

3.1 DETERMINACIÓN DE LA IMPEDANCIA POR MEDIO DE LA SOLUCIÓN DE LAS ECUACIONES DE ESTADO ESTACIONARIO

En este grupo están los algoritmos de protección que, para la determinación de la impedancia utilizan las ecuaciones de estado estacionario, es decir parten de señales de entradas sin armónicos, senoidales puras; por ello son caracterizadas normalmente como algoritmos senoidales. Los cuatro algoritmos más conocidos son:

- Gilbert/Shovlin
- Método T2 de Lobos
- Mann/Morrison
- Gilcrest/Rockefeller

3.1.1 Algoritmo de Gilber/Shovlin

Señales de entrada: $i(t)$ y $u(t)$ senoidales. Se supone la corriente con fase cero, sin con ello limitar la generalidad.

$$i(t) = \hat{I} \sin(\omega t) \quad (3.1.1)$$

$$u(t) = \hat{U} \sin(\omega t + \varphi) \quad (3.1.2)$$

En forma compleja:

$$\underline{I} = j \frac{\hat{I}}{\sqrt{2}} \quad (3.1.3)$$

$$\underline{U} = j \frac{\hat{U}}{\sqrt{2}} e^{j\varphi} \quad (3.1.4)$$

$$\underline{Z} = \frac{\underline{U}}{\underline{I}} = \frac{\hat{U}}{\hat{I}} e^{j\varphi} = \frac{\hat{U}}{\hat{I}} \cos \varphi + j \frac{\hat{U}}{\hat{I}} \sin \varphi = R + jX \quad (3.1.5)$$

La comparación de coeficientes resulta en que:

$$R = \frac{\hat{U}}{\hat{I}} \cos(\varphi) \quad (3.1.6)$$

$$X = \frac{\hat{U}}{\hat{I}} \sin(\varphi) \quad (3.1.7)$$

Para poder estimar los valores picos y el ángulo de fase de la corriente y la tensión, el algoritmo necesita correspondientemente 3 muestras equidistantes de $i(t)$ y $u(t)$ (ventana de datos $n=3$).

$$i_1 = \hat{I} \sin(\omega t_3 - 2\Delta) \quad (3.1.8a)$$

$$i_2 = \hat{I} \sin(\omega t_3 - \Delta) \quad (3.1.8b)$$

$$i_3 = \hat{I} \sin(\omega t_3) \quad (3.1.8c)$$

$$u_1 = \hat{u} \sin(\omega t_3 + \varphi - 2\Delta) \quad (3.1.9a)$$

$$u_2 = \hat{u} \sin(\omega t_3 + \varphi - \Delta) \quad (3.1.9b)$$

$$u_3 = \hat{u} \sin(\omega t_3 + \varphi) \quad (3.1.9c)$$

con $\Delta = \omega \Delta t$

Con esto se tiene 6 ecuaciones no lineales con las 4 incógnitas $\hat{I}, \hat{u}, t_3$ y φ . Se trata de un sistema sobredimensionado. Una posible solución, a partir de la eliminación de t_3 mediante complicadas operaciones trigonométricas, sería:

$$R = \frac{u_2 i_2 - 0.5 \cdot (u_3 i_1 + u_1 i_3)}{i_2^2 - i_1 i_3} \quad (3.1.10)$$

$$X = \frac{u_2 i_3 - u_3 i_2}{i_2^2 - i_1 i_3} \cdot \sin(\Delta) \quad (3.1.11)$$

Las pruebas realizadas con este algoritmo han mostrado que el mismo, excepto para las oscilaciones de la frecuencia, reacciona con alta sensibilidad a las señales de entrada no ideales, mostrando sobre- y subalcances para todos los ángulos de fase de la línea. Por ello, la aplicación de este algoritmo tiene sentido con filtros de banda angosta.

3.1.2 Método T2 de Lobos

Este método necesita 2 muestras de corriente y tensión. Partiendo del modelo de línea de 1er. orden se tiene:

$$u(t) = R i(t) + L \frac{di(t)}{dt} \quad (3.1.12)$$

Con ello se da, para las muestras en estado estacionario:

$$i_1 = \hat{i} \sin(\omega t_1) \quad (3.1.13a)$$

$$i_2 = \hat{i} \sin(\omega t_1 + \Delta) \quad (3.1.13b)$$

$$u_1 = R \hat{i} \sin(\omega t_1) + L \omega \hat{i} \cos(\omega t_1) = \hat{i} \cdot (R \sin(\omega t_1) + X \cos(\omega t_1)) \quad (3.1.14a)$$

$$u_2 = \hat{i} \cdot (R \sin(\omega t_1 + \Delta) + X \cos(\omega t_1 + \Delta)) \quad (3.1.14b)$$

con $\Delta = \omega \Delta t$

Dividiendo se obtiene:

$$\frac{u_1}{i_1} = R + X \cdot \frac{\cos(\omega t_1)}{\sin(\omega t_1)} \quad (3.1.15a)$$

$$\frac{u_2}{i_2} = R + X \cdot \frac{\cos(\omega t_1 + \Delta)}{\sin(\omega t_1 + \Delta)} \quad (3.1.15b)$$

Por medio de transformaciones de las ecuaciones 3.15 a y b se llega a las siguientes ecuaciones para el cálculo de la resistencia y la reactancia.

$$R = \frac{u_1 \cdot (i_2 \cos(\Delta) - i_1) + u_2 \cdot (i_1 \cos(\Delta) - i_2)}{i_1 \cdot (i_2 \cos(\Delta) - i_1) + i_2 \cdot (i_1 \cos(\Delta) - i_2)} \quad (3.1.16)$$

$$X = \frac{u_2 i_1 - u_1 i_2}{i_1 \cdot (i_2 \cos(\Delta) - i_1) + i_2 \cdot (i_1 \cos(\Delta) - i_2)} \cdot \sin(\Delta) \quad (3.1.17)$$

Este método es, respecto a las armónicas superiores, algo más sensible que el algoritmo de Gilbert/Shovlin. Para ambos métodos existe una correspondencia más lineal entre la amplitud de los armónicos de orden superior y del error resultante de los mismos. Este

método reacciona de forma más sensible a la componente unidireccional, error de los transformadores y oscilaciones de frecuencia que el algoritmo de Gilbert/Shovlin.

3.1.3 Algoritmo de Mann/Morrison

Punto de partida del cálculo son nuevamente las ecuaciones de corriente y tensión de estado estacionario:

$$i(t) = \hat{i} \sin(\omega t) \quad (3.1.18)$$

$$u(t) = \hat{u} \sin(\omega t + \varphi) \quad (3.1.19)$$

Si se forman las derivadas primeras de la corriente y la tensión, luego se obtiene:

$$\dot{i}(t) = \omega \hat{i} \cos(\omega t) \quad (3.1.20)$$

$$\dot{u}(t) = \omega \hat{u} \cos(\omega t + \varphi) \quad (3.1.21)$$

Elevando al cuadrado y sumando las ecuaciones 3.1.18 a 3.1.21 se obtiene:

$$\hat{u}^2 = u^2 + \left(\frac{\dot{u}}{\omega}\right)^2 \quad (3.1.22)$$

$$\hat{i}^2 = i^2 + \left(\frac{\dot{i}}{\omega}\right)^2 \quad (3.1.23)$$

Con ello, para el módulo de la impedancia vale:

$$|Z| = \sqrt{\frac{\hat{u}^2}{\hat{i}^2}} = \sqrt{\frac{u^2 + \left(\frac{\dot{u}}{\omega}\right)^2}{i^2 + \left(\frac{\dot{i}}{\omega}\right)^2}} \quad (3.1.24)$$

Para la determinación del ángulo de fase φ , se dividen de a pares las ec. 3.1.20 y 3.1.19, así como las ec. 3.1.19 y 3.1.21:

$$\tan(\omega t) = \frac{i}{I/\omega} \quad (3.1.25)$$

$$\tan(\omega t + \varphi) = \frac{u}{\dot{u}/\omega} \quad (3.1.26)$$

De allí resulta:

$$\varphi = \arctan\left(\frac{u}{\dot{u}/\omega}\right) - \arctan\left(\frac{i}{I/\omega}\right) \quad (3.1.27)$$

$$\varphi = \arctan\left(\frac{\omega u}{\dot{u}}\right) - \arctan\left(\frac{\omega i}{I}\right) \quad (3.1.28)$$

Los valores instantáneos de corriente y tensión así como los valores instantáneos de sus derivadas primeras se pueden aproximar, a partir de dos valores muestreados, medidos en el intervalo Δt , de acuerdo a las siguientes ecuaciones:

$$u \approx \frac{u_1 + u_2}{2} \quad (3.1.29)$$

$$\dot{u} \approx \frac{u_1 - u_2}{\Delta t} \quad (3.1.30)$$

De allí, se obtiene para el módulo y ángulo de la impedancia:

$$|Z| = \sqrt{\frac{\left(\frac{u_1 + u_2}{2}\right)^2 + \left(\frac{u_1 - u_2}{\Delta}\right)^2}{\left(\frac{i_1 + i_2}{2}\right)^2 + \left(\frac{i_1 - i_2}{\Delta}\right)^2}} \quad (3.1.31)$$

$$\varphi = \arctan\left(\frac{\Delta}{2} \cdot \frac{u_1 + u_2}{u_1 - u_2}\right) - \arctan\left(\frac{\Delta}{2} \cdot \frac{i_1 + i_2}{i_1 - i_2}\right) \quad (3.1.32)$$

con $\Delta = \omega \Delta t$

Para la determinación de la impedancia alcanzan luego 2 valores muestreados de corriente y tensión. En contraposición con el método de Gilbert/Shovlin no es necesario aquí ninguna sobredeterminación para calcular la impedancia.

La resistencia y la reactancia se obtienen a partir de la ec. 3.1.31 y 3.1.32 como:

$$R = |Z| \cdot \cos(\varphi) \quad (3.1.33)$$

$$X = |Z| \cdot \sin(\varphi) \quad (3.1.34)$$

Como para los algoritmos de Gilbert/Shovlin y T2 de Lobos, todas las señales de entradas no ideales conducen a fuertes sobre- y subfunciones. Aquí son los errores considerablemente más grandes que para los algoritmos tratados. El fundamento para la sensibilidad es la aproximación de la diferenciación por medio de diferencias. Dado que se parte de señales de entrada senoidales puras, se amplifican adicionalmente todavía más las componentes de alta frecuencia.

3.1.4 Algoritmo de Gilcrest/Rockefeller

A diferencia del algoritmo de Mann/Morrison, en este método se aplica, para el cálculo de la resistencia y reactancia, las derivadas primera y segunda de las señales de entrada, para eliminar la componente continua.

$$\dot{u}(t) = \omega \hat{u} \cos(\omega t + \varphi) \quad (3.1.35)$$

$$\ddot{u}(t) = -\omega^2 \hat{u} \cos(\omega t + \varphi) \quad (3.1.36)$$

$$\dot{i}(t) = \omega \hat{i} \cos(\omega t) \quad (3.1.37)$$

$$\ddot{i}(t) = -\omega^2 \hat{i} \cos(\omega t) \quad (3.1.38)$$

Elevando al cuadrado y sumando se tiene:

$$\omega^2 \hat{u}^2 = \dot{u}^2 + \left(\frac{\ddot{u}}{\omega}\right)^2 \quad (3.1.39)$$

$$\omega^2 \hat{i}^2 = \dot{i}^2 + \left(\frac{\ddot{i}}{\omega}\right)^2 \quad (3.1.40)$$

Con esto se puede describir el módulo y ángulo de fase en forma análoga al algoritmo de Mann/Morrison, por medio de las siguientes ecuaciones:

$$|\underline{Z}| = \sqrt{\frac{\hat{u}^2}{\hat{i}^2}} = \sqrt{\frac{\dot{u}^2 + \left(\frac{\ddot{u}}{\omega}\right)^2}{\dot{i}^2 + \left(\frac{\ddot{i}}{\omega}\right)^2}} \quad (3.1.41)$$

$$\varphi = \arctan\left(\frac{\omega \dot{u}}{\ddot{u}}\right) - \arctan\left(\frac{\omega \dot{i}}{\ddot{i}}\right) \quad (3.1.42)$$

Las derivadas de las tensiones se aproximan como sigue:

$$\dot{u} \approx \frac{u_3 - u_1}{2 \Delta t} \quad (3.1.43)$$

$$\ddot{u} \approx \frac{\frac{u_3 - u_2}{\Delta t} - \frac{u_2 - u_1}{\Delta t}}{\Delta t} = \frac{u_3 - 2u_2 + u_1}{\Delta t^2} \quad (3.1.44)$$

En forma análoga se pueden calcular las derivadas de la corriente. Con $\Delta = \omega \Delta t$ resulta, para el módulo de la impedancia, ángulo de fase, R y X:

$$|\underline{Z}| = \sqrt{\frac{\Delta^2 \cdot (u_3 - u_1)^2 + 4 \cdot (u_3 - 2u_2 + u_1)^2}{\Delta^2 \cdot (i_3 - i_1)^2 + 4 \cdot (i_3 - 2i_2 + i_1)^2}} \quad (3.1.45)$$

$$\varphi = \arctan\left(\frac{\Delta}{2} \cdot \frac{u_3 - u_1}{u_3 - 2u_2 + u_1}\right) - \arctan\left(\frac{\Delta}{2} \cdot \frac{i_3 - i_1}{i_3 - 2i_2 + i_1}\right) \quad (3.1.46)$$

$$R = |\underline{Z}| \cdot \cos(\varphi) \quad (3.1.47)$$

$$X = |\underline{Z}| \cdot \sin(\varphi) \quad (3.1.48)$$

A raíz de la diferenciación simple y doble de la corriente y tensión, el algoritmo muestra una sensibilidad todavía mayor a los armónicos superiores que el algoritmo de Mann/Morrison (aplicación de la primera derivada). Con esto, las componentes armónicas, las tensiones de arco, procesos transitorios en cortocircuitos y la saturación de transformadores de medición, conducen a errores extremadamente grandes. La única

ventaja de la aplicación de ambas derivadas es el amortiguamiento de la influencia de las componentes de corriente continua, que por otro lado resulta también insuficiente. Luego, es necesario para la aplicación de este método, todavía un mayor prefiltrado de la tensión y corriente, que para los algoritmos tratados anteriormente.

3.2 DETERMINACIÓN DE LA IMPEDANCIA POR MEDIO DE LA SOLUCIÓN DE LAS ECUACIONES DIFERENCIALES (ED)

Si, a diferencia de los algoritmos senoidales, debe tenerse en cuenta procesos transitorios, se puede modelar la red con un modelo con ecuaciones diferenciales. Mientras mayor precisión de la modelación se requiera, se obtiene ecuaciones diferenciales de mayor orden, lo cual aumenta también el esfuerzo de cálculo.

La solución de las ecuaciones diferenciales se puede realizar en principio de dos maneras. Por un lado, las diferenciaciones se pueden aproximar por medio de la formación de diferencias (ecuaciones de diferencia), por otro lado, una solución puede realizarse por medio de un procedimiento de integración numérica.

En la literatura se proponen exclusivamente algoritmos que parten de redes modelos de primer y segundo orden, donde se utilizan tanto la diferenciación numérica como también la integración. Modelos de mayor orden conducen a esfuerzos de cálculo desmedidamente grandes y provocan problemas numéricos en la solución. Por lo tanto, se tratará solamente los siguientes algoritmos conocidos de la literatura:

- Método A3 de Lobos (ED de 1er. orden, ecuaciones en diferencia)
- Método A4 de Lobos (ED de 1er. orden, ecuaciones en diferencia)
- Bornard/Bastide (ED de 1er. orden, ecuaciones en diferencia)
- McInnes/Morrison (ED de 1er. orden, integración)
- Ranjbar/Cory (ED de 1er. orden, integración)
- **Smolinski**
-

3.2.1 Método A3 de Lobos

Este método parte de que la línea cortocircuitada puede ser descrita por una ED de primer orden.

$$u(t) = R i(t) + L \frac{di(t)}{dt} \quad (3.2.1)$$

Luego, $di(t)/dt = i$ puede aproximarse por medio de cocientes de diferencias $\Delta i/\Delta t$. La diferencia de corrientes Δi se calcula a partir de los valores muestreados vecinos en t_1 y t_2 :

$$i_{12} \approx \frac{i_2 - i_1}{t_2 - t_1} = \frac{i_2 - i_1}{\Delta t} \quad (3.2.2)$$

Para el cálculo de la corriente y tensión se aplican los valores medios:

$$u_{12} \approx \frac{u_1 + u_2}{2} \quad (3.2.3)$$

$$i_{12} \approx \frac{i_1 + i_2}{2} \quad (3.2.4)$$

Si parte de corrientes y tensiones senoidales, luego para el instante $t_1 + \Delta t/2$ vale:

$$u_{12} = \hat{u} \sin(\omega t_2 - \frac{\Delta}{2}) \quad (3.2.5)$$

con $\Delta = \omega \Delta t$

De la aproximación (ec. 3.2.3) se obtiene:

$$u_{12} \approx \frac{1}{2} \hat{u} \cdot (\sin(\omega t_2 - \Delta) + \sin(\omega t_2)) \approx \hat{u} \cdot \sin(\omega t_2 - \frac{\Delta}{2}) \cdot \cos(\frac{\Delta}{2}) \quad (3.2.6)$$

Las ec. 3.2.5 y 3.2.6 se diferencian por el factor $\cos(\Delta/2)$. Así, se puede calcular a partir del valor aproximado, el valor exacto por corrección.

$$u_{12} = \frac{u_1 + u_2}{2 \cdot \cos(\frac{\Delta}{2})} \quad (3.2.7)$$

En forma análoga para la corriente:

$$i_{12} = \frac{i_1 + i_2}{2 \cdot \cos(\frac{\Delta}{2})} \quad (3.2.8)$$

El verdadero valor de la derivada de la corriente en el instante $t_1 + \Delta/2$ se puede derivar de la misma manera:

$$\dot{i}_{12} = \frac{d}{dt} [\hat{i} \cdot \sin(\omega t_2 - \frac{\Delta}{2})] = \omega \hat{i} \cdot \cos(\omega t_2 - \frac{\Delta}{2}) \quad (3.2.9)$$

La formación de las diferencias dan como resultado sin embargo:

$$\begin{aligned} \dot{i}_{12} &= \frac{i_2 - i_1}{\Delta t} = \frac{\hat{i} \cdot (\sin(\omega t_2) - \sin(\omega t_2 - \Delta))}{\Delta t} \\ &= \hat{i} \cdot \cos(\omega t_2 - \frac{\Delta}{2}) \cdot \frac{2 \cdot \sin(\frac{\Delta}{2})}{\Delta t} \end{aligned} \quad (3.2.10)$$

Si se incorpora la ec. 3.2.9, queda luego:

$$\dot{i}_{12} = \frac{i_{12}}{\omega} \cdot \frac{2 \cdot \sin(\frac{\Delta}{2})}{\Delta t}$$

Con esto, el valor exacto de la derivada es:

$$\dot{i}_{12} = \frac{i_2 - i_1}{\Delta t} \cdot \frac{\omega \Delta t}{2 \cdot \sin\left(\frac{\Delta}{2}\right)} = \frac{i_2 - i_1}{\Delta t} \cdot \frac{\frac{\Delta}{2}}{\sin\left(\frac{\Delta}{2}\right)} \quad (3.2.11)$$

Para los valores calculados a partir de las muestras 1 y 2 se obtiene, según la ec. 3.2.1:

$$u_{12} = R i_{12} + L \dot{i}_{12} \quad (3.2.12)$$

De la misma forma se puede proceder para los instantes de tiempo t_2 y t_3 :

$$u_{23} = R i_{23} + L \dot{i}_{23} \quad (3.2.13)$$

Las ec. 3.2.12 y 3.2.13 se resuelven según R y $X = \omega L$. Si se incorporan las ec. 3.2.7, 3.2.8 y 3.2.11, luego resulta:

$$R = \frac{i_1 \cdot (u_2 + u_3) - i_2 \cdot (u_1 + 2 \cdot u_2 + u_3) + i_3 \cdot (u_1 + u_2)}{2 \cdot (i_1 \cdot i_3 - i_2^2)} \quad (3.2.14)$$

$$X = \frac{i_1 \cdot (u_2 + u_3) - i_2 \cdot (u_1 - u_3) - i_3 \cdot (u_1 + u_2)}{2 \cdot (i_1 \cdot i_3 - i_2^2)} \cdot \tan\left(\frac{\Delta}{2}\right) \quad (3.2.15)$$

Este método es, para la mayoría de las señales de entrada, más robusto a los errores que los algoritmos senoidales. Correspondientemente a su aplicación, este reacciona relativamente insensible a la componente unidireccional, para las oscilaciones de frecuencia de la red prácticamente no presenta ningún error. Los armónicos superiores conducen, a raíz de los términos de corrección, a una zona de sobre y subfunción aproximadamente de igual tamaño. Las distorsiones de las señales de entrada producida

por los transformadores provocan principalmente errores de subfunción de 30%. Para su aplicación en redes con cables de alta tensión, el algoritmo no es adecuado sin filtro adicional.

3.2.2 Método A4 de Lobos

En el método A3 de Lobos deben formarse los valores medios para la tensión y la corriente, para obtener la corriente, tensión y la derivada de la corriente en el mismo instante de tiempo. En lugar de las tres muestras, el método A4 aplica 4 valores muestreados de corriente y tensión. Luego, deben aproximarse todavía las derivadas, los valores de corriente y tensión corresponden directamente a los valores muestreados. Se aplica aquí la ecuación diferencial de primer orden de la línea. Las derivadas se aproximan por medio de la formación de diferencias sobre dos intervalos de muestreo.

$$\dot{i}_2 \approx \frac{i_3 - i_1}{2 \Delta t} \quad (3.2.16a)$$

$$\dot{i}_3 \approx \frac{i_4 - i_2}{2 \Delta t} \quad (3.2.16b)$$

Para señales senoidales puras, estas pueden ser corregidas de la misma forma:

$$\begin{aligned} \dot{i}_2 &= \frac{\hat{i} \cdot (\sin(\omega t_2 + \Delta) - \sin(\omega t_2 - \Delta))}{2 \Delta t} \\ &= \omega \hat{i} \cdot \cos(\omega t_2) \cdot \frac{\sin(\Delta)}{\Delta} \end{aligned} \quad (3.2.17)$$

Así resulta para señales senoidales puras, en forma análoga que para las ec. 3.29 y 3.2.11

$$\dot{i}_2 = \omega \hat{i} \cdot \cos(\omega t_2) = \frac{i_3 - i_1}{\Delta t} \cdot \frac{\Delta}{2 \sin(\Delta)} \quad (3.2.18)$$

Si se reemplaza las derivadas $\dot{I}_2$ e $\dot{I}_3$ obtenidas de la ec. 3.2.18 y los correspondientes valores de corriente y tensión, en la ec. diferencial 3.2.1 y se despeja R y X , luego resulta:

$$R = \frac{u_2 \cdot (i_4 - i_2) - u_3 \cdot (i_3 - i_1)}{i_2 \cdot (i_4 - i_2) - i_3 \cdot (i_3 - i_1)} \quad (3.2.19)$$

$$X = \frac{u_2 i_2 - u_3 i_3}{i_2 \cdot (i_4 - i_2) - i_3 \cdot (i_3 - i_1)} \cdot 2 \sin(\Delta) \quad (3.2.20)$$

Este método requiere, debido a la aproximación errónea para los valores intermedios de la corriente y la tensión, a diferencia del método A3, un menor esfuerzo de cálculo, pero, para casi todas las señales de entrada no ideales, entrega zonas de sub y sobrefunción el doble de grandes comparado con el método A3. Por ello se prefiere el método A3.

3.2.3 Algoritmo de Bornard/Bastide

Si se parte nuevamente del modelo de línea de 1er orden, pero esta vez considerando las capacidades transversales, resistencia de la falla, etc..mediante el cálculo de una tensión de falla $e(t)$, luego puede describirse la línea cortocircuitada por medio de la siguiente ecuación:

$$u(t) + e(t) = R i(t) + L \frac{i(t) - i(t - \Delta t)}{\Delta t} \quad (3.2.21)$$

En esta ec., la derivada vale para el instante $t - \Delta/2$. Esto hace necesario una corrección, que es introducida más tarde. Para la simplificación de la forma de escritura, se lleva a cabo las siguientes substituciones:

$$R = a_1 + a_2 \quad (3.2.22a)$$

$$L = -a_2 \Delta t \quad (3.2.22b)$$

De la ec. 3.2.21 se obtiene luego la expresión:

$$u(t) + e(t) = a_1 i(t) + a_2 i(t - \Delta t) \quad (3.2.23)$$

Si se pone en función de los valores de muestra de la corriente y tensión, luego vale para el instante t_2 :

$$u_2 + e_2 = a_1 i_2 + a_2 i_1 \quad (3.2.24)$$

Para determinar los coeficientes a_1 y a_2 , la ec. 3.2.24 se considera hasta $n-1$ instantes de muestreo.

$$U + E = I \cdot A \quad (3.2.25)$$

$$\text{mit } U = \begin{bmatrix} u_2 \\ \vdots \\ u_n \end{bmatrix}; \quad E = \begin{bmatrix} e_2 \\ \vdots \\ e_n \end{bmatrix}; \quad I = \begin{bmatrix} i_2 & i_1 \\ \vdots & \vdots \\ i_n & i_{n-1} \end{bmatrix}; \quad A = \begin{bmatrix} a_1 \\ a_2 \end{bmatrix}.$$

Para la determinación clara de ambas magnitudes a_1 y a_2 son suficientes 2 ecuaciones ($n=3$). Si hay a disposición más puntos de muestreo, luego puede realizarse un cálculo de compensación. Para ello se forma la suma del cuadrado de los errores como función objetivo:

$$F(A) = \sum_{i=2}^n e_i^2 = E^T \cdot E = [(I \cdot A)^T - U^T] \cdot [I A - U]$$

$$= A^T I^T I A - 2 A^T I^T U + U^T U \quad (3.2.26)$$

Si se minimiza el error, luego debe satisfacerse que:

$$\frac{\partial F(A)}{\partial A} = 0 \quad (3.2.27)$$

De allí sigue que: $I^T I A = I^T U \quad (3.2.28)$

En conexión se realizan nuevas substituciones:

$$P = I^T I \quad (3.2.29)$$

$$B = I^T U \quad (3.2.30)$$

Ya que solo hay que calcular 2 magnitudes para a_1 y a_2 , luego P resulta una matriz de 2x2 y B un vector de 2 componentes.

$$\begin{bmatrix} p_{11} & p_{12} \\ p_{21} & p_{22} \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \end{bmatrix} \quad (3.2.31)$$

Para los coeficientes a_1 y a_2 se da, con $p_{12} = p_{21}$:

$$a_1 = \frac{p_{22} b_1 - p_{12} b_2}{p_{11} p_{22} - p_{12}^2} \quad (3.2.32)$$

$$a_2 = \frac{p_{11} b_2 - p_{12} b_1}{p_{11} p_{22} - p_{12}^2} \quad (3.2.33)$$

Si se consideran las matrices P y B, así puede describirse sus elementos con la ayuda de las ec. 3.2.25, 3.2.29 y 3.2.30, como sigue:

$$p_{11} = \sum_{j=2}^n i_j^2 \quad (3.2.34a)$$

$$b_1 = \sum_{j=2}^n i_j u_j \quad (3.2.35a)$$

$$p_{12} = \sum_{j=2}^n i_j \cdot i_{j-1} \quad (3.2.34b)$$

$$b_2 = \sum_{j=2}^n i_{j-1} u_j \quad (3.2.35a)$$

$$p_{21} = p_{12} \quad (3.2.34c)$$

$$p_{22} = \sum_{j=2}^n i_{j-1}^2 \quad (3.2.34d)$$

Con esto se da para R y L de la línea:

$$\begin{aligned} R = a_1 + a_2 &= \frac{p_{22} b_1 - p_{12} b_2 + p_{11} b_2 - p_{12} b_1}{p_{11} p_{22} - p_{12}^2} \\ &= \frac{b_1 \cdot (p_{22} - p_{12}) + b_2 \cdot (p_{11} - p_{12})}{p_{11} p_{22} - p_{12}^2} \\ &= \frac{\sum i_j u_j \cdot (\sum i_{j-1}^2 - \sum i_j i_{j-1}) + \sum i_{j-1} u_j \cdot (\sum i_j^2 - \sum i_j i_{j-1})}{\sum i_j^2 \cdot \sum i_{j-1}^2 - (\sum i_j i_{j-1})^2} \quad (3.2.36) \end{aligned}$$

$$\begin{aligned} X = -a_2 \Delta t &= \frac{p_{12} b_1 - p_{11} b_2}{p_{11} p_{22} - p_{12}^2} \cdot \Delta t \\ &= \frac{\sum i_j u_j \cdot \sum i_j i_{j-1} - \sum i_{j-1} u_j \cdot \sum i_j^2}{\sum i_j^2 \cdot \sum i_{j-1}^2 - (\sum i_j i_{j-1})^2} \cdot \Delta t \quad (3.2.37) \end{aligned}$$

Con: $\sum x = \sum_{l=2}^n x$

Como para los métodos A3 y A4, puede corregirse la R y la L. Para señales senoidales puras, resulta a partir del modelo de línea descrito por la ec. diferencial del 1er orden:

$$\hat{u} \sin(\omega t + \varphi) = R_O \hat{I} \sin(\omega t) + \omega L_O \hat{I} \cos(\omega t) \quad (3.2.38)$$

mit R_O : Valor exacto de la R

L_O : Valor exacto de la L

Correspondientemente a la ec. 3.2.21, se obtiene como aproximación:

$$\begin{aligned} \hat{u} \sin(\omega t + \varphi) &= R \hat{I} \sin(\omega t) + L \hat{I} \cdot \frac{2 \sin(\frac{\Delta}{2}) \cdot \cos(\omega t - \frac{\Delta}{2})}{\Delta t} \\ &= R \hat{I} \sin(\omega t) + L \hat{I} \cdot \frac{2 \sin(\frac{\Delta}{2})}{\Delta t} \cdot \left(\cos(\omega t) \cos(\frac{\Delta}{2}) + \sin(\omega t) \sin(\frac{\Delta}{2}) \right) \\ &= \hat{I} \sin(\omega t) \cdot \left(R + L \cdot \frac{2 \sin(\frac{\Delta}{2})}{\Delta t} \cdot \sin(\frac{\Delta}{2}) \right) \\ &\quad + \omega \hat{I} \cos(\omega t) \cdot \left(L \cdot \frac{2 \sin(\frac{\Delta}{2})}{\Delta t} \cdot \cos(\frac{\Delta}{2}) \right) \end{aligned} \quad (3.2.39)$$

Con $\Delta = \omega \Delta t$

La comparación de coeficientes entre las ec. 3.2.28 y 3.2.29 dan:

$$L_O = L \cdot \frac{2 \sin(\frac{\Delta}{2}) \cdot \cos(\frac{\Delta}{2})}{\Delta} = L \cdot \frac{\sin(\Delta)}{\Delta} \quad (3.2.40)$$

$$\begin{aligned} R_O &= R + L \cdot \frac{2 \sin(\frac{\Delta}{2}) \cdot \sin(\frac{\Delta}{2})}{\Delta t} \\ &= R + L_O \cdot \frac{\sin(\frac{\Delta}{2})}{\cos(\frac{\Delta}{2})} \cdot \frac{\Delta}{\Delta t} \\ &= R + L_O \cdot \tan(\frac{\Delta}{2}) \cdot \omega \end{aligned} \quad (3.2.41)$$

Este algoritmo es, a raíz del cálculo de compensación de falla, más sensible a las perturbaciones que los algoritmos senoidales y los métodos A3 y A4.

3.2.4 Algoritmo de McInnes/Morrison

Para la determinación de la R y de X, en este algoritmo se soluciona la ec. diferencial de 1er orden por medio de integración sobre 2 intervalos de tiempo de igual tamaño desplazados uno a otro

[t1....t6] y [t6....t11]:

$$\int_{t_1}^{t_6} u \, dt = R \cdot \int_{t_1}^{t_6} i \, dt + L \cdot \int_{i_1}^{i_6} di = R \cdot \int_{t_1}^{t_6} i \, dt + L \cdot (i_6 - i_1) \quad (3.2.42a)$$

$$\int_{t_6}^{t_{11}} u \, dt = R \cdot \int_{t_6}^{t_{11}} i \, dt + L \cdot \int_{i_6}^{i_{11}} di = R \cdot \int_{t_6}^{t_{11}} i \, dt + L \cdot (i_6 - i_{11}) \quad (3.2.42b)$$

Las cuatro integrales definidas pueden calcularse en forma numérica con el procedimiento del trapecio y la cuerda

Luego la solución de R y L queda:

$$R = \frac{(i_6 - i_1) \cdot \int_{t_6}^{t_{11}} u \, dt - (i_{11} - i_6) \cdot \int_{t_1}^{t_6} u \, dt}{(i_6 - i_1) \cdot \int_{t_6}^{t_{11}} i \, dt - (i_{11} - i_6) \cdot \int_{t_1}^{t_6} i \, dt} \quad (3.2.43)$$

$$L = \frac{\int_{t_1}^{t_6} u \, dt \cdot \int_{t_6}^{t_{11}} i \, dt - \int_{t_6}^{t_{11}} u \, dt \cdot \int_{t_1}^{t_6} i \, dt}{(i_6 - i_1) \cdot \int_{t_6}^{t_{11}} i \, dt - (i_{11} - i_6) \cdot \int_{t_1}^{t_6} i \, dt} \quad (3.2.44)$$

Este algoritmo también muestra, en comparación con el A3 y A4, una sensibilidad a las perturbaciones sensiblemente más pequeña en relación con la 3. y 5. armónica. Aquí supera el procedimiento de compensación de error de Bornard/Bastide. El motivo de ese comportamiento está en la supresión parcial de esas componentes espectrales por medio

del efecto alizante de la integración. Correspondientemente pequeños son también los errores en el caso de cortocircuitos de arcos y de cables.

3.2.5 Algoritmo de Ranjbar/Cory

Igual que el algoritmo de McInnes/Morrison, este procedimiento soluciona la ecuación diferencial de 1er. Orden por medio de integración numérica, donde los armónicos superiores son filtrados /18/. Con ello, este método representa una transición hacia los algoritmos de filtrado tratados más adelante. La aplicación del método a la ecuación diferencial de la línea da:

$$\sum_1 \int u \, dt = R \cdot \sum_1 \int i \, dt + L \cdot \sum_1 \int di \quad (3.2.45)$$

Cada integral se calcula por medio de la regla del trapezio. La integral de di se convierte en la diferencia de los límites de integración. Para poder determinar R y L, la ecuación diferencial debe solucionarse nuevamente en dos intervalos de suma.

$$R = \frac{\sum_2 \int u \, dt \cdot \sum_1 \int di - \sum_1 \int u \, dt \cdot \sum_2 \int di}{\sum_2 \int i \, dt \cdot \sum_1 \int di - \sum_1 \int i \, dt \cdot \sum_2 \int di} \quad (3.2.46)$$

$$L = \frac{\sum_1 \int u \, dt \cdot \sum_2 \int i \, dt - \sum_2 \int u \, dt \cdot \sum_1 \int i \, dt}{\sum_2 \int i \, dt \cdot \sum_1 \int di - \sum_1 \int i \, dt \cdot \sum_2 \int di} \quad (3.2.47)$$

El efecto de filtrado por medio de la suma de integrales en diferentes intervalos de tiempo elimina la 2da y 3er armónica.

$$\sum \int x \, dx = \int_0^{2\frac{\pi}{3}} x \, dx + \int_{\frac{\pi}{2}}^{2\frac{\pi}{3} + \frac{\pi}{2}} x \, dx \quad (3.2.48)$$

Los ángulos dados como límites de integración están referidos a la oscilación fundamental de 50Hz. Ambas integrales superan un área de ángulos de $2/3\pi$. Este es el período de la componente de mayor frecuencia (3er armónico en la fig. 3.2.43b). Con ello no se origina para la suma, ninguna componente a partir de la 3er armónica. El desplazamiento de una integral respecto a la otra se selecciona en $\pi/2$ de tal forma que se originan una 2da armónica en ambas integrales F1 y F2 las cuales se compensan y con ello se elimina la segunda armónica.

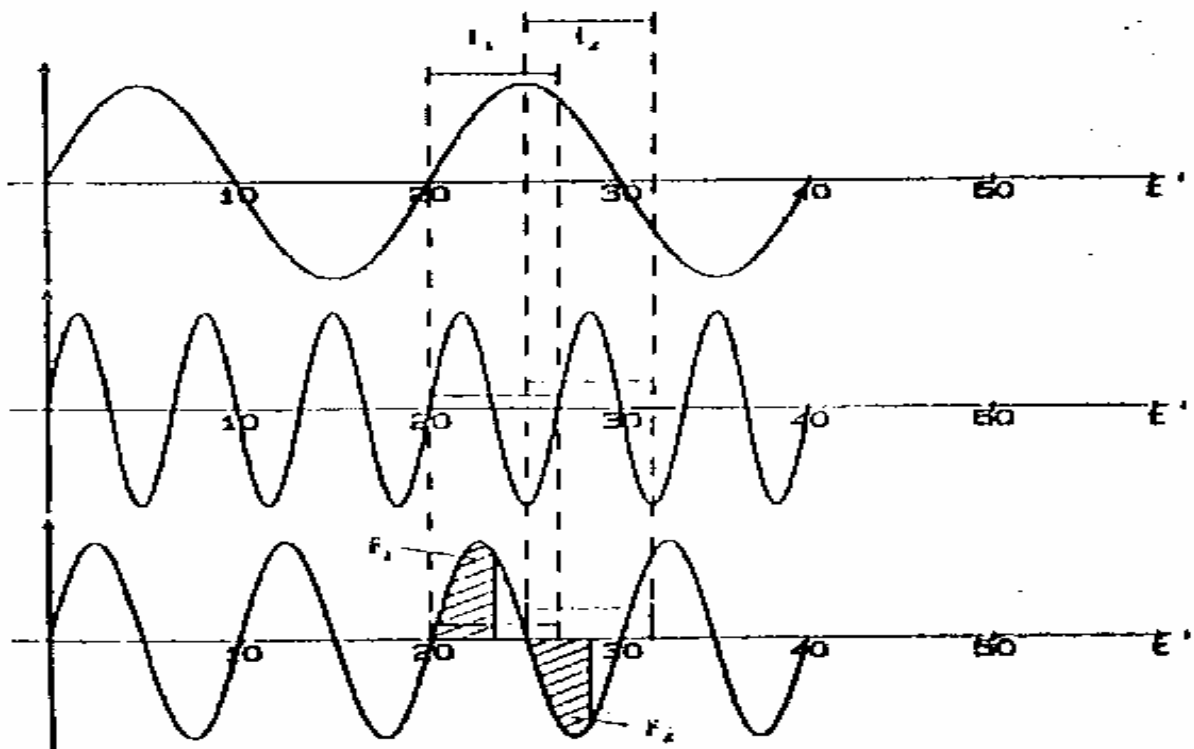


Fig. 3.2.43 Supresión de las componentes de armónicos superiores por medio del desplazamiento de los intervalos de integración

- a) Oscilación fundamental
- b) 3er armónico
- c) 2do armónico

Para juzgar la capacidad de actuación de este algoritmo, a continuación se selecciona la suma de integrales de tal forma que se supriman la 3., 5. y 7. armónicas. Para esta suma se da:

$$\sum \int x \, dx = \int_0^{2\frac{\pi}{7}} x \, dx + \int_{\frac{\pi}{3}}^{2\frac{\pi}{7} + \frac{\pi}{3}} x \, dx + \int_{\frac{\pi}{5}}^{2\frac{\pi}{7} + \frac{\pi}{5}} x \, dx + \int_{\frac{\pi}{3} + \frac{\pi}{5}}^{2\frac{\pi}{7} + \frac{\pi}{3} + \frac{\pi}{5}} x \, dx \quad (3.2.49)$$

Hay que tener en cuenta que un número mayor de componentes de oscilaciones de orden superior tenidas en cuenta, conduce a una ventana de datos más amplia. Además debe aplicarse, correspondientemente a los diferentes intervalos de integración, una división de ángulos más fina y con ello una frecuencia de muestreo más elevada. Para tener en cuenta las 3., 5. y 7. armónicas se da una frecuencia de muestreo de $h=60$ valores por período con un ancho de ventana de 13 ms.

Ya que debido a la alta frecuencia de muestreo, la integración numérica es muy precisa, luego no aparecen errores de cálculo para señales senoidales puras.

3.2.6 Algoritmo de Smolinski

A diferencia de los algoritmos tratados hasta aquí que solucionan la ec. diferencial de primer orden que parten del modelo de línea de primer orden, este algoritmo tiene en cuenta además las capacidades de la línea en su propuesta [20]. Para ello se modela la línea por medio de un equivalente Π . En el caso de cortocircuito se puede luego describir a la línea con una ecuación diferencial de 2do orden. Para la interrelación entre la corriente y tensión en el lugar de instalación del relé vale la siguiente ecuación:

$$u(t) = R_L i(t) + L_L \frac{di}{dt} - R_L C_L \frac{du}{dt} - L_L C_L \frac{d^2 u}{dt^2} \quad (3.2.50)$$

Si se reemplaza las diferenciales por diferencias luego resulta el valor de la tensión muestreada u_k :

$$u_k = R_L i_k + L_L \frac{i_{k+1} - i_{k-1}}{2 \Delta t} - R_L C_L \frac{u_{k+1} - u_{k-1}}{2 \Delta t} - L_L C_L \frac{\frac{u_{k+1} - u_k}{\Delta t} - \frac{u_k - u_{k-1}}{\Delta t}}{\Delta t} \quad (3.2.51)$$

Si se considera la ec. (3.2.51) en cuatro puntos de tiempos de muestreo consecutivos, luego se obtienen el siguiente sistema de cuatro ecuaciones:

$$\begin{bmatrix} u_2 \\ u_3 \\ u_4 \\ u_5 \end{bmatrix} = \begin{bmatrix} i_2 & \frac{i_3 - i_1}{2 \Delta t} & -\frac{u_3 - u_1}{2 \Delta t} & -\frac{u_3 - 2u_2 + u_1}{\Delta t^2} \\ i_3 & \frac{i_4 - i_2}{2 \Delta t} & -\frac{u_4 - u_2}{2 \Delta t} & -\frac{u_4 - 2u_3 + u_2}{\Delta t^2} \\ i_4 & \frac{i_5 - i_3}{2 \Delta t} & -\frac{u_5 - u_3}{2 \Delta t} & -\frac{u_5 - 2u_4 + u_3}{\Delta t^2} \\ i_5 & \frac{i_6 - i_4}{2 \Delta t} & -\frac{u_6 - u_4}{2 \Delta t} & -\frac{u_6 - 2u_5 + u_4}{\Delta t^2} \end{bmatrix} \cdot \begin{bmatrix} R_L \\ L_L \\ C_L R_L \\ C_L L_L \end{bmatrix} \quad (3.2.52)$$

Aquí es de interés solamente las magnitudes R_L y L_L . Los productos $C_L R_L$ y $C_L L_L$ no se calculan. A raíz de la relación lineal entre R_L y $C_L R_L$ así como L_L y $C_L L_L$ se puede simplificar la solución del sistema de ecuaciones (3.2.52) por medio de la introducción de sub-matrices.

$$\begin{bmatrix} U_1 \\ U_2 \end{bmatrix} = \begin{bmatrix} A & B \\ C & D \end{bmatrix} \cdot \begin{bmatrix} P \\ C_L P \end{bmatrix} \quad (3.2.53)$$

$$\text{mit } U_1 = \begin{bmatrix} u_2 \\ u_3 \end{bmatrix}; \quad U_2 = \begin{bmatrix} u_4 \\ u_5 \end{bmatrix}; \quad P = \begin{bmatrix} R_L \\ L_L \end{bmatrix};$$

$$A = \begin{bmatrix} i_2 & \frac{i_3 - i_1}{2 \Delta t} \\ i_3 & \frac{i_4 - i_2}{2 \Delta t} \end{bmatrix}; \quad B = \begin{bmatrix} -\frac{u_3 - u_1}{2 \Delta t} & -\frac{u_3 - 2u_2 + u_1}{2} \\ -\frac{u_4 - u_2}{2 \Delta t} & -\frac{u_4 - 2u_3 + u_2}{\Delta t^2} \end{bmatrix};$$

$$C = \begin{bmatrix} i_4 & \frac{i_5 - i_3}{2 \Delta t} \\ i_5 & \frac{i_6 - i_4}{2 \Delta t} \end{bmatrix}; \quad D = \begin{bmatrix} -\frac{u_5 - u_3}{2 \Delta t} & -\frac{u_5 - 2u_4 + u_3}{\Delta t^2} \\ -\frac{u_6 - u_4}{2 \Delta t} & -\frac{u_6 - 2u_5 + u_4}{\Delta t^2} \end{bmatrix};$$

A partir de los coeficientes calculados por medio de mediciones A, B, C, D, luego se puede eliminar primero la capacidad de la línea CL en la ec. 3.2.50, y luego se puede calcular el vector deseado P.

$$U_1 = A \cdot P + C_L \cdot B \cdot P \quad (3.2.54)$$

$$U_2 = C \cdot P + C_L \cdot D \cdot P \quad (3.2.55)$$

La ec. 3.2.55 se multiplica a la izquierda con $B \cdot D^{-1}$. Aquí hay tener en cuenta que CL es un escalar.

$$B \cdot D^{-1} \cdot U_2 = B \cdot D^{-1} \cdot (C \cdot P + C_L \cdot D \cdot P) \quad (3.2.56)$$

Restando (3.2.56) y (3.2.54) se obtiene:

$$U_1 - B \cdot D^{-1} \cdot U_2 = (A - B \cdot D^{-1} \cdot C) \cdot P \quad (3.2.57)$$

$$P = \begin{bmatrix} R_L \\ L_L \end{bmatrix} = (A - B \cdot D^{-1} \cdot C)^{-1} \cdot (U_1 - B \cdot D^{-1} \cdot U_2) \quad (3.2.58)$$

De esta forma se reduce la inversión inicial necesaria de la matriz 4x4 a la inversión de dos matrices 2x2.

Para señales de entrada sinusoidales el algoritmo se torna inestable debido a la dependencia lineal del sistema de ecuaciones (3.2.52).

Para la consideración del comportamiento de la actuación para el caso de señales de entrada con perturbaciones, se aplica según /19/ una frecuencia de muestreo de $h = 60$ con un ancho de ventana de datos de solo 2 ms.

3.3 DETERMINACIÓN DE LA IMPEDANCIA POR MEDIO DE LA ESTIMACION DE LOS FASORES DE FRECUENCIA FUNDAMENTAL APLICANDO FILTROS DIGITALES

En los algoritmos de filtrado se determinan primero a partir de los valores muestreados, los fasores de 50 Hz de la corriente y tensión. A partir de la parte real e imaginaria de tales fasores se puede calcular la impedancia compleja y de esa forma la resistencia y reactancia hasta el punto de falla en la línea.

3.3.1 Algoritmo de Phadke/Ibrahim

El algoritmo de Padke/Ibrahin, /20/ calcula los fasores complejos de 50Hz por medio del análisis de Fourier a partir de los valores muestreados en un período completo. Para la parte real e imaginaria de la tensión vale:

$$\text{Re}\{\underline{U}\} = \frac{2}{n} \sum_{k=0}^{n-1} u_k \sin \frac{k 2\pi}{n} \quad (3.3.1a)$$

$$\text{Im}\{\underline{U}\} = \frac{2}{n} \sum_{k=0}^{n-1} u_k \cos \frac{k 2\pi}{n} \quad (3.3.1b)$$

Luego para la corriente:

$$\text{Re}\{\underline{I}\} = \frac{2}{n} \sum_{k=0}^{n-1} i_k \sin \frac{k 2\pi}{n} \quad (3.3.2a)$$

$$\text{Im}\{\underline{I}\} = \frac{2}{n} \sum_{k=0}^{n-1} i_k \cos \frac{k 2\pi}{n} \quad (3.3.2b)$$

Aquí se pueden calcular los términos senos y cosenos off-line. Si se forma el cociente entre U e I complejos luego resulta la impedancia compleja Z:

$$\underline{Z} = \frac{\underline{U}}{\underline{I}} = \frac{\text{Re}\{\underline{U}\} \cdot \text{Re}\{\underline{I}\} + \text{Im}\{\underline{U}\} \cdot \text{Im}\{\underline{I}\}}{(\text{Re}\{\underline{I}\})^2 + (\text{Im}\{\underline{I}\})^2} + j \cdot \frac{\text{Im}\{\underline{U}\} \cdot \text{Re}\{\underline{I}\} - \text{Re}\{\underline{U}\} \cdot \text{Im}\{\underline{I}\}}{(\text{Re}\{\underline{I}\})^2 + (\text{Im}\{\underline{I}\})^2} \quad (3.3.3)$$

Luego para la resistencia R y la reactancia X vale:

$$R = \frac{\text{Re}\{\underline{U}\} \cdot \text{Re}\{\underline{I}\} + \text{Im}\{\underline{U}\} \cdot \text{Im}\{\underline{I}\}}{(\text{Re}\{\underline{I}\})^2 + (\text{Im}\{\underline{I}\})^2} \quad (3.3.4)$$

$$X = \frac{\text{Im}\{\underline{U}\} \cdot \text{Re}\{\underline{I}\} - \text{Re}\{\underline{U}\} \cdot \text{Im}\{\underline{I}\}}{(\text{Re}\{\underline{I}\})^2 + (\text{Im}\{\underline{I}\})^2} \quad (3.3.5)$$

Como frecuencia de muestreo es adecuado para este algoritmo $h = 20$ valores por período. El ancho de ventana de datos está en alrededor de $b = 20$ ms, es decir un período. Estos valores pueden variar según el comportamiento dinámico de las señales de tensión y corriente.

3.3.2 Algoritmo de Slemon/Robertson

La forma de calcular del algoritmo de Slemon/Robertson [21] es prácticamente igual al de del de Phadke/Ibrahim. También para este algoritmo se calcula primero los fasores de la fundamental de corriente y tensión por medio del análisis discreto de Fourier. Aquí sin embargo se calcula el módulo y el ángulo de la impedancia.

$$|\underline{Z}| = \sqrt{\frac{(\operatorname{Re}\{\underline{U}\})^2 + (\operatorname{Im}\{\underline{U}\})^2}{(\operatorname{Re}\{\underline{I}\})^2 + (\operatorname{Im}\{\underline{I}\})^2}} \quad (3.3.6)$$

$$\varphi = \arctan\left(\frac{\operatorname{Im}\{\underline{U}\}}{\operatorname{Re}\{\underline{U}\}}\right) - \arctan\left(\frac{\operatorname{Im}\{\underline{I}\}}{\operatorname{Re}\{\underline{I}\}}\right) \quad (3.3.7)$$

Luego se obtiene:

$$R = |\underline{Z}| \cdot \cos \varphi \quad (3.3.8)$$

$$X = |\underline{Z}| \cdot \sin \varphi \quad (3.3.9)$$

De [21] se recomienda también para este algoritmo una frecuencia de muestreo de $h = 20$ por período y una ventana de datos de $b = 20$ ms.

El comportamiento del algoritmo es idéntico al de Phadke/Ibrahim. Una desventaja es la aplicación de las funciones $\sin(x)$, $\cos(x)$, $\arctan(x)$ y $\sqrt{x}$ que traen consigo un mayor esfuerzo de cálculo en comparación de las sumas y multiplicaciones simples del algoritmo de Phadke/Ibrahim, y por ello se prefiere este.

3.3.3 Algoritmo de Horton

Este algoritmo [22] aplica, para la determinación de la impedancia, el análisis de Walsh, el cual muestra parecidos con el análisis de Fourier introducidos por el de Phadke /Ibrahim y Slemon/Robertson. En este procedimiento, para la determinación de los fasores complejos de la fundamental de la tensión y corriente, se forma las señales de entrada por medio de una suma de funciones rectángulo.

Para la serie de Fourier vale:

$$i(t) = \frac{A_0}{2} + \sum_{v=1}^{\infty} (A_v \cos(v\omega_0 t) + B_v \sin(v\omega_0 t)) \quad (3.3.10)$$

$$A_v = \frac{2}{T} \int_0^T i(t) \cos(v\omega_0 t) dt \quad (3.3.11)$$

$$B_v = \frac{2}{T} \int_0^T i(t) \sin(v\omega_0 t) dt \quad (3.3.12)$$

Con esto se describe cualquiera de las funciones periódicas $i(t)$ con período T por medio de una serie infinita de coeficientes A_v y B_v . Para el cálculo de la función i en el lugar t hay que multiplicar esos coeficientes con los factores $\cos(v\omega_0 t)$ y $\sin(v\omega_0 t)$ específicos para ese desarrollo en serie y finalmente sumarlos.

Al contrario que con las funciones senos y cosenos, las funciones de Walsh representan trayectorias rectangulares (figura 3.3.10). Las mismas se pueden derivar de las funciones de Rademacher /23/.

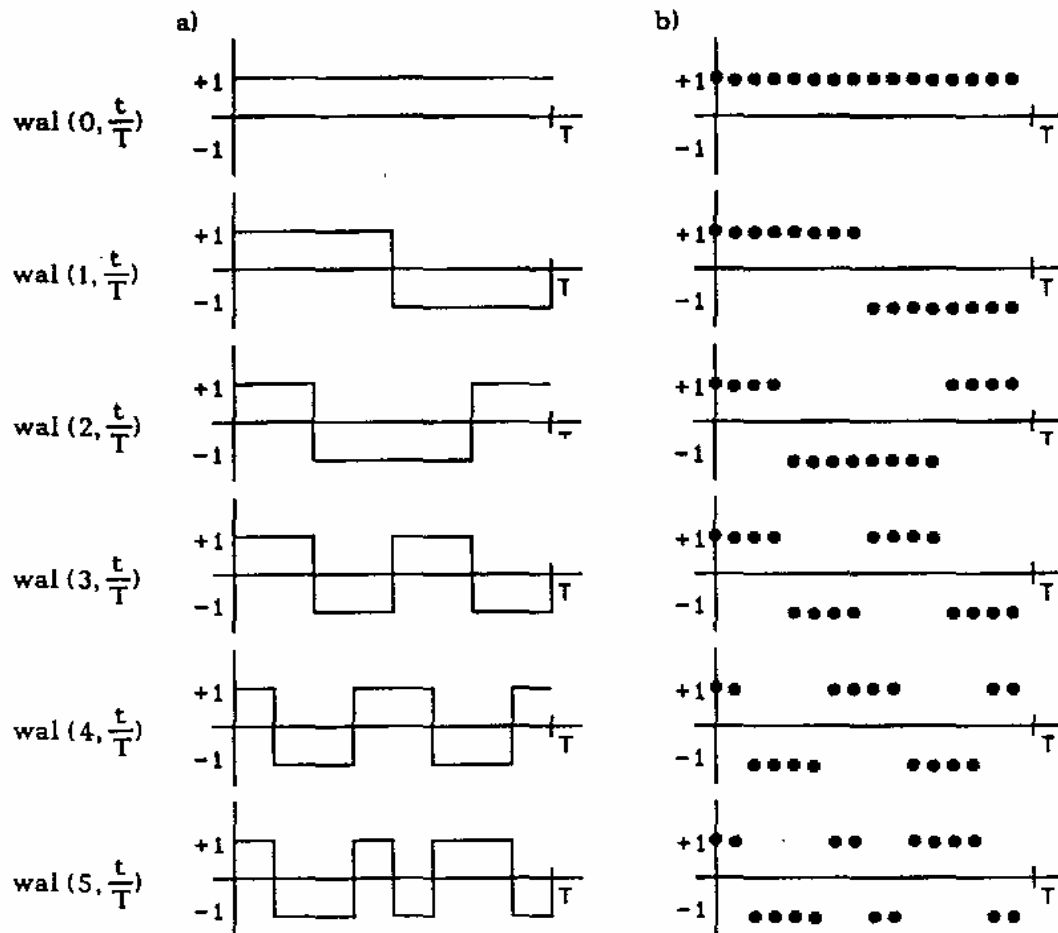


Fig. 3.3.10 a) Funciones Walsh
b) Walsh determinantes

Si se ponderan esas funciones rectangulares con factores adecuados W_μ , luego se puede describir cualquier función periódica con período T , en forma análoga que para la serie de Fourier.

$$i(t) = \sum_{\mu=1}^{\infty} W_{\mu} \text{wal}(\mu, \frac{t}{T}) \quad (3.3.13)$$

$$W_{\mu} = \frac{1}{T} \int_0^T i(t) \text{wal}(\mu, \frac{t}{T}) dt \quad (3.3.14)$$

Las funciones de Walsh se pueden determinar en forma simple por medio de los valores de muestreo /24/. Para ello hay que seleccionar la frecuencia de muestreo correspondiente a la frecuencia rectangular más elevada. Las funciones en la fig. 3.3.10 a transcurren en una secuencia de ± 1 (fig. 3.3.10b). Las ecs. (3.3.13) y (3.3.14) permiten la determinación de los valores de corriente i en determinados instantes de muestreo. En la precisión de ese valor de corriente i entra la frecuencia de muestreo por medio del valor W_{μ} .

Los coeficientes de W_{μ} permiten también la determinación de los coeficientes A_1 y B_1 de la corriente $i(t)$, que corresponden a la parte real e imaginaria del fasor complejo de la fundamental.

Para una frecuencia de muestreo de $h = 16$ vale según /23/:

$$A_1 = 0.9 \cdot W_2 + 0.373 \cdot W_6 - 0.074 \cdot W_{10} \quad (3.3.15)$$

$$B_1 = 0.9 \cdot W_1 - 0.373 \cdot W_5 - 0.074 \cdot W_9 \quad (3.3.16)$$

A partir de aquí se puede determinar la impedancia de la línea en forma correspondiente con los procedimientos descriptos en los apartados 3.3.1 y 3.3.2.

La ventaja de este procedimiento está en que las multiplicaciones $A \nu * \cos(\nu \omega t)$ se suprimen y hay que realizar sumas a la reproducción de la señal. Al respecto de dan algunas multiplicaciones adicionales de las ecuaciones (3.3.15) y (3.3.16). Además la precisión de los resultados es afectada por el análisis de los rectángulo. En principio no hay que esperar resultados distintos que con los algoritmos de Phadke/Ibrahim y Slemon/Robertson (apartados 3.3.1 y 3.3.2). Por ese motivo no vale la pena mayores investigaciones.

El algoritmo de /25/ trabaja también con funciones rectangulares. Dado el parecido con el procedimiento de Horton no se realizarán otros tratamientos.

3.3.4 Algoritmo de Sachdev/Baribeau

También en el algoritmo de Sachdev/Baribeau /26/ tiene lugar la determinación de la impedancia por medio de la parte real e imaginaria de la corriente y tensión. Para ello se

aplican filtros de regresión para la determinación de los fasores de la fundamental. Punto de partida de esos filtros es la descripción de la corriente de acuerdo a la siguiente ecuación:

$$i(t) = \hat{i}_g e^{-t/T_g} + \sum_{v=1}^n \hat{i}_v \sin(v\omega t + \varphi_v) \quad (3.3.17)$$

La función exponencial se aproxima por medio de una serie de Taylor con los dos primeros términos. Si se tiene en cuenta solo el tercer armónico, luego resulta la siguiente expresión:

$$\begin{aligned} i(t) &= \hat{i}_g \left(1 - \frac{t}{T}\right) + \hat{i}_1 \sin(\omega t + \varphi_1) + \hat{i}_3 \sin(3\omega t + \varphi_3) \\ &= \hat{i}_g \left(1 - \frac{t}{T}\right) + \hat{i}_1 \cos(\varphi_1) \sin(\omega t) + \hat{i}_1 \sin(\varphi_1) \cos(\omega t) \\ &\quad + \hat{i}_3 \cos(\varphi_3) \sin(3\omega t) + \hat{i}_3 \sin(\varphi_3) \cos(3\omega t) \end{aligned} \quad (3.3.18)$$

Para el valor de la tensión en el instante t_1 , se introducen solo apócopies:

$$a_{11} = 1 \quad (3.3.19a) \quad x_1 = \hat{i}_g \quad (3.3.20a)$$

$$a_{12} = \sin(\omega t_1) \quad (3.3.19b) \quad x_2 = \hat{i}_1 \cos(\varphi_1) = \text{Re}\{\underline{I}\} \quad (3.3.20b)$$

$$a_{13} = \cos(\omega t_1) \quad (3.3.19c) \quad x_3 = \hat{i}_1 \sin(\varphi_1) = \text{Im}\{\underline{I}\} \quad (3.3.20c)$$

$$a_{14} = \sin(3\omega t_1) \quad (3.3.19d) \quad x_4 = \hat{i}_3 \cos(\varphi_3) \quad (3.3.20d)$$

$$a_{15} = \cos(3\omega t_1) \quad (3.3.19e) \quad x_5 = \hat{i}_3 \sin(\varphi_3) \quad (3.3.20e)$$

$$a_{16} = t_1 \quad (3.3.19f) \quad x_6 = -\hat{i}_g / T_g \quad (3.3.20f)$$

Luego de la ec. (3.3.18):

$$i(t_1) = a_{11} x_1 + a_{12} x_2 + \dots + a_{16} x_6 \quad (3.3.21)$$

Si se muestrean las magnitudes de entrada con 6 pasos de muestreo, resulta luego un sistema de 6 ecuaciones con 6 incógnitas, que se puede solucionar como:

$$\mathbf{A} \cdot \mathbf{X} = \mathbf{I} \rightarrow \mathbf{X} = \mathbf{A}^{-1} \cdot \mathbf{I} \quad (3.3.22)$$

Si se aplica más de 6 valores de muestreo luego el sistema queda sobredeterminado. Para n puntos de apoyo la matriz A en la ec. 3.3.22 tiene la forma rectangular $6 * n$ y no se puede invertir. Pero es posible un cálculo de compensación de error.

$$\mathbf{X} = (\mathbf{A}^T \cdot \mathbf{A})^{-1} \cdot \mathbf{A}^T : \mathbf{I} = \mathbf{B} \cdot \mathbf{I} \quad (3.3.23)$$

Aquí B es la pseudo-inversa de la matriz A que se puede determinar solo con mucho esfuerzo computacional. Dado que la matriz A no contiene sin embargo ninguna variable, la pseudo-inversa se puede calcular fuera de línea. Además son de interés solamente en 2do y 3er elemento del vector X para la determinación de la impedancia. Que corresponden a la parte real e imaginaria del valor complejo de la fundamental de la corriente. Esos elementos se pueden calcular a partir del producto de la 2da y 3er fila de la matriz B con el vector de datos I.

La resistencia R y la reactancia X se obtienen luego en forma análoga al procedimiento de Phadke/Ibrahim (ecuaciones (3.3.4) y (3.3.5)):

$$R = \frac{\operatorname{Re}\{\underline{U}\} \cdot \operatorname{Re}\{\underline{I}\} + \operatorname{Im}\{\underline{U}\} \cdot \operatorname{Im}\{\underline{I}\}}{(\operatorname{Re}\{\underline{I}\})^2 + (\operatorname{Im}\{\underline{I}\})^2} = \frac{x_{2u} \cdot x_{2i} + x_{3u} \cdot x_{3i}}{x_{2i}^2 + x_{3i}^2} \quad (3.3.24)$$

$$X = \frac{\operatorname{Im}\{\underline{U}\} \cdot \operatorname{Re}\{\underline{I}\} - \operatorname{Re}\{\underline{U}\} \cdot \operatorname{Im}\{\underline{I}\}}{(\operatorname{Re}\{\underline{I}\})^2 + (\operatorname{Im}\{\underline{I}\})^2} = \frac{x_{3u} \cdot x_{2i} - x_{2u} \cdot x_{3i}}{x_{2i}^2 + x_{3i}^2} \quad (3.3.25)$$

Como frecuencias de muestreo se recomiendan, de /24/, para este algoritmo h = 12 valores por período. La ventana de datos se da con n = 9 valores de muestreo, de tal forma que resulta para 6 incógnitas una sobredeterminación de 3.

3.3.5 Algoritmo de Carr/Jackson

Para la determinación de los fasores complejos de la fundamental de la corriente y tensión se aplica en este algoritmo de correlación cruzada discreta /27/. Aquí se correlaciona los valores de muestreo de la tensión u_k con un seno y un coseno. La componente de la fundamental de la corriente entrega luego un señal, las restantes componentes se amortiguan más o menos bien. Para lo 4 valores de muestreo por período dado en /27/ vale:

$$z_k = \sum_{n=1}^4 u_{k+n-4} \cdot \cos\left(\frac{\pi(n+k)}{2}\right) \quad (3.3.26)$$

$$y_k = \sum_{n=1}^4 u_{k+n-4} \cdot \sin\left(\frac{\pi(n+k)}{2}\right) \quad (3.3.27)$$

$$c_k = z_k - j y_k \quad (3.3.28)$$

Los valores de la secuencia $\{z_k\}$ representan luego la componente real y los valores de $\{y_k\}$ la componente imaginaria de la señal de entrada. Como ejemplo, para el instante de muestreo $k = 0$ vale:

$$\begin{aligned}
 z_0 &= \sum_{n=1}^4 u_{n-4} \cdot \cos\left(\frac{\pi n}{2}\right) \\
 &= u_{-3} \cos\left(\frac{\pi}{2}\right) + u_{-2} \cos(\pi) + u_{-1} \cos\left(\frac{3\pi}{2}\right) + u_0 \cos(2\pi) \\
 &= -u_{-2} + u_0
 \end{aligned} \tag{3.3.29}$$

$$\begin{aligned}
 y_0 &= \sum_{n=1}^4 u_{n-4} \cdot \sin\left(\frac{\pi n}{2}\right) \\
 &= u_{-3} \sin\left(\frac{\pi}{2}\right) + u_{-2} \sin(\pi) + u_{-1} \sin\left(\frac{3\pi}{2}\right) + u_0 \sin(2\pi) \\
 &= u_{-3} - u_{-1}
 \end{aligned} \tag{3.3.30}$$

Correspondientemente vale para los instantes de muestreo $k = 1 \dots 3$:

$$\underline{u}_0 = z_0 - jy_0 = (-u_{-2} + u_0) - j(u_{-3} - u_{-1}) \tag{3.3.31a}$$

$$\underline{u}_1 = (-u_{-2} + u_0) - j(-u_{-1} + u_1) \tag{3.3.31b}$$

$$\underline{u}_2 = (u_0 - u_2) - j(-u_{-1} + u_1) \tag{3.3.31c}$$

$$\underline{u}_3 = (u_0 - u_2) - j(u_1 - u_3) \tag{3.3.31d}$$

Si no coinciden exactamente las frecuencias de la señal de entrada y de las funciones seno y coseno, muestra la señal de salida un "zumbido" superpuesto, cuya frecuencia resulta de la suma de las señales que forman parte de en la DKK. De allí se realiza adicionalmente la formación de un valor medio deslizante para dos valores de salida consecutivos. Para los

instantes de muestreo $k = -1$ y $k = 0$ vale (nuevamente tomando como ejemplo la tensión):

$$\underline{U}_0' = \frac{\underline{U}_1 + \underline{U}_2}{2} = \frac{1}{2} \left((u_{-4} - 2u_{-2} + u_0) - j(2u_{-3} - 2u_{-1}) \right) \quad (3.3.32a)$$

Entsprechend folgt f (Ir die drel darauf folgenden Werte:

$$\underline{U}_1' = \frac{1}{2} \left((-2u_{-2} + 2u_0) - j(u_{-3} - 2u_{-1} + u_1) \right) \quad (3.3.32b)$$

$$\underline{U}_2' = \frac{1}{2} \left((u_{-2} + 2u_0 - u_2) - j(-2u_{-1} - 2u_1) \right) \quad (3.3.32c)$$

$$\underline{U}_3' = \frac{1}{2} \left((2u_0 - 2u_2) - j(-u_{-1} + 2u_1 - u_3) \right) \quad (3.3.32d)$$

La determinación del fasor complejo de la fundamental de la corriente tienen lugar en forma análoga que para la tensión.

La determinación de la impedancia tienen lugar por medio de una división simple de los fasores complejos de la tensión y de la corriente. Para el instante de tiempo $k = 0$ vale:

$$\underline{Z} = \frac{\underline{U}_0'}{\underline{I}_0'} = \frac{(u_{-4} - 2u_{-2} + u_0) - j(2u_{-3} - 2u_{-1})}{(i_{-4} - 2i_{-2} + i_0) - j(2i_{-3} - 2i_{-1})} \quad (3.3.33)$$

Como todos los algoritmos de filtrado presentados hasta aquí, el algoritmo de Carr/Jackson también calcula la impedancia, para señales de entrada puramente senoidales, sin errores. Dado que resulta, para la frecuencia de muestreo extremadamente baja de $h = 4$ valores por período, una frecuencia límite de 100 Hz, no tiene sentido realizar consideraciones referidas a la sensibilidad a las oscilaciones de alta frecuencia.

3.3.6 Algoritmo de Girgis

En el algoritmo de Girgis [30] se aplican filtros de Kallman de distinto orden para el cálculo de las componentes real e imaginaria de la tensión y la corriente. Aquí se realiza la siguiente postulación para la trayectoria temporal de la corriente:

$$i(t) = x_1 \cos(\omega_n t) - x_2 \sin(\omega_n t) + x_3 \quad (3.3.64)$$

$$\text{mit } x_1 = \sqrt{2} \operatorname{Re}\{I\}$$

$$x_2 = \sqrt{2} \operatorname{Im}\{I\}$$

x_3 : Componente continua exponencialmente decreciente

Si se tiene en cuenta adicionalmente un ruido de medición, luego se obtiene para un valor de muestreo de la corriente en el "instante de tiempo $(k+1) \cdot \Delta t$ " (acuación de medición):

$$i_{k+1} = \begin{bmatrix} \cos(\omega_N(k+1)\Delta t) & -\sin(\omega_N(k+1)\Delta t) & 1 \end{bmatrix} \cdot \begin{bmatrix} x1_{k+1} \\ x2_{k+1} \\ x3_{k+1} \end{bmatrix} + v_{k+1}$$

$$i_{k+1} = C_{k+1} X_{k+1} + v_{k+1} \quad (3.3.65)$$

mit v_{k+1} : Ruido de medición

En la ecuación (3.3.65), $x1_{k+1}$ y $x2_{k+1}$ son los parámetros constantes a estimar; $x3_{k+1}$ es una función del tiempo que no se sigue evaluando. De allí resulta para la ecuación de estado:

$$\begin{bmatrix} x1_{k+1} \\ x2_{k+1} \\ x3_{k+1} \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & e^{-\beta \Delta t} \end{bmatrix} \cdot \begin{bmatrix} x1_k \\ x2_k \\ x3_k \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ w_k \end{bmatrix}$$

$$X_{k+1} = A X_k + W_k \quad (3.3.66)$$

mit w_k : Ruido del sistema

Aquí w_k es un ruido, que, adicionalmente al ruido de medición, afecta a la componente continua y del mismo modo se atenúa en forma exponencial.

Para el modelo descrito por medio de las ecuaciones (3.3.65) y (3.3.66) se puede calcular a partir de los valores de muestreo de la corriente i_k , los valores estimados $\hat{X}_k$ de las magnitudes estado X_k . Para ello se aplica un filtro de Kalman. Las deducciones en /31/ valen para el sistema:

$$X_{k+1} = A X_k + W_k \quad (3.3.67)$$

$$Y_{k+1} = C_{k+1} X_{k+1} + V_{k+1} \quad (3.3.68)$$

La ecuación de estado coincide aquí con la ec. (3.3.66), el vector Y_{k+1} es en el ejemplo anterior el escalar i_{k+1} de la ecuación (3.3.65). Para el filtro, en la representación general, resulta:

mit $\hat{P}_{k+1}$:	Valor estimado de la matriz de error
P_k :	Matriz de error
Q_{k+1} :	Matriz de covarianza del ruido del sistema
K_{k+1} :	Factor de amplificación de Kalman
C_{k+1} :	Matriz de parámetros de medición
R_{k+1} :	Matriz de covarianza del ruido de medición
X_{k+1} :	Valor estimado del vector de magnitudes de estado
Y_{k+1} :	Vector de valores medidos

Estas ecuaciones de validez general se pueden aplicar según Girgis /30/ al modelo descrito de la corriente por medio de las ecuaciones (3.3.65) y (3.3.66). La ecuación (3.3.71) que caracteriza el algoritmo de Girgis se representa en la figura 3.3.45.

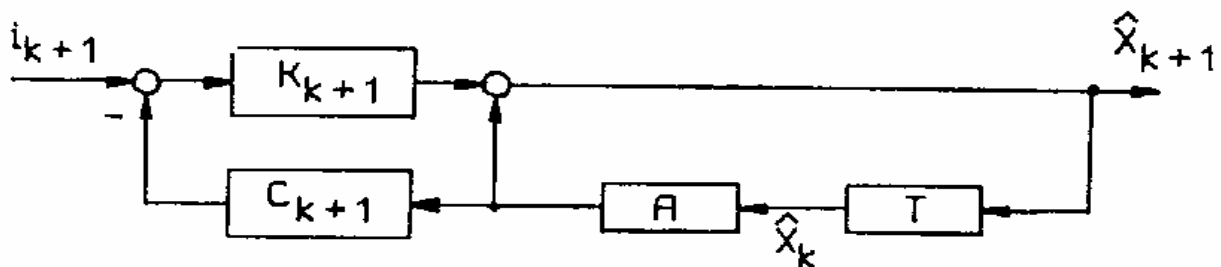


Fig. 3.3.45 Estructura del filtro de Kalman para el algoritmo de Girgis; T: Componente de tiempo muerto Δt

En la fig. 3.3.45 aparecen las siguientes magnitudes:

$$A = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & e^{-\beta \Delta t} \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0.96 \end{bmatrix}; \quad \text{Matriz del sistema}$$

mit $\beta = 1/8.5 \text{ ms}$ und $\Delta t = 0.33 \text{ ms}$;

$$C_k^T = \begin{bmatrix} \cos(\omega_N k \Delta t) \\ \sin(\omega_N k \Delta t) \\ 1 \end{bmatrix}; \quad \text{Parámetros de medición}$$

$$X_k = \begin{bmatrix} \sqrt{2} \operatorname{Re}\{I\} \\ \sqrt{2} \operatorname{Im}\{I\} \\ i_{gk} \end{bmatrix}; \quad \text{Valor estimado de las magnitudes de estado}$$

Para la explicación del filtro de Kalman se parte de una medición exacta. Luego, la salida del bloque con el vector de amplificación de Kalman K_{k+1} es cero, de tal forma que se mantiene la señal $\hat{x}_{k+1}$ sobre la componente de tiempo muerto. La matriz A es, en el caso de parámetros a estimar en forma constante, la matriz unidad. Por medio del vector C_{k+1} se calcula en cada paso de muestreo el valor estimado $\hat{i}_{k+1}$, el cual coincide con el valor medido i_{k+1} para una medición más exacta. En el caso de errores de medición se forma una diferencia, que conduce, por medio del vector de amplificación de Kalman K_{k+1} , a una corrección del valor estimado. Ese vector se obtiene a través de un algoritmo en el que, excepto en el modelo del sistema, realiza asunciones sobre el ruido en el sistema y señal de medición. Para la determinación del vector de amplificación de Kalman K_{k+1} se procede según los pasos descritos a continuación. De la ec. (3.3.70) hay que determinar un valor estimado $\hat{p}_{k+1}$ para la matriz de errores, el cual representa una medida de la precisión de los valores de medición i_k .

$$\hat{P}_{k+1} = A \hat{P}_k A^T + Q_{k+1} \quad (3.3.73)$$

Ya que en el primer paso no hay ningún valor para $\mathbf{P}_k$, debe darse como valor de partida a $\hat{\mathbf{P}}_0$. Se propone:

$$\hat{\mathbf{P}}_0 = \sigma_{w0}^2 \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 7.56 & 0 & 0 \\ 0 & 7.56 & 0 \\ 0 & 0 & 7.56 \end{bmatrix}$$

mit σ_{w0}^2 : Desviación estándar del ruido del sistema

La desviación estándar σ_{w0}^2 debe determinarse aquí por medio simulaciones y es dependiente de la distancia a la falla. En el ejemplo tratado resulta un valor de $\sigma_{w0}^2 = 7.56$.

En todos los otros pasos de muestreo se determina en forma correspondiente nuevamente la matriz de covarianza del ruido del sistema $\mathbf{Q}_k$. Los elementos no diagonales son aquí nulos dado que se incorpora ruido blanco. De otra forma resultan elementos no diagonales distintos a cero.

$$\mathbf{Q}_k = \sigma_{wk}^2 \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix}: \text{ Matriz de covarianza del ruido del sistema}$$

Las componentes de perturbación luego de la ocurrencia de un cortocircuito se representan como exponencialmente decrecientes:

$$\sigma_{wk}^2 = \frac{1}{4} \sigma_{w0}^2 e^{-k\Delta t/T_1} \quad (3.3.74)$$

$$\text{mit } T_1 = \frac{L_L}{2(R_L + R_F)} = 8.5 \text{ ms}$$

Dado que las matrices en la ec. (3.3.73) son matrices diagonales, está justificada la aplicación de $\hat{\mathbf{P}}_0$ como matriz diagonal. Junto al ruido del sistema se aplica un ruido de medición $\mathbf{v}_k$ en la ec. (3.3.70), para el que valen las mismas condiciones que para el ruido del sistema.

Ya que el vector de valores de medición se compone en el caso tratado aquí de solamente el valor escalar de muestreo de la corriente y se reduce con ello la matriz de parámetros de medición a un vector fila, se transforma también la matriz de covarianza del ruido de medición en un valor escalar:

$$R_k = \frac{1}{4} \sigma_{v0}^2 e^{-k\Delta t/T_1} \quad (3.3.75)$$

$$\text{mit } \sigma_{v0}^2 = 0.03$$

El ruido de medición y la matriz de errores entran según la ec. (3.3.70) directamente en la determinación del vector de amplificación de Kalman:

$$K_k = \hat{P}_k C_k^T (C \hat{P}_k C_k^T + R_k)^{-1} \quad (3.3.76)$$

YA que la expresión entre paréntesis en la ec. (3.3.76), a partir de los motivos ya mencionados, de la misma forma es solo un escalar, luego se reduce la inversión de matriz a una simple división.

Con ello, quedan fijos todas las magnitudes para el cálculo de la ec. de Kalman (3.3.71) representada en la fig. 3.3.45. Para el valor estimado del vector de estado se obtiene:

$$\hat{X}_{k+1} = A \hat{X}_k + K_{k+1} (I_{k+1} - C_{k+1} A \hat{X}_k) \quad (3.3.77)$$

Para la preparación del próximo paso de tiempo debe calcularse ahora la matriz actual de errores según la ec. (3.3.72).

$$P_{k+1} = \hat{F}_{k+1} - K_{k+1} C_{k+1} \hat{P}_{k+1} \quad (3.3.78)$$

Para el filtro de tensión se procede en forma análoga que para la corriente, sin embargo sin tener en cuenta la componente continua decadente en el tiempo. Con ello resulta un filtro de Kalman de 2do orden. La desviación estándar para el ruido del sistema toma un

valor de $\sigma_{w0}^2 = 16$, para el decaimiento exponencial de las componentes de perturbación se aplica la misma constante de tiempo que para la corriente.

A partir de los valores estimados de las componentes real e imaginaria de la corriente y tensión, pueden calcularse en forma análoga a la ec. (3.3.4) y (3.3.5), la resistencia R y la reactancia X:

$$R = \frac{\operatorname{Re}\{\underline{U}\} \cdot \operatorname{Re}\{\underline{I}\} + \operatorname{Im}\{\underline{U}\} \cdot \operatorname{Im}\{\underline{I}\}}{(\operatorname{Re}\{\underline{I}\})^2 + (\operatorname{Im}\{\underline{I}\})^2} \quad (3.3.79)$$

$$X = \frac{\operatorname{Im}\{\underline{U}\} \cdot \operatorname{Re}\{\underline{I}\} - \operatorname{Re}\{\underline{U}\} \cdot \operatorname{Im}\{\underline{I}\}}{(\operatorname{Re}\{\underline{I}\})^2 + (\operatorname{Im}\{\underline{I}\})^2} \quad (3.3.80)$$

La frecuencia de muestreo del algoritmo de Girgis se aplica según /30/ con $h = 60$ valores por período de la red.