



Ingeniería Energética

E-ISSN: 1815-5901

orestes@cipel.ispjae.edu.cu

Instituto Superior Politécnico José

Antonio Echeverría

Cuba

Bañuelos Cabral, Eduardo Salvador; Gutiérrez Robles, José Alberto; Naredo Villagrán, José Luis; Sotelo-Castañón, Julián; Galván Sánchez, Verónica Adriana; García-Sánchez, Jorge Luis

Identificación de parámetros con métodos numéricos para el modelado de sistemas eléctricos con dependencia frecuencial

Ingeniería Energética, vol. XXXVI, núm. 2, mayo-agosto, 2015, pp. 155-167

Instituto Superior Politécnico José Antonio Echeverría

La Habana, Cuba

Disponible en: <http://www.redalyc.org/articulo.oa?id=329138826005>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

redalyc.org

Sistema de Información Científica

Red de Revistas Científicas de América Latina, el Caribe, España y Portugal

Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso abierto



Identificación de parámetros con métodos numéricos para el modelado de sistemas eléctricos con dependencia frecuencial

Identification of parameters with numerical methods for the modelling of electrical systems with frequency dependence

Eduardo Salvador – Bañuelos-Cabral
José Alberto – Gutiérrez Robles
José Luís – Naredo-Villagrán

Julián – Sotelo-Castañón
Verónica Adriana – Galván Sánchez
Jorge Luis – García-Sánchez

Recibido: febrero de 2014

Aprobado: diciembre de 2014

Resumen/ Abstract

El artículo presenta una descripción detallada de las técnicas más utilizadas para hacer ajuste racional de funciones del dominio de la frecuencia. Las técnicas son: Aproximación Asintótica de Bode, Mínimos Cuadrados Ordinarios, Mínimos Cuadrados Iterativamente Reponderados, Ajuste Vectorial y Levenberg-Marquardt. Estas técnicas se comparan aproximando una función analítica. Después, se aplican al ajuste racional de parámetros dependientes de la frecuencia de una línea de transmisión monofásica. El efecto del ajuste se evalúa considerando los casos de línea abierta, corto circuito y perfectamente acoplada. La transformada numérica de Laplace (NLT) se utiliza como referencia para la evaluación. Se concluye que la adecuada implementación de cada técnica depende de varios factores; del tipo de función a ajustar, del rango de ajuste, y otros. Además, no es posible garantizar que una de las técnicas siempre converja al mejor resultado. El artículo propone algunas guías para seleccionar la técnica más adecuada para una aplicación particular.

Palabras clave: Bode, Mínimos Cuadrados, Ajuste Vectorial, Levenberg-Marquardt, y Líneas de Transmisión.

This paper provides a detailed description of the rational fitting techniques that are most used to approximate frequency domain functions. The techniques are: Bode asymptotic approximations, Ordinary Least-Squares, Iteratively Reweighted Least-Squares, Vector Fitting and Levenberg-Marquardt. These techniques are compared by approximating of an analytic function. Then, the techniques are applied to the rational-fitting of the frequency-dependent parameters corresponding to a single-phase transmission line. The effect of the rational representations is evaluated considering transients on cases of open-ended, short-circuited and perfectly matched lines. The numerical Laplace transform technique (NLT) is used as reference for the evaluations. It follows that the proper implementation of each fitting technique depends on various factors; the type of function being adjusted, the desired frequency range, an others. Moreover, one cannot guarantee that there is one of these techniques that will always yield the best fitting. The paper provides guidelines for selecting the most convenient one for a particular application.

Key words: Bode, Least-Squares, Vector Fitting, Levenberg-Marquardt, and Transmission Lines.

INTRODUCCIÓN

En varios campos de la ingeniería, para propósitos de análisis, diseño o simulación a menudo es necesario tener representaciones precisas de sistemas físicos con funciones de transferencia.

Específicamente, para el estudio de fenómenos electromagnéticos en sistemas eléctricos de potencia, los componentes deben ser modelados en un amplio rango de frecuencia.

El problema de identificación en el dominio de la frecuencia se basa en la estimación de una función racional, con coeficientes reales (a_n , b_m) para ajustar una función compleja $F(s)$ con la estructura algebraica mostrada en la ecuación (1).

$$F(s) \cong \frac{N(s)}{D(s)} = \frac{a_0 + a_1s + a_2s^2 + \dots + a_ns^n}{1 + b_1s + b_2s^2 + \dots + b_ms^m} \quad (1)$$

Donde es el vector de frecuencia compleja y usualmente $n < m$.

Existen diferentes procesos de estimación de parámetros para la ecuación (1). A continuación se presentan algunos antecedentes. En 1959 Levi [1], sugirió un procedimiento de linealización y presentó un método de ajuste para curvas en el dominio de la frecuencia. El método se basó en mínimos cuadrados y en el uso de derivadas parciales para minimizar el error cuadrático de una función. En 1963 Sanathanan y Koerner [2], utilizaron el método de Levi introduciendo un procedimiento iterativo. Consideraron el error como la diferencia entre las magnitudes de la función a ajustar y una función ajustada, esta última se representa como la relación de dos polinomios. Los coeficientes de los polinomios son el resultado de minimizar la suma del error al cuadrado. Por otro lado, J. R. Martí [3], presentó en 1982 un método basado en ajuste asintótico utilizando la técnica de Bode. Los parámetros de la aproximación son determinados automáticamente por la rutina, como una función de aproximación que "*libremente*" se adapta por sí misma a la forma de la función aproximada. Más recientemente, Gustavsen y Semlyen en 1999 [4], desarrollaron una metodología general para el ajuste de respuestas en el dominio de la frecuencia. Conocida como Ajuste Vectorial o "*Vector Fitting*" se ha posicionado como una de las más populares para el ajuste racional. Esta técnica ha sido utilizada en diferentes campos de la ingeniería para el modelado de sistemas eléctricos [5-7]. Por otro lado, Noda [8], presentó en 2005 un algoritmo que particiona el rango completo de frecuencia para evitar el mal condicionamiento del sistema, además de utilizar una matriz de ponderación para mejorar iterativamente la aproximación. Este algoritmo ha sido utilizado en [9] para modelar el comportamiento de una línea de transmisión.

Por último, Pordanjani *et al.* [10], en 2011 propusieron un método basado en un algoritmo genético, el cual garantiza pasividad.

Técnicas de ajuste racional

En esta sección se hace la descripción detallada de las técnicas de ajuste racional más utilizadas como son: Aproximación Asintótica de Bode, Mínimos Cuadrados Ordinarios, Mínimos Cuadrados Iterativamente Reponderados, Ajuste Vectorial y Levenberg-Marquardt.

Aproximación asintótica de bode

La técnica por aproximación asintótica conocida como diagramas de Bode [11], fue introducida en 1945, y posteriormente aplicada por J. R. Martí [3], para el ajuste racional de modelos de líneas de transmisión. Los diagramas consisten en dos gráficas, una muestra cómo varía la amplitud de $F(s)$ con la frecuencia y la otra muestra cómo lo hace el ángulo de fase. Para simplificar el desarrollo de la técnica se considera que todos los polos y ceros de $F(s)$ son reales, distintos y de primer orden, así factorizando la ecuación (1) se llega a la ecuación (2).

$$F(s) = \frac{k(s + z_1)(s + z_2) \dots (s + z_n)}{s(s + p_1)(s + p_2) \dots (s + p_m)} \quad (2)$$

Si se denota como ($s = j\omega$), entonces la ecuación (2) se representa como la ecuación (3),

$$F(s) = \frac{k(j\omega + z_1)(j\omega + z_2) \dots (j\omega + z_n)}{j\omega(j\omega + p_1)(j\omega + p_2) \dots (j\omega + p_m)} \quad (3)$$

Representando la función en su "*forma estándar*", se tiene la ecuación (4).

$$F(j\omega) = \frac{k \left(1 + \frac{j\omega}{z_1}\right) \left(1 + \frac{j\omega}{z_2}\right) \dots \left(1 + \frac{j\omega}{z_n}\right) z_1 z_2 \dots z_n}{j\omega \left(1 + \frac{j\omega}{p_1}\right) \left(1 + \frac{j\omega}{p_2}\right) \dots \left(1 + \frac{j\omega}{p_m}\right) p_1 p_2 \dots p_m} \quad (4)$$

Si se hace la asignación $K_0 = kz_1 z_2 \dots z_n / p_1 p_2 \dots p_m$ entonces la ecuación (4) en su forma polar se determina por la ecuación (5).

$$F(j\omega) = \frac{K_0 \left| 1 + \frac{j\omega}{z_1} \right| \angle \alpha_1 \left| 1 + \frac{j\omega}{z_2} \right| \angle \alpha_2 \cdots \left| 1 + \frac{j\omega}{z_n} \right| \angle \alpha_n}{\left| \omega \right| \angle 90^\circ \left| 1 + \frac{j\omega}{p_1} \right| \angle \beta_1 \left| 1 + \frac{j\omega}{p_2} \right| \angle \beta_2 \cdots \left| 1 + \frac{j\omega}{p_m} \right| \angle \beta_m}. \quad (5)$$

El método de aproximación asintótica solo ajusta la amplitud de la ecuación (5), la cual se expresa como, la ecuación (6).

$$|F(j\omega)| = \frac{K_0 \left| 1 + \frac{j\omega}{z_1} \right| \left| 1 + \frac{j\omega}{z_2} \right| \cdots \left| 1 + \frac{j\omega}{z_n} \right|}{\left| \omega \right| \left| 1 + \frac{j\omega}{p_1} \right| \left| 1 + \frac{j\omega}{p_2} \right| \cdots \left| 1 + \frac{j\omega}{p_m} \right|}. \quad (6)$$

La amplitud involucra multiplicación y división de factores asociados con los polos y ceros. Para reducir la ecuación (6) a sumas y restas de factores, ésta se expresa en términos de un valor logaritmo base 10, el decibel. Como su nombre indica, un decibel es la décima parte de un bel, escala utilizada por A. G. Bell para medir la ganancia de potencia en un sistema. Debido a que se tienen 10 decibeles en un bel, la relación de potencia en decibeles se define como: ver ecuación (7).

$$10 \log_{10} \frac{P_{\text{salida}}}{P_{\text{entrada}}}. \quad (7)$$

Como la potencia involucra valores de voltaje o corriente al cuadrado y asumiendo impedancias de carga y fuente iguales se tiene, la ecuación (8).

$$10 \log_{10} \frac{V_{\text{salida}}^2}{V_{\text{entrada}}^2} = 20 \log_{10} \frac{V_{\text{salida}}}{V_{\text{entrada}}}. \quad (8)$$

En la mayoría de las funciones de transferencia se utiliza la relación entre voltajes y/o corrientes. Por lo tanto, para el caso de la aproximación asintótica se utiliza la ecuación (8) para expresar la ecuación (6) en decibeles, como se indica en la ecuación (9).

$$A_{dB} = 20 \log_{10} |F(j\omega)|. \quad (9)$$

Así, utilizando propiedades de logaritmos se tiene, la ecuación (10).

$$A_{dB} = 20 \log_{10} |K_0| + 20 \log_{10} \left| 1 + j\omega/z_1 \right| + 20 \log_{10} \left| 1 + j\omega/z_n \right| - 20 \log_{10} |\omega| - 20 \log_{10} \left| 1 + j\omega/p_1 \right| - \cdots - 20 \log_{10} \left| 1 + j\omega/p_m \right|. \quad (10)$$

Para realizar el diagrama de la ecuación (10), se esboza cada término de forma separada y después se combinan las gráficas. En la expresión se pueden identificar 3 tipos de términos, la forma de aproximar cada uno de ellos se explica a continuación.

- Término $20 \log_{10} |K_0|$. - Este se aproxima con una línea horizontal debido a que K_0 es constante.
- Término $-20 \log_{10} |\omega|$. - Este se aproxima con una línea recta de -20 dB/década, la cual hace intercepción con 0 dB en $\omega = 1$.
- Términos $20 \log_{10} |1 + j\omega/z_n|$ ó $-20 \log_{10} |1 + j\omega/p_m|$. - Sin considerar el signo, ambos términos tienen el siguiente comportamiento: $20 \log_{10} |1 + j\omega/z_n| \rightarrow 0$ si $\omega \rightarrow 0$ - $20 \log_{10} |1 + j\omega/z_n| \rightarrow 20 \log_{10} (\omega/z_n)$ si $\omega \rightarrow \infty$

Es posible aproximar el segundo término por dos líneas rectas, asíntota de baja frecuencia y asíntota de alta frecuencia, respectivamente.

La asíntota de alta frecuencia que aproxima el término $20 \log_{10} (\omega/z_n)$ tiene una pendiente positiva de 20 dB/década, esta recta hace intercepción con 0 dB en $\omega = z_n$, este valor se conoce como "*frecuencia de ruptura*". Entonces, el término $20 \log_{10} |1 + j\omega/z_n|$ se aproxima con dos líneas rectas de pendiente 0 y 20 dB/década, y análogamente el término $-20 \log_{10} |1 + j\omega/p_m|$ se aproxima con dos líneas rectas de pendiente 0 y -20 dB/década de acuerdo a la explicación anterior. La adaptación de este método para el ajuste racional consiste en su implementación numérica, comparando la función a ajustar con la suma de los segmentos de líneas rectas, e introduciendo un polo o un cero donde sea necesario, a fin de mantener la gráfica de la aproximación dentro de un margen de error establecido [3].

Mínimos cuadrados ordinarios

El cálculo de los coeficientes para los polinomios $N(s)$ y $D(s)$ en la ecuación (1), puede ser formulado como un sistema lineal de mínimos cuadrados multiplicando ambos miembros de la ecuación por el denominador. De esta forma se obtiene un sistema sobre determinado $\mathbf{Ax}=\mathbf{b}$. Donde el k -ésimo renglón de $\mathbf{A}$ y los vectores $\mathbf{x}$ y $\mathbf{b}$ están dados por: ver ecuación (11a, 11b, 11c).

$$\mathbf{A}_k = \begin{bmatrix} 1 & s_k & \cdots & s_k^n & -s_k F(s_k) & \cdots & -s_k^m F(s_k) \end{bmatrix} \quad (11a)$$

$$\mathbf{x} = [a_0 \quad a_1 \quad \cdots \quad a_n \quad b_1 \quad b_2 \quad \cdots \quad b_m]^T \quad (11b)$$

$$\mathbf{b} = [F(s_1) \quad \cdots \quad F(s_k)]^T. \quad (11c)$$

En términos matriciales, dado un vector $\mathbf{b} \in \mathbf{R}^k$ y una matriz $\mathbf{A} \in \mathbf{R}^{k \times l}$, se desea encontrar un vector $\mathbf{x} \in \mathbf{R}^l$, tal que el producto $\mathbf{Ax}$ sea la mejor aproximación de $\mathbf{b}$. Dicho en otras palabras, el problema lineal de mínimos cuadrados se acerca a la optimización si se encuentra un vector $\mathbf{x}$ que en algún sentido sea la mejor aproximación entre $\mathbf{Ax}$ y $\mathbf{b}$ [12]. Así, se debe definir el sentido de cercanía es decir, la norma que mide el vector residual $\mathbf{r}$. La norma más utilizada es la euclidiana (l_2). Así, la solución en mínimos cuadrados de la ecuación (11), es el vector $\mathbf{x}$ que hace de $\|\mathbf{r}\|_2 = \|\mathbf{b} - \mathbf{Ax}\|_2$ un mínimo, como se expresa en la ecuación (12).

$$\min_{\mathbf{x} \in \mathbf{R}^n} \|\mathbf{b} - \mathbf{Ax}\|_2^2. \quad (12)$$

La expresión (12), es la función objetivo a minimizar. Como las normas no pueden ser negativas, minimizar una norma equivale a minimizar su cuadrado.

Ecuaciones normales

La propiedad más significativa del vector residual del problema de mínimos cuadrados es que está enteramente en el espacio nulo de la matriz $\mathbf{A}^T$. En términos algebraicos, una solución $\mathbf{x}$ satisface la ecuación (13), [12]:

$$\mathbf{A}^T (\mathbf{b} - \mathbf{Ax}) = 0 \quad (13)$$

o bien la ecuación (14),

$$\mathbf{A}^T \mathbf{Ax} = \mathbf{A}^T \mathbf{b}. \quad (14)$$

El sistema de ecuaciones dado por la ecuación (14) se conoce como de ecuaciones normales. La matriz $\mathbf{A}^T \mathbf{A}$, es semidefinida positiva si y solo si las columnas de $\mathbf{A}$ son linealmente independientes. Las ecuaciones normales son siempre compatibles, es decir, existe siempre una solución para la ecuación (14), aunque $\mathbf{A}^T \mathbf{A}$ sea singular [12]. La solución encontrada por las ecuaciones normales está dada por la ecuación (15).

$$\mathbf{x} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{b}. \quad (15)$$

La matriz $\mathbf{A}$, formada de acuerdo con la ecuación (11a), generalmente está mal condicionada debido al amplio rango de frecuencia que s puede tomar. Por este mal condicionamiento, la técnica presenta un mayor error para valores pequeños del rango de frecuencia. Es posible mejorar el mal condicionamiento de la matriz realizando un escalamiento en sus columnas, éste consiste en la multiplicación de la matriz por una matriz diagonal $\mathbf{D}$ de acuerdo a la ecuación (16), [12]:

$$\mathbf{D} = \text{diag}\{1/\|a_k\|_2\} \quad (16)$$

Donde $\|a_k\|_2$, es la norma euclidiana de la k -ésima columna de $\mathbf{A}$. El sistema lineal $\mathbf{Ax}=\mathbf{b}$, ahora queda determinado por la ecuación (17):

$$\mathbf{ADD}^{-1}\mathbf{x} = \mathbf{b}. \quad (17)$$

Denotando como $\mathbf{A}' = \mathbf{AD}$ y $\mathbf{x}' = \mathbf{D}^{-1}\mathbf{x}$, se tiene la ecuación (18),

$$\mathbf{A}'\mathbf{x}' = \mathbf{b}. \quad (18)$$

Se desea que los coeficientes (a_n , b_m) para los polinomios $N(s)$ y $D(s)$ en la ecuación (1), sean cantidades reales. Para resolver la ecuación (18), en este sentido, ésta se separa en sus partes real e imaginaria de acuerdo con la ecuación (19):

$$\begin{bmatrix} \mathbf{A}'_r \\ \mathbf{A}'_i \end{bmatrix} \mathbf{x}' = \begin{bmatrix} \mathbf{b}_r \\ \mathbf{b}_i \end{bmatrix} \quad (19)$$

Donde los sub-índices r e i denotan las partes real e imaginaria respectivamente. Con lo anterior se asegura que los coeficientes (a_n, b_m) sean cantidades reales. Una vez que se obtiene la solución del sistema de ecuaciones dado por (19), se calcula el vector $\mathbf{x}$ que contiene los coeficientes de la función racional.

Mínimos Cuadrados Iterativamente Preponderados

El método se utiliza para resolver ciertos problemas de optimización, en los cuales, la función objetivo a minimizar está dada por la ecuación (20), [12]:

$$\gamma(\mathbf{x}) = \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{W}(\mathbf{b} - \mathbf{A}\mathbf{x})\|_2^2 \quad (20)$$

Donde $\mathbf{W}$ es una matriz diagonal de pesos, $\text{diag}(w_1, w_m, \dots, w_m) > 0$. Las ecuaciones normales para el método se muestran en la ecuación (21):

$$\mathbf{A}^T \mathbf{W} \mathbf{A} \mathbf{x} = \mathbf{A}^T \mathbf{W} \mathbf{b}. \quad (21)$$

La solución para la ecuación (21), está dada por la ecuación (22):

$$\mathbf{x} = (\mathbf{A}^T \mathbf{W} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{W} \mathbf{b}. \quad (22)$$

Es posible llevar la ecuación (22) hasta un procedimiento iterativo utilizando la estructura algebraica mostrada en la ecuación (23):

$$\mathbf{x}^{(i+1)} = (\mathbf{A}^T \mathbf{W}^{(i)} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{W}^{(i)} \mathbf{b} \quad (23)$$

Donde i representa el número de iteración. En la solución iterativa de las ecuaciones normales ponderadas se deben tener en cuenta el escalamiento y las cantidades reales revisadas en la sección anterior. La matriz $\mathbf{W}$, de acuerdo con la ecuación (23), se debe actualizar de forma iterativa. Para actualizar la matriz de pesos se debe considerar la función de error presentada en la ecuación (24):

$$e_k = F(s_k) - \hat{F}(s_k, \mathbf{x}) \quad (24)$$

Donde $\hat{F}(s_k, \mathbf{x})$ es la respuesta en frecuencia de la función racional identificada. Cada elemento de la matriz de pesos se actualiza de acuerdo con la ecuación (25):

$$w_k^{(i)} = w_k^{(i-1)} |e_k^{(i-1)}|. \quad (25)$$

Para el primer paso de iteración, $\mathbf{W}$ es la matriz identidad. Si la función objetivo a minimizar cumple con $\gamma(\mathbf{x}^{(i)}) \leq \gamma(\mathbf{x}^{(i+1)})$, el vector $\mathbf{x}^{(i)}$ es la solución final.

Ajuste Vectorial

El Ajuste Vectorial o “*Vector Fitting*” es una metodología para el ajuste de respuestas en el dominio de la frecuencia con aproximaciones por funciones racionales. El ajuste se realiza reemplazando un juego de “*polos de arranque o polos iniciales*” por un juego de polos recalculados a través de un método iterativo [4]. La metodología realiza la aproximación en dos etapas, ambas con polos conocidos. La primera etapa comienza con la distribución de los polos iniciales, reales o complejos, sobre todo el rango de frecuencia de interés. En esta etapa un método iterativo mejora la posición de los polos. En la segunda etapa, una vez representado el sistema en fracciones parciales, se realiza el cálculo de los residuos en una sola iteración. Si se considera la aproximación racional dada por la ecuación (26):

$$F(s) \cong \sum_{n=1}^N \frac{c_n}{s - a_n} + d + sh \quad (26)$$

Donde c_n son los residuos, a_n son los polos, d es un término constante y h es un término proporcional. El problema del ajuste racional es encontrar todos los parámetros de la ecuación (26) para obtener una aproximación de $F(s)$ sobre un rango de frecuencia dado. Se observa que la ecuación (26), corresponde a un sistema no lineal, debido a que a_n aparece en el denominador.

Identificación de los polos

El inicio de la metodología consiste en especificar un conjunto de polos de arranque $\overline{a_n}$ en la ecuación (26). Además, se multiplica $F(s)$ por una función desconocida $\sigma(s)$. Por tanto es necesario introducir una función de aproximación para $\sigma(s)$. Como resultado se tienen las ecuaciones (27) y (28).

$$\sigma(s)F(s) \cong \sigma F_{fit}(s) = \sum_{n=1}^N \frac{c_n}{s - \overline{a_n}} + d + sh \quad (27)$$

$$\sigma(s) \cong \sigma_{fit}(s) = \sum_{n=1}^N \frac{\tilde{c}_n}{s - \overline{a_n}} + 1 \quad (28)$$

donde: $\sigma F_{fit}(s)$ es el ajuste de $\sigma(s)F(s)$ y $\sigma_{fit}(s)$ es el ajuste de $\sigma(s)$. Si se despeja $F(s)$ de la ecuación (27), se llega a la ecuación (29),

$$F(s) \cong \frac{\sigma F_{fit}(s)}{\sigma(s)}. \quad (29)$$

Ahora sustituyendo (27) y (28) en (29), se tiene la ecuación (30),

$$F(s) \cong \frac{h \prod_{n=1}^{N+1} s - z_n}{\prod_{n=1}^N s - \tilde{z}_n} \bigg/ \frac{\prod_{n=1}^N s - \overline{a_n}}{\prod_{n=1}^N s - \tilde{z}_n} = h \frac{\prod_{n=1}^{N+1} s - z_n}{\prod_{n=1}^N s - \tilde{z}_n}. \quad (30)$$

La ecuación (30), indica que los ceros de $\sigma(s)$ son una aproximación de los polos de $F(s)$, así el problema inicia con el cálculo de los ceros de la función auxiliar, $\sigma(s)$. Se debe notar en las ecuación (27) y (28), que tanto la función de aproximación racional para $\sigma(s)$ $F(s)$ como la función para $\sigma(s)$ tienen los mismos polos. Ahora multiplicando la ecuación (28), por $F(s)$ se obtiene la ecuación (31):

$$\sigma(s)F(s) \cong \left(\sum_{n=1}^N \frac{\tilde{c}_n}{s - \overline{a_n}} + 1 \right) F(s). \quad (31)$$

Se igualan las ecuaciones (27) y (31), para obtener la ecuación (32):

$$\sum_{n=1}^N \frac{c_n}{s - \overline{a_n}} + d + sh = \left(\sum_{n=1}^N \frac{\tilde{c}_n}{s - \overline{a_n}} + 1 \right) F(s). \quad (32)$$

Manipulando algebraicamente la ecuación anterior se llega a la ecuación (33),

$$\sum_{n=1}^N \frac{c_n}{s - \overline{a_n}} + d + sh - \sum_{n=1}^N \frac{\tilde{c}_n F(s)}{s - \overline{a_n}} = F(s). \quad (33)$$

La ecuación (33), indica que se tiene un problema no lineal donde $F(s)$ depende de $F(s)$. Una forma de solución adecuada es aproximar los polos mediante la determinación de los ceros de $\sigma(s)$, los cuales se obtienen de $\tilde{c}_n$. Así se utiliza el sistema de ecuaciones dado por la ecuación (33), planteándolo de acuerdo a la ecuación (34)

$$\mathbf{A}_k \mathbf{x} = \mathbf{b} \quad (34)$$

Donde cada termino se expresa de acuerdo a las ecuaciones (35a), (35b), (35c).

$$\mathbf{A}_k = \begin{bmatrix} \frac{1}{s_k - \overline{a_1}} & \cdots & \frac{1}{s_k - \overline{a_N}} & 1 & s_k & -\frac{F(s_k)}{s_k - \overline{a_1}} & \cdots & -\frac{F(s_k)}{s_k - \overline{a_N}} \end{bmatrix} \quad (35a)$$

$$\mathbf{x} = [c_1 \quad \cdots \quad c_N \quad d \quad h \quad \tilde{c}_1 \quad \cdots \quad \tilde{c}_N]^T \quad (35b)$$

$$\mathbf{b} = [F(s_1) \quad \cdots \quad F(s_k)]^T. \quad (35c)$$

Representando la ecuación (35a,35b,35c), para un rango amplio de frecuencias se obtiene un sistema sobre determinado, finalmente se trabaja en cantidades reales de la misma forma que se plantea en la sección de mínimos cuadrados. La ecuación (34), expresada para un rango amplio de frecuencias, entrega como resultado un vector de

residuos $\boldsymbol{\gamma}$ de la función auxiliar; el cálculo de los ceros de esta función corresponden a los valores propios de la matriz mostrada en la ecuación (36),

$$\mathbf{H} = \mathbf{A} - \mathbf{b}\boldsymbol{\gamma}^T \quad (36)$$

Donde $\mathbf{A}$ es una matriz diagonal que contiene los polos iniciales y $\mathbf{b}$ es un vector columna elementos unidad. $\boldsymbol{\gamma}^T$, es un vector fila que contiene los residuos de $\sigma(s)$.

Identificación de residuos

Una vez que se han obtenido los polos para el ajuste, los residuos se pueden obtener resolviendo el problema original en la ecuación (26), empleando los ceros de $\sigma_{fit}(s)$ como los nuevos polos a_n . Con esto se obtiene nuevamente un sistema lineal sobre determinado $\mathbf{Ax}=\mathbf{b}$, donde el vector solución es $\mathbf{x}=[c_1 \cdots c_N \ d \ h]^T$, el cual se resuelve como la ecuación (34), donde la matriz ahora no tiene los términos negativos [4].

Método de Levenberg-Marquardt

Levenberg-Marquardt actúa más como el de gradiente descendiente (Apéndice A) cuando los parámetros están lejos de su valor óptimo, pero más como el de Gauss-Newton (Apéndice B) cuando los parámetros están cerca de su valor óptimo. Primero Levenberg y luego unos años más tarde Marquardt propusieron este método de acuerdo a la ecuación (37), [13]:

$$(\mathbf{J}^T \mathbf{J} + \lambda \mathbf{I}) \mathbf{h}_{LM} = -\mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (37)$$

El término λ es una cantidad positiva. Es fácil notar que para valores pequeños de λ el método se comporta más como Gauss-Newton y para valores grandes de λ se parece más al método del gradiente. Por último, también se ha sugerido la modificación mostrada en la ecuación (38):

$$(\mathbf{J}^T \mathbf{J} + \lambda \text{diag}(\mathbf{J}^T \mathbf{J})) \mathbf{h}_{LM} = -\mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (38)$$

La solución a las ecuaciones normales para el método de Levenberg-Marquardt queda expresada en la ecuación (39):

$$\mathbf{h}_{LM} = -(\mathbf{J}^T \mathbf{J} + \lambda \text{diag}(\mathbf{J}^T \mathbf{J}))^{-1} \mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (39)$$

Existen diferentes métodos de cálculo para λ . En [13], se realiza una revisión de los mismos; sin embargo el más común consiste en evaluar la función objetivo y, si $\gamma(\mathbf{x}+\mathbf{h}) \leq \gamma(\mathbf{x})$, entonces $\mathbf{x}$ es remplazado por $\mathbf{x}+\mathbf{h}$, y λ es reducido por un factor v . De otra forma λ se incrementa por un factor u sin actualizar $\mathbf{x}$ y el algoritmo continúa a la siguiente iteración.

Análisis de casos

Una función con dependencia frecuencial

puede tener como origen una medición experimental. De esta forma se obtiene un modelo analítico a partir de esos datos. Otro origen puede ser una función analítica dependiente de la frecuencia, aquí la técnica reemplaza a esa función analítica por una función racional. Esto es adecuado siempre y cuando la función racional sea más simple de resolver que la función original y se mantengan todas las características cualitativas y cuantitativas. Con el fin de explorar la precisión y versatilidad de las técnicas se presentan dos casos de estudio; el primero de ellos es ajustar una función analítica caótica. El segundo caso de estudio es analizar el efecto final en la simulación de un transitorio en una línea monofásica, aquí lo que se hace es ajustar la ecuación de los parámetros de la línea y reemplazarla por una función racional [14-17].

Caso de prueba 1

Para el primer caso se toma la siguiente función analítica de orden 5.ver ecuación (40):

$$F(s) = \frac{-81.42 + 18.62e6 s + 27.71e6 s^2 + 62.84e3 s^3 + 42.54 s^4 + 5e - 7 s^5}{1 + 54.68e5 s + 36.42e7 s^2 + 83.37e4 s^3 + 90.28 s^4 + 36.35e - 3 s^5} \quad (40)$$

Se discretiza la función en el rango de frecuencias dado por $f = [1 \times 10^{-2} \ 1 \times 10^8]$. Y se utilizan los 5 métodos descritos para realizar el ajuste de $F(s)$. Cada uno de ellos arroja como resultado una función analítica distinta, la cual, dependiendo del método se tienen polos-ceros, coeficientes de polinomios o representación en fracciones parciales. La figura 1, muestra el comportamiento de todas estas funciones en comparación con la función original. Aunque en algunos casos las diferencias son relativamente pequeñas, la figura 2, muestra con claridad los errores porcentuales de cada ajuste utilizando la formulación de la ecuación (41):

$$e\% = \left| \frac{f_{Original} - f_{Ajustada}}{f_{Original}} \right| \times 100. \quad (41)$$

Del análisis de resultados del ajuste de esta función analítica, se puede concluir que; Ajuste Vectorial, Mínimos Cuadrados Iterativamente Reponderados y Levenberg-Marquardt ajustan de manera adecuada esta función analítica, sin embargo cabe señalar que tanto Bode como mínimos cuadrados la ajustan adecuadamente por secciones; esta es la razón por la cual no es correcto concluir que unas técnicas son mejores que otras, sino que se concluye que para esta prueba específica unas técnicas son más adecuadas que otras. La tabla 1, muestra las funciones resultantes de cada ajuste.

Tabla 1. Resultados de los ajustes de la función analítica propuesta.	
Metodología	$F(s)$ resultante
Bode	$F_{Bode}(s) = \frac{1183.12(s + 0.3141)(s + 0.4398)(s + 640.88)(s + 1130.97)(s + 50265.48)}{(s + 6911.5)(s + 6283.18)(s + 4398.22)(s + 2981.37)(s + 0.157)(s + 0.0628)}$
Mínimos Cuadrados	$F_{MC}(s) = \frac{0.084 + 2.25 \times 10^{-4}s + 1.68 \times 10^{-7}s^2 - 3.01 \times 10^{-14}s^3 - 3.78 \times 10^{-21}s^4 - 4 \times 10^{-29}s^5}{1 + 0.0032s + 3.36 \times 10^{-7}s^2 + 1.43 \times 10^{-10}s^3 - 2.74 \times 10^{-17}s^4 - 2.90 \times 10^{-24}s^5}$
Mínimos Cuadrados Ponderados	$F_{MCIR}(s) = \frac{4.83 \times 10^5 - 1.06 \times 10^{11}s - 1.57 \times 10^{11}s^2 - 3.58 \times 10^8s^3 - 2.42 \times 10^5s^4 - 0.0028s^5}{1 - 3.11 \times 10^{10}s - 2.07 \times 10^{12}s^2 - 4.75 \times 10^9s^3 - 5.14 \times 10^5s^4 - 2.07 \times 10^2s^5}$
Ajuste Vectorial	$F_{AV}(s) = \frac{10.17}{s + 455.25} + \frac{0.049}{s + 0.015} - \frac{1.55e-5}{s + 1.82e-7} + \frac{580 - i0.33}{s + (1.01 + i4.58e3)} + \frac{580 + i0.33}{s + (1.01 - i4.58e3)}$
Levenberg-Marquardt	$F_{LM}(s) = \frac{-8.06 \times 10^2 + 1.77 \times 10^8s + 2.64 \times 10^8s^2 + 6 \times 10^5s^3 + 4.06 \times 10^2s^4 + 4.77 \times 10^{-6}s^5}{1 + 5.22 \times 10^7s + 3.47 \times 10^9s^2 + 7.96 \times 10^6s^3 + 8.62 \times 10^2s^4 + 0.34s^5}$

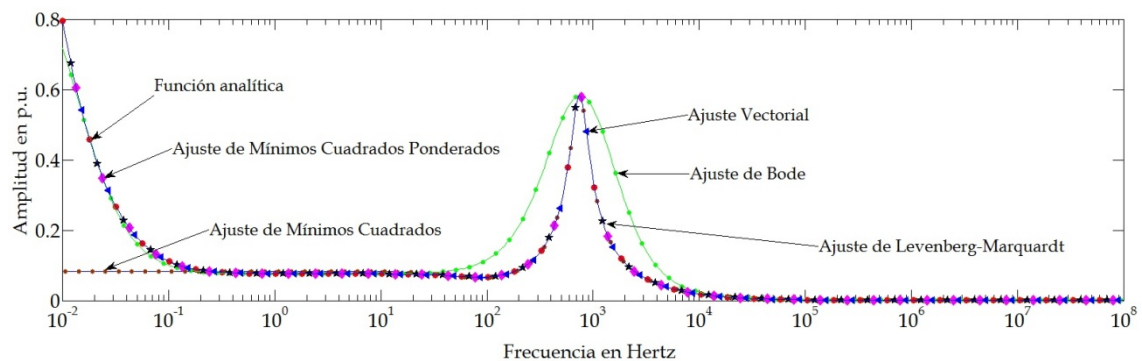


Fig. 1. Gráfica de los ajustes.

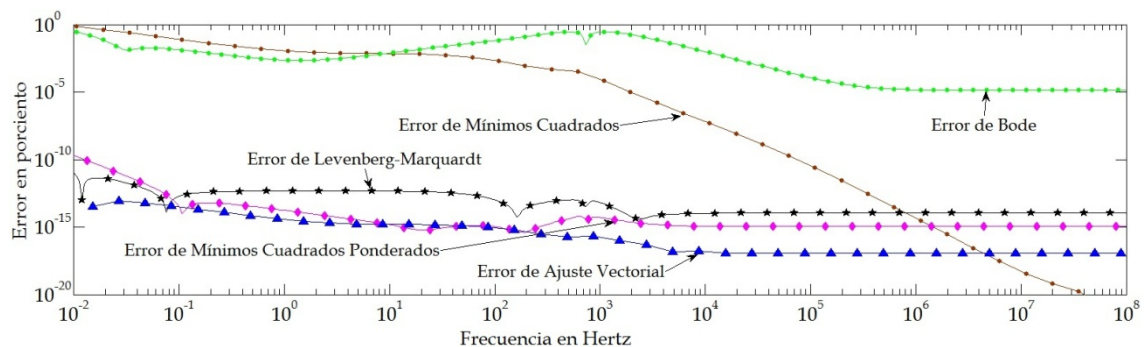


Fig. 2. Gráfica de los errores de los ajustes en porcentaje.

Caso de prueba 2

La simulación de un evento con simetría en un sistema de potencia, se puede hacer con su secuencia positiva en por unidad, es decir, como un sistema monofásico como el mostrado en la figura 3. Una longitud normal en un sistema de potencia es de 100 km, por lo tanto se utiliza esta longitud en la simulación. Los demás datos utilizados son: el radio del conductor es 0.0272 m, $R_1 = 1.058 \Omega$ es la resistencia interna del generador, $L = 70$ mH es la inductancia interna del

generador, R_2 es variable para representar la línea con terminación abierta, perfectamente acoplada o con terminación en corto circuito, la resistividad del terreno se toma como $100 \Omega\cdot\text{m}$ y la fuente como senoidal. Para realizar esta simulación en dominio de la frecuencia, se sustituye Y y Z por sus modelos ajustados como funciones racionales del tipo dado por la ecuación (1). La figura 4, muestra el comportamiento de Y y Z en función de la frecuencia.

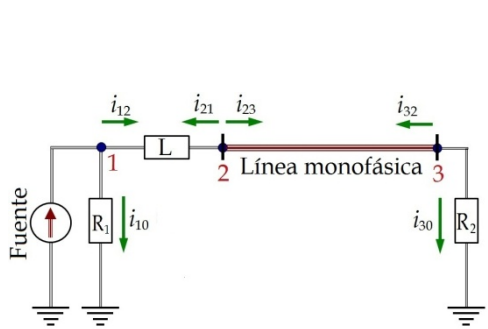


Fig. 3. Diagrama de la línea monofásica para el caso de prueba 2.

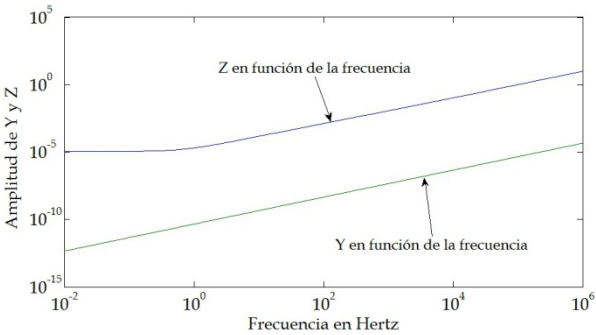


Fig. 4. Amplitudes de la admitancia e impedancia de La línea monofásica en función de la frecuencia.

Las figuras 5 y 6, muestran los errores porcentuales de todas las funciones, en estas figura se nota con claridad cuál es el mejor o peor ajuste, lo cual no significa que para otra función se vayan a comportar de la misma manera. Los resultados obtenidos con los ajustes se comparan con los que se obtienen utilizando las funciones Y y Z originales. Se utiliza la transformada numérica de Laplace para realizar las simulaciones [18-20]. Los resultados se muestran en las figuras 7, 8, 9, 10, 11 y 12.

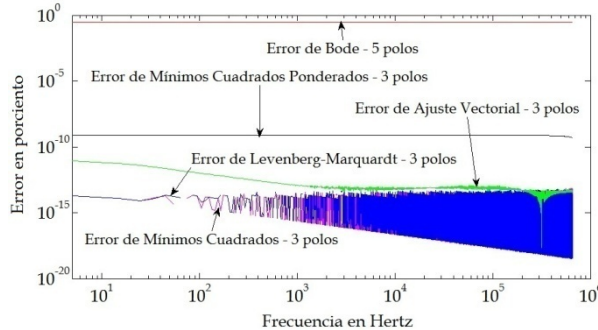


Fig. 5. Error del ajuste de Y.

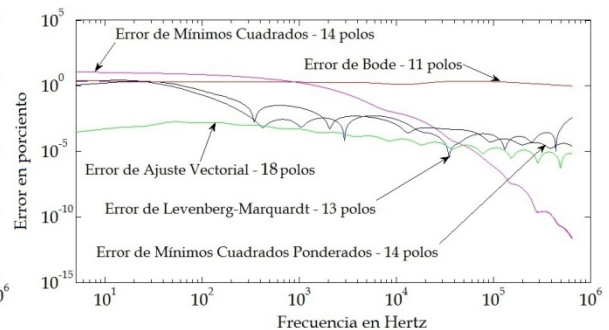


Fig. 6. Error del ajuste de Z.

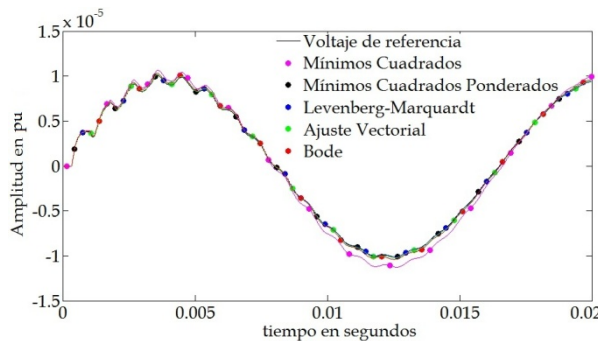


Fig. 7. Prueba de corto circuito con $R_2=0.001 \Omega$.

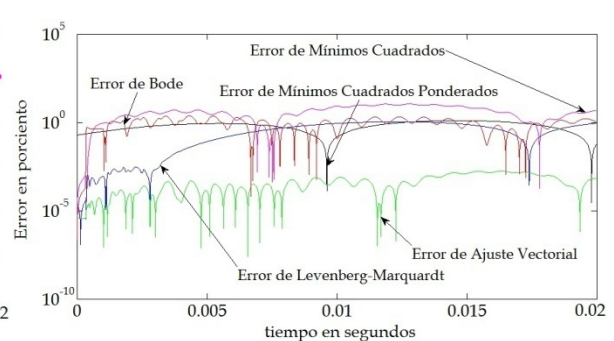
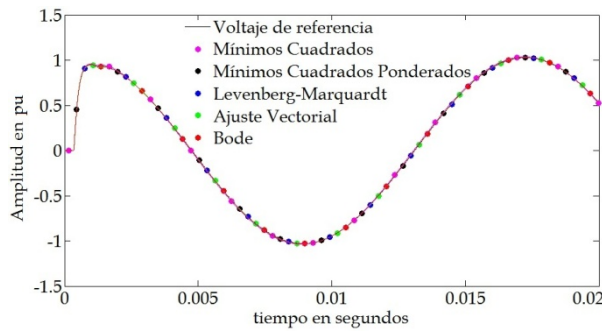
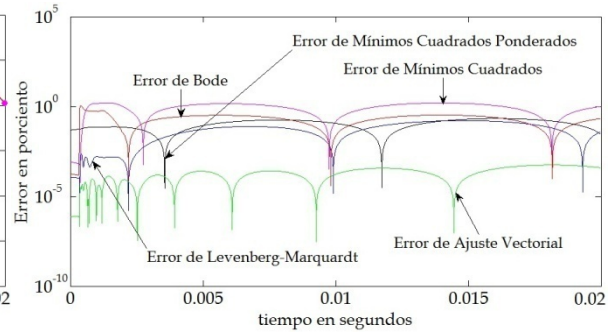
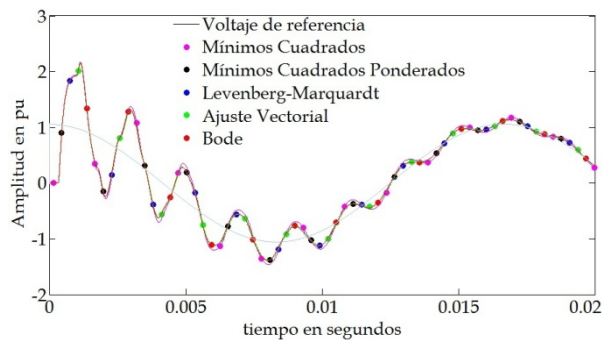
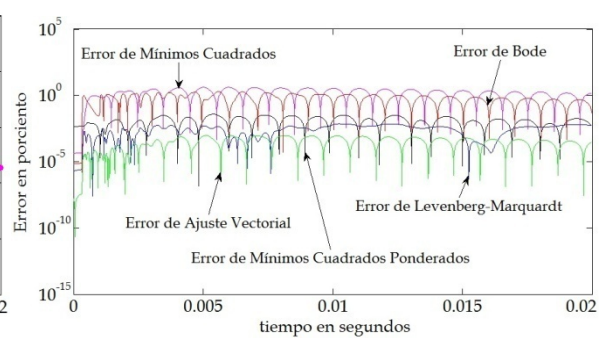


Fig. 8. Error para el caso de $R_2=0.001 \Omega$.

Fig. 9. Línea perfectamente acoplada con $R_2=Z_c$.Fig. 10. Error para el caso de $R_2=Z_c$.Fig. 11. Línea acoplada con $R_2=10000 \Omega$.Fig. 12. Error para el caso de $R_2=10000 \Omega$.

CONCLUSIONES

Aproximación Asintótica de Bode, Mínimos Cuadrados Ordinarios, Mínimos Cuadrados Iterativamente Reponderados, Ajuste Vectorial y Levenberg-Marquardt fueron descritos a detalle y aplicados en la identificación de parámetros para el modelado de sistemas eléctricos.

Dos casos de prueba fueron presentados para la evaluación de las técnicas de ajuste, primero se utilizaron para aproximar una función analítica, después se aplicaron en el ajuste racional de los parámetros Y y Z de una línea de transmisión monofásica.

Las principales conclusiones son las siguientes:

- 1) Para el caso del ajuste racional de la función analítica, se puede notar con claridad que dependiendo del perfil de la función a ajustar unas técnicas son más adecuadas que otras. Es decir, si la función analítica se presenta de la forma de las ecuaciones (1), (2) ó (26), esto incrementa o disminuye la complejidad para cada técnica.
- 2) En el caso del modelado de la línea de transmisión monofásica tres simulaciones en el dominio del tiempo fueron realizadas; línea en corto circuito, línea perfectamente acoplada y línea en circuito abierto. Analizando los ajustes para Y y Z , se puede notar que no siempre la misma técnica da el mejor resultado. También se concluye que el error en el ajuste para Y y Z se acarrea al error de la simulación en el dominio del tiempo (siendo para este caso específico menor para la técnica de Ajuste Vectorial). Sin embargo, debido a que esta técnica fue la que más polos utilizó para hacer el ajuste, la relación precisión-tiempo de ejecución no necesariamente será la más adecuada y siempre dependerá del usuario cuál de estas características tiene mayor peso.
- 3) A través de la aplicación de todas las técnicas fue posible obtener modelos para los parámetros Y y Z de la línea de transmisión, todos estos modelos mostraron su valía al presentar errores muy cercanos entre sí, esto fue posible determinarlo haciendo uso de la TNL.
- 4) Al final, y como punto final, no es posible determinar cuál técnica tendrá mejores resultados. Esto depende, entre otros factores, del tipo de función a ser ajustada, del grado del ajuste deseado, de la complejidad misma de la función, del rango de frecuencia deseado para el ajuste, del número de muestras usado e incluso de la precisión computacional.

Apéndice a- método del gradiente descendente

El gradiente de una función en un punto indica la dirección y sentido de la máxima variación de ésta en dicho punto. El método actualiza los parámetros en dirección opuesta al gradiente de la función objetivo. En este sentido, el gradiente de la función objetivo se expresa con la ecuación (42):

$$\frac{\partial}{\partial \mathbf{x}} \gamma(\mathbf{x}) = (\hat{\mathbf{F}}(\mathbf{s}, \mathbf{x}) - \mathbf{F}(\mathbf{s}))^T \frac{\partial}{\partial \mathbf{x}} (\hat{\mathbf{F}}(\mathbf{s}, \mathbf{x}) - \mathbf{F}(\mathbf{s})) \quad (42)$$

Si se asigna a la variable $\mathbf{J}$, llamada Jacobiano, la función derivada parcial del lado derecho de la ecuación (A.1), se tiene la ecuación (43):

$$\frac{\partial}{\partial \mathbf{x}} \gamma(\mathbf{x}) = (\hat{\mathbf{F}}(\mathbf{s}, \mathbf{x}) - \mathbf{F}(\mathbf{s}))^T \mathbf{J} \quad (43)$$

El Jacobiano representa la sensibilidad de la función a los parámetros $\mathbf{x}$. Por lo tanto, la perturbación $\mathbf{h}$ calculada de acuerdo al método del gradiente que mueve los parámetros en la dirección descendente más pronunciada se calcula de acuerdo a la ecuación (44):

$$\mathbf{h}_{GD} = -\mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (44)$$

Donde por simplicidad se ha omitido la dependencia de las funciones.

Apéndice B- método de gauss-newton

El método supone que la función objetivo es cuadrática en relación a los parámetros cerca de la solución óptima. La función evaluada con los parámetros del modelo perturbado puede ser aproximada a través de la serie de Taylor de primer orden de acuerdo a la ecuación (45):

$$\hat{\mathbf{F}}(\mathbf{s}, \mathbf{x} + \mathbf{h}) \approx \hat{\mathbf{F}}(\mathbf{s}, \mathbf{x}) + \frac{\partial \hat{\mathbf{F}}(\mathbf{s}, \mathbf{x})}{\partial \mathbf{x}} \mathbf{h} = \hat{\mathbf{F}}(\mathbf{s}, \mathbf{x}) + \mathbf{J} \mathbf{h} \quad (45)$$

Por otro lado, la función objetivo a minimizar se expresa con la ecuación (46),

$$\gamma(\mathbf{x}) = \frac{1}{2} \hat{\mathbf{F}}^T \hat{\mathbf{F}} - \hat{\mathbf{F}}^T \mathbf{F} + \frac{1}{2} \mathbf{F}^T \mathbf{F} \quad (46)$$

Donde por simplificación se omite la dependencia de las funciones. Sustituyendo las ecuaciones (45) en la ecuación (46) se llega a la ecuación (47):

$$\gamma(\mathbf{x} + \mathbf{h}) = \frac{1}{2} \hat{\mathbf{F}}^T \hat{\mathbf{F}} - \hat{\mathbf{F}}^T \mathbf{F} + \frac{1}{2} \mathbf{F}^T \mathbf{F} + (\hat{\mathbf{F}} - \mathbf{F})^T \mathbf{J} \mathbf{h} + \frac{1}{2} \mathbf{h}^T \mathbf{J}^T \mathbf{J} \mathbf{h} \quad (47)$$

La ecuación anterior muestra que $\gamma(\mathbf{x})$ es aproximadamente cuadrática para la perturbación $\mathbf{h}$ y que el Hessiano para $\gamma(\mathbf{x})$ es aproximado por $\mathbf{J}^T \mathbf{J}$. La perturbación $\mathbf{h}$ que minimiza la función objetivo se encuentra con la ecuación (48a).

$$\frac{\partial \gamma(\mathbf{x} + \mathbf{h})}{\partial \mathbf{h}} = 0 \quad (48a)$$

Donde la derivada de la ecuación (47), es la ecuación (48b):

$$\frac{\partial \gamma(\mathbf{x} + \mathbf{h})}{\partial \mathbf{h}} = (\hat{\mathbf{F}} - \mathbf{F})^T \mathbf{J} + \mathbf{h}^T \mathbf{J}^T \mathbf{J} \quad (48b)$$

Así, las ecuaciones normales para el método de Gauss-Newton están dadas por la ecuación (49):

$$(\mathbf{J}^T \mathbf{J}) \mathbf{h}_{GN} = -\mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (49)$$

Por lo tanto, la perturbación $\mathbf{h}$ calculada de acuerdo al método Gauss-Newton que actualiza los parámetros está dada por la ecuación (50):

$$\mathbf{h}_{GN} = -(\mathbf{J}^T \mathbf{J})^{-1} \mathbf{J}^T (\hat{\mathbf{F}} - \mathbf{F}) \quad (50)$$

REFERENCIAS

- [1]. LEVY E. C., "Complex curve fitting". *IRE Transactions on Automatic Control*. 1959, vol. AC-4, núm. 1, p. 37-44, DOI: [10.1109/TAC.1959.6429401](https://doi.org/10.1109/TAC.1959.6429401), ISSN 0096-199X.
- [2]. SANATHANAN C. K., KOERNER J., "Transfer function synthesis as a ratio of two complex polynomials". *IEEE Transactions on Automatic Control*. 1963, vol. 8, núm. 1, p. 56-58, DOI: [10.1109/TAC.1963.1105517](https://doi.org/10.1109/TAC.1963.1105517), ISSN 0018-9286.
- [3]. MARTI J. R., "Accurate modeling of frequency-dependent transmission lines in electromagnetic transient simulations". *IEEE Transactions Power Apparatus and Systems*. 1982, vol. PAS-101, n. 1, p. 147-157, DOI: [10.1109/TPAS.1982.317332](https://doi.org/10.1109/TPAS.1982.317332), ISSN 0018-9510.
- [4]. GUSTAVSEN B., SEMLYEN A., "Rational approximation of frequency domain responses by vector fitting". *IEEE Transactions on Power Delivery*. 1999, vol. 14, n. 3, p. 1052-1061, DOI: [10.1109/61.772353](https://doi.org/10.1109/61.772353), ISSN 0885-8977.
- [5]. SCHNEIDER H., WEILAND T., "Transient equivalent surface field excitations via rational approximation by vector fitting". *International Conference on Electromagnetics in Advanced Applications (ICEAA)*. 2014, p. 44-47, Palm Beach, Aruba, 3-8 Agosto, DOI: [10.1109/ICEAA.2014.6903826](https://doi.org/10.1109/ICEAA.2014.6903826), ISBN 978-1479973255.
- [6]. TANG M., SUN J., GUO J., LI H., ZHANG W., LIANG D., "Time-Domain Modeling and simulation of partial discharge on medium-voltage cables by vector fitting method". *IEEE Transactions on Magnetics*. 2014, vol. 50, n. 2, p. 993-996, DOI: [10.1109/TMAG.2013.2283729](https://doi.org/10.1109/TMAG.2013.2283729), ISSN 0018-9464.
- [7]. GORNIK P., BANDURSKI W., "A new universal approach to time-domain modeling and simulation of UWB channel containing convex obstacles using vector fitting algorithm". *IEEE Transactions on Antennas and Propagation*. 2014, vol. 62, n. 12, p. 6394-6405, DOI: [10.1109/TAP.2014.2359203](https://doi.org/10.1109/TAP.2014.2359203), ISSN 0018-926X.
- [8]. NODA T., "Identification of a multiphase network equivalent for electromagnetic transient calculations using partitioned frequency response". *IEEE Transactions on Power Delivery*. 2005, vol. 20, n. 2, p. 1134-1142, DOI: [10.1109/TPWRD.2004.834347](https://doi.org/10.1109/TPWRD.2004.834347), ISSN 0885-8977.
- [9]. NODA T., "Application of frequency-partitioning fitting to the phase-domain frequency-dependent modeling of overhead transmission lines". *IEEE Transactions on Power Delivery*. 2014, vol. 30, n. 1, p. 174-183, DOI: [10.1109/TPWRD.2014.2329532](https://doi.org/10.1109/TPWRD.2014.2329532), ISSN 0885-8977.
- [10]. PORDANJANI I. R., CHUNG C. Y., MAZIN H. E., XU W., "A method to construct equivalent circuit model from frequency responses with guaranteed passivity". *IEEE Transactions on Power Delivery*. 2011, vol. 26, n. 1, p. 400-409, DOI: [10.1109/TPWRD.2010.2080323](https://doi.org/10.1109/TPWRD.2010.2080323), ISSN 0885-8977.
- [11]. NILSSON J. A., "RIEDEL S. A., *Electric circuits*". 8a ed. Prentice Hall, 2008. 880 p., p. 799-815, ISBN 978-0131989252.
- [12]. BJÖRCK A., "Numerical methods for least squares problems". 1a ed. Philadelphia, PA: SIAM, 1996. 408 p., p. 165-175, ISBN 978-0898713602.
- [13]. FAN J., "The modified Levenberg-Marquardt method for nonlinear equations with cubic convergence". *Mathematics of Computation*. 2011, vol. 81, n. 277, p. 447-466, DOI: <http://dx.doi.org/10.1090/S0025-5718-2011-02496-8#sthash.CDCFAJOK.dpuf>, ISSN 1088-6842.
- [14]. RAMOS-LEANOS O., NAREDO J. L., MAHSEREDJIAN J., DUFOUR C., GUTIERREZ-ROBLES J. A., KOCAR I., "A wideband line/cable model for real-time simulations of power system transients". *IEEE Transactions on Power Delivery*. 2012, vol. 27, n. 4, p. 2211-2218, DOI: [10.1109/TPWRD.2012.2206833](https://doi.org/10.1109/TPWRD.2012.2206833), ISSN 0885-8977.
- [15]. RAMOS-LEANOS O., NAREDO J. L., GUTIERREZ-ROBLES J. A., "An advanced transmission line and cable model in MATLAB for the simulation of power-system transients", Chapter 12 (p. 269-304) of the book *MATLAB - A Fundamental Tool for Scientific Computing and Engineering Applications*. 2012, ISBN 978-953-51-0750-7.
- [16]. NAREDO J. L., MAHSEREDJIAN J., GUTIERREZ-ROBLES J. A., RAMOS-LEANOS O., DUFOUR C., BELANGER J., "Improving numerical performance of transmission line models in EMTP". *Proceedings of the 9th International Conference on Power Systems Transients, IPST2011*. 2011, vol. N/A, n. N/A, p. N/A, Delft, The Netherlands, 14-17 June. ISBN N/A.
- [17]. SEMLYEN A., GUSTAVSEN B., "Phase-domain transmission-line modeling with enforcement of symmetry via the propagated characteristic admittance matrix". *IEEE Transactions on Power Delivery*. 2012, vol. 27, n. 2, p. 626-631, DOI: [10.1109/TPWRD.2011.2180405](https://doi.org/10.1109/TPWRD.2011.2180405), ISSN 0885-8977.
- [18]. GOMEZ P., ESPINO-CORTES F.P., URIBE F.A., "Accurate computation of transient profiles along multiconductor transmission systems by means of the numerical Laplace transform". *IEEE Transactions on Power Delivery*. 2014, vol. 29, n. 5, p. 2385-2393, DOI: [10.1109/TPWRD.2014.2313526](https://doi.org/10.1109/TPWRD.2014.2313526), ISSN 0885-8977.
- [19]. GOMEZ P., ESPINO-CORTES F.P., "Computation of transient voltage profiles along transmission lines by successive application of the numerical Laplace transform". *North American Power Symposium (NAPS)*. 2013, vol. N/A, n. N/A, p. N/A, Manhattan, KS, 22-24 Sep., ISBN N/A.
- [20]. SHI-PENG B., JING-SHENG J., "Modeling and calculation for conductive coupling caused by lightning over-voltage in substation based on numerical inverse Laplace transform". *Power and Energy Engineering Conference (APPEEC), 2012 Asia-Pacific*. 2012, p. 1-4 Shanghai, 27-29 March, DOI: [10.1109/APPEEC.2012.6307361](https://doi.org/10.1109/APPEEC.2012.6307361), ISBN 978-1457705458.

AUTORES

Eduardo Salvador Bañuelos Cabral

Recibió los grados de Licenciatura y Maestría en la Universidad de Guadalajara, México en 2006 y 2009 respectivamente. Actualmente es estudiante de Doctorado en Cinvestav-Guadalajara. México.

e-mail: ebanuelos@gdl.cinvestav.mx

José Alberto Gutiérrez Robles

Recibió su grado de Licenciatura y Maestría en Ciencias en Ingeniería Eléctrica de la Universidad de Guadalajara, México, en 1993 y 1998, respectivamente, y su grado de Doctor en Ciencias de Cinvestav-Guadalajara en 2002. Actualmente es Profesor de tiempo completo en el Departamento de Matemáticas de la Universidad de Guadalajara. México.

e-mail: jose.gutierrez@cucei.udg.mx

José Luis Naredo Villagrán

Recibió su grado de Doctor en Ciencias de UBC-Canadá en 1997. Actualmente es Profesor de tiempo completo en Cinvestav-Unidad Guadalajara, México.

e-mail: jlnaredo@gdl.cinvestav.mx

Julián Sotelo Castañón

Recibió los grados de Licenciatura y Maestría en la Universidad de Guadalajara, México en 2006 y 2009 respectivamente. Actualmente es estudiante de Doctorado en Cinvestav-Guadalajara. México.

e-mail: jsotelo@gdl.cinvestav.mx

Verónica Adriana Galván Sánchez

Recibió su grado de Licenciatura y Maestría en la Universidad de Guadalajara, México en 2008 y 2011 respectivamente. Actualmente es estudiante de Doctorado en Cinvestav-Guadalajara. México.

e-mail: vgalvan@gdl.cinvestav.mx

Jorge Luis García Sánchez

Recibió los grados de Licenciatura y Maestría en la Universidad de Guadalajara, México en 2006 y 2009 respectivamente. Actualmente es estudiante de Doctorado en Cinvestav-Guadalajara. México.

e-mail: jlgarcia@gdl.cinvestav.mx