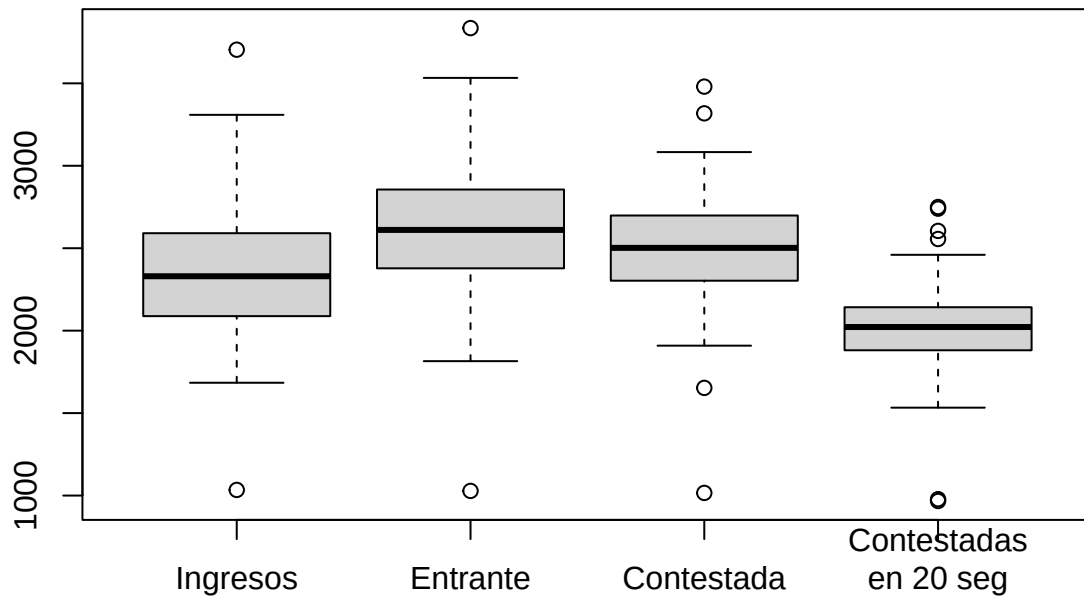


Figure 6: Grafico de caja y bigotes por categoria de llamadas

Por otro lado, en la figura 6 encontramos un grafico de caja y bigotes con el fin de caracterizar la mediana y los cuartiles de las variables que se analizaron para determinar las llamadas entrantes. Encontramos que la mayoría posee una media similar situada entre las 2000 y 2500 aproximadamente y que cuentan con algunos datos atípicos. De igual forma se presentan ciertas medidas descriptivas a continuación:

Table 2: Medidas descriptivas de las variables

Ingresos	Entrante	Contestada	Abandonada	Contestadas en 20 seg	Transferidas	Nivel Atencion	Nivel Servicio	Nivel Abandono
Min. :1034	Min. :1028	Min. :1016	Min. : 12.0	Min. : 967	Min. : 1.0	Min. :0.8920	Min. :0.5330	Min. :0.0110
1st Qu.:2088	1st Qu.:2378	1st Qu.:2303	1st Qu.: 69.0	1st Qu.:1881	1st Qu.: 223.5	1st Qu.:0.9465	1st Qu.:0.7135	1st Qu.:0.0280
Median :2330	Median :2611	Median :2502	Median : 94.0	Median :2022	Median : 252.0	Median :0.9630	Median :0.7900	Median :0.0370
Mean :2351	Mean :2612	Mean :2498	Mean :113.4	Mean :2012	Mean : 261.7	Mean :0.9580	Mean :0.7778	Mean :0.0420
3rd Qu.:2591	3rd Qu.:2856	3rd Qu.:2698	3rd Qu.:140.5	3rd Qu.:2142	3rd Qu.: 280.0	3rd Qu.:0.9720	3rd Qu.:0.8495	3rd Qu.:0.0535
Max. :3704	Max. :3835	Max. :3480	Max. :355.0	Max. :2750	Max. :1604.0	Max. :0.9890	Max. :0.9660	Max. :0.1080

Adicionalmente a lo anterior se realizó un breve análisis descriptivo con las 4 variables que mantienen un comportamiento similar y que se relación con la variable de interés encontrando que su mediana se encuentra entre 2000 y 2600 aproximadamente, también se destaca que poseen distribuciones simétricas y que todas presentan datos atípicos. Las llamadas contestadas en menos de 20 segundos logran conseguir un menor rango intercuartílico respecto a las demás variables, y por la forma de los bigotes se puede decir que los datos se encuentran agrupados o más cercanos al tercer cuartil en comparación con las variables de ingresos y entrante

4.2 Desarrollo

4.2.1 Modelo Arima

Una vez se han procesado y caracterizado todos los datos la investigación comienza diseñando un modelo ARIMA, teniendo en cuenta las características y especificaciones dadas previamente para aplicarlas con dicha informacion. Cabe aclarar que en primera instancia se busca pronosticar la cantidad de llamadas durante los periodos de tiempo que se ejerce la operación en base a los datos del pasado, por lo que no se tendrán en cuenta en este momento otras variables para asociar al momento de crear el modelo.

4.2.1.1 Transformacion de la serie Para la selección de un modelo adecuado para la serie temporal, es necesario realizar ciertas transformaciones para garantizar que este tenga un buen ajuste al momento de seleccionar el orden más adecuado. Por lo tanto, es necesario determinar ciertos comportamientos, para lo cual se realizó el grafico de la serie en primer lugar.

De acuerdo con el grafico de la figura 7 se observa una serie estacionaria, sin aparente tendencia ni señal de estacionalidad únicamente se ve una caída finalizando el mes de diciembre. Sin embargo, este valor atípico se puede explicar debido a la asociación a las festividades de fin de año, fuera de ello se mantiene en el mismo rango entre las 2000 y

3700 llamadas aproximadamente a lo largo del tiempo

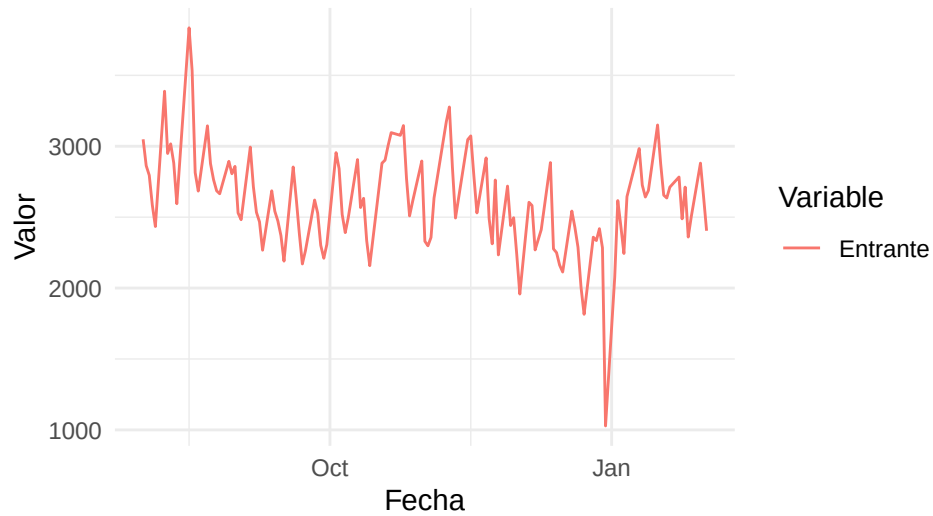


Figure 7: Grafico historico de los registros de llamadas entrantes

Respecto a la función de autocorrelación que se muestra a continuación en la figura 8 encontramos que existe un decaimiento lento hacia cero por lo que se reafirma la falta de estacionariedad de esta y se observan bastantes rezagos significativos. Se podría sugerir un modelo de orden MA dada la representación de la FAC, sin embargo, dada la forma de la función de autocorrelación se buscaría diferenciar por lo menos una vez la serie para obtener una serie estacionaria

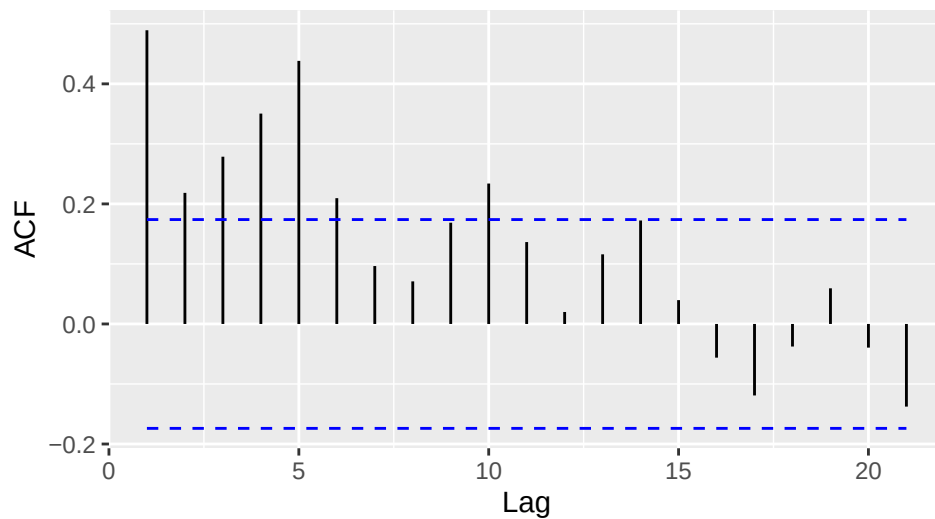


Figure 8: Funcion de autocorrelacion de las llamadas entrantes

Se encuentra una leve formación de ondas en la Figura 9 indicando posibles patrones de estacionalidad en la serie, sin embargo, no son tan claras en ciertos rezagos. De igual forma se observan bastantes rezagos significativos y un

decaimiento rápido hacia cero por lo que también sería viable un modelo de orden AR.

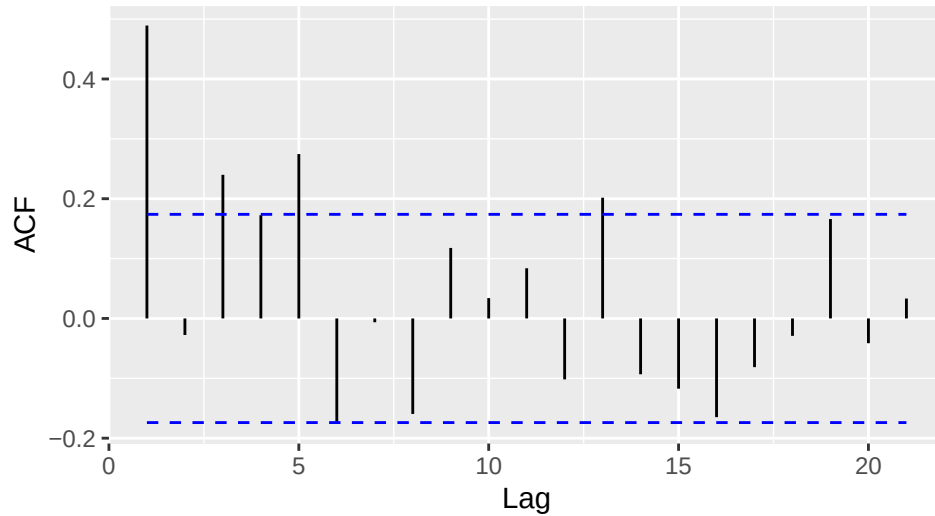


Figure 9: Funcion de autocorrelacion parcial de las llamadas entrantes

Una vez observada tanto la FAC como la PACF, se busca suavizar un poco la serie y estabilizar la varianza a traves de transformaciones a pesar de que encontramos estacionariedad, por lo cual se toman los cuadrados de la serie. De esta formar si la varianza se mantiene constante será posible mejorar la calidad de los pronósticos aumentando la dispersión de los valores atípicos y por lo cual se podría obtener un mejor ajuste y los supuestos del modelo ARIMA validarían todos los supuestos. Se efectúa la prueba para verificar el cumplimiento de estacionariedad en la serie para esto nos apoyaremos en la prueba de Dickey Fuller.

```
##
## Augmented Dickey-Fuller Test
##
## data: calls2
## Dickey-Fuller = -2.9694, Lag order = 5, p-value = 0.1736
## alternative hypothesis: stationary
```

Encontramos que con un p-valor de 0.1736, y un nivel de significancia del 5% existe suficiente evidencia estadística para afirmar que la serie NO es estacionaria. Sin embargo, al aplicar la prueba de phillips perron nos encontramos que con un p valor de 0.01 existe suficiente evidencia estadística para afirmar que al serie es estacionaria.

```
##
## Phillips-Perron Unit Root Test
##
## data:  calls2
## Dickey-Fuller Z(alpha) = -68.866, Truncation lag parameter = 4, p-value
## = 0.01
## alternative hypothesis: stationary
```

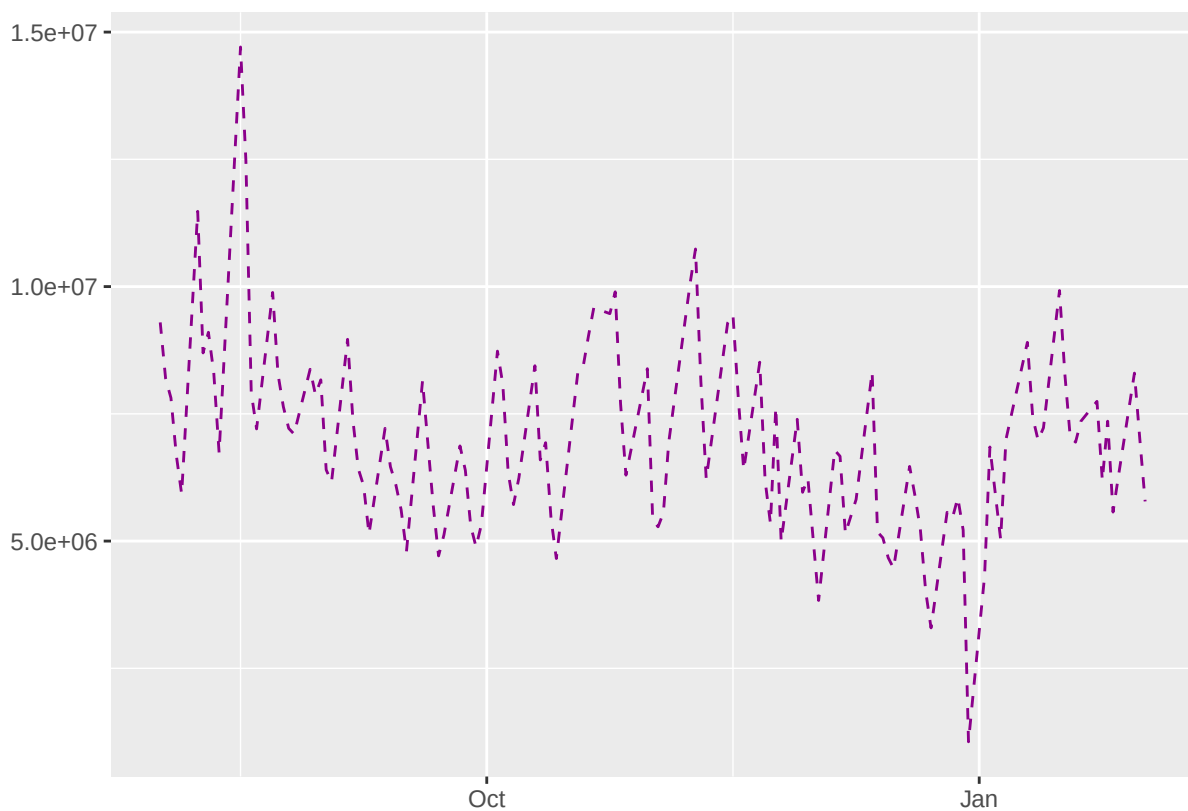


Figure 10: Grafico historico de los registros de llamadas entrantes con transformacion cuadratica

En la figura 10 podemos observar el comportamiento de la serie encontrando valores extremos significativos en los meses de agosto y enero, sin embargo, dichos valores extremos se encuentran ubicados tanto en un mínimo como en su máximo por lo que podemos corroborar que la series NO es del todo estacionaria.

Ahora bien, queda cierta incertidumbre ya que la prueba de Dickey Fuller nos puede dar indicios de algún patrón que se

está repitiendo a través del tiempo el cual deberá ser modelado. Para ello se realizará una transformación para asegurar una mayor estabilidad en la media.

```
##  
## Augmented Dickey-Fuller Test  
##  
## data: dcall1  
## Dickey-Fuller = -6.4573, Lag order = 4, p-value = 0.01  
## alternative hypothesis: stationary  
  
##  
## Phillips-Perron Unit Root Test  
##  
## data: calls2  
## Dickey-Fuller Z(alpha) = -68.866, Truncation lag parameter = 4, p-value  
## = 0.01  
## alternative hypothesis: stationary
```

Nuevamente aplicamos ambos test, y en esta oportunidad nos encontramos que la serie cumple con la estacionariedad, por lo que ya es estable en media y varianza. Una vez ajustado graficamos la función de autocorrelacion y la función de autocorrelacion parcial para establecer los posibles modelos que mejor se puedan ajustar.

Si comparáramos Figura 10 respecto a la Figura 11 existe una diferencia significativa resaltando el suavizamiento de los datos a típicos por lo que se podría mejorar la calidad del pronóstico, únicamente se observa cierto ruido al comienzo de la serie transformada.

4.2.1.2 Seleccíon del modelo Luego de la transformación de la serie se empiezan a evaluar ciertos modelos que se ajusten de la mejor manera a las características de la nueva serie temporal en este caso se tomó un modelo arima (4,1,5) en base a la Figura 12 y 13 en las cuales destaca el quinto y cuarto rezago respectivamente. A continuación, se presentan todos los coeficientes del modelo, así como su criterio de informacion y su sigma cuadrado

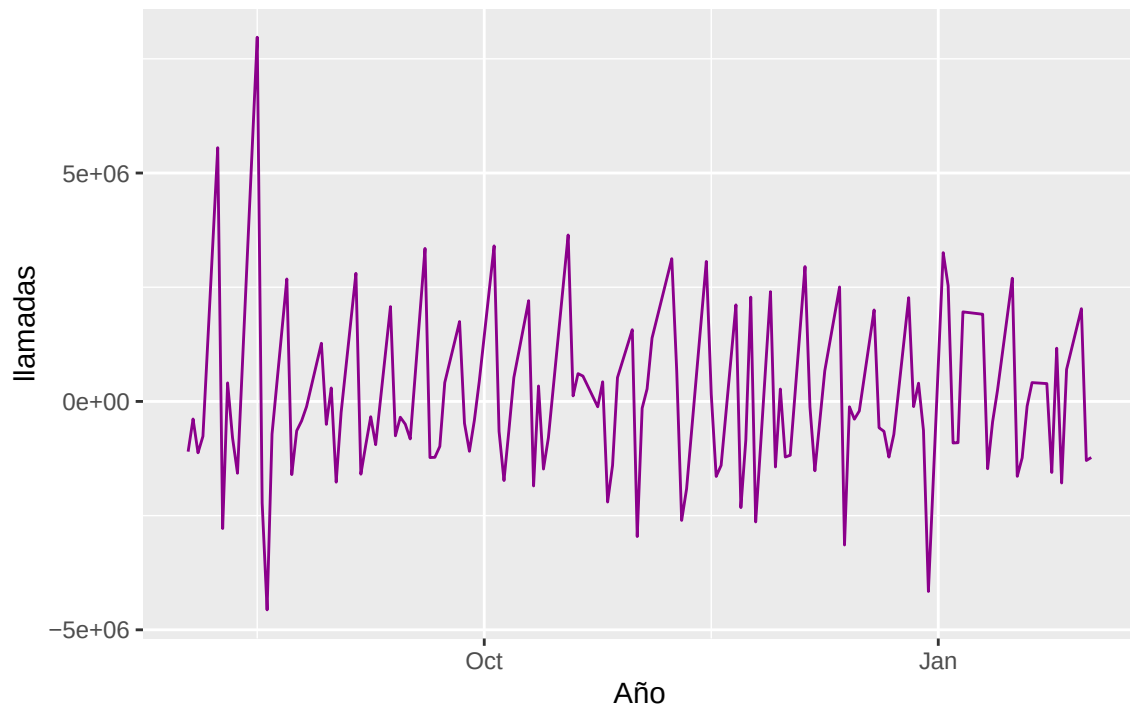


Figure 11: Grafico historico de las llamadas entrantes con transformacion cuadratica y diferenciacion

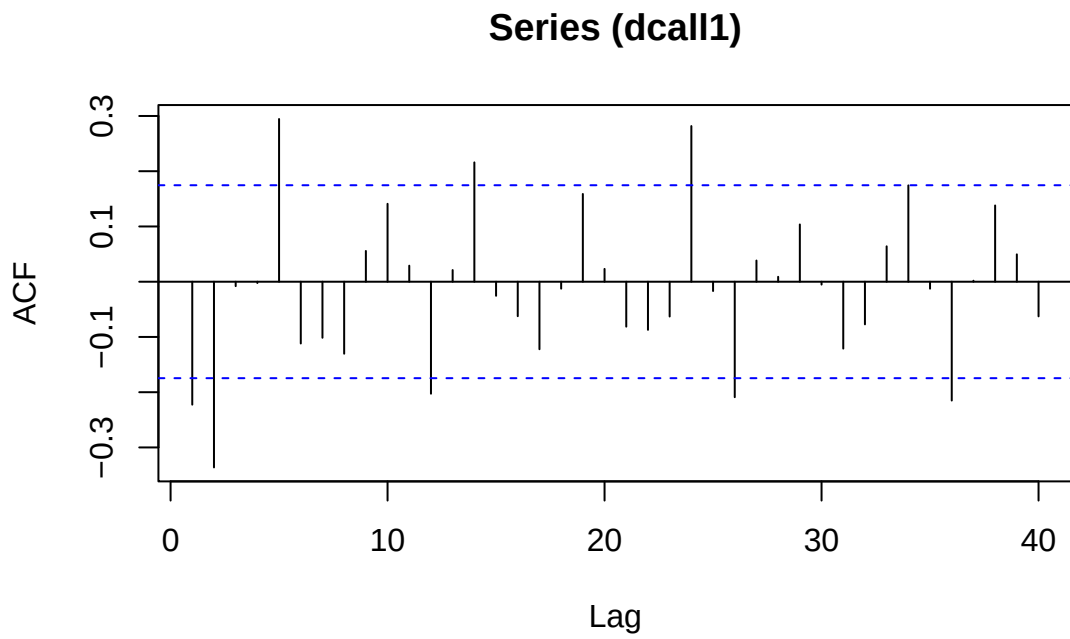


Figure 12: Funcion de autocorrelacion y funcion de autocorrelacion parcial despues de su transformacion

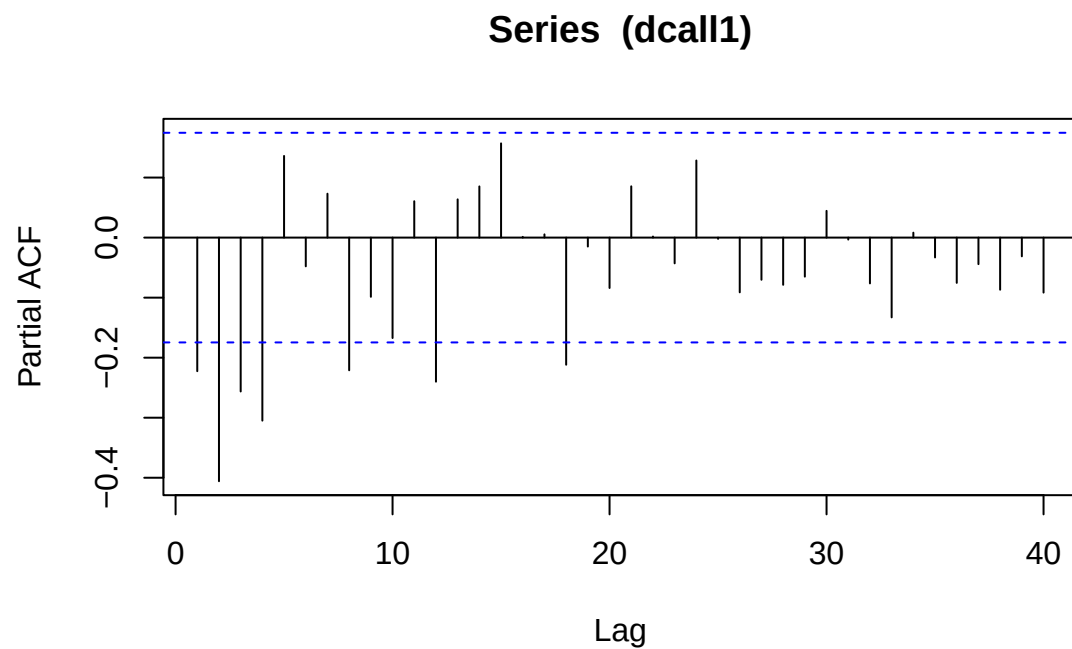


Figure 13: Funcion de autocorrelacion y funcion de autocorrelacion parcial despues de su transformacion

- Coeficientes del modelo

```
##          ar1          ar2          ar3          ar4          ma1          ma2
## -1.16427929 -0.90741039 -1.29384172 -0.73128777  0.78856818 -0.02282861
##          ma3          ma4          ma5
##  0.72136225 -0.07978086 -0.70892842
```

- AIC

```
## [1] 3936.023
```

- Sigma Cuadrado

```
## [1] 1.7311e+12
```

Una vez planteado el modelo verificando la función de autocorrelacion parcial y la función de autocorrelacion nos

Table 3: Comparacion de modelos por varianza y criterio de informacion de Akaike

#	Orden (p,d,q)	AIC	σ^2
m1	(2,1,5)	3947.24	2.010e+12
m2	(4,1,5)	3936.02	1.731e+12
m3	(2,1,2)	3947.75	2.154e+12
m4	(1,1,2)	3946.13	2.193e+12

apoyamos en la función autoarima dado que a traves de esta es posible evaluar múltiples modelos y compararlos por sus criterios de informacion.

Se plantearon 3 modelos basados en los cortes observados en la función de autocorrelacion y autocorrelacion parcial, dado su decaimiento hacia cero y determinando los rezagos significativos. Así mismo, nos apoyamos de la función de autoarima como se mencionó anteriormente para determinar un posible cuarto modelo para evaluar cual podría ser más viable, cabe aclarar que en esta oportunidad se selecciona el modelo de menor criterio de informacion respecto a las posibles combinaciones que se realizaron previamente. Estos fueron los modelos propuestos:

Al momento de comparar los modelos podemos decir que el mejor modelo dado el AIC corresponde a el modelo arima número 2, el cual también posee la menor varianza por lo que podría explicar de una mejor manera como se distribuyen los datos, así como un menor ruido por lo que podría capturar de una mejor manera la estructura de la serie temporal. A continuación, se presenta la ecuación del modelo:

$$y_t = -1.1643y_{t-1} - 0.9074y_{t-2} - 1.2938y_{t-3} - 0.7313y_{t-4} + 0.7886\epsilon_{t-1} - 0.00228\epsilon_{t-2}*$$

$$*0.7214\epsilon_{t-3} - 0.0798\epsilon_{t-4} - 0.7089\epsilon_{t-5} + \epsilon_t$$

$$\epsilon_t \sim N(0, \sigma^2)$$

4.2.1.3 Supuestos del modelo

Correlacion Una vez definido el modelo se procede a evaluar los supuestos para determinar qué tan viable es el modelo y en que partes puede llegar a tener falencias durante la etapa de pronósticos, en primera instancia se realizara un gráfico de los residuales para determinar si hay algún tipo de distribución en los datos

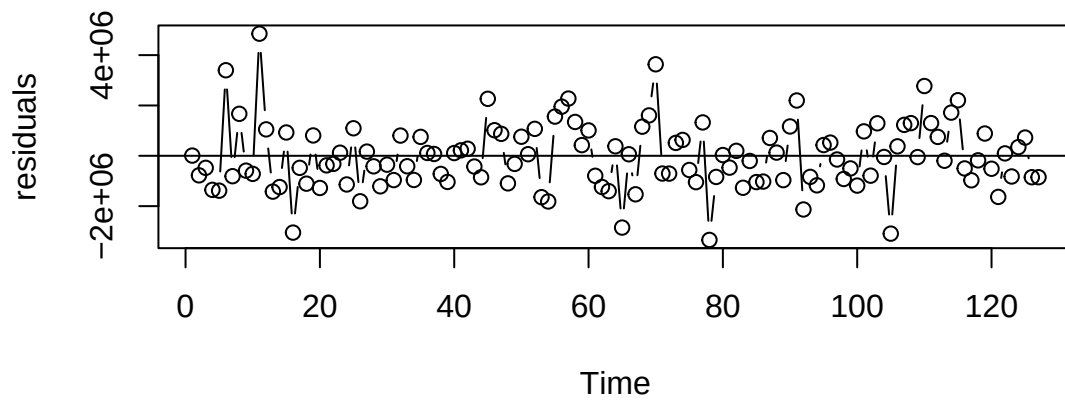


Figure 14: Residuales del modelo aplicado

Encontramos que los datos no siguen una tendencia definida, ni tampoco alguna clase de patrón por lo que se puede decir que el modelo está capturando la información de forma adecuada. Por lo que se procede a evaluar los residuales en su función de autocorrelación y autocorrelación parcial

Sample ACF of Residuals from ARIMA(4,1,5) Modelo Calls

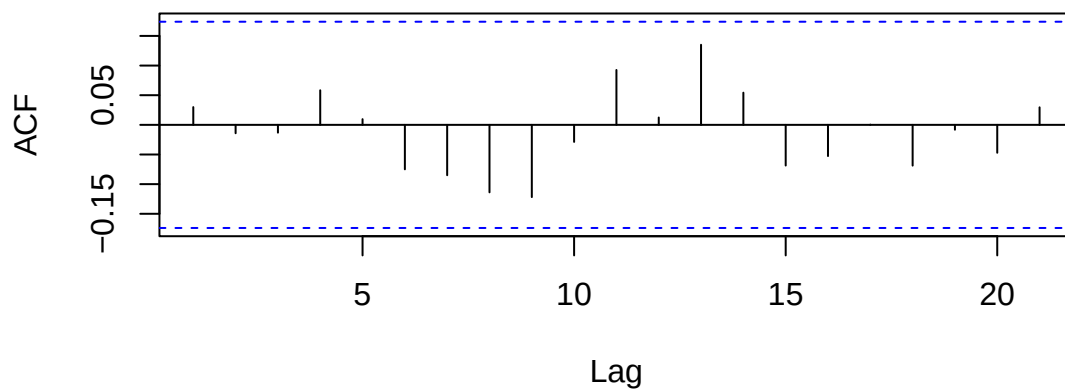


Figure 15: Funcion de autocorrelacion y autocorrelacion parcial de los residuales

Sample PACF of Residuals from ARIMA(4,1,5) Modelo Calls

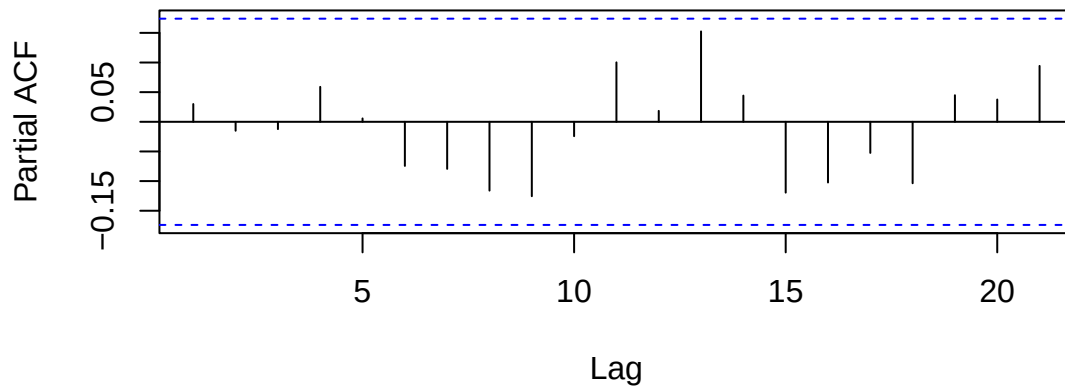


Figure 16: Funcion de autocorrelacion y autocorrelacion parcial de los residuales

Se puede observar que no hay rezagos significativos tanto en la FAC como FACP de los residuales como se muestra en la figura 15 y 16, por lo que esto nos puede dar indicios que el modelo se ajusta de manera correcta en su parte autorregresiva y de media móvil

```
##  
## Box-Ljung test  
##  
## data: residuals(m1)  
## X-squared = 6.342, df = 10, p-value = 0.7858
```

Con el fin de descartar correlación en los residuales se realizó el test de Box-Ljung este consiste en la evaluación de rezagos para identificar correlaciones muestrales con las correlaciones esperadas encontrando que un p-valor de 0.7229 existe suficiente evidencia estadística para afirmar que los residuales del modelo no están correlacionados

Normalidad Respecto al supuesto de normalidad se realizaron a traves de pruebas de hipotesis y observacion del grafico QQ-plot para determinar si existe la presencia de esta en los residuales.

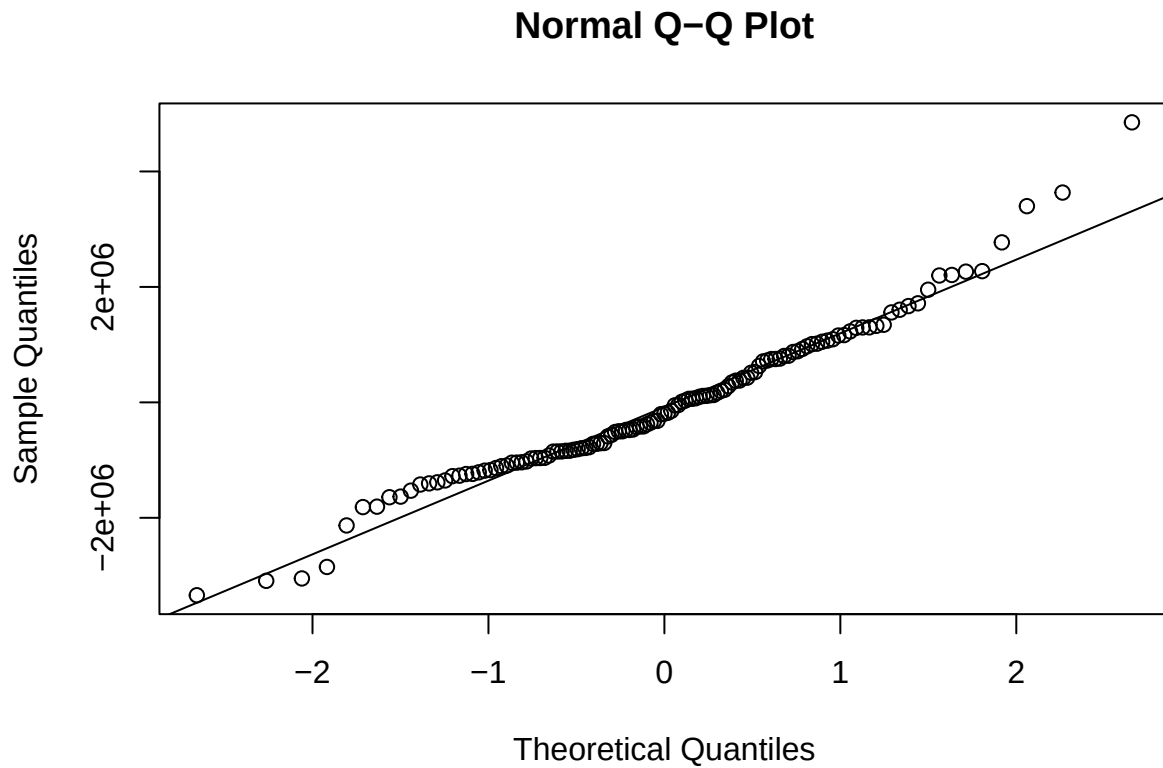


Figure 17: Grafico Cuantil Cuantil de los residuales del modelo

Existe cierta incertidumbre en el supuesto de normalidad dado que los residuales tienen un buen ajuste salvo por algunos datos atípicos que se pueden apreciar en el gráfico. Sin embargo, se aplica un segundo test en este caso la prueba de kolmogorov - Smirnov encontrando que con un p-valor de 0.5662 existiría normalidad en los residuales.

```
##
## Asymptotic one-sample Kolmogorov-Smirnov test
##
## data: residuals(m1)
## D = 0.069795, p-value = 0.5662
## alternative hypothesis: two-sided
```

Media cero Aplicamos una prueba sencilla de t-student para evaluar el supuesto de media cero en los residuales, encontrando que con un p-valor de 0.6733 no hay suficiente evidencia estadística para afirmar que la media de los

residuales es cero. Esto es un buen indicio, ya que nos asegura que no hay algún tipo de tendencia o sesgo

```
##  
## One Sample t-test  
##  
## data: residuals(m1)  
## t = -0.42263, df = 126, p-value = 0.6733  
## alternative hypothesis: true mean is not equal to 0  
## 95 percent confidence interval:  
## -280189.7 181575.3  
## sample estimates:  
## mean of x  
## -49307.21
```

4.2.1.4 Pronosticos del modelo Una vez se han validado los supuestos del modelo, pasamos a la etapa de uso, para ello se realizó el pronóstico para los próximos 32 periodos de la serie, en este caso se realizaron 32 pronósticos ya que esta fue la cantidad de datos de entrenamiento que se tomó con la base de datos empleada para dicho estudio.

```
## $pred  
## Time Series:  
## Start = 128  
## End = 159  
## Frequency = 1  
## [1] 5665453 6204136 7168679 7273315 6259330 5703046 6430076 7323805 7084798  
## [10] 6018232 5788878 6679387 7405455 6828766 5856897 5921082 6943416 7374062  
## [19] 6572664 5745266 6130978 7154647 7269386 6312926 5715865 6381862 7301835  
## [28] 7098347 6075400 5773708 6643711 7376883  
##
```

```

## $se

## Time Series:

## Start = 128

## End = 159

## Frequency = 1

## [1] 1328176 1557472 1567274 1682684 1702965 1790025 1841740 1893924 1940186

## [10] 1976828 2044606 2086164 2140563 2168498 2220321 2263539 2315750 2349057

## [19] 2385049 2425450 2472147 2513605 2544932 2578802 2617685 2662200 2696913

## [28] 2727204 2758523 2798391 2837549 2869489

## Time Series:

## Start = 128

## End = 159

## Frequency = 1

## [1] 2380.221 2490.810 2677.439 2696.908 2501.865 2388.105 2535.759 2706.253

## [9] 2661.728 2453.208 2406.009 2584.451 2721.297 2613.191 2420.103 2433.327

## [17] 2635.036 2715.522 2563.721 2396.929 2476.081 2674.817 2696.180 2512.554

## [25] 2390.787 2526.235 2702.191 2664.272 2464.833 2402.854 2577.540 2716.042

```

Ahora bien, se procede a transformar la serie nuevamente a sus valores originales para determinar el mejor ajuste y los valores reales.

Por ultimo, calculamos el error cuadratico medio de la serie encontrando un valor de 230.090 mientras que el error cuadratico medio porcentual fue de 8,1430%

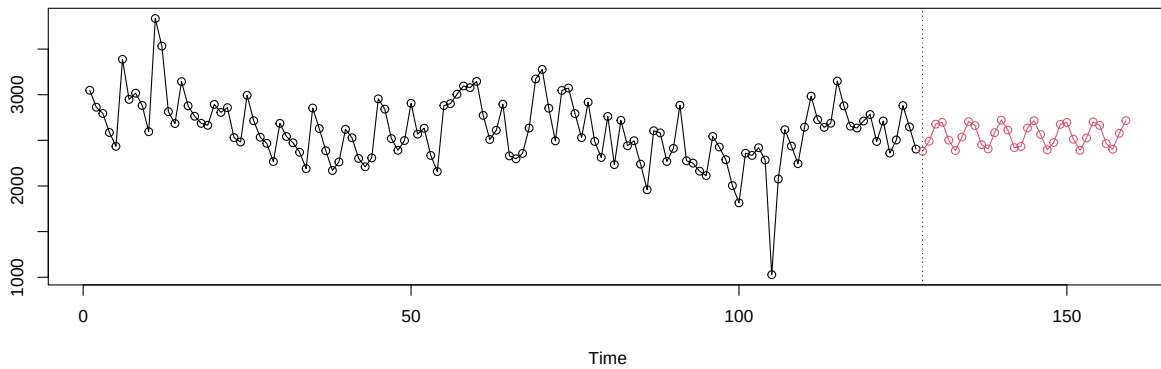


Figure 18: Grafico serie temporal llamadas entrantes mas pronosticos realizados

4.2.2 Random Forest

Se inicia el desarrollo del modelo Random Forest dividiendo los datos en conjuntos de entrenamiento y prueba, con una proporción de 80 - 20 respectivamente. El primer grupo se utiliza para entrenar el modelo, mientras que el segundo se utiliza para evaluar la calidad de las predicciones y determinar su precisión.

En este primer modelo de Random Forest, la variable respuesta seleccionada es el número de llamadas entrantes, mientras que las variables explicativas incluyen el día de la semana y los retrasos de las llamadas. A continuación, se presenta la tabla y el grafico correspondientes a estos datos.

En este primer random forest se tomo como variable respuesta las llamadas entrantes y como variables explicativa el dia de la semana y los retardos de la misma obteniendo la siguiente tabla y grafico

Table 4: Comparacion datos de prueba vs predicciones modelo Random Forest con unicamente los rezagos

Numero	Prediccion	Datos de prueba
1	2351.73	2505.00
2	2502.22	2424.00
3	2855.07	3074.00
4	2888.17	2690.00
5	2561.86	2736.00
6	2559.55	2464.00
7	2434.03	2304.00
8	2698.07	2930.00
9	2651.79	2609.00
10	2533.32	2815.00
11	2570.08	2509.00
12	2493.26	2669.00
13	2945.63	3059.00
14	2875.22	2845.00
15	2658.35	2551.00
16	2433.28	2402.00
17	2428.47	2405.00
18	2837.26	3009.00
19	2839.66	3049.00
20	2941.24	2550.00
21	2433.28	2599.00
22	2574.23	2442.00
23	2899.56	2996.00
24	2789.24	2652.00
25	2553.81	2525.00
26	2428.72	2667.00
27	2547.89	2532.00
28	2891.22	3160.00
29	2891.83	2957.00
30	2851.09	3102.00
31	2855.45	2978.00
32	2699.37	2895.00

En la figura 19 se presenta una comparacion entre las predicciones realizadas vs los valores de prueba para identificar o reconocer la calidad de los pronosticos en cuestion.

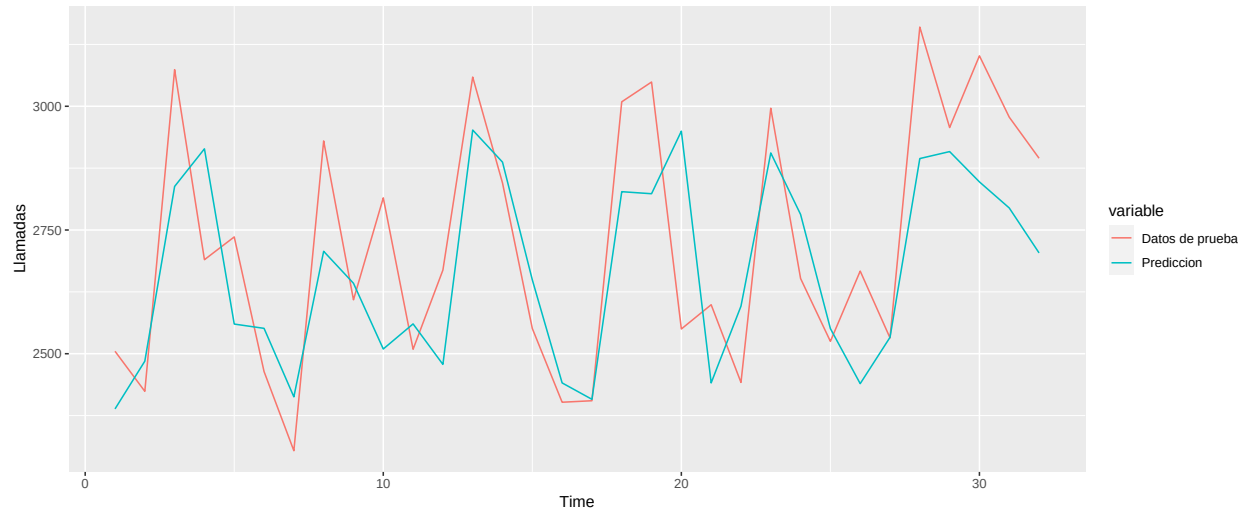


Figure 19: Comparacion prediccion y datos de prueba con Random forest aplicado con un retardo de la serie

Después de realizar el pronóstico, notamos que el modelo presenta cierto ruido en las predicciones, lo que se refleja en datos atípicos. Por lo tanto, resulta interesante evaluar otras variables explicativas. Calculamos el error absoluto medio de 145.20159 y un MAPE (porcentaje de error absoluto medio) de 5,2574%. Para seleccionar las variables que se agregaran al modelo, realizamos una matriz de correlación que se presenta en la figura 20. De esta manera, buscamos verificar cuales variables explican de mejor manera el modelo.

Para la seleccion de las variables que se van agregar al modelo se realizo una matriz de correlacion, de esta forma se buscar verificar cuales explican de mejor manera el modelo

```
## corrplot 0.92 loaded
```

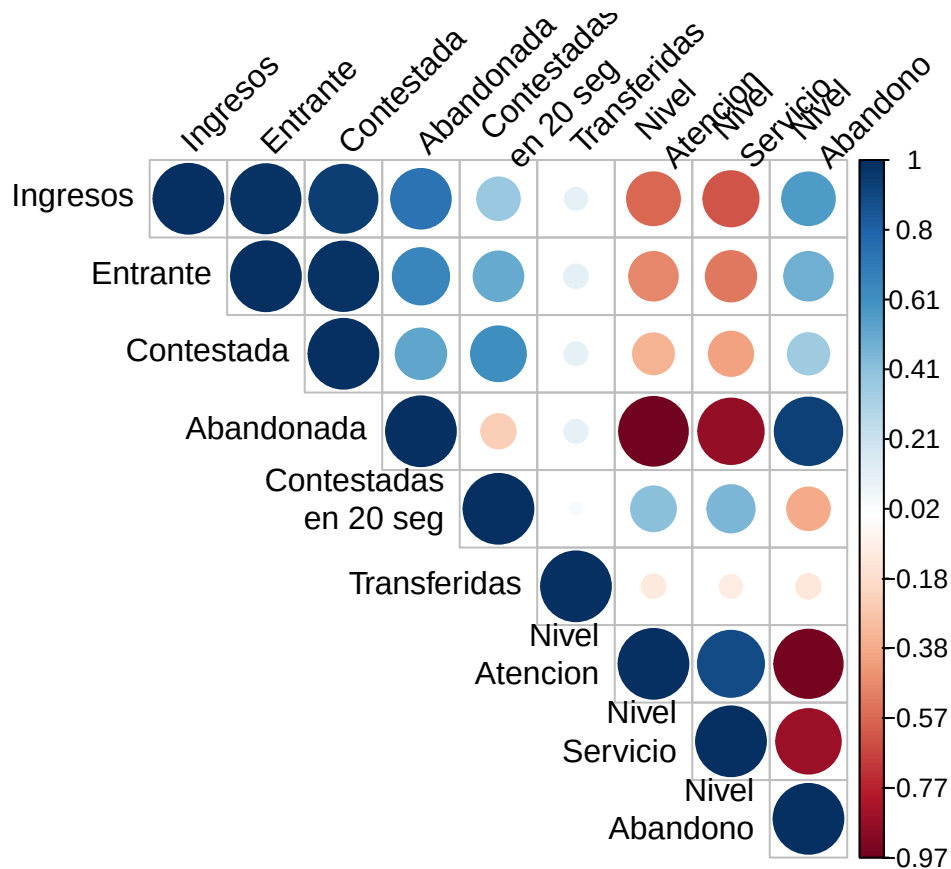


Figure 20: Grafico de correlaciones entre las variables de la base de datos

A partir de aquí, se observa que las llamadas entrantes muestran una correlación significativa con las llamadas contestadas, las llamadas abandonadas y los ingresos. Por otro lado, existe una correlación negativa entre los niveles de abandono y atención. Además, no se encuentra ninguna correlación entre las llamadas entrantes y las llamadas transferidas o las llamadas contestadas en menos de 20 segundos

A continuación, llevamos a cabo el mismo proceso anterior teniendo en cuenta las variables mencionadas anteriormente, basándonos en su correlación. Los datos resultantes son los siguientes: el modelo random forest muestra un ajuste menor en comparación con el modelo anterior, aunque se observa un aumento en el error en las últimas predicciones. En este caso, se calcula el error absoluto medio de 166.6695 y un MAPE de 5.9997%.

En la figura 21 y en la tabla 5 se presenta una breve comparativa de los resultados al momento de realizar la predicción encontrando un mejor ajuste en los mínimos de la gráfica, sin embargo si se observa un mayor error en los valores máximos de la serie

Table 5: Comparacion datos de prueba vs predicciones modelo Random Forest con rezagos y variables explicativas adicionales

Numero	Prediccion	Datos de prueba
1	2383.69	2505.00
2	2549.48	2424.00
3	2887.17	3074.00
4	2900.28	2690.00
5	2684.06	2736.00
6	2628.34	2464.00
7	2380.65	2304.00
8	2668.97	2930.00
9	2808.93	2609.00
10	2796.63	2815.00
11	2607.15	2509.00
12	2598.02	2669.00
13	2849.20	3059.00
14	2879.92	2845.00
15	2668.02	2551.00
16	2452.90	2402.00
17	2327.24	2405.00
18	2570.67	3009.00
19	2784.28	3049.00
20	2898.25	2550.00
21	2738.72	2599.00
22	2535.57	2442.00
23	2912.92	2996.00
24	2891.44	2652.00
25	2642.07	2525.00
26	2500.94	2667.00
27	2627.63	2532.00
28	2878.32	3160.00
29	2879.06	2957.00
30	2625.46	3102.00
31	2814.60	2978.00
32	2622.69	2895.00

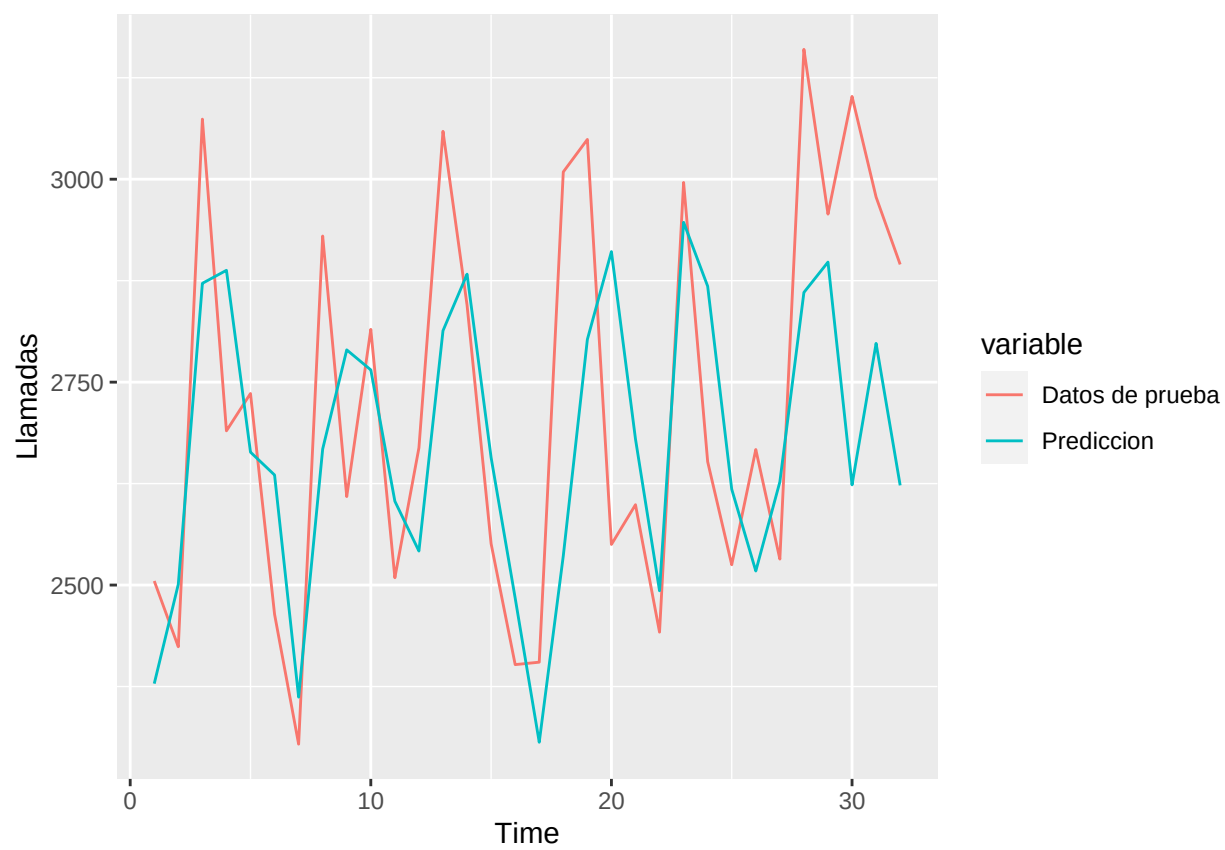


Figure 21: Comparacion prediccion y datos de prueba con Random forest a traves de variables incluidas en el modelo por correlacion

En ambos modelos se utilizó un bosque aleatorio compuesto por un total de 500 árboles de decisión. Para evaluar la convergencia de las predicciones a medida que aumenta la cantidad de árboles, se realizó un gráfico comparativo (Figura22). Se observó que a partir del árbol numero 50 la cantidad de llamadas pronosticadas se estabiliza, lo que indica que se podría reducir significativamente la cantidad de árboles utilizados sin comprometer la precisión del modelo

Finalmente se generó un gráfico que muestra los datos de prueba y las predicciones realizadas por ambos modelos de random forest (Figura 23). Los resultados obtenidos fueron muy similares entre los dos modelos, a pesar de que el primero contiene menos variables. Sin embargo, se observó que el primer modelo presenta un error aproximadamente de 0.3% menor en comparación con el segundo modelo random forest

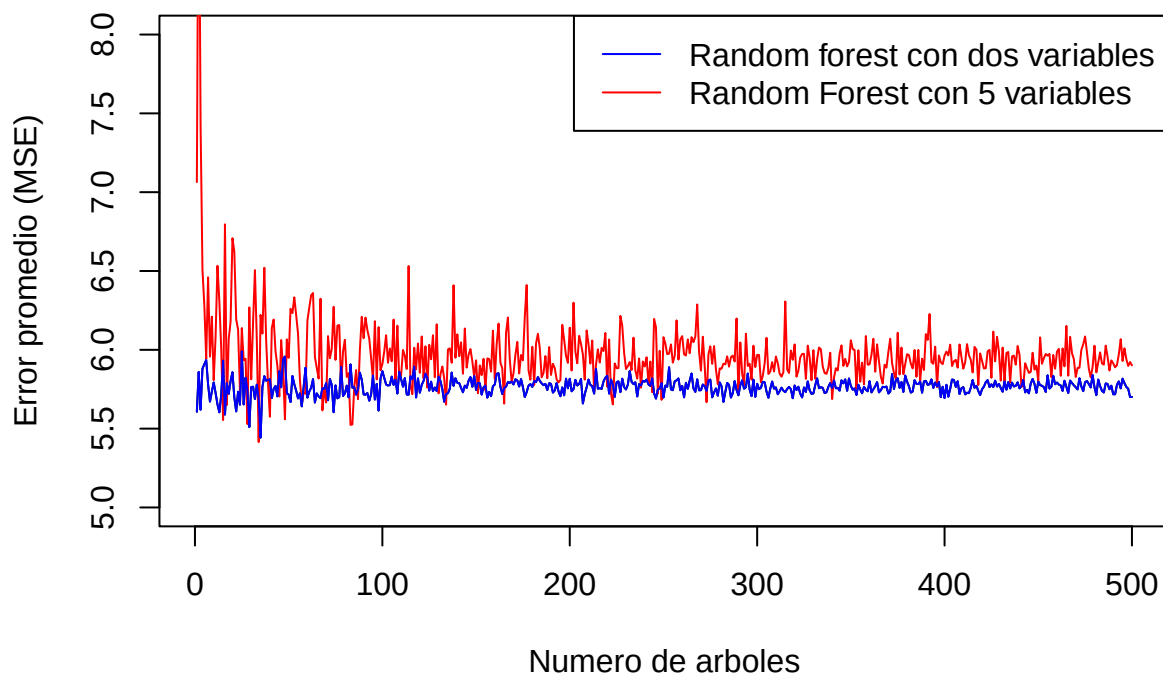


Figure 22: Convergencia de random forest al incrementar la cantidad de arboles de decision

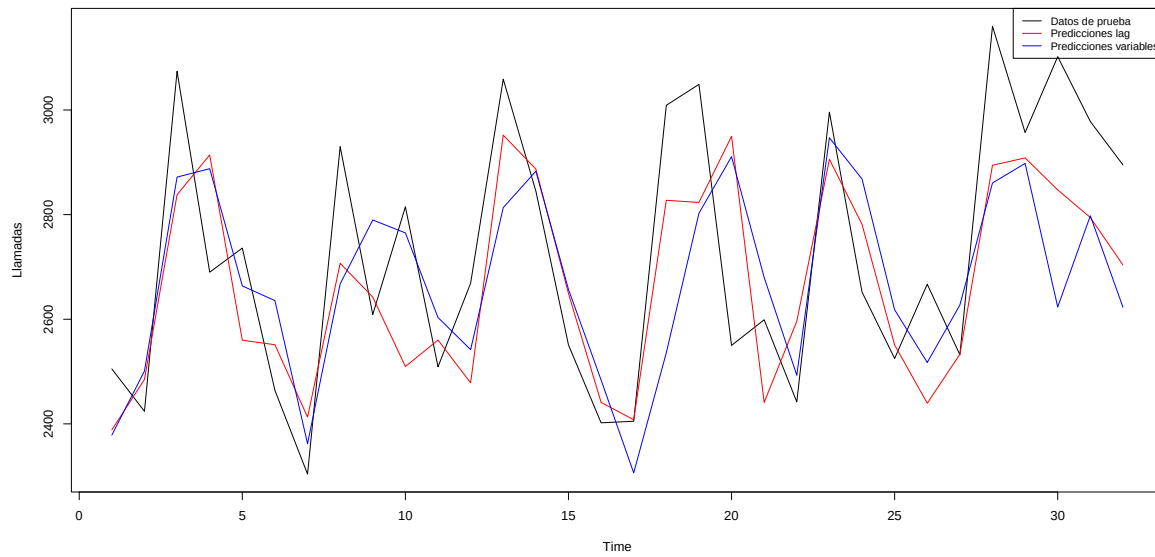


Figure 23: Comparacion de las predicciones entre los datos de prueba con todos los modelos random forest

4.2.3 Redes Neuronales

En este caso, se tomaron los datos y se dividió en conjuntos de entrenamiento y prueba para utilizarlos en el entrenamiento de una red neuronal de tipo feedforward o de propagación hacia adelante. La estructura de la red consiste en 3 capas ocultas, donde la primera capa tiene 2 unidades de procesamiento, la segunda capa tiene 1 unidad, y la última capa tiene una unidad de procesamiento que proporciona la respuesta correspondiente.

Antes de ejecutar la red neuronal, se realizó la normalización de las variables tanto en el conjunto de entrenamiento como en el conjunto de prueba. Esta técnica permite facilitar el proceso de aprendizaje de la red neuronal y evita el sesgo causado por valores extremos durante el proceso de aprendizaje, asegurando que no haya un aumento desproporcionado en los pesos de las unidades de procesamiento

En cuanto a la función de activación, se utiliza la función sigmoideal o logística como referencia. Continuación se presenta la estructura de la red en a la figura 24

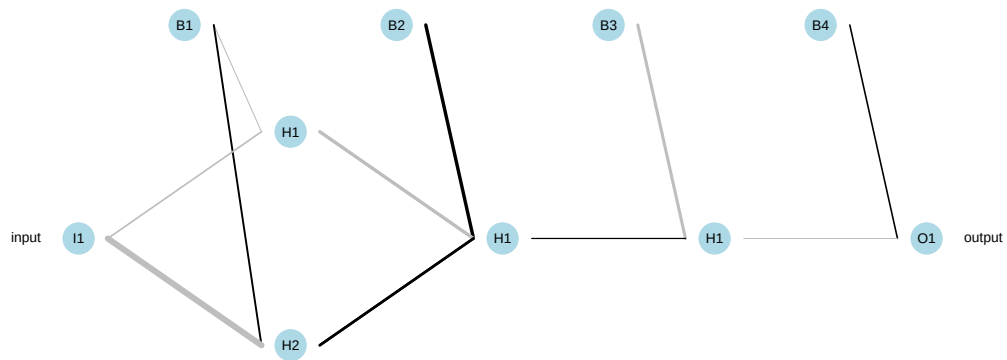


Figure 24: Estructura de la red neuronal

Cada union de capas representa el peso al cual se encuentra sujeta cada unidad de procesamiento, por lo que entre mas intenso sea el color significa que hay un peso mayor en estos nodos generando asi una mayor influencia al momento del aprendizaje de la red asi como su pronostico. Estos fueron los pesos encontrados en la red:

```
## [[1]]
## [[1]][[1]]
##          [,1]      [,2]
## [1,] -0.4542610  0.7938997
## [2,] -0.7105004 -2.4590090
##
## [[1]][[2]]
##          [,1]
## [1,]  1.397828
## [2,] -1.263208
## [3,]  1.055544
```



```
##
## [[1]][[3]]
##          [,1]
## [1,] -1.2719192
## [2,]  0.4735837
##
## [[1]][[4]]
##          [,1]
## [1,]  0.6704034
## [2,] -0.3696637
```

Finalmente, se presenta una tabla con las predicciones y un gráfico correspondiente (figura 25). En el gráfico, podemos observar que el modelo muestra ajustes y errores similares a los modelos anteriores aplicados. Es importante destacar que la red neuronal tiene dificultades para pronosticar los datos que se encuentran en los mínimos de la serie, lo cual es notable.

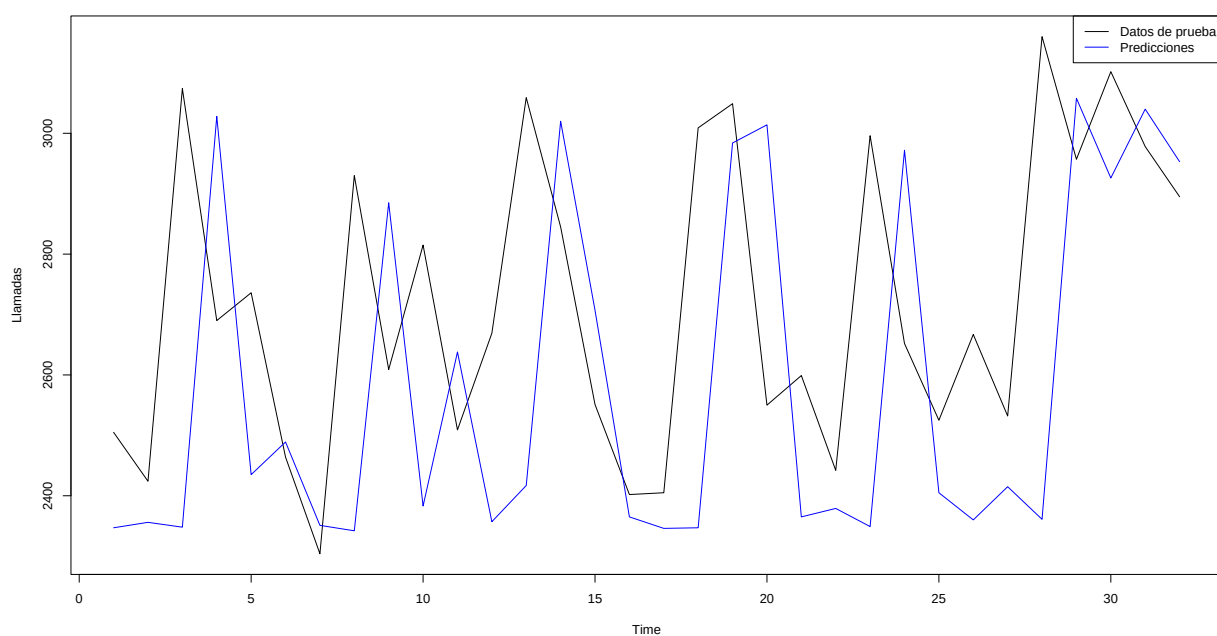


Figure 25: Comparacion prediccion y datos de prueba con red neuronal aplicado con un retardo de la serie

Table 6: Comparacion datos de prueba vs predicciones meidante red neuronal con rezagos de la serie

Numero	Prediccion	Datos de prueba
1	2347.04	2505.00
2	2356.79	2424.00
3	2348.36	3074.00
4	3028.32	2690.00
5	2435.17	2736.00
6	2489.21	2464.00
7	2351.77	2304.00
8	2342.83	2930.00
9	2885.24	2609.00
10	2383.26	2815.00
11	2638.16	2509.00
12	2357.39	2669.00
13	2417.37	3059.00
14	3020.36	2845.00
15	2707.69	2551.00
16	2365.27	2402.00
17	2346.92	2405.00
18	2347.10	3009.00
19	2984.63	3049.00
20	3014.43	2550.00
21	2365.04	2599.00
22	2379.36	2442.00
23	2349.75	2996.00
24	2972.43	2652.00
25	2405.49	2525.00
26	2360.02	2667.00
27	2415.87	2532.00
28	2361.31	3160.00
29	3058.45	2957.00
30	2926.49	3102.00
31	3040.64	2978.00
32	2953.08	2895.00

4.3 Comparacion de los modelos aplicados

Después de ejecutar todos los modelos establecidos previamente, se realiza una comparación visual y utilizando medidas descriptivas para determinar cual modelo puede tener el mejor ajuste en el pronóstico de series temporales y en qué condiciones se deben realizar para garantizar la reproducibilidad en los datos a lo largo del tiempo. Esta comparación nos permitirá identificar el modelo que se ajuste mejor a nuestros datos y que pueda proporcionar pronósticos más precisos y confiables en función de las características y patrones temporales presentes en los datos. A continuación, se muestra en la tabla 7, la serie original de datos de prueba y sus predicciones por cada uno de los métodos usados (series de tiempo, Random forest y redes neuronales). Adicionalmente, con estas series se hace un gráfico comparativo

(Figura 26) en donde se puede observar que la serie temporal no logra capturar tanto la variabilidad de los datos dado que en los picos de la serie pronosticada existe ciertas diferencias respecto a los demás modelos planteado, para el caso de la redes neuronales nos encontramos con tendencias parecidas sin embargo conservan estructuras casi iguales a la serie original sin retardos por lo que a simple vista el mejor modelo sería el de Random Forest

Table 7: Comparativo de las predicciones realizadas por todos los modelos aplicados

Datos de prueba	Prediccion Serie temporal	Prediccion Random Forest	Prediccion Red neuronal
2505.00	2380.22	2361.47	2347.04
2424.00	2490.81	2495.88	2356.79
3074.00	2677.44	2865.19	2348.36
2690.00	2696.91	2896.62	3028.32
2736.00	2501.87	2563.10	2435.17
2464.00	2388.11	2557.09	2489.21
2304.00	2535.76	2444.63	2351.77
2930.00	2706.25	2703.46	2342.83
2609.00	2661.73	2658.74	2885.24
2815.00	2453.21	2535.34	2383.26
2509.00	2406.01	2568.98	2638.16
2669.00	2584.45	2488.45	2357.39
3059.00	2721.30	2938.35	2417.37
2845.00	2613.19	2883.03	3020.36
2551.00	2420.10	2659.41	2707.69
2402.00	2433.33	2428.24	2365.27
2405.00	2635.04	2431.66	2346.92
3009.00	2715.52	2852.11	2347.10
3049.00	2563.72	2832.66	2984.63
2550.00	2396.93	2954.49	3014.43
2599.00	2476.08	2428.24	2365.04
2442.00	2674.82	2577.62	2379.36
2996.00	2696.18	2916.92	2349.75
2652.00	2512.55	2787.52	2972.43
2525.00	2390.79	2564.09	2405.49
2667.00	2526.23	2414.95	2360.02
2532.00	2702.19	2544.60	2415.87
3160.00	2664.27	2891.67	2361.31
2957.00	2464.83	2902.68	3058.45
3102.00	2402.85	2855.19	2926.49
2978.00	2577.54	2847.76	3040.64
2895.00	2716.04	2705.63	2953.08

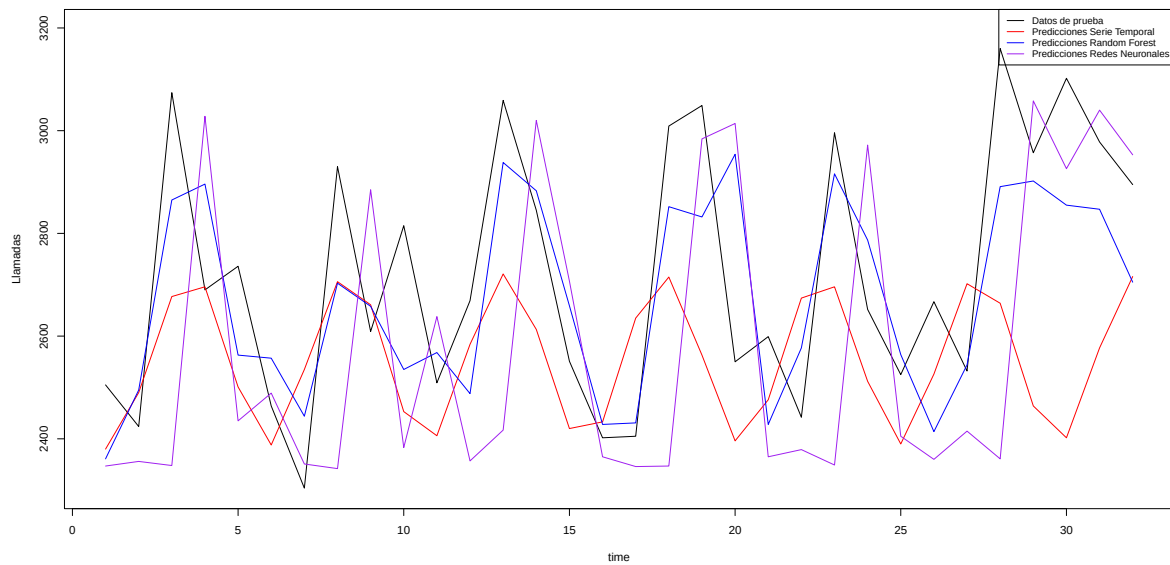


Figure 26: Comparacion entre serie temporal,random forest y red neuronal

Ahora otra alternativa es comparar los modelos respecto a su Error Cuadratico Medio, MAPE y Varianza como se observa en la tabla 8

Medida	Serie temporal	Random Forest	Red neuronal
MSE	77863	28986	125065
MAPE	8.14	5.24	9.57
Var	14556	35947	82941

Se observa que el modelo con mejor ajuste en términos de Error Porcentual Absoluto Medio es el Random Forest, con un valor de 5,24. Esto indica que este modelo se acerca más tanto en términos visuales como en el error al realizar las predicciones, sin sufrir de sobreajuste. Por otro lado, aunque el error en las series temporales es menor en comparación con la red neuronal, se encuentra que los pronósticos no explican adecuadamente el comportamiento de la serie debido a patrones periódicos. Esto refleja una menor varianza en las series temporales, ya que los datos no presentan una dispersión tan amplia debido a su comportamiento. además, se observa que el error cuadrático medio sigue la misma tendencia, lo que refuerza la idea de que el Random forest tuvo un mejor ajuste en comparación con los otros modelos utilizados

5. Conclusiones

- En este estudio se evaluaron diferentes modelos de aprendizaje y series temporales, y se encontró que el modelo Random Forest simple brindo la mejor respuesta, ya que logro reducir significativamente el error. además, se observó que este modelo converge rápidamente, lo que no genera una influencia computacional que dificulte su estimación
- En ninguno de los modelos se encontró posible asociar variables explicativas distintas a los retardos de la serie. Por lo tanto, se obtuvieron mejores resultados al entrenar los modelos más simples al momento de realizar las predicciones.
- La transformación de la serie favoreció el pronóstico de las llamadas en varios modelos, lo que permitió disminuir el impacto de valore a típicos, como se observó en modelos como la red neuronal o la serie temporal
- Para futuros trabajos, se recomienda identificar o detallar nuevas variables que puedan explicar el modelo, ya que contra únicamente con los retardos anteriores de la serie limita el desarrollo de nuevos modelos para el pronóstico. Seria beneficioso explorar la posibilidad de utilizar series multivariada o emplear otro tipo de red neuronal

6. Referencias

- Ali, J. (2012). International Journal of Computer Science
- Benavides, A. M. (Enero de 2022). Modelo para la prediccion de interacciones recibidas en un centro de experiencia telefonica.
- Boulin, J. M. (2010). CALL CENTER DEMAND FORECASTING: IMPROVING SALES CALLS. Massachusetts.
- Breiman, L. (2001). Random Forests. Obtenido de <https://link.springer.com/article/10.1023/A:1010933404324>
- Chen, X. (09 de 02 de 2012). Random forests for genomic data analysis. Science Direct. Obtenido de <https://doi.org/10.1016/j.ygeno.2012.04.003>

- Eduardo Raffo, L. R. (2012). Medición de la atención en un call center usando box-jenkins.
- Hurtado, P. A. (2003). Analisis de Capacidad del Call Center de una.
- Rigatti, S. J. (2017). Random Forest. Journal of Insurance Medicine.
- Saavedra, F. I. (s.f.). Redes Neuronales Artificiales. Chile.