

Qucs

Technical Papers

Stefan Jahn
Michael Margraf
Vincent Habchi
Raimund Jacob

Copyright © 2003, 2004, 2005, 2006, 2007 Stefan Jahn <stefan@lkcc.org>

Copyright © 2003, 2004, 2005, 2006, 2007 Michael Margraf <michael.margraf@alumni.tu-berlin.de>

Copyright © 2004, 2005 Vincent Habchi, F5RCS <10.50@free.fr>

Copyright © 2005 Raimund Jacob <raimi@lkcc.org>

Permission is granted to copy, distribute and/or modify this document under the terms of the GNU Free Documentation License, Version 1.1 or any later version published by the Free Software Foundation. A copy of the license is included in the section entitled "GNU Free Documentation License".

Contents

Chapter 1

Scattering parameters

1.1 Introduction and definition

Voltage and current are hard to measure at high frequencies. Short and open circuits (used by definitions of most n-port parameters) are hard to realize at high frequencies. Therefore, microwave engineers work with so-called scattering parameters (S parameters), that uses waves and matched terminations (normally 50Ω). This procedure also minimizes reflection problems.

A (normalized) wave is defined as ingoing wave $\underline{a}$ or outgoing wave $\underline{b}$:

$$\underline{a} = \underbrace{\frac{\underline{u} + \underline{Z}_0 \cdot \underline{i}}{2}}_{\underline{U}_{forward}} \cdot \frac{1}{\sqrt{|\operatorname{Re}\underline{Z}_0|}} \quad \underline{b} = \underbrace{\frac{\underline{u} - \underline{Z}_0^* \cdot \underline{i}}{2}}_{\underline{U}_{backward}} \cdot \frac{1}{\sqrt{|\operatorname{Re}\underline{Z}_0|}} \quad (1.1)$$

where $\underline{u}$ is (effective) voltage, $\underline{i}$ (effective) current flowing into the device and $\underline{Z}_0$ reference impedance. The waves are related to power in the following way.

$$P = (|\underline{a}|^2 - |\underline{b}|^2) \quad (1.2)$$

Sometimes waves are defined with peak voltages and peak currents. The only difference that appears then is the relation to power:

$$P = \frac{1}{2} \cdot (|\underline{a}|^2 - |\underline{b}|^2) \quad (1.3)$$

Now, characterizing an n-port is straight-forward:

$$\begin{pmatrix} \underline{b}_1 \\ \vdots \\ \underline{b}_n \end{pmatrix} = \begin{pmatrix} \underline{S}_{11} & \cdots & \underline{S}_{1n} \\ \vdots & \ddots & \vdots \\ \underline{S}_{n1} & \cdots & \underline{S}_{nn} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_1 \\ \vdots \\ \underline{a}_n \end{pmatrix} \quad (1.4)$$

One final note: The reference impedance $\underline{Z}_0$ can be arbitrary chosen. It normally is real, and there is no urgent reason to use a complex one. The definitions in equation ??, however, are made from complex impedances. These ones stem from [?], where they are named "power waves". These power waves are a useful way to define waves with complex reference impedances, but they differ from the waves introduced in the following chapter. For real reference impedances both definitions equal each other.

1.2 Waves on Transmission Lines

This section should derive the existence of the voltage and current waves on a transmission line. This way, it also proves that the definitions from the last section make sense.

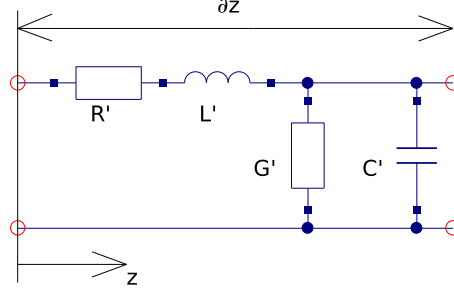


Figure 1.1: Infinite short piece of transmission line

Figure ?? shows the equivalent circuit of an infinite short piece of an arbitrary transmission line. The names of the components all carry a single quotation mark which indicates a per-length quantity. Thus, the units are ohms/m for R' , henry/m for L' , siemens/m for G' and farad/m for C' . Writing down the change of voltage and current across a piece with length ∂z results in the transmission line equations.

$$\frac{\partial u}{\partial z} = -R' \cdot i(z) - L' \cdot \frac{\partial i}{\partial t} \quad (1.5)$$

$$\frac{\partial i}{\partial z} = -G' \cdot u(z) - C' \cdot \frac{\partial u}{\partial t} \quad (1.6)$$

Transforming these equations into frequency domain leads to:

$$\frac{\partial \underline{U}}{\partial z} = -\underline{I}(z) \cdot (R' + j\omega L') \quad (1.7)$$

$$\frac{\partial \underline{I}}{\partial z} = -\underline{U}(z) \cdot (G' + j\omega C') \quad (1.8)$$

Taking equation ?? and setting it into the first derivative of equation ?? creates the wave equation:

$$\frac{\partial^2 \underline{U}}{\partial z^2} = \underline{\gamma}^2 \cdot \underline{U} \quad (1.9)$$

with $\underline{\gamma}^2 = (\alpha + j\beta)^2 = (R' + j\omega L') \cdot (G' + j\omega C')$. The complete solution of the wave equation is:

$$\underline{U}(z) = \underbrace{\underline{U}_1 \cdot \exp(-\underline{\gamma} \cdot z)}_{\underline{U}_f(z)} + \underbrace{\underline{U}_2 \cdot \exp(\underline{\gamma} \cdot z)}_{\underline{U}_b(z)} \quad (1.10)$$

As can be seen, there is a voltage wave $\underline{U}_f(z)$ traveling forward (in positive z direction) and there is a voltage wave $\underline{U}_b(z)$ traveling backwards (in negative z direction). By setting equation ?? into equation ??, it becomes clear that the current behaves in the same way:

$$\underline{I}(z) = \underbrace{\frac{\gamma}{R' + j\omega L'}}_{\underline{Y}_L} \cdot (\underline{U}_f(z) - \underline{U}_b(z)) =: \underline{I}_f(z) + \underline{I}_b(z) \quad (1.11)$$

Note that both current waves are counted positive in positive z direction. In literature, the backward flowing current wave $\underline{I}_b(z)$ is sometime counted the otherway around which would avoid the negative sign within some of the following equations.

Equation ?? introduces the characteristic admittance $\underline{Y}_L$. The propagation constant $\underline{\gamma}$ and the characteristic impedance $\underline{Z}_L$ are the two fundamental properties describing a transmission line.

$$\underline{Z}_L = \frac{1}{\underline{Y}_L} = \frac{\underline{U}_f}{\underline{I}_f} = -\frac{\underline{U}_b}{\underline{I}_b} = \sqrt{\frac{R' + j\omega L'}{G' + j\omega C'}} \approx \sqrt{\frac{L'}{C'}} \quad (1.12)$$

Note that $\underline{Z}_L$ is a real value if the line loss (due to R' and G') is small. This is often the case in reality. A further very important quantity is the reflexion coefficient $\underline{r}$ which is defined as follows:

$$\underline{r} = \frac{\underline{U}_b}{\underline{U}_f} = -\frac{\underline{I}_b}{\underline{I}_f} = \frac{\underline{Z}_e - \underline{Z}_L}{\underline{Z}_e + \underline{Z}_L} \quad (1.13)$$

The equation shows that a part of the voltage and current wave is reflected back if the end of a transmission line is not terminated by an impedance that equals $\underline{Z}_L$. The same effect occurs in the middle of a transmission line, if its characteristic impedance changes.

$\underline{U} = \underline{U}_f + \underline{U}_b$	$\underline{I} = \underline{I}_f + \underline{I}_b$
$\underline{U}_f = \frac{1}{2} \cdot (\underline{U} + \underline{I} \cdot \underline{Z}_L)$	$\underline{I}_f = \frac{1}{2} \cdot (\underline{U} / \underline{Z}_L + \underline{I})$
$\underline{U}_b = \frac{1}{2} \cdot (\underline{U} - \underline{I} \cdot \underline{Z}_L)$	$\underline{I}_b = \frac{1}{2} \cdot (\underline{I} - \underline{U} / \underline{Z}_L)$

1.3 Computing with S-parameters

1.3.1 S-parameters in CAE programs

The most common task of a simulation program is to compute the S parameters of an arbitrary network that consists of many elementary components connected to each other. To perform this, one can build a large matrix containing the S parameters of all components and then use matrix operations to solve it. However this method needs heavy algorithms. A more elegant possibility was published in [?]. Each step computes only one connection and so unites two connected components to a single S parameter block. This procedure has to be done with every connection until there is only one block left whose S parameters therefore are the simulation result.

Connecting port k of circuit ($\underline{S}$) with port l of circuit ($\underline{T}$), the new S-parameters are

$$\underline{S}'_{ij} = \underline{S}_{ij} + \frac{\underline{S}_{kj} \cdot \underline{T}_{ll} \cdot \underline{S}_{ik}}{1 - \underline{S}_{kk} \cdot \underline{T}_{ll}} \quad (1.14)$$

with i and j both being ports of ($\underline{S}$). Furthermore, it is

$$\underline{S}'_{mj} = \frac{\underline{S}_{kj} \cdot \underline{T}_{ml}}{1 - \underline{S}_{kk} \cdot \underline{T}_{ll}} \quad (1.15)$$

with m being a port of the circuit ($\underline{T}$). If two ports of the same circuit ($\underline{S}$) are connected, the new S-parameters are

$$\underline{S}'_{ij} = \underline{S}_{ij} + \frac{\underline{S}_{kj} \cdot \underline{S}_{il} \cdot (1 - \underline{S}_{lk}) + \underline{S}_{lj} \cdot \underline{S}_{ik} \cdot (1 - \underline{S}_{kl}) + \underline{S}_{kj} \cdot \underline{S}_{ll} \cdot \underline{S}_{ik} + \underline{S}_{lj} \cdot \underline{S}_{kk} \cdot \underline{S}_{il}}{(1 - \underline{S}_{kl}) \cdot (1 - \underline{S}_{lk}) - \underline{S}_{kk} \cdot \underline{S}_{ll}}. \quad (1.16)$$

If more than two ports are connected at a node, one have to insert one or more ideal tee components. Its S-parameters write as follows.

$$(\underline{S}) = \frac{1}{3} \cdot \begin{pmatrix} -1 & 2 & 2 \\ 2 & -1 & 2 \\ 2 & 2 & -1 \end{pmatrix} \quad (1.17)$$

For optimisation reasons it may be desirable to insert a cross if at least four components are connected at one node. Its S-parameters write as follows.

$$(\underline{S}) = \frac{1}{2} \cdot \begin{pmatrix} -1 & 1 & 1 & 1 \\ 1 & -1 & 1 & 1 \\ 1 & 1 & -1 & 1 \\ 1 & 1 & 1 & -1 \end{pmatrix} \quad (1.18)$$

The formulas (??), (??) and (??) were obtained using the “nontouching-loop” rule being an analytical method for solving a flow graph. A few basic definitions have to be understood.

A “path” is a series of branches into the same direction with no node touched more than once. A paths value is the product of the coefficients of the branches. A “loop” is formed when a path starts and finishes at the same node. A “first-order” loop is a path coming to closure with no node passed more than once. Its value is the product of the values of all branches encountered on the route. A “second-order” loop consists of two first-order loops not touching each other at any node. Its value is calculated as the product of the values of the two first-order loops. Third- and higher-order loops are three or more first-order loops not touching each other at any node.

The nontouching-loop rule can be applied to solve any flow graph. In the following equation in symbolic form T represents the ratio of the dependent variable in question and the independent variable.

$$T = \frac{P_1 \cdot (1 - \Sigma L_1^{(1)} + \Sigma L_2^{(1)} - \Sigma L_3^{(1)} + \dots) + P_2 \cdot (1 - \Sigma L_1^{(2)} + \Sigma L_2^{(2)} - \Sigma L_3^{(2)} + \dots) + P_3 \cdot (1 - \Sigma L_1^{(3)} + \Sigma L_2^{(3)} - \Sigma L_3^{(3)} + \dots) + P_4 \cdot (1 - \dots) + \dots}{1 - \Sigma L_1 + \Sigma L_2 - \Sigma L_3 + \dots} \quad (1.19)$$

In eq. (??) ΣL_1 stands for the sum of all first-order loops, ΣL_2 is the sum of all second-order loops, and so on. P_1, P_2, P_3 etc., stand for the values of all paths that can be found from the independent variable to the dependent variable. $\Sigma L_1^{(1)}$ denotes the sum of those first-order loops which do not touch (hence the name) the path of P_1 at any node, $\Sigma L_2^{(1)}$ denotes then the sum of those second-order loops which do not touch the path P_1 at any point, $\Sigma L_1^{(2)}$ consequently denotes the sum of those first-order loops which do not touch the path of P_2 at any point. Each path is multiplied by the factor in parentheses which involves all the loops of all orders that the path does not touch.

When connecting two different networks the signal flow graph in fig. ?? is used to compute the new S-parameters. With equally reference impedances on port k and port l the relations $\underline{a}_k = \underline{b}_l$ and $\underline{a}_l = \underline{b}_k$ are satisfied.

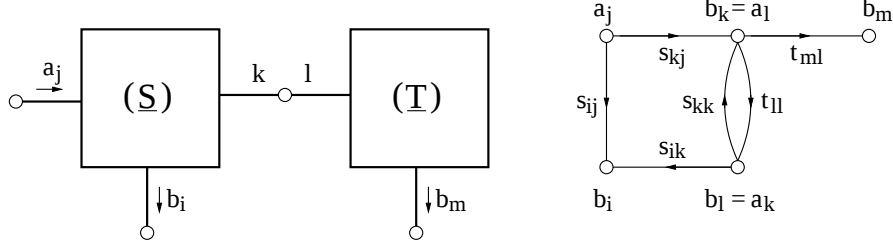


Figure 1.2: signal flow graph of a joint between ports k and l on different networks

There is only one first-order loop (see fig. ??) within this signal flow graph. This loops value yields to

$$L_{11} = \underline{S}_{kk} \cdot \underline{T}_{ll} \quad (1.20)$$

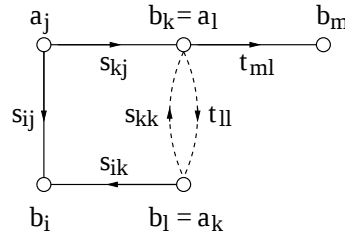


Figure 1.3: loops in the signal flow graph when connecting ports k and l on different networks

The paths that can be found from the independent variable $\underline{a}_j$ to the dependent variable $\underline{b}_i$ (as depicted in fig. ??) can be written as

$$P_1 = \underline{S}_{kj} \cdot \underline{T}_{ll} \cdot \underline{S}_{ik} \quad (1.21)$$

$$P_2 = \underline{S}_{ij} \quad (1.22)$$

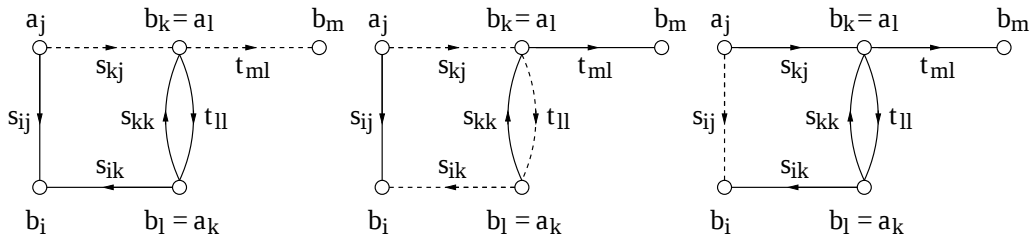


Figure 1.4: paths in the signal flow graph when connecting ports k and l on different networks

Applying the nontouching-loop rule, i.e. eq. (??), gives the new S-parameter $\underline{S}'_{ij}$

$$\begin{aligned}\underline{S}'_{ij} &= \frac{\underline{b}_i}{\underline{a}_j} = \frac{P_1 \cdot (1 - L_{11}) + P_2 \cdot 1}{1 - L_{11}} \\ &= \frac{\underline{S}_{ij} \cdot (1 - \underline{S}_{kk} \cdot \underline{T}_{ll}) + \underline{S}_{kj} \cdot \underline{T}_{ll} \cdot \underline{S}_{ik}}{1 - \underline{S}_{kk} \cdot \underline{T}_{ll}} = \underline{S}_{ij} + \frac{\underline{S}_{kj} \cdot \underline{T}_{ll} \cdot \underline{S}_{ik}}{1 - \underline{S}_{kk} \cdot \underline{T}_{ll}}\end{aligned}\quad (1.23)$$

The only path that can be found from the independent variable $\underline{a}_j$ to the dependent variable $\underline{b}_m$ (as depicted in fig. ??) can be written as

$$P_1 = \underline{S}_{kj} \cdot \underline{T}_{ml} \quad (1.24)$$

Thus the new S-parameter $\underline{S}'_{mj}$ yields to

$$\underline{S}'_{mj} = \frac{\underline{b}_m}{\underline{a}_j} = \frac{P_1 \cdot 1}{1 - L_{11}} = \frac{\underline{S}_{kj} \cdot \underline{T}_{ml}}{1 - \underline{S}_{kk} \cdot \underline{T}_{ll}} \quad (1.25)$$

When connecting the same network the signal flow graph in fig. ?? is used to compute the new S-parameters. With equally reference impedances on port k and port l the relations $\underline{a}_k = \underline{b}_l$ and $\underline{a}_l = \underline{b}_k$ are satisfied.

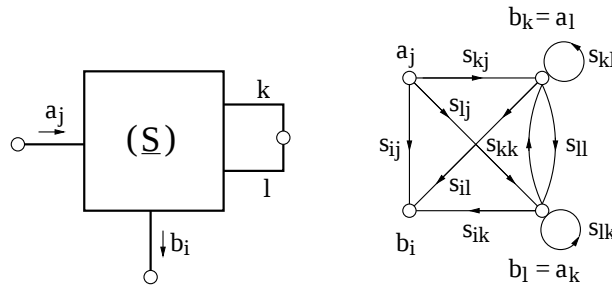


Figure 1.5: signal flow graph of a joint between ports k and l on the same network

There are three first-order loops and a second-order loop (see fig. ??) within this signal flow graph. These loops' values yield to

$$L_{11} = \underline{S}_{kk} \cdot \underline{S}_{ll} \quad (1.26)$$

$$L_{12} = \underline{S}_{kl} \quad (1.27)$$

$$L_{13} = \underline{S}_{lk} \quad (1.28)$$

$$L_{21} = L_{12} \cdot L_{13} = \underline{S}_{kl} \cdot \underline{S}_{lk} \quad (1.29)$$

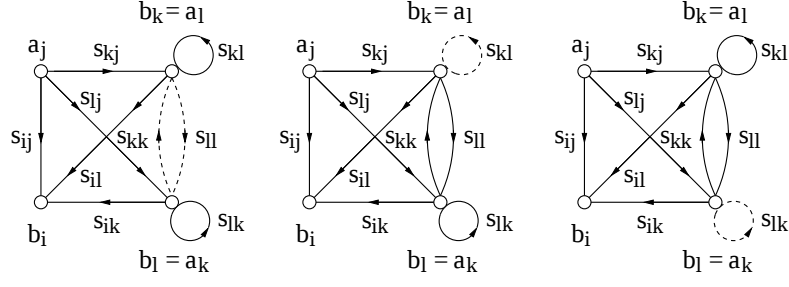


Figure 1.6: loops in the signal flow graph when connecting ports k and l on the same network

There are five different paths that can be found from the independent variable $\underline{a}_j$ to the dependent variable $\underline{b}_i$ (as depicted in fig. ??) which can be written as

$$P_1 = \underline{S}_{kj} \cdot \underline{S}_{ll} \cdot \underline{S}_{ik} \quad (1.30)$$

$$P_2 = \underline{S}_{kj} \cdot \underline{S}_{il} \quad (1.31)$$

$$P_3 = \underline{S}_{lj} \cdot \underline{S}_{ik} \quad (1.32)$$

$$P_4 = \underline{S}_{ij} \quad (1.33)$$

$$P_5 = \underline{S}_{lj} \cdot \underline{S}_{kk} \cdot \underline{S}_{il} \quad (1.34)$$

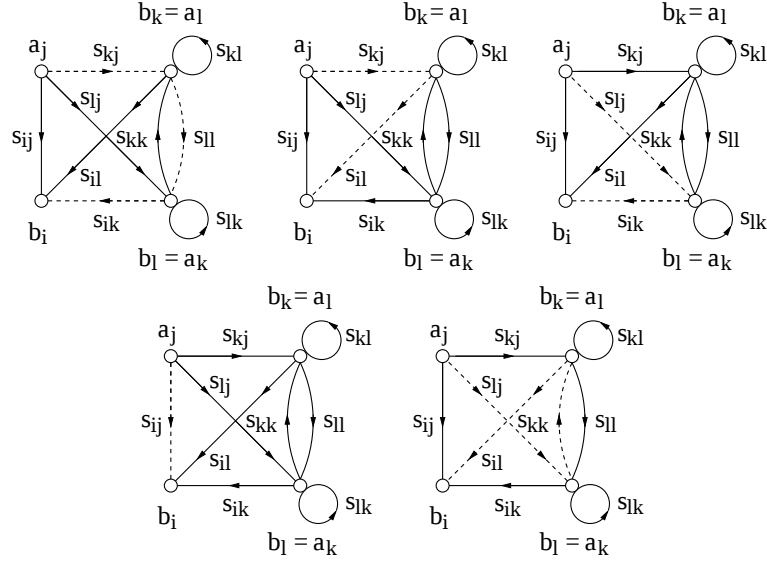


Figure 1.7: paths in the signal flow graph when connecting ports k and l on the same network

Thus the new S-parameter $\underline{S}'_{ij}$ yields to

$$\begin{aligned}
\underline{S}'_{ij} &= \frac{P_1 + P_2 \cdot (1 - L_{13}) + P_3 \cdot (1 - L_{12}) + P_4 \cdot (1 - (L_{11} + L_{12} + L_{13}) + L_{21}) + P_5}{1 - (L_{11} + L_{12} + L_{13}) + L_{21}} \\
&= P_4 + \frac{P_1 + P_2 \cdot (1 - L_{13}) + P_3 \cdot (1 - L_{12}) + P_5}{1 - (L_{11} + L_{12} + L_{13}) + L_{21}} \\
&= \underline{S}_{ij} + \frac{\underline{S}_{kj} \cdot \underline{S}_{ll} \cdot \underline{S}_{ik} + \underline{S}_{kj} \cdot \underline{S}_{il} \cdot (1 - \underline{S}_{lk}) + \underline{S}_{lj} \cdot \underline{S}_{ik} \cdot (1 - \underline{S}_{kl}) + \underline{S}_{lj} \cdot \underline{S}_{kk} \cdot \underline{S}_{il}}{1 - (\underline{S}_{kk} \cdot \underline{S}_{ll} + \underline{S}_{kl} + \underline{S}_{lk}) + \underline{S}_{kl} \cdot \underline{S}_{lk}} \\
&= \underline{S}_{ij} + \frac{\underline{S}_{kj} \cdot \underline{S}_{ll} \cdot \underline{S}_{ik} + \underline{S}_{kj} \cdot \underline{S}_{il} \cdot (1 - \underline{S}_{lk}) + \underline{S}_{lj} \cdot \underline{S}_{ik} \cdot (1 - \underline{S}_{kl}) + \underline{S}_{lj} \cdot \underline{S}_{kk} \cdot \underline{S}_{il}}{(1 - \underline{S}_{kl}) \cdot (1 - \underline{S}_{lk}) - \underline{S}_{kk} \cdot \underline{S}_{ll}}
\end{aligned} \tag{1.35}$$

This short introduction to signal flow graphs and their solution using the nontouching-loop rule verifies the initial formulas used to compute the new S-parameters for the reduced subnetworks.

1.3.2 Differential S-parameter ports

The implemented algorithm for the S-parameter analysis calculates S-parameters in terms of the ground node. In order to allow differential S-parameters as well it is necessary to insert an ideal impedance transformer with a turns ratio of 1:1 between the differential port and the device under test.

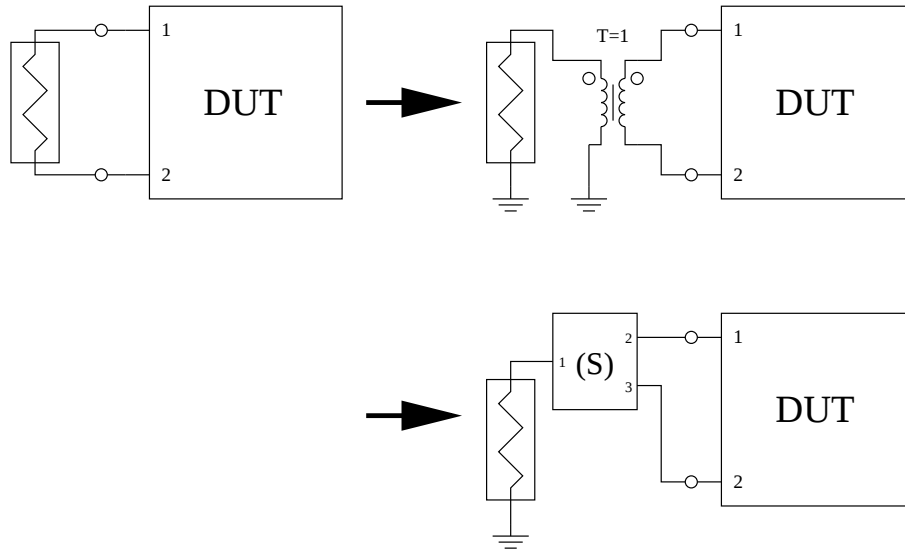


Figure 1.8: transformation of differential port into single ended port

The S-parameter matrix of the inserted ideal transformer being a three port device can be written as follows.

$$(\underline{S}) = \frac{1}{3} \cdot \begin{pmatrix} 1 & 2 & -2 \\ 2 & 1 & 2 \\ -2 & 2 & 1 \end{pmatrix} \tag{1.36}$$

This transformation can be applied to each S-parameter port in a circuit regardless whether it is actually differential or not.

It is also possible to do the impedance transformation within this step (for S-parameter ports with impedances different than 50Ω). This can be done by using a transformer with an impedance ration of

$$r = T^2 = \frac{50\Omega}{Z} \quad (1.37)$$

With Z being the S-parameter port impedance. The S-parameter matrix of the inserted ideal transformer now writes as follows.

$$(\underline{S}) = \frac{1}{2 \cdot Z_0 + Z} \cdot \begin{pmatrix} 2 \cdot Z_0 - Z & 2 \cdot \sqrt{Z_0 \cdot Z} & -2 \cdot \sqrt{Z_0 \cdot Z} \\ 2 \cdot \sqrt{Z_0 \cdot Z} & Z & 2 \cdot Z_0 \\ -2 \cdot \sqrt{Z_0 \cdot Z} & 2 \cdot Z_0 & Z \end{pmatrix} \quad (1.38)$$

With Z being the new S-parameter port impedance and Z_0 being 50Ω .

1.4 Applications

1.4.1 Stability

A very important task in microwave design (especially for amplifiers) is the question, whether the circuit tends to unwanted oscillations. A two-port oscillates if, despite of no signal being fed into it, AC power issues from at least one of its ports. This condition can be easily expressed in terms of RF quantities, so a circuit is stable if:

$$|\underline{r}_1| < 1 \quad \text{and} \quad |\underline{r}_2| < 1 \quad (1.39)$$

with $\underline{r}_1$ being reflexion coefficient of port 1 and $\underline{r}_2$ the one of port 2.

A further question can be asked: What conditions must be fulfilled to have a two-port be stable for all combinations of passive impedance terminations at port 1 and port 2? Such a circuit is called unconditionally stable. [?] is one of the best discussions dealing with this subject.

A circuit is unconditionally stable if the following two relations hold:

$$K = \frac{1 - |S_{11}|^2 - |S_{22}|^2 + |\Delta|^2}{2 \cdot |S_{12} \cdot S_{21}|} > 1 \quad (1.40)$$

$$|\Delta| = |S_{11} \cdot S_{22} - S_{12} \cdot S_{21}| < 1 \quad (1.41)$$

with Δ being the determinant of the S parameter matrix of the two port. K is called Rollet stability factor. Two relations must be fulfilled to have a necessary and sufficient criterion.

A more practical criterion (necessary and sufficient) for unconditional stability is obtained with the μ -factor:

$$\mu = \frac{1 - |S_{11}|^2}{|S_{22} - S_{11}^* \cdot \Delta| + |S_{12} \cdot S_{21}|} > 1 \quad (1.42)$$

Because of symmetry reasons, a second stability factor must exist that also gives a necessary and sufficient criterion for unconditional stability:

$$\mu' = \frac{1 - |S_{22}|^2}{|S_{11} - S_{22}^* \cdot \Delta| + |S_{12} \cdot S_{21}|} > 1 \quad (1.43)$$

For conditional stable two-ports it is interesting which load and which source impedance may cause instability. This can be seen using stability circles [?]. A disadvantage of this method is that the radius of the below-mentioned circles can become infinity. (A circle with infinite radius is a line.)

Within the reflexion coefficient plane of the load (r_L -plane), the stability circle is:

$$\underline{r}_{center} = \frac{S_{22}^* - S_{11} \cdot \Delta^*}{|S_{22}|^2 - |\Delta|^2} \quad (1.44)$$

$$\text{Radius} = \frac{|S_{12}| \cdot |S_{21}|}{|S_{22}|^2 - |\Delta|^2} \quad (1.45)$$

If the center of the r_L -plane lies within this circle and $|S_{11}| \leq 1$ then the circuit is stable for all reflexion coefficients inside the circle. If the center of the r_L -plane lies outside the circle and $|S_{11}| \leq 1$ then the circuit is stable for all reflexion coefficients outside the circle.

Very similar is the situation for reflexion coefficients in the source plane (r_S -plane). The stability circle is:

$$\underline{r}_{center} = \frac{S_{11}^* - S_{22} \cdot \Delta^*}{|S_{11}|^2 - |\Delta|^2} \quad (1.46)$$

$$\text{Radius} = \frac{|S_{12}| \cdot |S_{21}|}{|S_{11}|^2 - |\Delta|^2} \quad (1.47)$$

If the center of the r_S -plane lies within this circle and $|S_{22}| \leq 1$ then the circuit is stable for all reflexion coefficients inside the circle. If the center of the r_S -plane lies outside the circle and $|S_{22}| \leq 1$ then the circuit is stable for all reflexion coefficients outside the circle.

1.4.2 Gain

Maximum available and stable power gain (only for unconditional stable 2-ports) [?]:

$$G_{max} = \left| \frac{S_{21}}{S_{12}} \right| \cdot \left(K - \sqrt{K^2 - 1} \right) \quad (1.48)$$

where K is Rollet stability factor.

The (bilateral) transmission power gain of a two-port can be split into three parts [?]:

$$G = G_S \cdot G_0 \cdot G_L \quad (1.49)$$

with

$$G_S = \frac{(1 - |r_S|^2) \cdot (1 - |r_1|^2)}{|1 - r_S \cdot r_1|^2} \quad (1.50)$$

$$G_0 = |S_{21}|^2 \quad (1.51)$$

$$G_L = \frac{1 - |r_L|^2}{|1 - r_L \cdot S_{22}|^2 \cdot (1 - |r_1|^2)} \quad (1.52)$$

where r_1 is reflexion coefficient of the two-port input.

The curves of constant gain are circles in the reflexion coefficient plane. The circle for the load-mismatched two-port with gain G_L is

$$r_{center} = \frac{(S_{22}^* - S_{11} \cdot \Delta^*) \cdot G_L}{G_L \cdot (|S_{22}|^2 - |\Delta|^2) + 1} \quad (1.53)$$

$$\text{Radius} = \frac{\sqrt{1 - G_L \cdot (1 - |S_{11}|^2 - |S_{22}|^2 + |\Delta|^2) + G_L^2 \cdot |S_{12} \cdot S_{21}|^2}}{G_L \cdot (|S_{22}|^2 - |\Delta|^2) + 1} \quad (1.54)$$

The circle for the source-mismatched two-port with gain G_S is

$$r_{center} = \frac{G_S \cdot r_1^*}{1 - |r_1|^2 \cdot (1 - G_S)} \quad (1.55)$$

$$\text{Radius} = \frac{\sqrt{1 - G_S} \cdot (1 - |r_1|^2)}{1 - |r_1|^2 \cdot (1 - G_S)} \quad (1.56)$$

with

$$r_1 = S_{11} + \frac{S_{12} \cdot S_{21} \cdot r_L}{1 - r_L \cdot S_{22}} \quad (1.57)$$

The available power gain G_A of a two-port is reached when the load is conjugately matched to the output port. It is:

$$G_A = \frac{|S_{21}|^2 \cdot (1 - |r_S|^2)}{|1 - S_{11} \cdot r_S|^2 - |S_{22} - \Delta \cdot r_S|^2} \quad (1.58)$$

with $\Delta = S_{11}S_{22} - S_{12}S_{21}$. The curves with constant gain G_A are circles in the source reflexion coefficient plane (r_S -plane). The center $r_{S,c}$ and the radius R_S are:

$$r_{S,c} = \frac{g_A \cdot C_1^*}{1 + g_A \cdot (|S_{11}|^2 - |\Delta|^2)} \quad (1.59)$$

$$R_S = \frac{\sqrt{1 - 2 \cdot K \cdot g_A \cdot |S_{12}S_{21}| + g_A^2 \cdot |S_{12}S_{21}|^2}}{|1 + g_A \cdot (|S_{11}|^2 - |\Delta|^2)|} \quad (1.60)$$

with $C_1 = S_{11} - S_{22}^* \cdot \Delta$, $g_A = G_A/|S_{21}|^2$ and K Rollet stability factor.

The operating power gain G_P of a two-port is the power delivered to the load divided by the input power of the amplifier. It is:

$$G_P = \frac{|S_{21}|^2 \cdot (1 - |r_L|^2)}{|1 - S_{22} \cdot r_L|^2 - |S_{11} - \Delta \cdot r_L|^2} \quad (1.61)$$

with $\Delta = S_{11}S_{22} - S_{12}S_{21}$. The curves with constant gain G_P are circles in the load reflexion coefficient plane (r_L -plane). The center $r_{L,c}$ and the radius R_L are:

$$r_{L,c} = \frac{g_P \cdot C_2^*}{1 + g_P \cdot (|S_{22}|^2 - |\Delta|^2)} \quad (1.62)$$

$$R_L = \frac{\sqrt{1 - 2 \cdot K \cdot g_P \cdot |S_{12}S_{21}| + g_P^2 \cdot |S_{12}S_{21}|^2}}{|1 + g_P \cdot (|S_{22}|^2 - |\Delta|^2)|} \quad (1.63)$$

with $C_2 = S_{22} - S_{11}^* \cdot \Delta$, $g_P = G_P/|S_{21}|^2$ and K Rollet stability factor.

1.4.3 Two-Port Matching

Obtaining concurrent power matching of input and output in a bilateral circuit is not such simple, due to the backward transmission S_{12} . However, in linear circuits, this task can be easily solved by the following equations:

$$\Delta = S_{11} \cdot S_{22} - S_{12} \cdot S_{21} \quad (1.64)$$

$$B = 1 + |S_{11}|^2 - |S_{22}|^2 - |\Delta|^2 \quad (1.65)$$

$$C = S_{11} - S_{22}^* \cdot \Delta \quad (1.66)$$

$$r_S = \frac{1}{2 \cdot C} \cdot \left(B - \sqrt{B^2 - |2 \cdot C|^2} \right) \quad (1.67)$$

Here r_S is the reflexion coefficient that the circuit needs to see at the input port in order to reach concurrently matched in- and output. For the reflexion coefficient at the output r_L the same equations hold by simply changing the indices (exchange 1 by 2 and vice versa).

Chapter 2

Noise Waves

2.1 Definition

In microwave circuits described by scattering parameters, it is advantageous to regard noise as noise waves [?]. The noise characteristics of an n -port is then defined completely by one outgoing noise wave $\underline{b}_{noise,n}$ at each port (see 2-port example in fig. ??) and the correlation between these noise sources. Therefore, mathematically, you can characterize a noisy n -port by its $n \times n$ scattering matrix ($\underline{S}$) and its $n \times n$ noise wave correlation matrix ($\underline{C}$).

$$\begin{aligned}
 (\underline{C}) &= \begin{pmatrix} \overline{\underline{b}_{noise,1} \cdot \underline{b}_{noise,1}^*} & \overline{\underline{b}_{noise,1} \cdot \underline{b}_{noise,2}^*} & \cdots & \overline{\underline{b}_{noise,1} \cdot \underline{b}_{noise,n}^*} \\ \overline{\underline{b}_{noise,2} \cdot \underline{b}_{noise,1}^*} & \overline{\underline{b}_{noise,2} \cdot \underline{b}_{noise,2}^*} & \cdots & \overline{\underline{b}_{noise,2} \cdot \underline{b}_{noise,n}^*} \\ \vdots & \vdots & \ddots & \vdots \\ \overline{\underline{b}_{noise,n} \cdot \underline{b}_{noise,1}^*} & \overline{\underline{b}_{noise,n} \cdot \underline{b}_{noise,2}^*} & \cdots & \overline{\underline{b}_{noise,n} \cdot \underline{b}_{noise,n}^*} \end{pmatrix} \\
 &= \begin{pmatrix} \underline{c}_{11} & \underline{c}_{12} & \cdots & \underline{c}_{1n} \\ \underline{c}_{21} & \underline{c}_{22} & \cdots & \underline{c}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \underline{c}_{n1} & \underline{c}_{n2} & \cdots & \underline{c}_{nn} \end{pmatrix}
 \end{aligned} \tag{2.1}$$

Where $\overline{x}$ is the time average of x and $\underline{x}^*$ is the conjugate complex of $\underline{x}$. Noise correlation matrices are hermitian matrices because the following equations hold.

$$\text{Im}(\underline{c}_{nn}) = \text{Im}(\overline{|\underline{b}_{noise,n}|^2}) = 0 \tag{2.2}$$

$$\underline{c}_{nm} = \underline{c}_{mn}^* \tag{2.3}$$

Where $\text{Im}(\underline{x})$ is the imaginary part of $\underline{x}$ and $|\underline{x}|$ is the magnitude of $\underline{x}$.

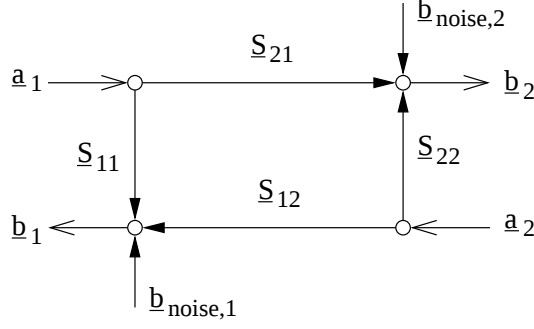


Figure 2.1: signal flow graph of a noisy 2-port

2.2 Noise Parameters

Having the noise wave correlation matrix, one can easily compute the noise parameters [?]. The following equations calculate them with regard to port 1 (input) and port 2 (output). (If one uses an n-port and want to calculate the noise parameters regarding to other ports, one has to replace the index numbers of S- and c-parameters accordingly. I.e. replace "1" with the number of the input port and "2" with the number of the output port.)

Noise figure:

$$F = 1 + \frac{\mathcal{C}_{22}}{k \cdot T_0 \cdot |\underline{S}_{21}|^2} \quad (2.4)$$

$$NF [\text{dB}] = 10 \cdot \lg F \quad (2.5)$$

Optimal source reflection coefficient (normalized according to the input port impedance):

$$\Gamma_{opt} = \eta_2 \cdot \left(1 - \sqrt{1 - \frac{1}{|\eta_2|^2}} \right) \quad (2.6)$$

With

$$\eta_1 = \mathcal{C}_{11} \cdot |\underline{S}_{21}|^2 - 2 \cdot \text{Re}(\mathcal{C}_{12} \cdot \underline{S}_{21} \cdot \underline{S}_{11}^*) + \mathcal{C}_{22} \cdot |\underline{S}_{11}|^2 \quad (2.7)$$

$$\eta_2 = \frac{1}{2} \cdot \frac{\mathcal{C}_{22} + \eta_1}{\mathcal{C}_{22} \cdot \underline{S}_{11} - \mathcal{C}_{12} \cdot \underline{S}_{21}} \quad (2.8)$$

Minimum noise figure:

$$F_{min} = 1 + \frac{\mathcal{C}_{22} - \eta_1 \cdot |\Gamma_{opt}|^2}{k \cdot T_0 \cdot |\underline{S}_{21}|^2 \cdot (1 + |\Gamma_{opt}|^2)} \quad (2.9)$$

$$NF_{min} = 10 \cdot \lg F_{min} \quad (2.10)$$

Equivalent noise resistance:

$$R_n = \frac{Z_{port,in}}{4 \cdot k \cdot T_0} \cdot \left(\mathcal{C}_{11} - 2 \cdot \text{Re} \left(\mathcal{C}_{12} \cdot \left(\frac{1 + \underline{S}_{11}}{\underline{S}_{21}} \right)^* \right) + \mathcal{C}_{22} \cdot \left| \frac{1 + \underline{S}_{11}}{\underline{S}_{21}} \right|^2 \right) \quad (2.11)$$

With $Z_{port,in}$ internal impedance of input port
 Boltzmann constant $k = 1.380658 \cdot 10^{-23}$ J/K
 standard temperature $T_0 = 290$ K

Calculating the noise wave correlation coefficients from the noise parameters is straightforward as well.

$$\underline{c}_{11} = k \cdot T_{min} \cdot (|S_{11}|^2 - 1) + K_x \cdot |1 - S_{11} \cdot \Gamma_{opt}|^2 \quad (2.12)$$

$$\underline{c}_{22} = |S_{21}|^2 \cdot (k \cdot T_{min} + K_x \cdot |\Gamma_{opt}|^2) \quad (2.13)$$

$$\underline{c}_{12} = \underline{c}_{21}^* = -S_{21}^* \cdot \Gamma_{opt}^* \cdot K_x + \frac{S_{11}}{S_{21}} \cdot \underline{c}_{22} \quad (2.14)$$

with

$$K_x = \frac{4 \cdot k \cdot T_0 \cdot R_n}{Z_0 \cdot |1 + \Gamma_{opt}|^2} \quad (2.15)$$

Once having the noise parameters, one can calculate the noise figure for every source admittance $Y_S = G_S + j \cdot B_S$, source impedance $Z_S = R_S + j \cdot X_S$, or source reflection coefficient r_S .

$$F = \frac{SNR_{in}}{SNR_{out}} = \frac{T_{equi}}{T_0} + 1 \quad (2.16)$$

$$= F_{min} + \frac{G_n}{R_S} \cdot ((R_S - R_{opt})^2 + (X_S - X_{opt})^2) \quad (2.17)$$

$$= F_{min} + \frac{G_n}{R_S} \cdot |\underline{Z}_S - \underline{Z}_{opt}|^2 \quad (2.18)$$

$$= F_{min} + \frac{R_n}{G_S} \cdot ((G_S - G_{opt})^2 + (B_S - B_{opt})^2) \quad (2.19)$$

$$= F_{min} + \frac{R_n}{G_S} \cdot |\underline{Y}_S - \underline{Y}_{opt}|^2 \quad (2.20)$$

$$= F_{min} + 4 \cdot \frac{R_n}{Z_0} \cdot \frac{|\underline{\Gamma}_{opt} - \underline{r}_S|^2}{(1 - |\underline{r}_S|^2) \cdot |1 + \underline{\Gamma}_{opt}|^2} \quad (2.21)$$

Where SNR_{in} and SNR_{out} are the signal to noise ratios at the input and output, respectively, T_{equi} is the equivalent (input) noise temperature. Note that G_n does not equal $1/R_n$.

All curves with constant noise figures are circles (in all planes, i.e. impedance, admittance and reflection coefficient). A circle in the reflection coefficient plane has the following parameters.

center point:

$$\underline{r}_{center} = \frac{\underline{\Gamma}_{opt}}{1 + N} \quad (2.22)$$

radius:

$$R = \frac{\sqrt{N^2 + N \cdot (1 - |\underline{\Gamma}_{opt}|^2)}}{1 + N} \quad (2.23)$$

with

$$N = \frac{Z_0}{4 \cdot R_n} \cdot (F - F_{min}) \cdot |1 + \Gamma_{opt}|^2 \quad (2.24)$$

2.3 Noise Wave Correlation Matrix in CAE

Due to the similar concept of S parameters and noise correlation coefficients, the CAE noise analysis can be performed quite alike the S parameter analysis (section ??). As each step uses the S parameters to calculate the noise correlation matrix, the noise analysis is best done step by step in parallel with the S parameter analysis. Performing each step is as follows: We have the noise wave correlation matrices ($(\underline{C})$, $(\underline{D})$) and the S parameter matrices ($(\underline{S})$, $(\underline{T})$) of two arbitrary circuits and want to know the correlation matrix of the special circuit resulting from connecting two circuits at one port.

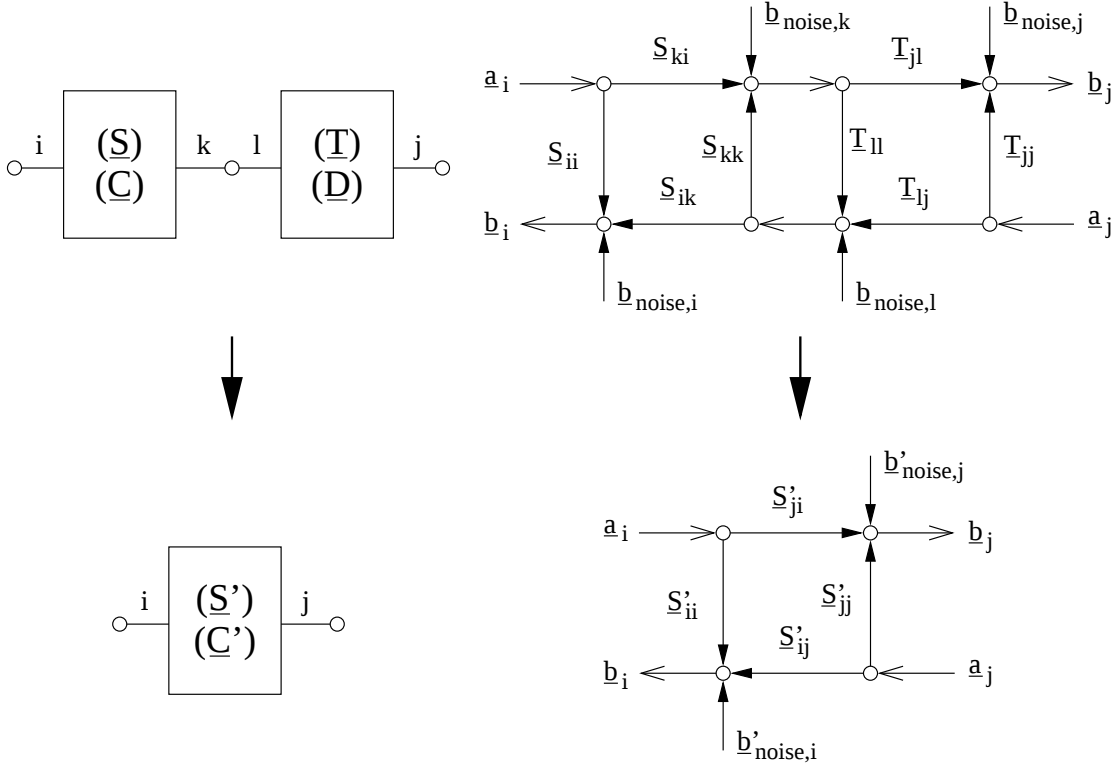


Figure 2.2: connecting two noisy circuits, scheme (left) and signal flow graph (right)

An example is shown in fig. ?? . What we have to do is to transform the inner noise waves $b_{noise,k}$ and $b_{noise,l}$ to the open ports. Let us look upon the example. According to the signal flow graph the resulting noise wave $b'_{noise,i}$ writes as follows:

$$b'_{noise,i} = b_{noise,i} + b_{noise,k} \cdot \frac{T_{il} \cdot S_{ik}}{1 - S_{kk} \cdot T_{il}} + b_{noise,l} \cdot \frac{S_{ik}}{1 - S_{kk} \cdot T_{il}} \quad (2.25)$$

The noise wave $b_{noise,j}$ does not contribute to $b'_{noise,i}$, because no path leads to port i . Calculating $b'_{noise,j}$ is quite alike:

$$b'_{noise,j} = b_{noise,j} + b_{noise,l} \cdot \frac{T_{jl} \cdot S_{kk}}{1 - S_{kk} \cdot T_{il}} + b_{noise,k} \cdot \frac{T_{jl}}{1 - S_{kk} \cdot T_{il}} \quad (2.26)$$

Now we can derive the first element of the new noise correlation matrix by multiplying eq. (??) with eq. (??).

$$\begin{aligned}
\mathcal{C}'_{ij} &= \overline{b'_{noise,i} \cdot b'^*_{noise,j}} \\
&= \overline{b_{noise,i} \cdot b^*_{noise,j}} \\
&\quad + \overline{b_{noise,i} \cdot b^*_{noise,l}} \cdot \left(\frac{T_{jl} \cdot S_{kk}}{1 - S_{kk} \cdot T_{ll}} \right)^* + \overline{b_{noise,i} \cdot b^*_{noise,k}} \cdot \left(\frac{T_{jl}}{1 - S_{kk} \cdot T_{ll}} \right)^* \\
&\quad + \overline{b_{noise,k} \cdot b^*_{noise,j}} \cdot \frac{T_{ll} \cdot S_{ik}}{1 - S_{kk} \cdot T_{ll}} \\
&\quad + \overline{b_{noise,k} \cdot b^*_{noise,l}} \cdot \frac{T_{ll} \cdot S_{ik} \cdot T_{jl}^* \cdot S_{kk}^*}{|1 - S_{kk} \cdot T_{ll}|^2} + \overline{b_{noise,k} \cdot b^*_{noise,k}} \cdot \frac{T_{ll} \cdot S_{ik} \cdot T_{jl}^*}{|1 - S_{kk} \cdot T_{ll}|^2} \\
&\quad + \overline{b_{noise,l} \cdot b^*_{noise,j}} \cdot \frac{S_{ik}}{1 - S_{kk} \cdot T_{ll}} \\
&\quad + \overline{b_{noise,l} \cdot b^*_{noise,l}} \cdot \frac{S_{ik} \cdot T_{jl}^* \cdot S_{kk}^*}{|1 - S_{kk} \cdot T_{ll}|^2} + \overline{b_{noise,l} \cdot b^*_{noise,k}} \cdot \frac{S_{ik} \cdot T_{jl}^*}{|1 - S_{kk} \cdot T_{ll}|^2}
\end{aligned} \tag{2.27}$$

The noise waves of different circuits are uncorrelated and therefore their time average product equals zero (e.g. $\overline{b_{noise,i} \cdot b^*_{noise,j}} = 0$). Thus, the final result is:

$$\begin{aligned}
\mathcal{C}'_{ij} &= (\mathcal{C}'_{ji})^* = (\mathcal{C}_{kk} \cdot T_{ll} + d_{ll} \cdot S_{kk}^*) \cdot \frac{S_{ik} \cdot T_{jl}^*}{|1 - S_{kk} \cdot T_{ll}|^2} \\
&\quad + \mathcal{C}_{ik} \cdot \left(\frac{T_{jl}}{1 - S_{kk} \cdot T_{ll}} \right)^* + d_{lj} \cdot \frac{S_{ik}}{1 - S_{kk} \cdot T_{ll}}
\end{aligned} \tag{2.28}$$

All other cases of connecting circuits can be calculated the same way using the signal flow graph. The results are listed below.

If index i and j are within the same circuit, it results in fig. ???. The following formula holds:

$$\begin{aligned}
\mathcal{C}'_{ij} &= (\mathcal{C}'_{ji})^* = \mathcal{C}_{ij} + (\mathcal{C}_{kk} \cdot |T_{ll}|^2 + d_{ll}) \cdot \frac{S_{ik} \cdot S_{jk}^*}{|1 - S_{kk} \cdot T_{ll}|^2} \\
&\quad + \mathcal{C}_{ik} \cdot \left(\frac{T_{ll} \cdot S_{jk}}{1 - S_{kk} \cdot T_{ll}} \right)^* + \mathcal{C}_{kj} \cdot \frac{T_{ll} \cdot S_{ik}}{1 - S_{kk} \cdot T_{ll}}
\end{aligned} \tag{2.29}$$

This equation is also valid, if i equals j .

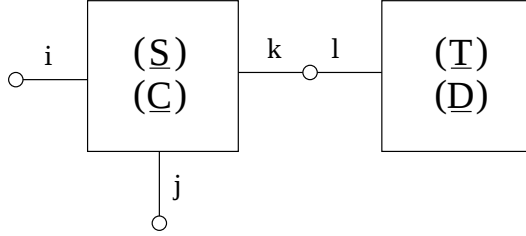


Figure 2.3: connecting two noisy circuits

If the connected ports k and l are from the same circuit, the following equations must be applied (see also fig. ??) to obtain the new correlation matrix coefficients.

$$M = (1 - \underline{S}_{kl}) \cdot (1 - \underline{S}_{lk}) - \underline{S}_{kk} \cdot \underline{S}_{ll} \quad (2.30)$$

$$K_1 = \frac{\underline{S}_{il} \cdot (1 - \underline{S}_{lk}) + \underline{S}_{il} \cdot \underline{S}_{ik}}{M} \quad (2.31)$$

$$K_2 = \frac{\underline{S}_{ik} \cdot (1 - \underline{S}_{kl}) + \underline{S}_{kk} \cdot \underline{S}_{il}}{M} \quad (2.32)$$

$$K_3 = \frac{\underline{S}_{jl} \cdot (1 - \underline{S}_{lk}) + \underline{S}_{il} \cdot \underline{S}_{jk}}{M} \quad (2.33)$$

$$K_4 = \frac{\underline{S}_{jk} \cdot (1 - \underline{S}_{kl}) + \underline{S}_{kk} \cdot \underline{S}_{jl}}{M} \quad (2.34)$$

$$\begin{aligned} \underline{c}'_{ij} = & \underline{c}_{ij} + \underline{c}_{kj} \cdot K_1 + \underline{c}_{lj} \cdot K_2 + K_3^* \cdot (\underline{c}_{ik} + \underline{c}_{kk} \cdot K_1 + \underline{c}_{lk} \cdot K_2) + \\ & K_4^* \cdot (\underline{c}_{il} + \underline{c}_{kl} \cdot K_1 + \underline{c}_{ll} \cdot K_2) \end{aligned} \quad (2.35)$$

These equations are also valid if i equals j .

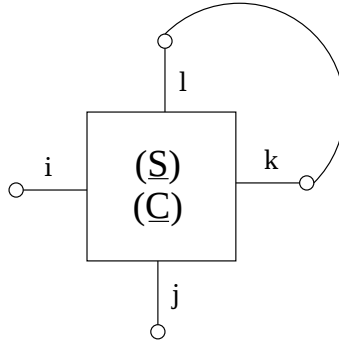


Figure 2.4: connection within a noisy circuits

The absolute values of the noise correlation coefficients are very small. To achieve a higher numerical precision, it is recommended to normalize the noise matrix with $k \cdot T_0$. After the simulation they do not have to be denormalized, because the noise parameters can be calculated by using equation (??) to (??) and omitting all occurrences of $k \cdot T_0$.

The transformer concept to deal with different port impedances and with differential ports (as described in section ??) can also be applied to this noise analysis.

2.4 Noise Correlation Matrix Transformations

The noise wave correlation matrix of a passive linear circuit generating thermal noise can simply be calculated using Bosma's theorem. The noise wave correlation matrices of active devices can be determined by forming the noise current correlation matrix and then transforming it to the equivalent noise wave correlation matrix.

The noise current correlation matrix (also called the admittance representation) $\underline{C}_Y$ is an $n \times n$ matrix.

$$\underline{C}_Y = \begin{pmatrix} \overline{i_1 \cdot i_1^*} & \overline{i_1 \cdot i_2^*} & \cdots & \overline{i_1 \cdot i_n^*} \\ \overline{i_2 \cdot i_1^*} & \overline{i_2 \cdot i_2^*} & \cdots & \overline{i_2 \cdot i_n^*} \\ \vdots & \vdots & \ddots & \vdots \\ \overline{i_n \cdot i_1^*} & \overline{i_n \cdot i_2^*} & \cdots & \overline{i_n \cdot i_n^*} \end{pmatrix} = \begin{pmatrix} \underline{C}_{11} & \underline{C}_{12} & \cdots & \underline{C}_{1n} \\ \underline{C}_{21} & \underline{C}_{22} & \cdots & \underline{C}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \underline{C}_{n1} & \underline{C}_{n2} & \cdots & \underline{C}_{nn} \end{pmatrix} \quad (2.36)$$

This definition is very likely the one made by eq. (??). The matrix has the same properties as well. Because in most transistor models the noise behaviour is expressed as the sum of effects of noise current sources it is easier to form this matrix representation.

2.4.1 Forming the noise current correlation matrix

Each element in the diagonal matrix is equal to the sum of the noise current of each element connected to the corresponding node. So the first diagonal element is the sum of noise currents connected to node 1, the second diagonal element is the sum of noise currents connected to node 2, and so on.

The off diagonal elements are the negative noise current of the element connected to the pair of corresponding node. Therefore a noise current source between nodes 1 and 2 goes into the matrix at location (1,2) and locations (2,1).

If a noise current source is grounded, it will only have contribute to one entry in the noise correlation matrix – at the appropriate location on the diagonal. If it is ungrounded it will contribute to four entries in the matrix – two diagonal entries (corresponding to the two nodes) and two off-diagonal entries.

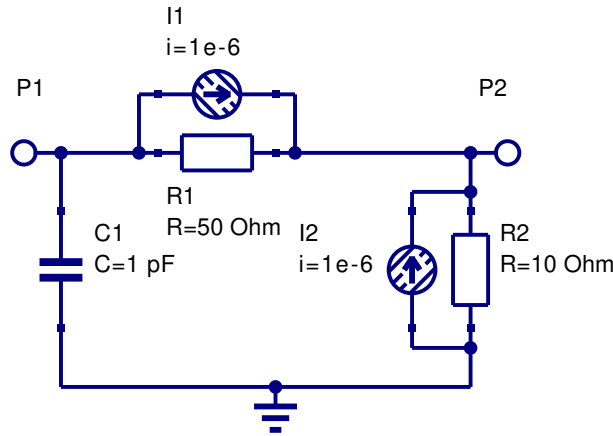


Figure 2.5: example circuit applied to noise analysis

Once having defined the spectral noise current densities of the noise currents within a transistor model the above rules for forming the $\underline{C}_Y$ matrix can be applied to the example circuit depicted in fig. ???. The noise current correlation matrix is accordingly

$$\underline{C}_Y = \begin{bmatrix} +\overline{i_1^2} & -\overline{i_1^2} \\ -\overline{i_1^2} & \overline{i_1^2} + \overline{i_2^2} \end{bmatrix} \quad (2.37)$$

2.4.2 Transformations

There are three usable noise correlation matrix representations for multiport circuits.

- admittance representation $\underline{C}_Y$ - based on noise currents
- impedance representation $\underline{C}_Z$ - based on noise voltages
- wave representation $\underline{C}_S$ - based on noise waves

According to Scott W. Wedge and David B. Rutledge [?] the transformations between these representations write as follows.

	$\underline{C}_Y$	$\underline{C}_Z$	$\underline{C}_S$
$\underline{C}_Y$	$\underline{C}_Y$	$Y \cdot \underline{C}_Z \cdot Y^+$	$(E + Y) \cdot \underline{C}_S \cdot (E + Y)^+$
$\underline{C}_Z$	$Z \cdot \underline{C}_Y \cdot Z^+$	$\underline{C}_Z$	$(E + Z) \cdot \underline{C}_S \cdot (E + Z)^+$
$\underline{C}_S$	$\frac{1}{4} (E + S) \cdot \underline{C}_Y \cdot (E + S)^+$	$\frac{1}{4} (E - S) \cdot \underline{C}_Z \cdot (E - S)^+$	$\underline{C}_S$

The signal as well as correlation matrices in impedance and admittance representations are assumed to be normalized in the above table. E denotes the identity matrix and the $^+$ operator indicates the transposed conjugate matrix (also called adjoint or adjugate).

Each noise correlation matrix transformation requires the appropriate signal matrix representation which can be obtained using the formulas given in section ?? on page ??.

2.5 Noise Wave Correlation Matrix of Components

Many components do not produce any noise. Every element of their noise correlation matrix therefore equals exactly zero. Examples are lossless, passive components, i.e. capacitors, inductors, transformers, circulators, phase shifters. Furthermore ideal voltage and current sources (without internal resistance) as well as gyrators also do not produce any noise.

If one wants to calculate the noise wave correlation matrix of a component, the most universal method is to take noise voltages and noise currents and then derive the noise waves by the use of equation (??). However, this can be very difficult.

A passive, linear circuit produces only thermal noise and thus its noise waves can be calculated with Bosma's theorem (assuming thermodynamic equilibrium).

$$(\underline{C}) = k \cdot T \cdot ((\underline{E}) - (\underline{S}) \cdot (\underline{S})^{*T}) \quad (2.38)$$

with $(\underline{S})$ being the S parameter matrix and $(\underline{E})$ identity matrix. Of course, this theorem can also be written with impedance and admittance representation of the noise correlation matrix:

$$\underline{C}_Z = 4 \cdot k \cdot T \cdot \text{Re}(\underline{Z}) \quad (2.39)$$

$$\underline{C}_Y = 4 \cdot k \cdot T \cdot \text{Re}(\underline{Y}) \quad (2.40)$$

Chapter 3

DC Analysis

3.1 Modified Nodal Analysis

Many different kinds of network element are encountered in network analysis. For circuit analysis it is necessary to formulate equations for circuits containing as many different types of network elements as possible. There are various methods for equation formulation for a circuit. These are based on three types of equations found in circuit theory:

- equations based on Kirchhoff's voltage law (KVL)
- equations based on Kirchhoff's current law (KCL)
- branch constitutive equations

The equations have to be formulated (represented in a computer program) automatically in a simple, comprehensive manner. Once formulated, the system of equations has to be solved. There are two main aspects to be considered when choosing algorithms for this purpose: accuracy and speed. The MNA, briefly for **M**odified **N**odal **A**alysis, has been proved to accomplish these tasks.

MNA applied to a circuit with passive elements, independent current and voltage sources and active elements results in a matrix equation of the form:

$$[A] \cdot [x] = [z] \quad (3.1)$$

For a circuit with N nodes and M independent voltage sources:

- The A matrix
 - is $(N+M) \times (N+M)$ in size, and consists only of known quantities
 - the $N \times N$ part of the matrix in the upper left:
 - * has only passive elements
 - * elements connected to ground appear only on the diagonal
 - * elements not connected to ground are both on the diagonal and off-diagonal terms
 - the rest of the A matrix (not included in the $N \times N$ upper left part) contains only 1, -1 and 0 (other values are possible if there are dependent current and voltage sources)
- The x matrix

- is an $(N+M) \times 1$ vector that holds the unknown quantities (node voltages and the currents through the independent voltage sources)
- the top N elements are the n node voltages
- the bottom M elements represent the currents through the M independent voltage sources in the circuit
- The z matrix
 - is an $(N+M) \times 1$ vector that holds only known quantities
 - the top N elements are either zero or the sum and difference of independent current sources in the circuit
 - the bottom M elements represent the M independent voltage sources in the circuit

The circuit is solved by a simple matrix manipulation:

$$[x] = [A]^{-1} \cdot [z] \quad (3.2)$$

Though this may be difficult by hand, it is straightforward and so is easily done by computer.

3.1.1 Generating the MNA matrices

The following section is an algorithmic approach to the concept of the Modified Nodal Analysis. There are three matrices we need to generate, the A matrix, the x matrix and the z matrix. Each of these will be created by combining several individual sub-matrices.

3.1.2 The A matrix

The A matrix will be developed as the combination of 4 smaller matrices, G , B , C , and D .

$$A = \begin{bmatrix} G & B \\ C & D \end{bmatrix} \quad (3.3)$$

The A matrix is $(M+N) \times (M+N)$ (N is the number of nodes, and M is the number of independent voltage sources) and:

- the G matrix is $N \times N$ and is determined by the interconnections between the circuit elements
- the B matrix is $N \times M$ and is determined by the connection of the voltage sources
- the C matrix is $M \times N$ and is determined by the connection of the voltage sources (B and C are closely related, particularly when only independent sources are considered)
- the D matrix is $M \times M$ and is zero if only independent sources are considered

Rules for making the G matrix

The G matrix is an $N \times N$ matrix formed in two steps.

1. Each element in the diagonal matrix is equal to the sum of the conductance (one over the resistance) of each element connected to the corresponding node. So the first diagonal element is the sum of conductances connected to node 1, the second diagonal element is the sum of conductances connected to node 2, and so on.

2. The off diagonal elements are the negative conductance of the element connected to the pair of corresponding node. Therefore a resistor between nodes 1 and 2 goes into the G matrix at location (1,2) and locations (2,1).

If an element is grounded, it will only have contribute to one entry in the G matrix – at the appropriate location on the diagonal. If it is ungrounded it will contribute to four entries in the matrix – two diagonal entries (corresponding to the two nodes) and two off-diagonal entries.

Rules for making the B matrix

The B matrix is an $N \times M$ matrix with only 0, 1 and -1 elements. Each location in the matrix corresponds to a particular voltage source (first dimension) or a node (second dimension). If the positive terminal of the i th voltage source is connected to node k , then the element (k,i) in the B matrix is a 1. If the negative terminal of the i th voltage source is connected to node k , then the element (k,i) in the B matrix is a -1. Otherwise, elements of the B matrix are zero.

If a voltage source is ungrounded, it will have two elements in the B matrix (a 1 and a -1 in the same column). If it is grounded it will only have one element in the matrix.

Rules for making the C matrix

The C matrix is an $M \times N$ matrix with only 0, 1 and -1 elements. Each location in the matrix corresponds to a particular node (first dimension) or voltage source (second dimension). If the positive terminal of the i th voltage source is connected to node k , then the element (i,k) in the C matrix is a 1. If the negative terminal of the i th voltage source is connected to node k , then the element (i,k) in the C matrix is a -1. Otherwise, elements of the C matrix are zero.

In other words, the C matrix is the transpose of the B matrix. This is not the case when dependent sources are present.

Rules for making the D matrix

The D matrix is an $M \times M$ matrix that is composed entirely of zeros. It can be non-zero if dependent sources are considered.

3.1.3 The x matrix

The x matrix holds our unknown quantities and will be developed as the combination of 2 smaller matrices v and j . It is considerably easier to define than the A matrix.

$$x = \begin{bmatrix} v \\ j \end{bmatrix} \quad (3.4)$$

The x matrix is $1 \times (M+N)$ (N is the number of nodes, and M is the number of independent voltage sources) and:

- the v matrix is $1 \times N$ and hold the unknown voltages
- the j matrix is $1 \times M$ and holds the unknown currents through the voltage sources

Rules for making the v matrix

The v matrix is an $1 \times N$ matrix formed of the node voltages. Each element in v corresponds to the voltage at the equivalent node in the circuit (there is no entry for ground – node 0).

For a circuit with N nodes we get:

$$v = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_N \end{bmatrix} \quad (3.5)$$

Rules for making the j matrix

The j matrix is an $1 \times M$ matrix, with one entry for the current through each voltage source. So if there are M voltage sources V_1, V_2 through V_M , the j matrix will be:

$$j = \begin{bmatrix} i_{V_1} \\ i_{V_2} \\ \vdots \\ i_{V_M} \end{bmatrix} \quad (3.6)$$

3.1.4 The z matrix

The z matrix holds our independent voltage and current sources and will be developed as the combination of 2 smaller matrices i and e. It is quite easy to formulate.

$$z = \begin{bmatrix} i \\ e \end{bmatrix} \quad (3.7)$$

The z matrix is $1 \times (M+N)$ (N is the number of nodes, and M is the number of independent voltage sources) and:

- the i matrix is $1 \times N$ and contains the sum of the currents through the passive elements into the corresponding node (either zero, or the sum of independent current sources)
- the e matrix is $1 \times M$ and holds the values of the independent voltage sources

Rules for making the i matrix

The i matrix is an $1 \times N$ matrix with each element of the matrix corresponding to a particular node. The value of each element of i is determined by the sum of current sources into the corresponding node. If there are no current sources connected to the node, the value is zero.

Rules for making the e matrix

The e matrix is an $1 \times M$ matrix with each element of the matrix equal in value to the corresponding independent voltage source.

3.1.5 A simple example

The example given in fig. ?? illustrates applying the rules for building the MNA matrices and how this relates to basic equations used in circuit analysis.

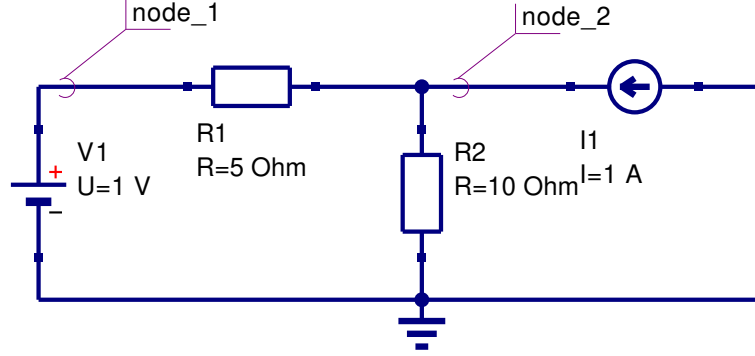


Figure 3.1: example circuit applied to modified nodal analysis

Going through the MNA algorithm

The G matrix is a 2×2 matrix because there are 2 different nodes apart from ground which is the reference node. On the diagonal you find the sum of the elements conductances connected to the nodes 1 and 2. The off-diagonal matrix entries contain the negative conductances of the elements connected between two nodes.

$$G = \begin{bmatrix} \frac{1}{R_1} & -\frac{1}{R_1} \\ -\frac{1}{R_1} & \frac{1}{R_1} + \frac{1}{R_2} \end{bmatrix} = \begin{bmatrix} 0.2 & -0.2 \\ -0.2 & 0.3 \end{bmatrix} \quad (3.8)$$

The B matrix (which is transposed to C) is a 1×2 matrix because there is one voltage source and 2 nodes. The positive terminal of the voltage source V_1 is connected to node 1. That is why

$$B = C^T = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \quad (3.9)$$

and the D matrix is filled with zeros only because there are no dependent (active and controlled) devices in the example circuit.

$$D = [0] \quad (3.10)$$

The x matrix is a 1×3 matrix. The MNA equations deliver a solution for the unknown voltages at each node in a circuit except the reference node and the currents through each voltage source.

$$x = \begin{bmatrix} v_1 \\ v_2 \\ i_{V_1} \end{bmatrix} \quad (3.11)$$

The z matrix is according to the rules for building it a 1×3 matrix. The upper two entries are the sums of the currents flowing into node 1 and node 2. The lower entry is the voltage value of the voltage source V_1 .

$$z = \begin{bmatrix} 0 \\ I_1 \\ U_1 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix} \quad (3.12)$$

According to the MNA algorithm the equation system is represented by

$$[A] \cdot [x] = [z] \quad (3.13)$$

which is equivalent to

$$\begin{bmatrix} G & B \\ C & D \end{bmatrix} \cdot [x] = [z] \quad (3.14)$$

In the example eq. (??) expands to:

$$\begin{bmatrix} \frac{1}{R_1} & -\frac{1}{R_1} & 1 \\ -\frac{1}{R_1} & \frac{1}{R_1} + \frac{1}{R_2} & 0 \\ 1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} v_1 \\ v_2 \\ i_{V_1} \end{bmatrix} = \begin{bmatrix} 0 \\ I_1 \\ U_1 \end{bmatrix} \quad (3.15)$$

The equation systems to be solved is now defined by the following matrix representation.

$$\begin{bmatrix} 0.2 & -0.2 & 1 \\ -0.2 & 0.3 & 0 \\ 1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} v_1 \\ v_2 \\ i_{V_1} \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix} \quad (3.16)$$

Using matrix inversion the solution vector x writes as follows:

$$[x] = [A]^{-1} \cdot [z] = \begin{bmatrix} v_1 \\ v_2 \\ i_{V_1} \end{bmatrix} = \begin{bmatrix} 1 \\ 4 \\ 0.6 \end{bmatrix} \quad (3.17)$$

The result in eq. (??) denotes the current through the voltage source V_1 is 0.6A, the voltage at node 1 is 1V and the voltage at node 2 is 4V.

How the algorithm relates to basic equations in circuit analysis

Expanding the matrix representation in eq. (??) to a set of equations denotes the following equation system consisting of 3 of them.

$$\text{I:} \quad 0 = \frac{1}{R_1} \cdot v_1 - \frac{1}{R_1} \cdot v_2 + i_{V_1} \quad \text{KCL at node 1} \quad (3.18)$$

$$\text{II:} \quad I_1 = -\frac{1}{R_1} \cdot v_1 + \left(\frac{1}{R_1} + \frac{1}{R_2} \right) \cdot v_2 \quad \text{KCL at node 2} \quad (3.19)$$

$$\text{III:} \quad U_1 = v_1 \quad \text{constitutive equation} \quad (3.20)$$

Apparently eq. I and eq. II conform to Kirchhoff's current law at the nodes 1 and 2. The last equation is just the constitutive equation for the voltage source V_1 . There are three unknowns (v_1 , v_2 and i_{V_1}) and three equations, thus the system should be solvable.

Equation III indicates the voltage at node 1 is 1V. Applying this result to eq. II and transposing it to v_2 (the voltage at node 2) gives

$$v_2 = \frac{I_1 + \frac{1}{R_1} \cdot U_1}{\frac{1}{R_1} + \frac{1}{R_2}} = 4V \quad (3.21)$$

The missing current through the voltage source V_1 can be computed using both the results $v_2 = 4V$ and $v_1 = 1V$ by transforming equation I.

$$i_{V_1} = \frac{1}{R_1} \cdot v_2 - \frac{1}{R_1} \cdot v_1 = 0.6A \quad (3.22)$$

The small example, shown in fig. ??, and the excursus into artless math verifies that the MNA algorithm and classic electrical handiwork tend to produce the same results.

3.2 Extensions to the MNA

As noted in the previous sections the D matrix is zero and the B and C matrices are transposed each other and filled with either 1, -1 or 0 provided that there are no dependent sources within the circuit. This changes when introducing active (and controlled) elements. Examples are voltage controlled voltage sources, transformers and ideal operational amplifiers. The models are depicted in section ?? and ??

3.3 Non-linear DC Analysis

Previous sections described using the modified nodal analysis solving linear networks including controlled sources. It can also be used to solve networks with non-linear components like diodes and transistors. Most methods are based on iterative solutions of a linearised equation system. The best known is the so called Newton-Raphson method.

3.3.1 Newton-Raphson method

The Newton-Raphson method is going to be introduced using the example circuit shown in fig. ?? having a single unknown: the voltage at node 1.

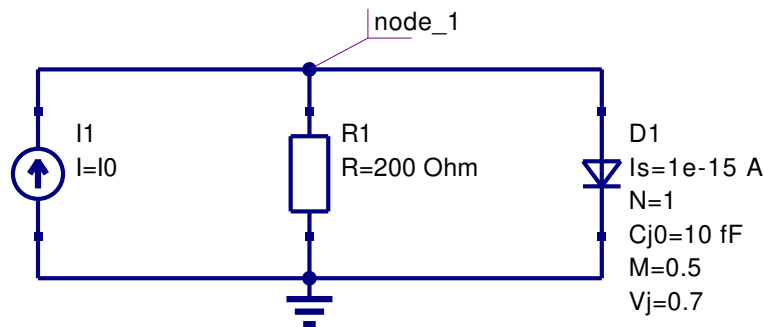


Figure 3.2: example circuit for non-linear DC analysis

The 1x1 MNA equation system to be solved can be written as

$$[G] \cdot [V_1] = [I_0] \quad (3.23)$$

whereas the value for G is now going to be explained. The current through a diode is simply determined by Shockley's approximation

$$I_d = I_S \cdot \left(e^{\frac{V_d}{V_T}} - 1 \right) \quad (3.24)$$

Thus Kirchhoff's current law at node 1 can be expressed as

$$I_0 = \frac{V}{R} + I_S \cdot \left(e^{\frac{V}{V_T}} - 1 \right) \quad (3.25)$$

By establishing eq. (??) it is possible to trace the problem back to finding the zero point of the function f .

$$f(V) = \frac{V}{R} + I_S \cdot \left(e^{\frac{V}{V_T}} - 1 \right) - I_0 \quad (3.26)$$

Newton developed a method stating that the zero point of a functions derivative (i.e. the tangent) at a given point is nearer to the zero point of the function itself than the original point. In mathematical terms this means to linearise the function f at a starting value $V^{(0)}$.

$$f(V^{(0)} + \Delta V) \approx f(V^{(0)}) + \left. \frac{\partial f(V)}{\partial V} \right|_{V^{(0)}} \cdot \Delta V \quad \text{with} \quad \Delta V = V^{(1)} - V^{(0)} \quad (3.27)$$

Setting $f(V^{(1)}) = 0$ gives

$$V^{(1)} = V^{(0)} - \frac{f(V^{(0)})}{\left. \frac{\partial f(V)}{\partial V} \right|_{V^{(0)}}} \quad (3.28)$$

or in the general case with m being the number of iteration

$$V^{(m+1)} = V^{(m)} - \frac{f(V^{(m)})}{\left. \frac{\partial f(V)}{\partial V} \right|_{V^{(m)}}} \quad (3.29)$$

This must be computed until $V^{(m+1)}$ and $V^{(m)}$ differ less than a certain barrier.

$$\left| V^{(m+1)} - V^{(m)} \right| < \varepsilon_{abs} + \varepsilon_{rel} \cdot \left| V^{(m)} \right| \quad (3.30)$$

With very small ε_{abs} the iteration would break too early and for little ε_{rel} values the iteration aims to a useless precision for large absolute values of V .

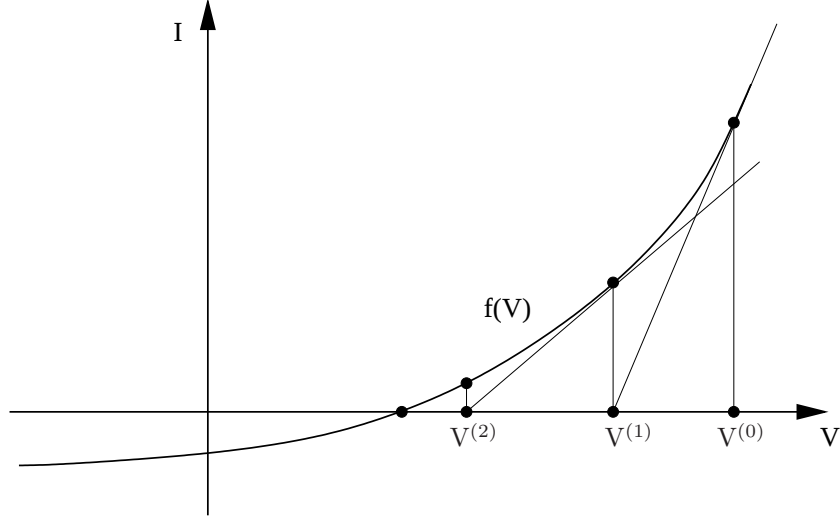


Figure 3.3: Newton-Raphson method for example circuit

With this theoretical background it is now possible to step back to eq. (??) being the determining equation for the example circuit. With

$$g_d^{(m)} = \left. \frac{\partial I_d}{\partial V} \right|_{V^{(m)}} = \frac{I_S}{V_T} \cdot e^{\frac{V^{(m)}}{V_T}} \quad (3.31)$$

and

$$\left. \frac{\partial f(V)}{\partial V} \right|_{V^{(m)}} = \frac{1}{R} + g_d^{(m)} \quad (3.32)$$

the eq. (??) can be written as

$$\left(g_d^{(m)} + \frac{1}{R} \right) \cdot V^{(m+1)} = I_0 - \left(I_d^{(m)} - g_d^{(m)} \cdot V^{(m)} \right) \quad (3.33)$$

when the expression

$$f(V^{(m)}) = \frac{1}{R} \cdot V^{(m)} + I_d^{(m)} - I_0 \quad (3.34)$$

based upon eq. (??) is taken into account. Comparing the introductory MNA equation system in eq. (??) with eq. (??) proposes the following equivalent circuit for the diode model.

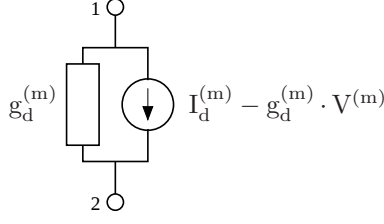


Figure 3.4: accompanied equivalent circuit for intrinsic diode

With

$$I_{eq} = I_d^{(m)} - g_d^{(m)} \cdot V^{(m)} \quad (3.35)$$

the MNA matrix entries can finally be written as

$$\begin{bmatrix} g_d & -g_d \\ -g_d & g_d \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} -I_{eq} \\ I_{eq} \end{bmatrix} \quad (3.36)$$

In analog ways all controlled current sources with non-linear current-voltage dependency built into diodes and transistors can be modeled. The left hand side of the MNA matrix (the $\mathbf{A}$ matrix) is called Jacobian matrix which is going to be build in each iteration step. For the solution vector x possibly containing currents as well when voltage sources are in place a likely convergence criteria as defined in eq. (??) must be defined for the currents.

Having understood the one-dimensional example, it is now only a small step to the general multi-dimensional algorithm: The node voltage becomes a vector $\mathbf{V}^{(m)}$, factors become the corresponding matrices and differentiations become Jacobian matrices.

The function whose zero must be found is the transformed MNA equation ??:

$$\mathbf{f}(\mathbf{V}^{(m)}) = \mathbf{G} \cdot \mathbf{V}^{(m)} - \mathbf{I}_0^{(m)} \quad (3.37)$$

The only difference to the linear case is that the vector $\mathbf{I}_0$ also contains the currents flowing out of the non-linear components. The iteration formula of the Newton-Raphson method writes:

$$\mathbf{V}^{(m+1)} = \mathbf{V}^{(m)} - \left(\frac{\partial \mathbf{f}(\mathbf{V})}{\partial \mathbf{V}} \bigg|_{\mathbf{V}^{(m)}} \right)^{-1} \cdot \mathbf{f}(\mathbf{V}^{(m)}) \quad (3.38)$$

Note that the Jacobian matrix is nothing else but the real part of the MNA matrix for the AC analysis:

$$\mathbf{J}^{(m)} = \frac{\partial \mathbf{f}(\mathbf{V})}{\partial \mathbf{V}} \bigg|_{\mathbf{V}^{(m)}} = \mathbf{G} - \frac{\partial \mathbf{I}_0}{\partial \mathbf{V}} \bigg|_{\mathbf{V}^{(m)}} = \mathbf{G} - \mathbf{J}_{nl}^{(m)} = \text{Re}(\mathbf{G}_{AC}) \quad (3.39)$$

where the index nl denotes only the non-linear terms. Putting equation ?? into equation ?? and multiplying it with the Jacobian matrix leads to

$$\mathbf{J}^{(m)} \cdot \mathbf{V}^{(m+1)} = \mathbf{J}^{(m)} \cdot \mathbf{V}^{(m)} - \mathbf{f}(\mathbf{V}^{(m)}) \quad (3.40)$$

$$= (\mathbf{G} - \mathbf{J}_{nl}^{(m)}) \cdot \mathbf{V}^{(m)} - \mathbf{G} \cdot \mathbf{V}^{(m)} + \mathbf{I}_0^{(m)} \quad (3.41)$$

$$= -\mathbf{J}_{nl}^{(m)} \cdot \mathbf{V}^{(m)} + \mathbf{I}_0^{(m)} \quad (3.42)$$

So, bringing the Jacobian back to the right side results in the new iteration formula:

$$\mathbf{V}^{(m+1)} = \left(\mathbf{J}^{(m)} \right)^{-1} \cdot \left(-\mathbf{J}_{nl}^{(m)} \cdot \mathbf{V}^{(m)} + \mathbf{I}_0^{(m)} \right) \quad (3.43)$$

The negative sign in front of $\mathbf{J}_{nl}$ is due to the definition of $\mathbf{I}_0$ flowing out of the component. Note that $\mathbf{I}_0^{(m)}$ still contains contributions of linear and non-linear current sources.

3.3.2 Convergence

Numerical as well as convergence problems occur during the Newton-Raphson iterations when dealing with non-linear device curves as they are used to model the DC behaviour of diodes and transistors.

Linearising the exponential diode eq. (??) in the forward region a numerical overflow can occur. The diagram in fig. ?? visualises this situation. Starting with $V^{(0)}$ the next iteration value gets $V^{(1)}$ which results in an indefinite large diode current. It can be limited by iterating in current instead of voltage when the computed voltage exceeds a certain value.

How this works is going to be explained using the diode model shown in fig. ?. When iterating in voltage (as normally done) the new diode current is

$$\hat{I}_d^{(m+1)} = g_d^{(m)} \left(\hat{V}^{(m+1)} - V^{(m)} \right) + I_d^{(m)} \quad (3.44)$$

The computed value $\hat{V}^{(m+1)}$ in iteration step $m+1$ is not going to be used for the following step when $V^{(m)}$ exceeds the critical voltage V_{CRIT} which gets explained in the below paragraphs. Instead, the value resulting from

$$I_d^{(m+1)} = I_S \cdot \left(e^{\frac{V^{(m+1)}}{nV_T}} - 1 \right) \quad (3.45)$$

is used (i.e. iterating in current). With

$$\hat{I}_d^{(m+1)} \stackrel{!}{=} I_d^{(m+1)} \quad \text{and} \quad g_d^{(m)} = \frac{I_S}{n \cdot V_T} \cdot e^{\frac{V^{(m)}}{n \cdot V_T}} \quad (3.46)$$

the new voltage can be written as

$$V^{(m+1)} = V^{(m)} + nV_T \cdot \ln \left(\frac{\hat{V}^{(m+1)} - V^{(m)}}{nV_T} + 1 \right) \quad (3.47)$$

Proceeding from Shockley's simplified diode equation the critical voltage is going to be defined. The explained algorithm can be used for all exponential DC equations used in diodes and transistors.

$$I(V) = I_S \cdot \left(e^{\frac{V}{nV_T}} - 1 \right) \quad (3.48)$$

$$y(x) = f(x) \quad (3.49)$$

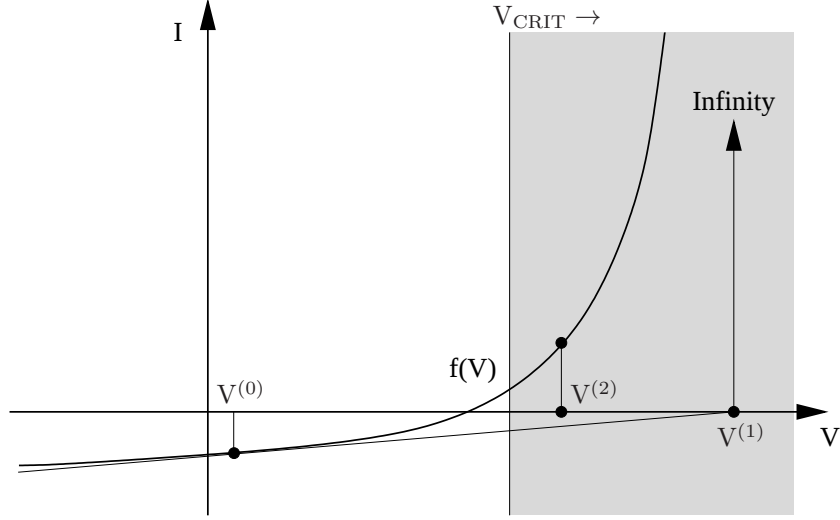


Figure 3.5: numerical problem with Newton-Raphson algorithm

The critical voltage V_{CRIT} is the voltage where the curve radius of eq. (??) has its minimum with I and V having equally units. The curve radius R for the explicit definition in eq. (??) can be written as

$$R = \left| \frac{\left(1 + \left(\frac{dy}{dx}\right)^2\right)^{3/2}}{\frac{d^2y}{dx^2}} \right| \quad (3.50)$$

Finding this equations minimum requires the derivative.

$$\frac{dR}{dx} = \frac{\frac{d^2y}{dx^2} \cdot \frac{3}{2} \left(1 + \left(\frac{dy}{dx}\right)^2\right)^{1/2} \cdot 2 \cdot \frac{dy}{dx} \cdot \frac{d^2y}{dx^2} - \left(1 + \left(\frac{dy}{dx}\right)^2\right)^{3/2} \cdot \frac{d^3y}{dx^3}}{\left(\frac{d^2y}{dx^2}\right)^2} \quad (3.51)$$

The diagram in fig. ?? shows the graphs of eq. (??) and eq. (??) with $n = 1$, $I_S = 100\text{nA}$ and $V_T = 25\text{mV}$.

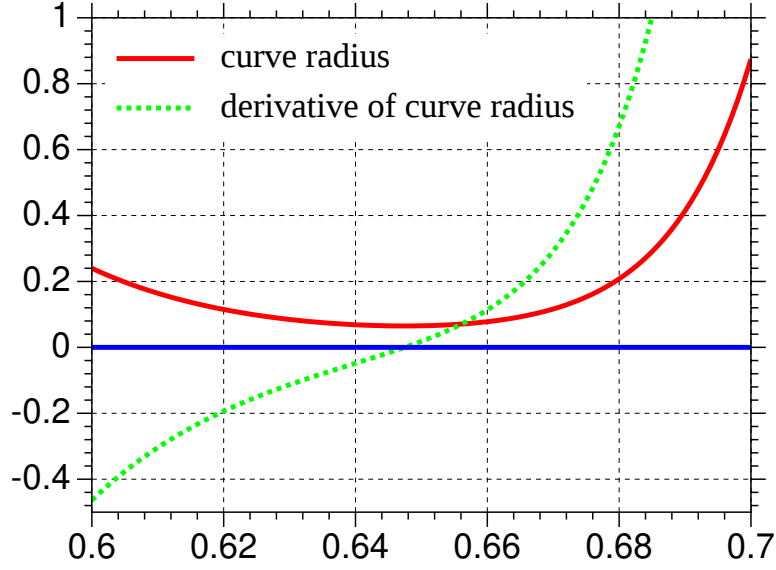


Figure 3.6: curve radius of exponential diode curve and its derivative

With the following higher derivatives of eq. (??)

$$\frac{dI(V)}{dV} = \frac{I_S}{nV_T} \cdot e^{\frac{V}{nV_T}} \quad (3.52)$$

$$\frac{d^2I(V)}{dV^2} = \frac{I_S}{n^2V_T^2} \cdot e^{\frac{V}{nV_T}} \quad (3.53)$$

$$\frac{d^3I(V)}{dV^3} = \frac{I_S}{n^3V_T^3} \cdot e^{\frac{V}{nV_T}} \quad (3.54)$$

the critical voltage results in

$$\frac{dR}{dx} \stackrel{!}{=} 0 = 3 - \frac{n^2V_T^2}{I_S^2} \cdot e^{-2\frac{V}{nV_T}} - 1 \quad \rightarrow \quad V_{CRIT} = nV_T \cdot \ln\left(\frac{nV_T}{I_S\sqrt{2}}\right) \quad (3.55)$$

In order to avoid numerical errors a minimum value of the pn-junction's derivative (i.e. the currents tangent in the operating point) g_{min} is defined. On the one hand this avoids very large deviations of the appropriate voltage in the next iteration step in the backward region of the pn-junction and on the other hand it avoids indefinite large voltages if g_d itself suffers from numerical errors and approaches zero.

The quadratic input I-V curve of field-effect transistors as well as the output characteristics of these devices can be handled in similar ways. The limiting (and thereby improving the convergence behaviour) algorithm must somehow ensure that the current and/or voltage deviation from one iteration step to the next step is not too a large value. Because of the wide range of existing variations how these curves are exactly modeled there is no standard strategy to achieve this. Anyway, the threshold voltage V_{Th} should play an important role as well as the direction which the current iteration step follows.

3.4 Overall solution algorithm for DC Analysis

In this section an overall solution algorithm for a DC analysis for linear as well as non-linear networks is given. With non-linear network elements at hand the Newton-Raphson (NR) algorithm is applied.

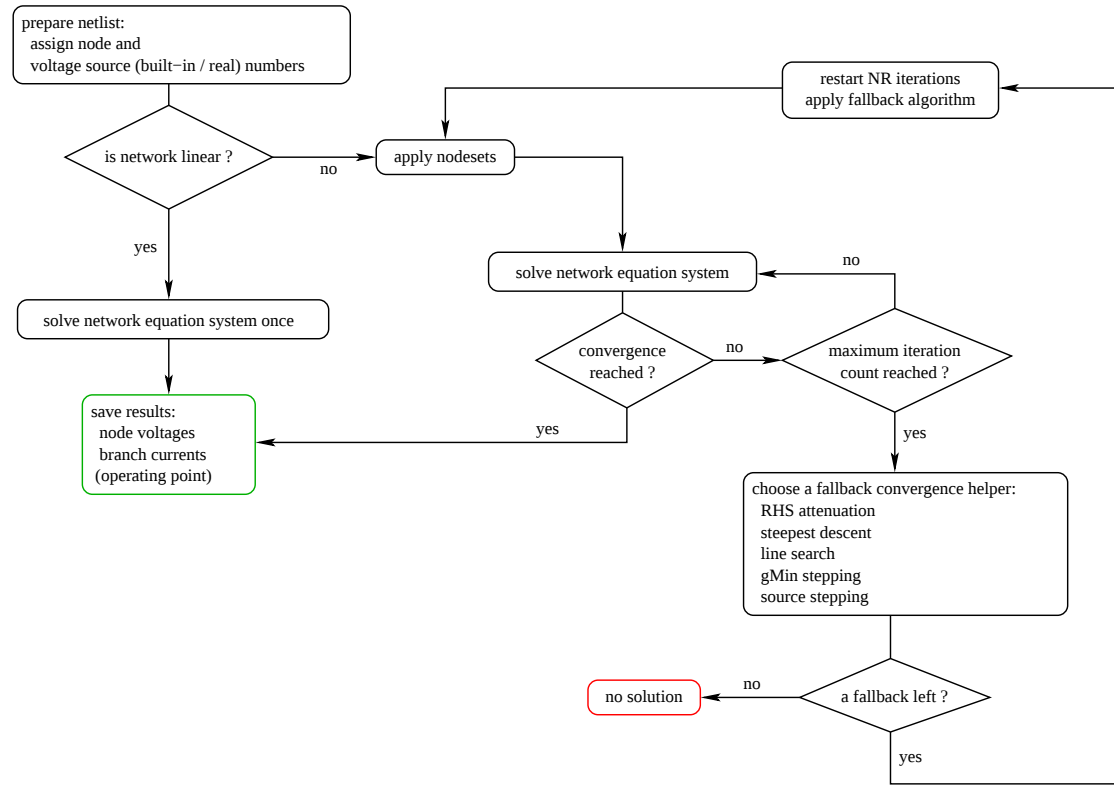


Figure 3.7: DC solution algorithm flow chart

The algorithm shown in fig. ?? has been proved to be able to find DC solutions for a large variety of networks. It must be said that the application of any of the fallback convergence helpers indicates a nearly or definitely singular equation system (e.g. floating nodes or overdetermining sources). The convergence problems are either due to an apparently “wrong” network topology or to the model implementation of non-linear components. For some of the problems also refer to the facts mentioned in section ?? on page ??. In some cases it may even occur that tiny numerical inaccuracies lead to non-convergences whereas the choice of a more accurate (but probably slower) equation system solver can help. With network topologies having more than a single stable solution (e.g. bistable flip-flops) it is recommended to apply nodesets, i.e. forcing the Newton-Raphson iteration into a certain direction by initial values.

When having problems to get a circuit have its DC solution the following actions can be taken to solve these problems.

- check circuit topology (e.g. floating nodes or overdetermining sources)

- check model parameters of non-linear components
- apply nodesets
- choose a more accurate equation system solver
- relax the convergence tolerances if possible
- increase the maximum iteration count
- choose the preferred fallback algorithm

The presented concepts are common to most circuit simulators each having to face the mentioned aspects. And probably facing it in a different manner with more or less big differences in their implementation details especially regarding the (fallback) convergence helpers. None of the algorithms based on Newton-Raphson ensures global convergence, thus very few documents have been published either for the complexity of the topic or for uncertainties in the detailed implementation each carrying the attribute “can help” or “may help”.

So for now the application of a circuit simulator to find the DC solution of a given network sometimes keeps being a task for people knowing what they want to achieve and what they can roughly expect.

Chapter 4

AC Analysis

The AC analysis is a small signal analysis in the frequency domain. Basically this type of simulation uses the same algorithms as the DC analysis (section ?? on page ??). The AC analysis is a linear modified nodal analysis. Thus no iterative process is necessary. With the Y-matrix of the components, i.e. now a complex matrix, and the appropriate extensions it is necessary to solve the equation system (??) similar to the (linear) DC analysis.

$$[A] \cdot [x] = [z] \quad \text{with} \quad A = \begin{bmatrix} Y & B \\ C & D \end{bmatrix} \quad (4.1)$$

Non-linear components have to be linearized at the DC bias point. That is, before an AC simulation with non-linear components can be performed, a DC simulation must be completed successfully. Then, the MNA stamp of the non-linear components equals their entries of the Jacobian matrix, which was already computed during the DC simulation. In addition to this real-valued elements, a further stamp has to be applied: The Jacobian matrix of the non-linear charges multiplied by $j\omega$ (see also section ??).

Chapter 5

AC Noise Analysis

5.1 Definitions

First some definition must be done:

Reciprocal Networks:

Two networks A and B are reciprocal to each other if their transimpedances have the following relation:

$$Z_{mn,A} = Z_{nm,B} \quad (5.1)$$

That means: Drive the current I into node n of circuit A and at node m the voltage $I \cdot Z_{mn,A}$ appears. In circuit B it is just the way around.

Adjoint Networks:

Network A and network B are adjoint to each other if the following equation holds for their MNA matrices:

$$[A]^T = [B] \quad (5.2)$$

5.2 The Algorithm

To calculate the small signal noise of a circuit, the AC noise analysis has to be applied [?]. This technique uses the principle of the AC analysis described in chapter ?? on page ?. In addition to the MNA matrix A one needs the noise current correlation matrix $\underline{C}_Y$ of the circuit, that contains the equivalent noise current sources for every node on its main diagonal and their correlation on the other positions.

The basic concept of the AC noise analysis is as follows: The noise voltage at node i should be calculated, so the voltage arising due to the noise source at node j is calculated first. This has to be done for every n nodes and after that adding all the noise voltages (by paying attention to their correlation) leads to the overall voltage. But that would mean to solve the MNA equation n times. Fortunately there is a more easy way. One can perform the above-mentioned n steps in one single step, if the reciprocal MNA matrix is used. This matrix equals the MNA matrix itself, if the network is reciprocal. A network that only contains resistors, capacitors, inductors, gyrators and transformers is reciprocal.

The question that needs to be answered now is: How to get the reciprocal MNA matrix for an arbitrary network? This is equivalent to the question: How to get the MNA matrix of the adjoint network. The answer is quite simple: Just transpose the MNA matrix!

For any network, calculating the noise voltage at node i is done by the following three steps:

$$1. \text{ Solving MNA equation: } [A]^T \cdot [x] = [A]^T \cdot \begin{bmatrix} v \\ j \end{bmatrix} = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ -1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \leftarrow i\text{-th row} \quad (5.3)$$

$$2. \text{ Creating noise correlation matrix: } (\underline{C}_Y) \quad (5.4)$$

$$3. \text{ Calculating noise voltage: } v_{noise,i} = \sqrt{[v]^T \cdot (\underline{C}_Y) \cdot [v]^*} \quad (5.5)$$

If the correlation between several noise voltages is also wanted, the procedure is straight forward: Perform step 1 for every desired node, put the results into a matrix and replace the vector $[v]$ in step 3 by this matrix. This results in the complete correlation matrix. Indeed, the above-mentioned algorithm is only a specialisation of transforming the noise correlation matrices (see section ??).

If the normal AC analysis has already be done with LU decomposition, then the most time consuming work of step 1 has already be done.

$$\text{instead of } Y = L \cdot U \quad \text{we have} \quad Y^T = U^T \cdot L^T \quad (5.6)$$

I.e. U^T becomes the new L matrix and L^T becomes the new U matrix, and the matrix equation do not need to be solved again, because only the right-hand side was changed. So altogether this is a quickly done task. (Note that in step 3, only the subvector $[v]$ of vector $[x]$ is used. See section ?? for details on this.)

If the noise voltage at another node needs to be known, only the right-hand side of step 1 changes. That is, a new LU decomposition is not needed.

Reusing the LU decomposed MNA matrix of the usual AC analysis is possible if there has been no pivoting necessary during the decomposition.

When using either Crout's or Doolittle's definition of the LU decomposition during the AC analysis the decomposition representation changes during the AC noise analysis as the matrix A gets transposed. This means:

$$A = L \cdot U \quad \text{with} \quad L = \begin{bmatrix} l_{11} & 0 & \dots & 0 \\ l_{21} & l_{22} & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ l_{n1} & \dots & \dots & l_{nn} \end{bmatrix} \quad \text{and} \quad U = \begin{bmatrix} 1 & u_{12} & \dots & u_{1n} \\ 0 & 1 & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{bmatrix} \quad (5.7)$$

becomes

$$A^T = U^T \cdot L^T \quad \text{with} \quad L = \begin{bmatrix} 1 & 0 & \dots & 0 \\ l_{21} & 1 & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ l_{n1} & \dots & \dots & 1 \end{bmatrix} \quad \text{and} \quad U = \begin{bmatrix} u_{11} & u_{12} & \dots & u_{1n} \\ 0 & u_{22} & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & u_{nn} \end{bmatrix} \quad (5.8)$$

Thus the forward substitution (as described in section ??) and the backward substitution (as described in section ??) must be slightly modified.

$$y_i = z_i - \sum_{k=1}^{i-1} y_k \cdot l_{ik} \quad i = 1, \dots, n \quad (5.9)$$

$$x_i = \frac{y_i}{u_{ii}} - \sum_{k=i+1}^n x_k \cdot \frac{u_{ik}}{u_{ii}} \quad i = n, \dots, 1 \quad (5.10)$$

Now the diagonal elements l_{ii} can be neglected in the forward substitution but the u_{ii} elements must be considered in the backward substitution.

5.2.1 A Simple Example

The network that is depicted in figure ?? is given. The MNA equation is (see chapter ??):

$$[A] \cdot [x] = \begin{bmatrix} 1/R_1 & 0 \\ G & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \quad (5.11)$$

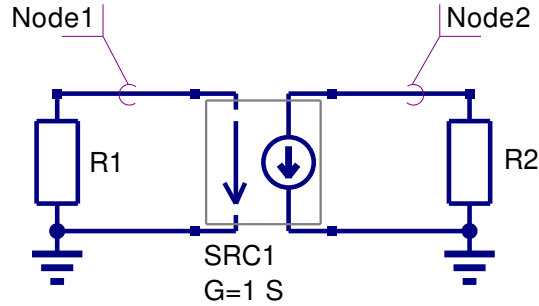


Figure 5.1: simple non-reciprocal network

Because of the controlled current source, the circuit is not reciprocal. The noise voltage at node 2 is the one to search for. Yes, this is very easy to calculate, because it is a simple example, but the algorithm described above should be used. This can be achieved by solving the equations

$$\begin{bmatrix} 1/R_1 & 0 \\ G & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} Z_{11} \\ Z_{21} \end{bmatrix} = \begin{bmatrix} -1 \\ 0 \end{bmatrix} \quad (5.12)$$

and

$$\begin{bmatrix} 1/R_1 & 0 \\ G & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} Z_{12} \\ Z_{22} \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \end{bmatrix} \quad (5.13)$$

So, the MNA matrix must be solved two times: First to get the transimpedance from node 1 to node 2 (i.e. Z_{21}) and second to get the transimpedance from node 2 to node 2 (i.e. Z_{22}). But why solving it two times, if only one voltage should be calculated? With every step transimpedances are calculated that are not need. Is there no more effective way?

Fortunately, there is Tellegen's Theorem: A network and its adjoint network are reciprocal to each other. That is, transposing the MNA matrix leads to the one of the reciprocal network. To check it out:

$$[A]^T \cdot [x] = \begin{bmatrix} 1/R_1 & G \\ 0 & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \quad (5.14)$$

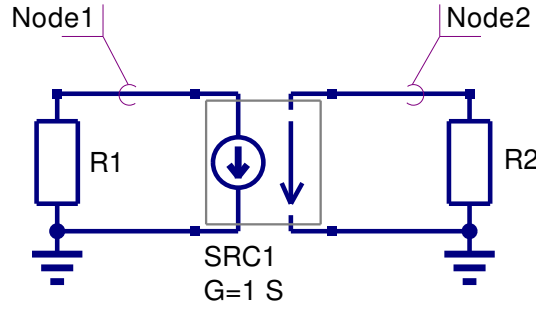


Figure 5.2: simple network to compare with adjoint network

Compare the transposed matrix with the reciprocal network in figure ???. It is true! But now it is:

$$\begin{bmatrix} 1/R_1 & G \\ 0 & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} Z_{12,reciprocal} \\ Z_{22,reciprocal} \end{bmatrix} = \begin{bmatrix} 1/R_1 & G \\ 0 & 1/R_2 \end{bmatrix} \cdot \begin{bmatrix} Z_{21} \\ Z_{22} \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \end{bmatrix} \quad (5.15)$$

Because Z_{21} of the original network equals Z_{12} of the reciprocal network, the one step delivers exactly what is needed. So the next step is:

$$([A]^T)^{-1} \cdot \begin{bmatrix} 0 \\ -1 \end{bmatrix} = \begin{bmatrix} R_1 & -G \cdot R_1 \cdot R_2 \\ 0 & R_2 \end{bmatrix} \cdot \begin{bmatrix} 0 \\ -1 \end{bmatrix} = \begin{bmatrix} G \cdot R_1 \cdot R_2 \\ -R_2 \end{bmatrix} = \begin{bmatrix} Z_{21} \\ Z_{22} \end{bmatrix} \quad (5.16)$$

Now, as the transimpedances are known, the noise voltage at node 2 can be computed. As there is no correlation, it writes as follows:

$$\langle v_{node2}^2 \rangle = \langle v_{R1,node2}^2 \rangle + \langle v_{R2,node2}^2 \rangle \quad (5.17)$$

$$= \langle i_{R1}^2 \rangle \cdot Z_{21} \cdot Z_{21}^* + \langle i_{R2}^2 \rangle \cdot Z_{22} \cdot Z_{22}^* \quad (5.18)$$

$$= \frac{4 \cdot k \cdot T \cdot \Delta f}{R_1} \cdot (G \cdot R_1 \cdot R_2)^2 + \frac{4 \cdot k \cdot T \cdot \Delta f}{R_2} \cdot (-R_2)^2 \quad (5.19)$$

$$= 4 \cdot k \cdot T \cdot \Delta f \cdot (R_1 \cdot (G \cdot R_2)^2 + R_2) \quad (5.20)$$

That's it. Yes, this could have be computed more easily, but now the universal algorithm is also clear.

5.3 Noise Current Correlation Matrix

The sections ?? and ?? show the noise current correlation matrices of noisy components. The equations are built for RMS noise currents with 1Hz bandwidth.

Chapter 6

Transient Analysis

The transient simulation is the calculation of a networks response on arbitrary excitations. The results are network quantities (branch currents and node voltages) as a function of time. Substantial for the transient analysis is the consideration of energy storing components, i.e. inductors and capacitors.

The relations between current and voltage of ideal capacitors and inductors are given by

$$V_C(t) = \frac{1}{C} \int I_C(t) \cdot dt \quad \text{and} \quad I_L(t) = \frac{1}{L} \int V_L(t) \cdot dt \quad (6.1)$$

or in terms of differential equations

$$I_C(t) = C \cdot \frac{dV_C}{dt} \quad \text{and} \quad V_L(t) = L \cdot \frac{dI_L}{dt} \quad (6.2)$$

To calculate these quantities in a computer program numerical integration methods are required. With the current-voltage relations of these components at hand it is possible to apply the modified nodal analysis algorithm in order to calculate the networks response. This means the transient analysis attempts to find an approximation to the analytical solution at discrete time points using numeric integration.

6.1 Integration methods

The following differential equation is going to be solved.

$$\frac{dx}{dt} = \dot{x}(t) = f(x, t) \quad (6.3)$$

This differential equation is transformed into an algorithm-dependent finite difference equation by quantizing and replacing

$$\dot{x}(t) = \lim_{h \rightarrow 0} \frac{x(t+h) - x(t)}{h} \quad (6.4)$$

by the following equation.

$$\dot{x}^n = \frac{x^{n+1} - x^n}{h^n} \quad (6.5)$$

There are several linear single- and multi-step numerical integration methods available, each having advantages and disadvantages concerning aspects of stability and accuracy. Integration methods can also be classified into implicit and explicit methods. Explicit methods are inexpensive per step but limited in stability and therefore not used in the field of circuit simulation to obtain a correct and stable solution. Implicit methods are more expensive per step, have better stability and therefore suitable for circuit simulation.

6.1.1 Properties of numerical integration methods

Beforehand some definitions and explanations regarding the terms often used in the following sections are made in order to avoid bigger confusions later on.

- step size

The step size is defined by the arguments difference of successive solution vectors, i.e. the time step h^n during transient analysis with n being the n -th integration step.

$$h^n = t^{n+1} - t^n \quad (6.6)$$

- order

The order k of an integration method is defined as follows: With two successive solution vectors x^{n+1} and x^n given, the successor x^{n+1} can be expressed by x^n by a finite Taylor series. The order of an integrations method equals the power of the step size up to which the approximate solution of the Taylor series differs less than x^n from the true solution x^{n+1} .

- truncation error

The truncation error ε_T depends on the order k of the integration method and results from the remainder term of the Taylor series.

- stability

In order to obtain an accurate network solution integration methods are required to be stable for a given step size h . Various stability definitions exist. This property is explained more in detail in the following sections. Basically it determines the usability of an integration algorithm.

- single- and multistep methods

Single step methods only use x^n in order to calculate x^{n+1} , multi step methods use x^i with $0 \leq i < n$.

- implicit and explicit methods

When using explicit integration methods the evaluation of the integration formula is sufficient for each integration step. With implicit methods at hand it is necessary to solve an equation system (with non-linear networks a non-linear equation system) because for the calculation of x^{n+1} , apart from x^n and $\dot{x}^n$, also $\dot{x}^{n+1}$ is used. For the transient analysis of electrical networks the implicit methods are better qualified than the explicit methods.

6.1.2 Elementary Methods

Implicit Euler Method (Backward Euler)

In the implicit Euler method the right hand side of eq. (??) is substituted by $f(x^{n+1}, t^{n+1})$ which yields

$$f(x, t) = f(x^{n+1}, t^{n+1}) \quad \rightarrow \quad x^{n+1} = x^n + h^n \cdot f(x^{n+1}, t^{n+1}) \quad (6.7)$$

The backward euler integration method is a first order single-step method.

Explicit Euler Method (Forward Euler)

In the explicit Euler method the right hand side of eq. (??) is substituted by $f(x^n, t^n)$ which yields

$$f(x, t) = f(x^n, t^n) \rightarrow x^{n+1} = x^n + h^n \cdot f(x^n, t^n) \quad (6.8)$$

The explicit Euler method has stability problems. The step size is limited by stability. In general explicit time marching integration methods are not suitable for circuit analysis where computation with large steps may be necessary when the solution changes slowly (i.e. when the accuracy does not require small steps).

Trapezoidal method

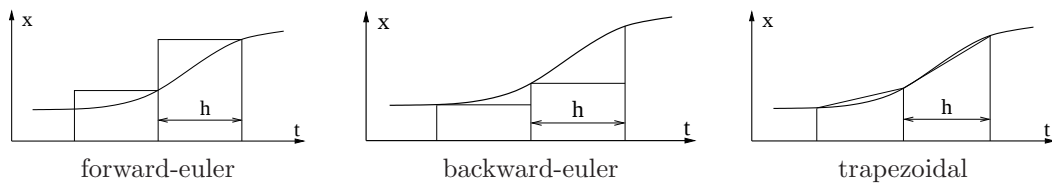
For the bilinear (also called trapezoidal) integration method $f(x, t)$ is substituted by

$$f(x, t) = \frac{1}{2} \cdot (f(x^{n+1}, t^{n+1}) + f(x^n, t^n)) \quad (6.9)$$

which yields

$$x^{n+1} = x^n + \frac{h^n}{2} \cdot (f(x^{n+1}, t^{n+1}) + f(x^n, t^n)) \quad (6.10)$$

In each integration step the average value of the intervals beginning and end is taken into account. The trapezoidal rule integration method is a second order single-step method. There is no more accurate second order integration method than the trapezoidal method.



6.1.3 Linear Multistep Methods

For higher order multi-step integration methods the general purpose method of resolution for the equation $\dot{x} = f(x, t)$

$$x^{n+1} = \sum_{i=0}^p a_i \cdot x^{n-i} + h \sum_{i=-1}^p b_i \cdot f(x^{n-i}, t^{n-i}) \quad (6.11)$$

is used. With $b_{-1} = 0$ the method is explicit and therefore not suitable for obtaining the correct and stable solution. When $b_{-1} \neq 0$ the method is implicit and suitable for circuit simulation, i.e. suitable for solving stiff problems. For differential equation systems describing electrical networks the eigenvalues strongly vary. These kind of differential equation systems are called stiff.

For a polynomial of order k the number of required coefficients is

$$2p + 3 \geq k + 1 \quad (6.12)$$

The $2p + 3$ coefficients are chosen to satisfy

$$x^{n+1} = x(t^{n+1}) \quad (6.13)$$

This can be achieved by the following equation system

$$\begin{aligned} \sum_{i=0}^p a_i &= 1 \\ \sum_{i=1}^p (-i)^j a_i + j \sum_{i=-1}^p (-i)^{j-1} b_i &= 1 \text{ for } j = 1 \dots k \end{aligned} \quad (6.14)$$

The different linear multistep integration methods which can be constructed by the equation system (??) vary in the equality condition corresponding with (??) and the choice of coefficients which are set to zero.

Gear

The Gear [?] formulae (also called BDF - backward differentiation formulae) have great importance within the multi-step integration methods used in transient analysis programs. The conditions

$$p = k - 1 \quad \text{and} \quad b_0 = b_1 = \dots = b_{k-1} = 0 \quad (6.15)$$

due to the following equation system

$$\begin{bmatrix} 0 & 1 & 1 & 1 & 1 \\ 1 & 0 & -1 & -2 & -3 \\ 2 & 0 & 1 & 4 & 9 \\ 3 & 0 & -1 & -8 & -27 \\ 4 & 0 & 1 & 16 & 81 \end{bmatrix} \cdot \begin{bmatrix} b_{-1} \\ a_0 \\ a_1 \\ a_2 \\ a_3 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} \quad (6.16)$$

for the Gear formulae of order 4. Order $k = 1$ yields the implicit Euler method. The example given in the equation system (??) results in the following integration formula.

$$\begin{aligned} x^{n+1} &= a_0 \cdot x^n + a_1 \cdot x^{n-1} + a_2 \cdot x^{n-2} + a_3 \cdot x^{n-3} + h \cdot b_{-1} \cdot f(x^{n+1}, t^{n+1}) \\ &= \frac{48}{25} \cdot x^n - \frac{36}{25} \cdot x^{n-1} + \frac{16}{25} \cdot x^{n-2} - \frac{3}{25} \cdot x^{n-3} + h \cdot \frac{12}{25} \cdot f(x^{n+1}, t^{n+1}) \end{aligned} \quad (6.17)$$

There is no more stable second order integration method than the Gear's method of second order. Only implicit Gear methods with order $k \leq 6$ are zero stable.

Adams-Bashford

The Adams-Bashford algorithm is an explicit multi-step integration method whence

$$p = k - 1 \quad \text{and} \quad a_1 = a_2 = \dots = a_{k-1} = 0 \quad \text{and} \quad b_{-1} = 0 \quad (6.18)$$

is set to satisfy the equation system (??). The equation system of the Adams-Bashford coefficients of order 4 is as follows.

$$\begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 1 \\ 0 & 0 & -2 & -4 & -6 \\ 0 & 0 & 3 & 12 & 27 \\ 0 & 0 & -4 & -32 & -108 \end{bmatrix} \cdot \begin{bmatrix} a_0 \\ b_0 \\ b_1 \\ b_2 \\ b_3 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} \quad (6.19)$$

This equation system results in the following integration formula.

$$\begin{aligned} x^{n+1} &= a_0 \cdot x^n + h \cdot b_0 \cdot f^n + h \cdot b_1 \cdot f^{n-1} + h \cdot b_2 \cdot f^{n-2} + h \cdot b_3 \cdot f^{n-3} \\ &= x^n + h \cdot \frac{55}{24} \cdot f^n - h \cdot \frac{59}{24} \cdot f^{n-1} + h \cdot \frac{37}{24} \cdot f^{n-2} - h \cdot \frac{9}{24} \cdot f^{n-3} \end{aligned} \quad (6.20)$$

The Adams-Bashford formula of order 1 yields the (explicit) forward Euler integration method.

Adams-Moulton

The Adams-Moulton algorithm is an implicit multi-step integration method whence

$$p = k - 2 \quad \text{and} \quad a_1 = a_2 = \dots = a_{k-2} = 0 \quad (6.21)$$

is set to satisfy the equation system (??). The equation system of the Adams-Moulton coefficients of order 4 is as follows.

$$\begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 1 \\ 0 & 2 & 0 & -2 & -4 \\ 0 & 3 & 0 & 3 & 12 \\ 0 & 4 & 0 & -4 & -32 \end{bmatrix} \cdot \begin{bmatrix} a_0 \\ b_{-1} \\ b_0 \\ b_1 \\ b_2 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} \quad (6.22)$$

This equation system results in the following integration formula.

$$\begin{aligned} x^{n+1} &= a_0 \cdot x^n + h \cdot b_{-1} \cdot f^{n+1} + h \cdot b_0 \cdot f^n + h \cdot b_1 \cdot f^{n-1} + h \cdot b_2 \cdot f^{n-2} \\ &= x^n + h \cdot \frac{9}{24} \cdot f^{n+1} + h \cdot \frac{19}{24} \cdot f^n - h \cdot \frac{5}{24} \cdot f^{n-1} + h \cdot \frac{1}{24} \cdot f^{n-2} \end{aligned} \quad (6.23)$$

The Adams-Moulton formula of order 1 yields the (implicit) backward Euler integration method and the formula of order 2 yields the trapezoidal rule.

6.1.4 Stability considerations

When evaluating the numerical formulations given for both implicit and explicit integration formulas once rounding errors are unavoidable. For small values of h the evaluation must be repeated very often and thus the rounding error possibly accumulates. With higher order algorithms it is possible to enlarge the step width and thereby reduce the error accumulation.

On the other hand it is questionable whether the construction of implicit algorithms is really valuable because of the higher computation effort caused by the necessary iteration (indices $n+1$ on both sides of the equation). In practice there is a class of differential equations which can be reasonably handled by implicit algorithms where explicit algorithms completely fail because of

the impracticable reduction of the step width. This class of differential equations are called stiff problems. The effect of stiffness causes for small variations in the actual solution to be computed very large deviations in the solution which get damped.

The numerical methods used for the transient analysis are required to be stiffly stable and accurate as well. The regions requirements in the complex plane are visualized in the following figure.

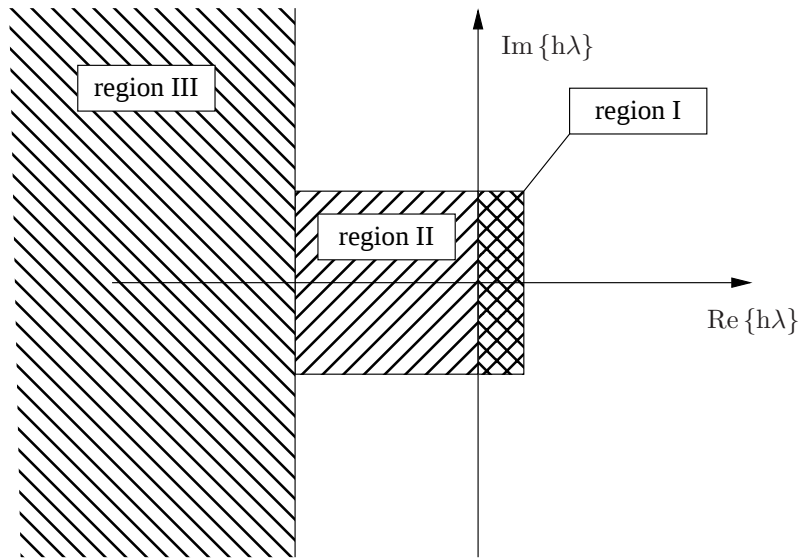


Figure 6.1: stability requirements for stiff differential equation systems

For values of $h\lambda$ in region II the numerical method must be stable and accurate, in region I accurate and in region III only stable. The area outside the specified regions are of no particular interest.

For the stability prediction of integration algorithms with regard to nonlinear differential equations and equation systems the simple and linear test differential equation

$$\dot{x} = \lambda x \quad \text{with} \quad \lambda \in \mathbb{C}, \operatorname{Re}\{\lambda\} < 0, x \geq 0 \quad (6.24)$$

is used. The condition $\operatorname{Re}\{\lambda\} < 0$ ensures the solution to be decreasing. The general purpose method of resolution given in (??) can be solved by the polynomial method setting

$$x^k = z^k \quad \text{with} \quad z \in \mathbb{C} \quad (6.25)$$

Thus we get the characteristic polynom

$$\varphi(z) = \varrho(z) + h\lambda \cdot \eta(z) = 0 \quad (6.26)$$

$$= \sum_{i=-1}^{n-1} a_i \cdot z^{n-i} + h\lambda \sum_{i=-1}^{n-1} b_i \cdot z^{n-i} \quad (6.27)$$

Because of the conditions defined by (??) the above eq. (??) can only be true for

$$|z| < 1 \quad (6.28)$$

which describes the inner unity circle on the complex plane. In order to compute the boundary of the area of absolute stability it is necessary to calculate

$$\mu(z) = h\lambda = -\frac{\varrho(z)}{\eta(z)} \quad \text{with} \quad z = e^{j\vartheta}, 0 \leq \vartheta \leq 2\pi \quad (6.29)$$

These equations describe closed loops. The inner of these loops describe the area of absolute stability. Because $\lambda \leq 0$ and $h \geq 0$ only the left half of the complex plane is of particular interest. An integration algorithm is call zero-stable if the stability area encloses $\mu = 0$. Given this condition the algorithm is as a matter of principle usable, otherwise not. If an algorithms stability area encloses the whole left half plane it is called A-stable. A-stable algorithms are stable for any h and all $\lambda < 0$. Any other kind of stability area introduces certain restrictions for μ .

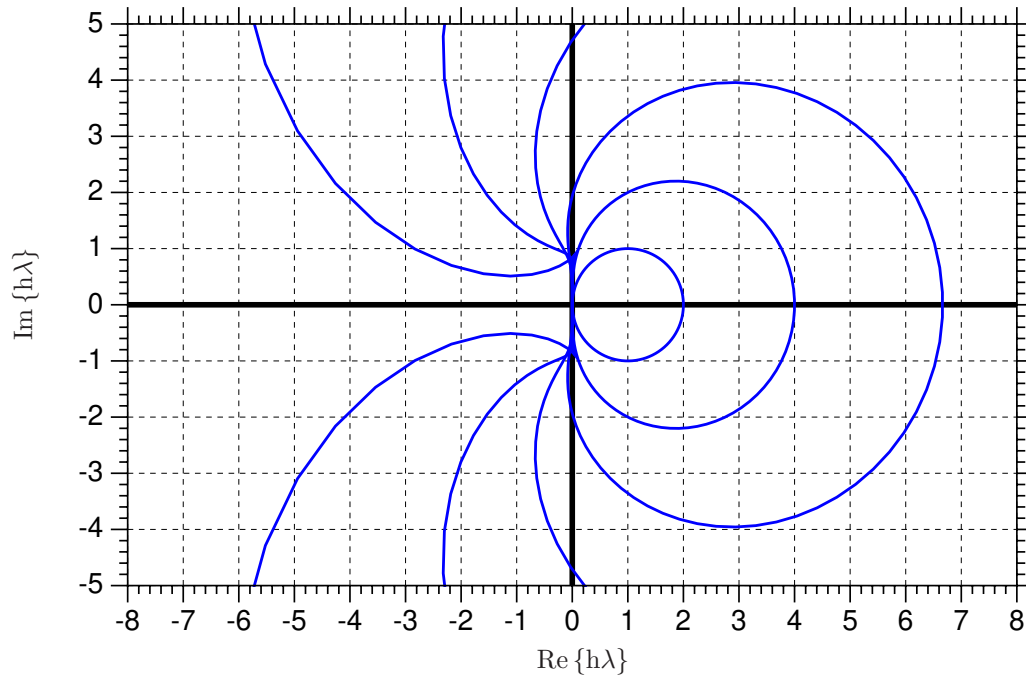


Figure 6.2: areas of absolute stability for order 1...6 Gear formulae

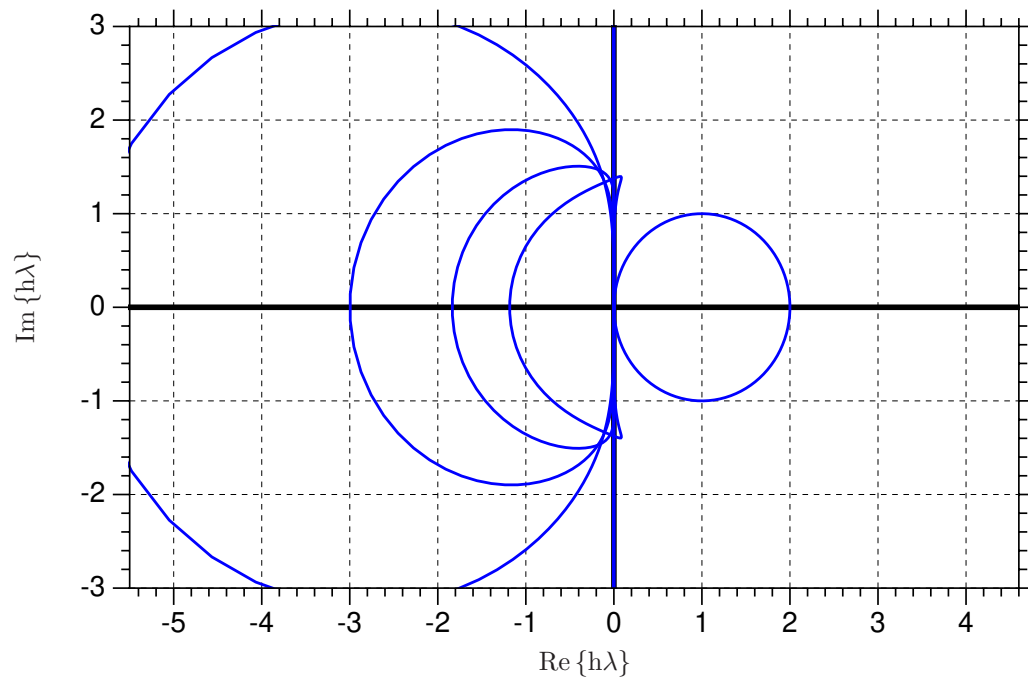


Figure 6.3: areas of absolute stability for order 1...6 Adams-Moulton formulae

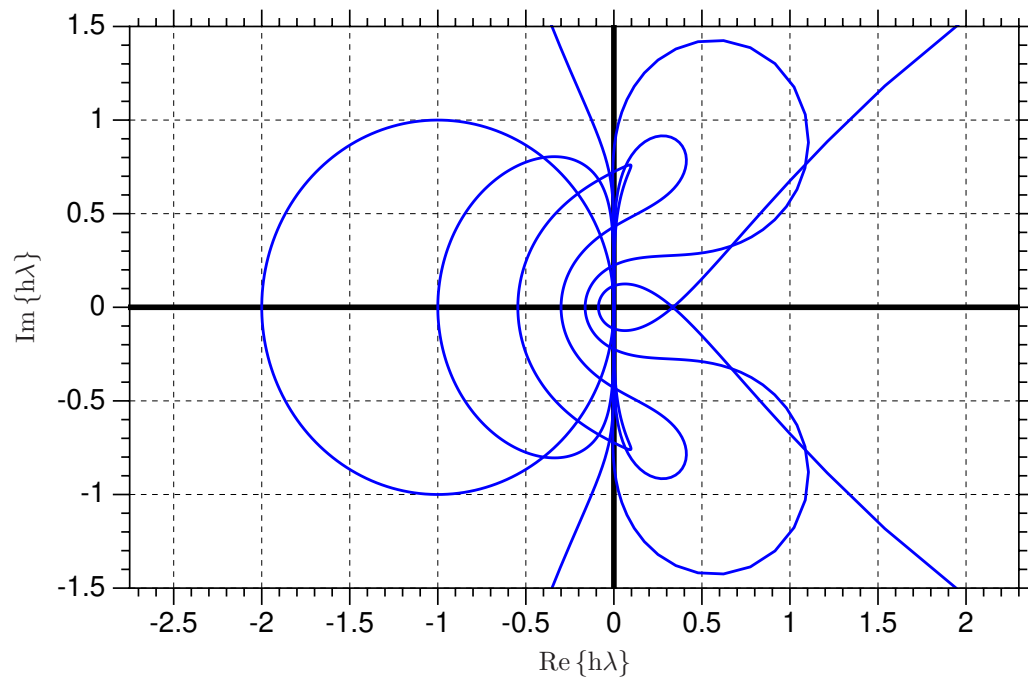


Figure 6.4: areas of absolute stability for order 1...6 Adams-Bashford formulae

The figures ??, ?? and ?? visualize the evaluation of eq. (??) for the discussed integration methods. All of the implicit formulae are zero-stable, thus principally usable. The (implicit) backward Euler, Gear order 2 and the trapezoidal integration methods are A-stable. Fig. ?? shows why the Gear formulae are of such great importance for the transient analysis of electrical networks. With least restrictions for μ they can be stabilized.

6.2 Predictor-corrector methods

In section ?? on pages ?? ff. various integration methods have been discussed. The elementary as well as linear multistep methods (in order to get more accurate methods) always assumed $a_{-1} = -1$ in its general form. Explicit methods were encountered by $b_{-1} = 0$ and implicit methods by $b_{-1} \neq 0$. Implicit methods have been shown to have a limited area of stability and explicit methods to have a larger range of stability. With increasing order k the linear multistep methods interval of absolute stability (intersection of the area of absolute stability in the complex plane with the real axis) decreases except for the implicit Gear formulae.

For these given reasons implicit methods can be used to obtain solutions of ordinary differential equation systems describing so called stiff problems. Now considering e.g. the implicit Adams-Moulton formulae of order 3

$$x^{n+1} = x^n + h \cdot \frac{5}{12} \cdot f^{n+1} + h \cdot \frac{8}{12} \cdot f^n - h \cdot \frac{1}{12} \cdot f^{n-1} \quad (6.30)$$

clarifies that f^{n+1} is necessary to calculate x^{n+1} (and the other way around as well). Every implicit integration method has this particular property. The above equation can be solved using iteration. This iteration is said to be convergent if the integration method is consistent and zero-stable. A linear multistep method that is at least first-order is called a consistent method. Zero-stability and consistency are necessary for convergence. The converse is also true.

The iteration introduces a second index m .

$$x^{n+1,m+1} = x^n + h \cdot \frac{5}{12} \cdot f^{n+1,m} + h \cdot \frac{8}{12} \cdot f^n - h \cdot \frac{1}{12} \cdot f^{n-1} \quad (6.31)$$

This iteration will converge for an arbitrary initial guess $x^{n+1,0}$ only limited by the step size h . In practice successive iterations are processed unless

$$|x^{n+1,m+1} - x^{n+1,m}| < \varepsilon_{abs} + \varepsilon_{rel} \cdot |x^{n+1,m}| \quad (6.32)$$

The disadvantage for this method is that the number of iterations until it converges is unknown. Alternatively it is possible to use a fixed number of correction steps. A cheap way of providing a good initial guess $x^{n+1,0}$ is using an explicit integration method, e.g. the Adams-Bashford formula of order 3.

$$x^{n+1,0} = x^n + h \cdot \frac{23}{12} \cdot f^n - h \cdot \frac{16}{12} \cdot f^{n-1} + h \cdot \frac{5}{12} \cdot f^{n-2} \quad (6.33)$$

Equation (??) requires no iteration process and can be used to obtain the initial guess. The combination of evaluating a single explicit integration method (the predictor step) in order to provide a good initial guess for the successive evaluation of an implicit method (the corrector step) using iteration is called predictor-corrector method. The motivation using an implicit integration method is its fitness for solving stiff problems. The explicit method (though possibly unstable) is used to provide a good initial guess for the corrector steps.

6.2.1 Order and local truncation error

The order of an integration method results from the truncation error ε_T which is defined as

$$\varepsilon_T = x(t^{n+1}) - x^{n+1} \quad (6.34)$$

meaning the deviation of the exact solution $x(t^{n+1})$ from the approximate solution x^{n+1} obtained by the integration method. For explicit integration methods with $b_{-1} = 0$ the local truncation error ε_{LTE} yields

$$\varepsilon_{LTE} = x(t^{n+1}) - x^{n+1} \quad (6.35)$$

and for implicit integration methods with $b_{-1} \neq 0$ it is

$$\varepsilon_{LTE} \approx x(t^{n+1}) - x^{n+1} \quad (6.36)$$

Going into equation (??) and setting $a_{-1} = -1$ the truncation error is defined as

$$\varepsilon_{LTE} = \sum_{i=-1}^p a_i \cdot x(t^{n-i}) + h \sum_{i=-1}^p b_i \cdot f(x(t^{n-i}), t^{n-i}) \quad (6.37)$$

With the Taylor series expansions

$$x(t^{n+i}) = x(t^n) + \frac{(ih)}{1!} \dot{x}(t^n) + \frac{(ih)^2}{2!} \ddot{x}(t^n) + \dots \quad (6.38)$$

$$f(x(t^{n+i}), t^{n+i}) = \dot{x}(t^{n+i}) = \dot{x}(t^n) + \frac{(ih)}{1!} \ddot{x}(t^n) + \frac{(ih)^2}{2!} \ddot{x}(t^n) + \dots \quad (6.39)$$

the local truncation error as defined by eq. (??) can be written as

$$\varepsilon_{LTE} = C_0 \cdot x(t^n) + C_1 h \cdot \dot{x}(t^n) + C_2 h^2 \cdot \ddot{x}(t^n) + \dots \quad (6.40)$$

The error terms C_0 , C_1 and C_2 in their general form can then be expressed by the following equation.

$$C_q = -\frac{1}{q!} \cdot \sum_{i=-1}^{p-1} a_i \cdot (p-i)^q - \frac{1}{(q-1)!} \sum_{i=-1}^{p-1} b_i \cdot (p-i)^{q-1} \quad (6.41)$$

A linear multistep integration method is of order k if

$$\varepsilon_{LTE} = C_{k+1} \cdot h^{k+1} \cdot x^{(k+1)}(t^n) + O(h^{k+2}) \quad (6.42)$$

The error constant C_{k+1} of an p -step integration method of order k is then defined as

$$C_{k+1} = -\frac{1}{(k+1)!} \cdot \sum_{i=-1}^{p-1} a_i \cdot (p-i)^{k+1} - \frac{1}{k!} \sum_{i=-1}^{p-1} b_i \cdot (p-i)^k \quad (6.43)$$

The practical computation of these error constants is now going to be explained using the Adams-Moulton formula of order 3 given by eq. (??). For this third order method with $a_{-1} = -1$, $a_0 = 1$,

$b_{-1} = 5/12$, $b_0 = 8/12$ and $b_1 = -1/12$ the following values are obtained using eq. (??).

$$C_0 = -\frac{1}{0!} \cdot (-1 \cdot 2^0 + 1 \cdot 1^0) = 0 \quad (6.44)$$

$$C_1 = -\frac{1}{1!} \cdot (-1 \cdot 2^1 + 1 \cdot 1^1) - \frac{1}{0!} \cdot \left(\frac{5}{12} 2^0 + \frac{8}{12} 1^0 - \frac{1}{12} 0^0 \right) = 0 \quad (6.45)$$

$$C_2 = -\frac{1}{2!} \cdot (-1 \cdot 2^2 + 1 \cdot 1^2) - \frac{1}{1!} \cdot \left(\frac{5}{12} 2^1 + \frac{8}{12} 1^1 - \frac{1}{12} 0^1 \right) = 0 \quad (6.46)$$

$$C_3 = -\frac{1}{3!} \cdot (-1 \cdot 2^3 + 1 \cdot 1^3) - \frac{1}{2!} \cdot \left(\frac{5}{12} 2^2 + \frac{8}{12} 1^2 - \frac{1}{12} 0^2 \right) = 0 \quad (6.47)$$

$$C_4 = -\frac{1}{4!} \cdot (-1 \cdot 2^4 + 1 \cdot 1^4) - \frac{1}{3!} \cdot \left(\frac{5}{12} 2^3 + \frac{8}{12} 1^3 - \frac{1}{12} 0^3 \right) = -\frac{1}{24} \quad (6.48)$$

In similar ways it can be verified for each of the discussed linear multistep integration methods that

$$C_p = 0 \quad \forall \quad 0 \leq p \leq k \quad (6.49)$$

The following table summarizes the error constants for the implicit Gear formulae (also called BDF - backward differentiation formulae).

	implicit Gear formulae (BDF)					
steps n	1	2	3	4	5	6
order k	1	2	3	4	5	6
error constant C_{k+1}	$-\frac{1}{2}$	$-\frac{2}{9}$	$-\frac{3}{22}$	$-\frac{12}{125}$	$-\frac{10}{137}$	$-\frac{20}{343}$

The following table summarizes the error constants for the explicit Gear formulae.

	explicit Gear formulae					
steps n	2	3	4	5	6	7
order k	1	2	3	4	5	6
error constant C_{k+1}	+1	+1	+1	+1	+1	+1

The following table summarizes the error constants for the explicit Adams-Bashford formulae.

	explicit Adams-Bashford					
steps n	1	2	3	4	5	6
order k	1	2	3	4	5	6
error constant C_{k+1}	$\frac{1}{2}$	$\frac{5}{12}$	$\frac{3}{8}$	$\frac{251}{720}$	$\frac{95}{288}$	$\frac{19087}{60480}$

The following table summarizes the error constants for the implicit Adams-Moulton formulae.

	implicit Adams-Moulton					
steps n	1	1	2	3	4	5
order k	1	2	3	4	5	6
error constant C_{k+1}	$-\frac{1}{2}$	$-\frac{1}{12}$	$-\frac{1}{24}$	$-\frac{19}{720}$	$-\frac{3}{160}$	$-\frac{863}{60480}$

6.2.2 Milne's estimate

The locale truncation error of the predictor of order k^* may be defined as

$$\varepsilon_{LTE}^* = C_{k^*+1}^* \cdot h^{k^*+1} \cdot x^{(k^*+1)}(t^n) + O(h^{k^*+2}) \quad (6.50)$$

and that of the corresponding corrector method of order k

$$\varepsilon_{LTE} = C_{k+1} \cdot h^{k+1} \cdot x^{(k+1)}(t^n) + O(h^{k+2}) \quad (6.51)$$

If a predictor and a corrector method with same orders $k = k^*$ are used the locale truncation error of the predictor-corrector method yields

$$\varepsilon_{LTE} \approx \frac{C_{k+1}}{C_{k+1}^* - C_{k+1}} \cdot (x^{n+1,m} - x^{n+1,0}) \quad (6.52)$$

This approximation is called Milne's estimate.

6.2.3 Adaptive step-size control

For all numerical integration methods used for the transient analysis of electrical networks the choice of a proper step-size is essential. If the step-size is too large, the results become inaccurate or even completely wrong when the region of absolute stability is left. And if the step-size is too small the calculation requires more time than necessary without raising the accuracy. Usually a chosen initial step-size cannot be used overall the requested time of calculation.

Basically a step-size h is chosen such that

$$\varepsilon_{LTE} < \varepsilon_{abs} + \varepsilon_{rel} \cdot |x^{n+1,m}| \quad (6.53)$$

Forming a step-error quotient

$$q = \frac{\varepsilon_{LTE}}{\varepsilon_{abs} + \varepsilon_{rel} \cdot |x^{n+1,m}|} \quad (6.54)$$

yields the following algorithm for the step-size control. The initial step size h^0 is chosen sufficiently small. After each integration step every step-error quotient gets computed and the largest q_{max} is then checked.

If $q_{max} > 1$, then a reduction of the current step-size is necessary. As new step-size the following expression is used

$$h^n = \left(\frac{\varepsilon}{q_{max}} \right)^{\frac{1}{k+1}} \cdot h^n \quad (6.55)$$

with k denoting the order of the corrector-predictor method and $\varepsilon < 1$ (e.g. ≈ 0.8). If necessary the process must be repeated.

If $q_{max} > 1$, then the calculated value in the current step gets accepted and the new step-size is

$$h^{n+1} = \left(\frac{\varepsilon}{q_{max}} \right)^{\frac{1}{k+1}} \cdot h^n \quad (6.56)$$

6.3 Energy-storage components

As already mentioned it is essential for the transient analysis to consider the energy storing effects of components. The following section describes how the modified nodal analysis can be used to take this into account.

6.3.1 Capacitor

The relation between current and voltage in terms of a differential equation for an ideal capacitor is

$$I_C(t) = C \cdot \frac{dV_C}{dt} \quad (6.57)$$

With

$$\frac{I_C(V, t)}{C} = \frac{dV_C}{dt} = f(x, t) \quad (6.58)$$

the discussed integration formulas (??), (??), (??) and (??) can be applied to the problem. Rewriting them in an explicit form regarding the next integration current results in

$$I_C^{n+1} = \frac{C}{h^n} V_C^{n+1} - \frac{C}{h^n} V_C^n \quad (\text{backward Euler}) \quad (6.59)$$

$$I_C^{n+1} = \frac{2C}{h^n} V_C^{n+1} - \frac{2C}{h^n} V_C^n - I_C^n \quad (\text{trapezoidal}) \quad (6.60)$$

$$I_C^{n+1} = \frac{C}{b_{-1} \cdot h^n} V_C^{n+1} - \frac{a_0 \cdot C}{b_{-1} \cdot h^n} V_C^n - \frac{a_1 \cdot C}{b_{-1} \cdot h^n} V_C^{n-1} - \dots - \frac{a_{k-1} \cdot C}{b_{-1} \cdot h^n} V_C^{n-k+1} \quad (6.61)$$

$$I_C^{n+1} = \underbrace{\frac{C}{b_{-1} \cdot h^n} V_C^{n+1}}_{g_{eq}} - \underbrace{\frac{a_0 \cdot C}{b_{-1} \cdot h^n} V_C^n - \frac{b_0}{b_{-1}} I_C^n - \frac{b_1}{b_{-1}} I_C^{n-1} - \dots - \frac{b_{k-2}}{b_{-1}} I_C^{n-k+2}}_{I_{eq}} \quad (6.62)$$

Each of these equations can be rewritten as

$$I_C^{n+1} = g_{eq} \cdot V_C^{n+1} + I_{eq} \quad (6.63)$$

which leads to the following companion model representing a current source with its accompanied internal resistance.

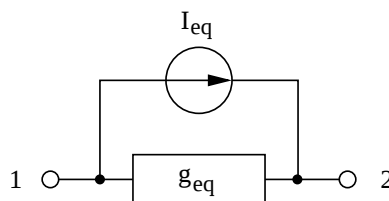


Figure 6.5: companion equivalent circuit of a capacitor during transient analysis

Thus the complete MNA matrix equation for an ideal capacitance writes as follows.

$$\begin{bmatrix} +g_{eq} & -g_{eq} \\ -g_{eq} & +g_{eq} \end{bmatrix} \cdot \begin{bmatrix} V_1^{n+1} \\ V_2^{n+1} \end{bmatrix} = \begin{bmatrix} -I_{eq} \\ +I_{eq} \end{bmatrix} \quad (6.64)$$

6.3.2 Inductor

The relation between current and voltage in terms of a differential equation for an ideal inductor can be written as

$$V_L(t) = L \cdot \frac{dI_L}{dt} \quad (6.65)$$

With

$$\frac{V_L(I, t)}{L} = \frac{dI_L}{dt} = f(x, t) \quad (6.66)$$

the discussed integration formulas (??), (??), (??) and (??) can be applied to the problem. Rewriting them in an explicit form regarding the next integration voltage results in

$$V_L^{n+1} = \frac{L}{h^n} I_L^{n+1} - \frac{L}{h^n} I_L^n \quad (6.67)$$

$$V_L^{n+1} = \frac{2L}{h^n} I_L^{n+1} - \frac{2L}{h^n} I_L^n - V_L^n \quad (6.68)$$

$$V_L^{n+1} = \frac{L}{b_{-1} \cdot h^n} I_L^{n+1} - \frac{a_0 \cdot L}{b_{-1} \cdot h^n} I_L^n - \frac{a_1 \cdot L}{b_{-1} \cdot h^n} I_L^{n-1} - \dots - \frac{a_{k-1} \cdot L}{b_{-1} \cdot h^n} I_L^{n-k+1} \quad (6.69)$$

$$V_L^{n+1} = \underbrace{\frac{L}{b_{-1} \cdot h^n} I_L^{n+1}}_{r_{eq}} - \underbrace{\frac{a_0 \cdot L}{b_{-1} \cdot h^n} I_L^n - \frac{b_0}{b_{-1}} V_L^n - \frac{b_1}{b_{-1}} V_L^{n-1} - \dots - \frac{b_{k-2}}{b_{-1}} V_L^{n-k+2}}_{V_{eq}} \quad (6.70)$$

Each of these equations can be rewritten as

$$V_L^{n+1} = r_{eq} \cdot I_L^{n+1} + V_{eq} \quad (6.71)$$

which leads to the following companion model representing a voltage source with its accompanied internal resistance.

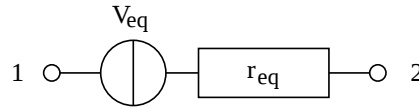


Figure 6.6: companion equivalent circuit of an inductor during transient analysis

Thus the complete MNA matrix equation for an ideal inductor writes as follows.

$$\begin{bmatrix} 0 & 0 & +1 \\ 0 & 0 & -1 \\ +1 & -1 & -r_{eq} \end{bmatrix} \cdot \begin{bmatrix} V_1^{n+1} \\ V_2^{n+1} \\ I_L^{n+1} \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ V_{eq} \end{bmatrix} \quad (6.72)$$

It is also possible to model the ideal inductor as a current source with an internal resistance which would yield a similar equivalent circuit as for the capacitor. But with the proposed model it is possible to use alike computation schemes for capacitors and inductors. Charges become fluxes, capacitances become inductances and finally voltages become currents and the other way around. Everything else (especially the coefficients in the integration formulas) can be reused.

6.3.3 Coupled Inductors

In a non-ideal transformer, there are two (or more) coupled inductors. The model for the transient simulation is not very different from the one of a single inductor. In addition to each coil, the mutual inductance has to be counted for.

$$V_{L1} = L_1 \cdot \frac{dI_{L1}}{dt} + M_{12} \cdot \frac{dI_{L2}}{dt} + I_{L1} \cdot R_1 \quad (6.73)$$

$$\text{with } M_{12} = k \cdot \sqrt{L_1 \cdot L_2} \quad (6.74)$$

$$\text{and } R_1 \text{ ohmic resistance of coil 1} \quad (6.75)$$

So it is:

$$V_{L1}^{n+1} = r_{eq11} \cdot I_{L1}^{n+1} + r_{eq12} \cdot I_{L2}^{n+1} + V_{eq}(I_{L1}^n, I_{L2}^n, \dots) \quad (6.76)$$

Note that r_{eq11} includes the ohmic resistance R_1 . For backward Euler, it therefore follows:

$$V_{L1}^{n+1} = \underbrace{\left(\frac{L_1}{h^n} + R_1\right)}_{r_{eq11}} \cdot I_{L1}^{n+1} + \underbrace{\frac{k \cdot \sqrt{L_1 \cdot L_2}}{h^n}}_{r_{eq12}} \cdot I_{L2}^{n+1} - \underbrace{\left(\frac{L_1}{h^n} + R_1\right) \cdot I_{L1}^n - \frac{k \cdot \sqrt{L_1 \cdot L_2}}{h^n} \cdot I_{L2}^n}_{V_{eq1}} \quad (6.77)$$

The voltage across the secondary coil V_{L2}^{n+1} goes likewise by just changing the indices. Finally, the MNA matrix writes (port numbers are according to figure ??):

$$\begin{bmatrix} 0 & 0 & 0 & 0 & +1 & 0 \\ 0 & 0 & 0 & 0 & 0 & +1 \\ 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ +1 & 0 & 0 & -1 & -r_{eq11} & -r_{eq12} \\ 0 & +1 & -1 & 0 & -r_{eq21} & -r_{eq22} \end{bmatrix} \cdot \begin{bmatrix} V_1^{n+1} \\ V_2^{n+1} \\ V_3^{n+1} \\ V_4^{n+1} \\ I_{L1}^{n+1} \\ I_{L2}^{n+1} \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ V_{eq1} \\ V_{eq2} \end{bmatrix} \quad (6.78)$$

These equations can also give an idea on how to model more than two coupled inductors. For three coupled inductors, the voltage across coil 1 writes:

$$V_{L1} = L_1 \cdot \frac{dI_{L1}}{dt} + M_{12} \cdot \frac{dI_{L2}}{dt} + M_{13} \cdot \frac{dI_{L3}}{dt} + I_{L1} \cdot R_1 \quad (6.79)$$

$$V_{L2} = L_2 \cdot \frac{dI_{L2}}{dt} + M_{12} \cdot \frac{dI_{L1}}{dt} + M_{23} \cdot \frac{dI_{L3}}{dt} + I_{L2} \cdot R_2 \quad (6.80)$$

$$V_{L3} = L_3 \cdot \frac{dI_{L3}}{dt} + M_{13} \cdot \frac{dI_{L1}}{dt} + M_{23} \cdot \frac{dI_{L2}}{dt} + I_{L3} \cdot R_3 \quad (6.81)$$

$$\text{with } M_{12} = k_{12} \cdot \sqrt{L_1 \cdot L_2} \quad (6.82)$$

$$\text{and } M_{13} = k_{13} \cdot \sqrt{L_1 \cdot L_3} \quad (6.83)$$

$$\text{and } M_{23} = k_{23} \cdot \sqrt{L_2 \cdot L_3} \quad (6.84)$$

This can be easily extended to an arbitrary number of coupled inductors.

6.3.4 Depletion Capacitance

For non-constant capacitances, especially depletion capacitance used in non-linear devices, instead of eq. (??) the following equation holds.

$$I_C(t) = \frac{dQ}{dt} \quad (6.85)$$

With

$$dQ = C \cdot dV_C \quad \text{and} \quad \left. \frac{dV_C}{dQ} \right|_{Q^{(m)}} = \frac{1}{C} \quad (6.86)$$

equation (??) can be written as

$$V_C^{(m+1)} = V_C^{(m)} - \frac{Q(V_C^{(m)})}{C^{(m)}} \quad (6.87)$$

$$\Rightarrow (V_C^{(m+1)} - V_C^{(m)}) \cdot C^{(m)} = -Q^{(m)} \quad (6.88)$$

yielding a similar iterative algorithm as already used for the non-linear DC analysis described in section ?? on page ?. The indices $^{(m)}$ indicated the m -th Newton-Raphson iteration. With this knowledge at hand it is possible to rewrite the explicit formula for the backward Euler integration (??), i.e. the next iteration step Q^{m+1} is replaced by the Newton-Raphson formula as follows.

$$\begin{aligned} I_C^{n+1,m+1} &= \frac{Q^{n+1,m+1} - Q^n}{h^n} \\ &= \frac{1}{h^n} \cdot (Q^{n+1,m} + C^{n+1,m} \cdot (V_C^{n+1,m+1} - V_C^{n+1,m}) - Q^n) \end{aligned} \quad (6.89)$$

The double indices now indicate the n -th integration step and the m -th Newton-Raphson iteration. The same can be done for the other integration formulas and results also in a similar equivalent companion model as shown in fig. ??.

The capacitance C and the charge Q within the above equations is computed according to the appropriate (non-linear) model formulations.

$$\begin{aligned} Q &= C_0 \cdot \left(+ \frac{V_J \cdot (1 - (1 - F)^{1-M})}{1 - M} \right. \\ &\quad + \frac{1 - F \cdot (1 + M)}{(1 - F)^{1+M}} \cdot (V_C - F \cdot V_J) \\ &\quad \left. + \frac{M}{2 \cdot V_J \cdot (1 - F)^{1+M}} \cdot (V_C^2 - F^2 \cdot V_J^2) \right) \end{aligned} \quad (6.90)$$

and

$$C = \frac{dQ}{dV_C} = \frac{C_0}{(1 - F)^M} \cdot \left(1 + \frac{M \cdot (V_C - F \cdot V_J)}{V_J \cdot (1 - F)} \right) \quad (6.91)$$

for a depletion capacitance with $V_C > F \cdot V_J$ and for $V_C < F \cdot V_J$ those capacitances yield

$$Q = \frac{C_0 \cdot V_J}{1 - M} \cdot \left(1 - \left(1 - \frac{V_C}{V_J} \right)^{1-M} \right) \quad (6.92)$$

with

$$C = \frac{dQ}{dV_C} = C_0 \cdot \left(1 - \frac{V_C}{V_J} \right)^{-M} \quad (6.93)$$

6.3.5 Diffusion Capacitance

The current through a diffusion capacitance can be approximated by

$$I_C(t) = \tau_D \frac{dI_D}{dt} \quad (6.94)$$

whence τ_D specifies the transit time through a pn-junction. The above formula can be rewritten as

$$I_C(t) = \tau_D \frac{dI_D}{dV_C} \cdot \frac{dV_C}{dt} = \tau_D \cdot g_D \cdot \frac{dV_C}{dt} \quad (6.95)$$

which means that eq. (??) can be used here, too. Also the formulas for the other integration methods can be easily rewritten and the equivalent companion model shown in fig. ?? is valid as well.

The capacitance C and the charge Q for a diffusion capacitance of a pn-junction according to the most model formulations write as follows.

$$Q = \tau_D \cdot I_D \quad (6.96)$$

$$C = \frac{dQ}{dV_C} = \tau_D \cdot g_D \quad (6.97)$$

6.3.6 MOS Gate Capacitances

The MOS gate capacitances are not constant values with respect to voltages (see section ?? on page ??). The capacitance values can best be described by the incremental capacitance:

$$C(V) = \frac{dQ(V)}{dV} \quad (6.98)$$

where $Q(V)$ is the charge on the capacitor and V is the voltage across the capacitor.

The formula for calculating the differential is difficult to derive (because not given in the Meyer capacitance model). Furthermore, the voltage is required as the accumulated capacitance over time. The timewise charge formula is:

$$Q(V) = \int_0^V C(V) \cdot dV \quad (6.99)$$

And for small intervalls:

$$Q(V) = \int_{V^n}^{V^{n+1}} C(V) \cdot dV \quad (6.100)$$

The integral has been approximated in SPICE by:

$$Q^{n+1} = (V^{n+1} - V^n) \cdot \frac{C(V^{n+1}) + C(V^n)}{2} \quad (6.101)$$

This last formula is the trapezoidal rule for integration over two points. The charge is approximated as the average capacitance times the change in voltage. If the capacitance is nonlinear, this approximation can be in error. To estimate the charge accurately, use Simpson's numerical integration rule. This method provides charge conservation control.

$$Q^{n+1} = (V^{n+1} - V^n) \cdot \frac{C(V^{n+1}) + 4C(V^n) + C(V^{n-1}))}{6} \quad (6.102)$$

6.4 Special time-domain models

6.4.1 AM modulated AC source

An AC voltage source in the time-domain is characterized by its frequency f , the initial phase ϕ and the amplitude A . During amplitude modulation the modulation level M must be considered. The output voltage of the source is determined by the following equation.

$$V_1(t) - V_2(t) = (1 + M \cdot V_3(t)) \cdot A \cdot \sin(\omega \cdot t + \phi) \quad (6.103)$$

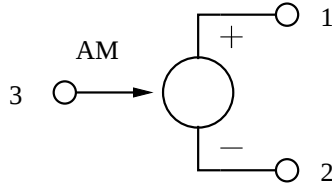


Figure 6.7: AM modulated AC source

The appropriate MNA matrix entries during the transient analysis describing a simple linear operation can be written as

$$\begin{bmatrix} \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & 0 \\ 1 & -1 & -M \cdot A \cdot \sin(\omega \cdot t + \phi) & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ J_1(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ I_3(t) \\ A \cdot \sin(\omega \cdot t + \phi) \end{bmatrix} \quad (6.104)$$

6.4.2 PM modulated AC source

The phase modulated AC source is also characterized by the frequency f , the amplitude A and by an initial phase ϕ . The output voltage in the time-domain is determined by the following equation

$$V_1(t) - V_2(t) = A \cdot \sin(\omega \cdot t + \phi + M \cdot V_3(t)) \quad (6.105)$$

whereas M denotes the modulation index and V_3 the modulating voltage.

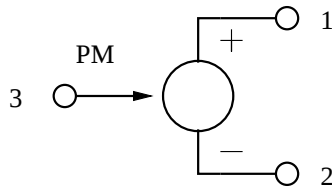


Figure 6.8: PM modulated AC source

The component is non-linear in the frequency- as well in the time-domain. In order to prepare the source for subsequent Newton-Raphson iterations the derivative

$$g = \frac{\partial(V_1 - V_2)}{\partial V_3} = M \cdot A \cdot \cos(\omega \cdot t + \phi + M \cdot V_3) \quad (6.106)$$

is required. With this at hand the MNA matrix entries of the PM modulated AC voltage source during the transient analysis can be written as

$$\begin{bmatrix} \cdot & \cdot & \cdot & +1 \\ \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & 0 \\ +1 & -1 & g & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ J_1(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ I_3(t) \\ g \cdot V_3 - A \cdot \sin(\omega \cdot t + \phi + M \cdot V_3) \end{bmatrix} \quad (6.107)$$

6.5 Components defined in the frequency domain

The time-domain simulation of components defined in the frequency-domain can be performed using an inverse Fourier transformation of the Y-parameters of the component (giving the impulse response) and an adjacent convolution with the prior node voltages (or branch currents) of the component.

This requires a memory of the node voltages and branch currents for each component defined in the frequency-domain. During a transient simulation the time steps are not equidistant and the maximum required memory length T_{end} of a component may not correspond with the time grid produced by the time step control (see section ?? on page ??) of the transient simulation. That is why an interpolation of exact values (voltage or current) at a given point in time is necessary.

Components defined in the frequency-domain can be divided into two major classes.

- Components with frequency-independent (non-dispersive) delay times and with or without constant losses.
- Components with frequency-dependent (dispersive) delay times and losses.

6.5.1 Components with frequency-independent delay times

Components with constant delay times are a special case. The impulse response corresponds to the node voltages and/or branch currents at some prior point in time optionally multiplied with a constant loss factor.

Voltage controlled current source

With no constant delay time the MNA matrix entries of a voltage controlled current source is determined by the following equations according to the node numbering in fig. ?? on page ??.

$$I_2 = -I_3 = G \cdot (V_1 - V_4) \quad (6.108)$$

The equations yield the following MNA entries during the transient analysis.

$$\begin{bmatrix} 0 & 0 & 0 & 0 \\ +G & 0 & 0 & -G \\ -G & 0 & 0 & +G \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \end{bmatrix} \quad (6.109)$$

With a constant delay time τ eq. (??) rewrites as

$$I_2(t) = -I_3(t) = G \cdot (V_1(t - \tau) - V_4(t - \tau)) \quad (6.110)$$

which yields the following MNA entries during the transient analysis.

$$\begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ V_4(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ -G \cdot (V_1(t - \tau) - V_4(t - \tau)) \\ +G \cdot (V_1(t - \tau) - V_4(t - \tau)) \\ I_4(t) \end{bmatrix} \quad (6.111)$$

Voltage controlled voltage source

The MNA matrix entries of a voltage controlled voltage source are determined by the following characteristic equation according to the node numbering in fig. ?? on page ??.

$$V_2 - V_3 = G \cdot (V_4 - V_1) \quad (6.112)$$

This equation yields the following augmented MNA matrix entries with a single extra branch equation.

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 \\ G & -1 & 1 & -G & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ J_1 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} \quad (6.113)$$

When considering an additional constant time delay τ eq. (??) must be rewritten as

$$V_2(t) - V_3(t) = G \cdot (V_4(t - \tau) - V_1(t - \tau)) \quad (6.114)$$

This representation requires a change of the MNA matrix entries which now yield the following matrix equation.

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & -1 & 1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ V_4(t) \\ J_1(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ I_3(t) \\ I_4(t) \\ G \cdot (V_4(t - \tau) - V_1(t - \tau)) \end{bmatrix} \quad (6.115)$$

Current controlled current source

With no time delay the MNA matrix entries of a current controlled current source are determined by the following equations according to the node numbering in fig. ?? on page ??.

$$I_2 = -I_3 = G \cdot I_1 = -G \cdot I_4 \quad (6.116)$$

$$V_1 = V_4 \quad (6.117)$$

These equations yield the following MNA matrix entries using a single extra branch equation.

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 1/G \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & -1/G \\ 1 & 0 & 0 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ J_1 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} \quad (6.118)$$

When additional considering a constant delay time τ eq. (??) must be rewritten as

$$I_2(t) = -I_3(t) = G \cdot I_1(t - \tau) = -G \cdot I_4(t - \tau) \quad (6.119)$$

Thus the MNA matrix entries change as well yielding

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 \\ 1 & 0 & 0 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ V_4(t) \\ J_1(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ -G \cdot J_1(t - \tau) \\ +G \cdot J_1(t - \tau) \\ I_4(t) \\ 0 \end{bmatrix} \quad (6.120)$$

Current controlled voltage source

The MNA matrix entries for a current controlled voltage source are determined by the following characteristic equations according to the node numbering in fig. ?? on page ??.

$$V_2 - V_3 = G \cdot I_2 = -G \cdot I_3 \quad (6.121)$$

$$V_1 = V_4 \quad (6.122)$$

These equations yield the following MNA matrix entries.

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 1 & -1 & 0 & G & 0 \\ 1 & 0 & 0 & -1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ J_1 \\ J_2 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \\ 0 \end{bmatrix} \quad (6.123)$$

With an additional time delay τ between the input current and the output voltage eq. (??) rewrites as

$$V_2(t) - V_3(t) = G \cdot I_2(t - \tau) = -G \cdot I_3(t - \tau) \quad (6.124)$$

Due to the additional time delay the MNA matrix entries must be rewritten as follows

$$\begin{bmatrix} 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 1 & -1 & 0 & 0 & 0 \\ 1 & 0 & 0 & -1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ V_4(t) \\ J_1(t) \\ J_2(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ I_3(t) \\ I_4(t) \\ 0 \\ G \cdot J_1(t - \tau) \end{bmatrix} \quad (6.125)$$

Ideal transmission line

The A-parameters of a transmission line (see eq (??) on page ??) are defined in the frequency domain. The equation system formed by these parameters write as

$$\text{I. } V_1 = V_2 \cdot \cosh(\gamma \cdot l) + I_2 \cdot Z_L \cdot \sinh(\gamma \cdot l) \quad (6.126)$$

$$\text{II. } I_1 = V_2 \cdot \frac{1}{Z_L} \sinh(\gamma \cdot l) + I_2 \cdot \cosh(\gamma \cdot l) \quad (6.127)$$

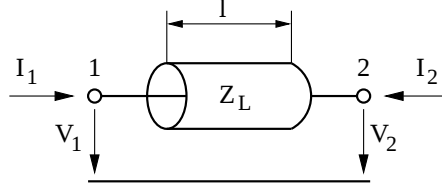


Figure 6.9: ideal transmission line

Applying $I + Z_L \cdot II$ and $I - Z_L \cdot II$ to the above equation system and using the following transformations

$$\cosh x + \sinh x = \frac{e^x + e^{-x}}{2} + \frac{e^x - e^{-x}}{2} = e^x \quad (6.128)$$

$$\cosh x - \sinh x = \frac{e^x + e^{-x}}{2} - \frac{e^x - e^{-x}}{2} = e^{-x} \quad (6.129)$$

yields

$$V_1 = V_2 \cdot e^{-\gamma \cdot l} + Z_L \cdot (I_1 + I_2 \cdot e^{-\gamma \cdot l}) \quad (6.130)$$

$$V_2 = V_1 \cdot e^{-\gamma \cdot l} + Z_L \cdot (I_2 + I_1 \cdot e^{-\gamma \cdot l}) \quad (6.131)$$

whereas γ denotes the propagation constant $\alpha + j\beta$, l the length of the transmission line and Z_L the line impedance.

These equations can be transformed from the frequency domain into the time domain using the inverse Fourier transformation. The frequency independent loss $\alpha \neq f(\omega)$ gives the constant factor

$$A = e^{-\alpha \cdot l} \quad (6.132)$$

The only remaining frequency dependent term is

$$e^{-j\beta \cdot l} = e^{-j\omega \cdot \tau} \quad \text{with} \quad \beta = \frac{\omega}{v_{ph}} = \frac{\omega}{c_0} = \frac{\omega \cdot \tau}{l} \quad (6.133)$$

which yields the following transformation

$$f(\omega) \cdot e^{-\gamma \cdot l} = A \cdot f(\omega) \cdot e^{-j\omega \cdot \tau} \iff A \cdot f(t - \tau) \quad (6.134)$$

All the presented time-domain models with a frequency-independent delay time are based on this simple transformation. It can be applied since the phase velocity $v_{ph} \neq f(\omega)$ is not a function of the frequency. This is true for all non-dispersive transmission media, e.g. air or vacuum. The given transformation can now be applied to the eq. (??) and eq. (??) defined in the frequency-domain to obtain equations in the time-domain.

The length T_{end} of the memory needed by the ideal transmission line can be easily computed by

$$T_{end} = \tau = \frac{l}{v_{ph}} = \frac{l}{c_0} \quad (6.135)$$

whereas c_0 denotes the speed of light in free space (since there is no dielectric involved during transmission) and l the physical length of the transmission line.

The MNA matrix for a lossy (or lossless with $\alpha = 0$) transmission line during the transient analysis is augmented by two new rows and columns in order to consider the following branch equations.

$$V_1(t) = Z_L \cdot I_1(t) + A \cdot (Z_L \cdot I_2(t - \tau) + V_2(t - \tau)) \quad (6.136)$$

$$V_2(t) = Z_L \cdot I_2(t) + A \cdot (Z_L \cdot I_1(t - \tau) + V_1(t - \tau)) \quad (6.137)$$

Thus the MNA matrix entries can be written as

$$\begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & -Z_L & 0 \\ 0 & 1 & 0 & -Z_L \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ J_1(t) \\ J_2(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ A \cdot (V_2(t - \tau) + Z_L \cdot J_2(t - \tau)) \\ A \cdot (V_1(t - \tau) + Z_L \cdot J_1(t - \tau)) \end{bmatrix} \quad (6.138)$$

with A denoting the loss factor derived from the constant (and frequency independent) line attenuation α and the transmission line length l .

$$A = e^{-\frac{\alpha}{2} \cdot l} \quad (6.139)$$

Ideal 4-terminal transmission line

The ideal 4-terminal transmission line is a two-port as well. It differs from the 2-terminal line as shown in fig. ?? in two new node voltages and branch currents.

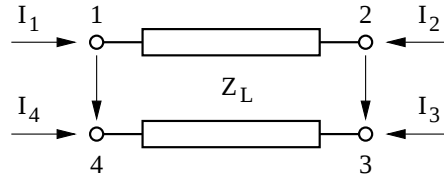


Figure 6.10: ideal 4-terminal transmission line

The differential mode of the ideal 4-terminal transmission line can be modeled by modifying the branch eqs. (??) and (??) of the 2-terminal line which yields

$$V_1(t) - V_4(t) = Z_L \cdot I_1(t) + A \cdot (Z_L \cdot I_2(t - \tau) + V_2(t - \tau) - V_3(t - \tau)) \quad (6.140)$$

$$V_2(t) - V_3(t) = Z_L \cdot I_2(t) + A \cdot (Z_L \cdot I_1(t - \tau) + V_1(t - \tau) - V_4(t - \tau)) \quad (6.141)$$

Two more conventions must be introduced

$$I_1(t) = -I_4(t) \quad (6.142)$$

$$I_2(t) = -I_3(t) \quad (6.143)$$

which is valid for the differential mode (i.e. the odd mode) of the transmission line and represents a kind of current mirror on each transmission line port.

According to these consideration the MNA matrix entries during transient analysis are

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & 1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 & 1 \\ \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ 1 & 0 & 0 & -1 & -Z_L & 0 \\ 0 & 1 & -1 & 0 & 0 & -Z_L \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ V_3(t) \\ V_4(t) \\ J_1(t) \\ J_2(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ I_3(t) \\ I_4(t) \\ A \cdot (V_2(t-\tau) - V_3(t-\tau) + Z_L \cdot J_2(t-\tau)) \\ A \cdot (V_1(t-\tau) - V_4(t-\tau) + Z_L \cdot J_1(t-\tau)) \end{bmatrix} \quad (6.144)$$

Logical devices

The analogue models of logical (digital) components explained in section ?? on page ?? do not include delay times. With a constant delay time τ the determining equations for the logical components yield

$$u_{out}(t) = f(V_{in,1}(t-\tau), V_{in,2}(t-\tau), \dots) \quad (6.145)$$

With the prior node voltages $V_{in,n}(t-\tau)$ known the MNA matrix entries in eq. (??) can be rewritten as

$$\begin{bmatrix} \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & 0 \\ 1 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_{out}(t) \\ V_{in,1}(t) \\ V_{in,2}(t) \\ I_{out}(t) \end{bmatrix} = \begin{bmatrix} I_0(t) \\ I_1(t) \\ I_2(t) \\ u_{out}(t) \end{bmatrix} \quad (6.146)$$

during the transient analysis. The components now appear to be simple linear components. The derivatives are not anymore necessary for the Newton-Raphson iterations. This happens to be because the output voltage does not depend directly on the input voltage(s) at exactly the same time point.

6.5.2 Components with frequency-dependent delay times and losses

In the general case a component with P ports which is defined in the frequency-domain can be represented by the following matrix equation.

$$\begin{bmatrix} Y_{11} & Y_{12} & \dots & Y_{1P} \\ Y_{21} & Y_{22} & & Y_{2P} \\ \vdots & & \ddots & \vdots \\ Y_{P1} & Y_{P2} & \dots & Y_{PP} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ \vdots \\ V_P \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ \vdots \\ I_P \end{bmatrix} \quad (6.147)$$

This matrix representation is the MNA representation during the AC analysis. With no specific time-domain model at hand the equation

$$[Y(j\omega)] \cdot [V(j\omega)] = [I(j\omega)] \quad (6.148)$$

must be transformed into the time-domain using a Fourier transformation.

The convolution integral

The multiplication in the frequency-domain is equivalent to a convolution in the time-domain after the transformation. It yields the following matrix equation

$$[H(t)] * [V(t)] = [I(t)] \quad (6.149)$$

whereas $H(t)$ is the impulse response based on the frequency-domain model and the $*$ operator denotes the convolution integral

$$H(t) * V(t) = \int_{-\infty}^{+\infty} H(\tau) \cdot V(t - \tau) d\tau \quad (6.150)$$

The lower bound of the given integral is set to zero since both the impulse response as well as the node voltages are meant to deliver no contribution to the integral. Otherwise the circuit appears to be unphysical. The upper limit should be bound to a maximum impulse response time T_{end}

$$H(t) * V(t) = \int_0^{T_{end}} H(\tau) \cdot V(t - \tau) d\tau \quad (6.151)$$

with

$$H(\tau) = 0 \quad \forall \quad \tau > T_{end} \quad (6.152)$$

Since there is no analytic representation for the impulse response as well as for the node voltages eq. (??) must be rewritten to

$$H(n \cdot \Delta t) * V(n \cdot \Delta t) = \sum_{k=0}^{N-1} H(k \cdot \Delta t) \cdot V((n - k) \cdot \Delta t) \quad (6.153)$$

with

$$\Delta t = \frac{T_{end}}{N} \quad (6.154)$$

whereas N denotes the number of samples to be used during numerical convolution. Using the current time step $t = n \cdot \Delta t$ it is possible to express eq. (??) as

$$I(t) = H(0) \cdot V(t) + \underbrace{\sum_{k=1}^{N-1} H(k \cdot \Delta t) \cdot V(t - k \cdot \Delta t)}_{I_{eq}} \quad (6.155)$$

With $G = H(0)$ the resulting MNA matrix equation during the transient analysis gets

$$[G] \cdot [V(t)] = [I(t)] - [I_{eq}] \quad (6.156)$$

This means, the component defined in the frequency-domain can be expressed with an equivalent DC admittance G and additional independent current sources in the time-domain. Each independent current source at node r delivers the following current

$$I_{eq_r} = \sum_{c=1}^P \sum_{k=1}^{N-1} H_{rc}(k \cdot \Delta t) \cdot V_c(t - k \cdot \Delta t) \quad (6.157)$$

whereas V_c denotes the node voltage at node c at some prior time and H_{rc} the impulse response of the component based on the frequency-domain representation. The MNA matrix equation during transient analysis can thus be written as

$$\begin{bmatrix} G_{11} & G_{12} & \dots & G_{1P} \\ G_{21} & G_{22} & & G_{2P} \\ \vdots & & \ddots & \vdots \\ G_{P1} & G_{P2} & \dots & G_{PP} \end{bmatrix} \cdot \begin{bmatrix} V_1(t) \\ V_2(t) \\ \vdots \\ V_P(t) \end{bmatrix} = \begin{bmatrix} I_1(t) \\ I_2(t) \\ \vdots \\ I_P(t) \end{bmatrix} - \begin{bmatrix} I_{eq_1} \\ I_{eq_2} \\ \vdots \\ I_{eq_P} \end{bmatrix} \quad (6.158)$$

Frequency- to time-domain transformation

With the number of samples N being a power of two it is possible to use the Inverse Fast Fourier Transformation (IFFT). The transformation to be performed is

$$Y(j\omega) \Leftrightarrow H(t) \quad (6.159)$$

The maximum impulse response time of the component is specified by T_{end} requiring the following transformation pairs.

$$Y(j\omega_i) \Leftrightarrow H(t_i) \quad \text{with} \quad i = 0, 1, 2, \dots, N-1 \quad (6.160)$$

with

$$t_i = 0, \Delta t, 2 \cdot \Delta t, \dots, (N-1) \cdot \Delta t \quad (6.161)$$

$$\omega_i = 0, \frac{1}{T_{end}}, \frac{1}{T_{end}}, \dots, \frac{N/2}{T_{end}} \quad (6.162)$$

The frequency samples in eq. (??) indicate that only half the values are required to obtain the appropriate impulse response. This is because the impulse response $H(t)$ is real valued and that is why

$$Y(j\omega) = Y^*(-j\omega) \quad (6.163)$$

The maximum frequency considered is determined by the maximum impulse response time T_{end} and the number of time samples N .

$$f_{max} = \frac{N/2}{2\pi \cdot T_{end}} = \frac{1}{4\pi \cdot \Delta t} \quad (6.164)$$

It could prove useful to weight the Y-parameter samples in the frequency-domain by multiplying them with an appropriate windowing function (e.g. Kaiser-Bessel).

Implementation considerations

For the method presented the Y-parameters of a component must be finite for $f \rightarrow 0$ as well as for $f \rightarrow f_{max}$. To obtain $G = H(0)$ the Y-parameters at $f = 0$ are required. This cannot be ensured for the general case (e.g. for an ideal inductor).

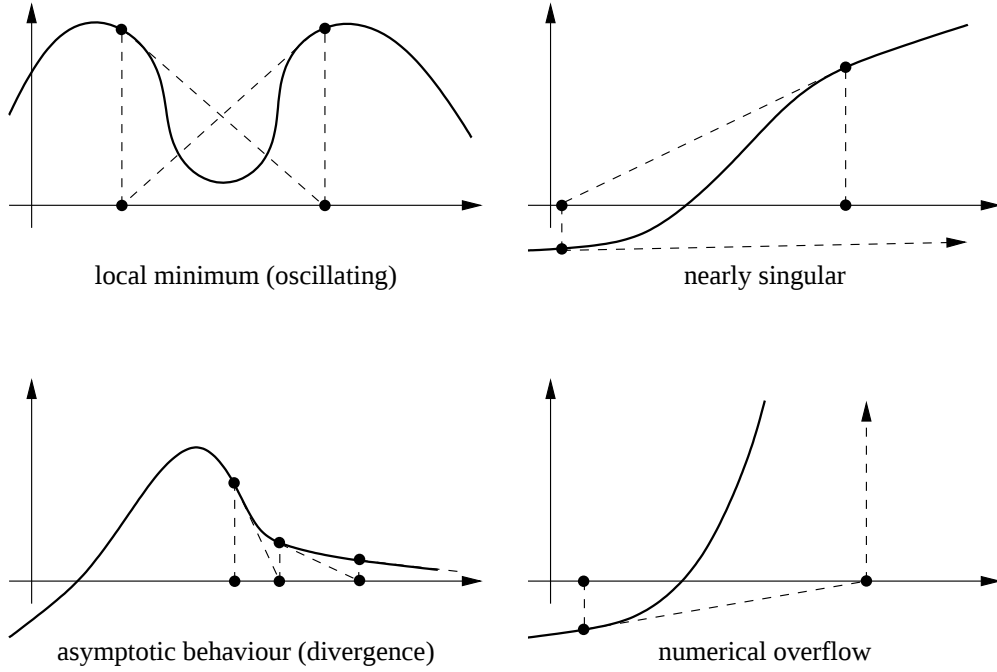
6.6 Convergence

Similar to the DC analysis convergence problems occur during the transient analysis (see section ?? on page ??) as well. In order to improve the overall convergence behaviour it is possible to improve the models on the one hand and/or to improve the algorithms on the other hand.

The implications during Newton-Raphson iterations solving the linear equation system

$$[A(x^k)] \cdot [x^{k+1}] = [z(x^k)] \quad (6.165)$$

are continuous device model equations (with continuous derivatives as well), floating nodes (make the Jacobian matrix A singular) and the initial guess x^0 . The convergence problems which in fact occur are local minimums causing the matrix A to be singular, nearly singular matrices and overflow problems.



6.6.1 Limiting schemes

The modified (damped) Newton-Raphson schemes are based on the limitation of the solution vector x^k in each iteration.

$$x^{k+1} = x^k + \alpha \cdot \Delta x^{k+1} \quad \text{with} \quad \Delta x^{k+1} = x^{k+1} - x^k \quad (6.166)$$

One possibility to choose a value for $\alpha \in [0, 1]$ is

$$\alpha = \frac{\gamma}{\|\Delta x^{k+1}\|_\infty} \quad (6.167)$$

This is a heuristic and does not ensure global convergence, but it can help solving some of the discussed problems. Another possibility is to pick a value α^k which minimizes the L_2 norm of the right hand side vector. This method performs a one-dimensional (line) search into Newton direction and guarantees global convergence.

$$x^{k+1} = x^k + \alpha^k \cdot \Delta x^{k+1} \quad \text{with an } \alpha^k \text{ which minimizes} \quad \|z(x^k + \alpha^k \cdot \Delta x^{k+1})\|_2 \quad (6.168)$$

The one remaining problem about that line search method for convergence improvement is its iteration into local minimums where the Jacobian matrix is singular. The damped Newton-Raphson method “pushes” the matrix into singularity as depicted in fig. ??.

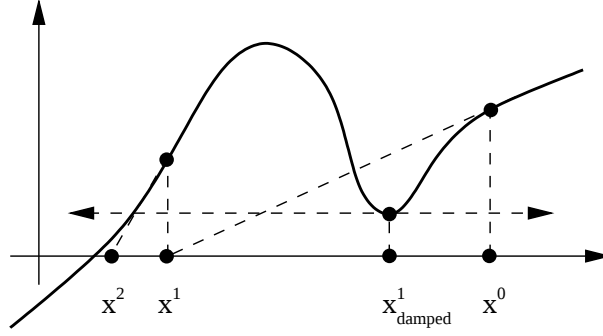


Figure 6.11: singular Jacobian problem

6.6.2 Continuation schemes

The basic idea behind this Newton-Raphson modification is to generate a sequence of problems such that a problem is a good initial guess for the following one, because Newton basically converges given a close initial guess.

The template algorithm for this modification is to solve the equation system

$$[A] \cdot [x] - [z] = 0 \quad (6.169)$$

$$F(x(\lambda), \lambda) = 0 \quad (6.170)$$

with the parameter $\lambda \in [0, 1]$ given that $x(\lambda)$ is sufficiently smooth. $F(x(0), 0)$ starts the continuation and $F(x(1), 1)$ ends the continuation. The algorithm outline is as follows: First solve the problem $F(x(0), 0)$, e.g. set $\lambda = \Delta\lambda = 0.01$ and try to solve $F(x(\lambda), \lambda)$. If Newton-Raphson converged then increase λ by $\Delta\lambda$ and double $\Delta\lambda = 2 \cdot \Delta\lambda$, otherwise half $\Delta\lambda = 0.5 \cdot \Delta\lambda$ and set $\lambda = \lambda_{prev} + \Delta\lambda$. Repeat this until $\lambda = 1$.

Source stepping

Applied to the solution of (non-linear) electrical networks one may think of $\alpha \in [0, 1]$ as a multiplier for the source vector S yielding $S(\alpha) = \alpha S$. Varying α from 0 to 1 and solve at each α . The actual circuit solution is done when $\alpha = 1$. This method is called source stepping. The solution vector $x(\alpha)$ is continuous in α (hence the name continuation scheme).

Minimal derivative g_{min} stepping

Another possibility to improve convergence of almost singular electrical networks is the so called g_{min} stepping, i.e. adding a tiny conductance to ground at each node of the Jacobian A matrix. The continuation starts e.g. with $g_{min} = 0.01$ and ends with $g_{min} = 0$ reached by the algorithm described in section ???. The equation system is slightly modified by adding the current g_{min} to each diagonal element of the matrix A .

Chapter 7

Harmonic Balance Analysis

Harmonic balance is a non-linear, frequency-domain, steady-state simulation. The voltage and current sources create discrete frequencies resulting in a spectrum of discrete frequencies at every node in the circuit. Linear circuit components are solely modeled in frequency domain. Non-linear components are modeled in time domain and Fourier-transformed before each solving step. The informations in this chapter are taken from [?] (chapter 3) which is a very nice and well-written publication on this topic.

The harmonic balance simulation is ideal for situations where transient simulation methods are problematic. These are:

- components modeled in frequency domain, for instance (dispersive) transmission lines
- circuit time constants large compared to period of simulation frequency
- circuits with lots of reactive components

Harmonic balance methods, therefore, are the best choice for most microwave circuits excited with sinusoidal signals (e.g. mixers, power amplifiers).

7.1 The Basic Concept

As the non-linear elements are still modeled in time domain, the circuit first must be separated into a linear and a non-linear part. The internal impedances Z_i of the voltage sources are put into the linear part as well. Figure ?? illustrates the concept. Let us define the following symbols:

M = number of (independent) voltage sources

N = number of connections between linear and non-linear subcircuit

K = number of calculated harmonics

L = number of nodes in linear subcircuit

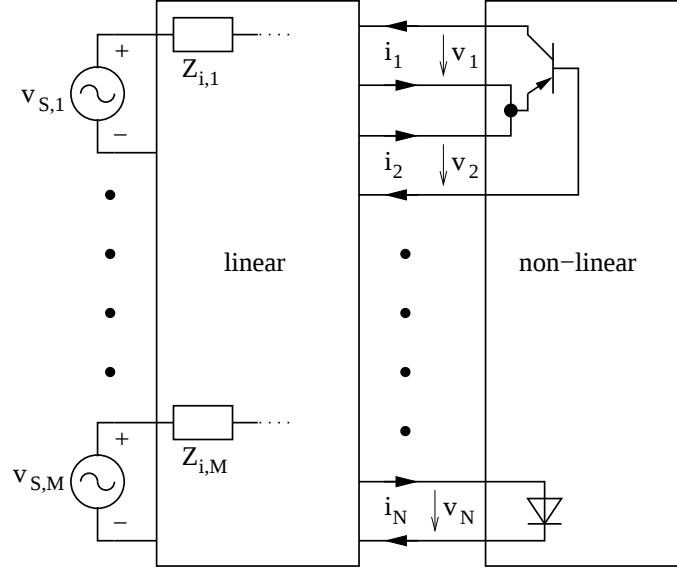


Figure 7.1: circuit partitioning in harmonic balance

The linear circuit is modeled by two transadmittance matrices: The first one $\tilde{\mathbf{Y}}$ relates the source voltages $v_{S,1} \dots v_{S,M}$ to the interconnection currents $i_1 \dots i_N$ and the second one $\hat{\mathbf{Y}}$ relates the interconnection voltages $v_1 \dots v_N$ to the interconnection currents $i_1 \dots i_N$. Taking both, we can express the current flowing through the interconnections between linear and non-linear subcircuit:

$$\mathbf{I} = \tilde{\mathbf{Y}}_{N \times M} \cdot \mathbf{V}_S + \hat{\mathbf{Y}}_{N \times N} \cdot \mathbf{V} = \mathbf{I}_S + \hat{\mathbf{Y}} \cdot \mathbf{V} \quad (7.1)$$

Because $\mathbf{V}_S$ is known and constant, the first term can already be computed to give $\mathbf{I}_S$. Taking the whole linear network as one block is called the "piecewise" harmonic balance technique.

The non-linear circuit is modeled by its current function $i(t) = f_g(v_1, \dots, v_P)$ and by the charge of its capacitances $q(t) = f_q(v_1, \dots, v_Q)$. These functions must be Fourier-transformed to give the frequency-domain vectors $\mathbf{Q}$ and $\mathbf{I}_G$, respectively.

A simulation result is found if the currents through the interconnections are the same for the linear and the non-linear subcircuit. This principle actually gave the harmonic balance simulation its name, because through the interconnections the currents of the linear and non-linear subcircuits have to be *balanced* at every *harmonic* frequency. To be precise the described method is called Kirchhoff's current law harmonic balance (KCL-HB). Theoretically, it would also be possible to use an algorithm that tries to balance the voltages at the subcircuit interconnections. But then the Z matrix (linear subcircuit) and current-dependent voltage laws (non-linear subcircuit) have to be used. That doesn't fit the need (see other simulation types).

So, the non-linear equation system that needs to be solved writes:

$$\mathbf{F}(\mathbf{V}) = \underbrace{(\mathbf{I}_S) + (\hat{\mathbf{Y}}) \cdot (\mathbf{V})}_{\text{linear}} + \underbrace{j \cdot \boldsymbol{\Omega} \cdot \mathbf{Q} + \mathbf{I}_G}_{\text{non-linear}} = \mathbf{0} \quad (7.2)$$

where matrix $\mathbf{\Omega}$ contains the angular frequencies on the first main diagonal and zeros anywhere else, $\mathbf{0}$ is the zero vector.

After each iteration step, the inverse Fourier transformation must be applied to the voltage vector $\mathbf{V}$. Then the time domain voltages $v_{0,1}...v_{K,N}$ are put into $i(t) = f_g(v_1, ..., v_P)$ and $q(t) = f_q(v_1, ..., v_Q)$ again. Now, a Fourier transformation gives the vectors $\mathbf{Q}$ and $\mathbf{I}_G$ for the next iteration step. After repeating this several times, a simulation result has hopefully be found.

Having found this result means having got the voltages $v_1...v_N$ at the interconnections of the two subcircuits. With these values the voltages at all nodes can be calculated: Forget about the non-linear subcircuit, put current sources at the former interconnections (using the calculated values) and perform a normal AC simulation. After that the simulation is complete.

A short note to the construction of the quantities: One big difference between the HB and the conventional simulation types like a DC or an AC simulation is the structure of the matrices and vectors. A vector used in a conventional simulation contains one value for each node. In an HB simulation there are many harmonics and thus, a vector contains K values for each node. This means that within a matrix, there is a $K \times K$ diagonal submatrix for each node. Using this structure, all equations can be written in the usual way, i.e. without paying attention to the special matrix and vector structure. In a computer program, however, a special matrix class is needed in order to not waste memory for the off-diagonal zeros.

7.2 Going through each Step

7.2.1 Creating Transadmittance Matrix

It needs several steps to get the transadmittance matrices $[\tilde{Y}]$ and $[\hat{Y}]$ mentioned in equation (??). First the MNA matrix of the linear subcircuit (figure ??) is created (chapter ??) without the voltage sources $S_1...S_M$ and without the non-linear components. Note that all nodes must appear in the matrix, even those where only non-linear components are connected. Then the transimpedance matrix is derived by exciting one by one the port nodes of the MNA matrix with unity current. After that the transadmittance matrix is calculated by inverting the transimpedance matrix. Finally the matrices $[\tilde{Y}]$ and $[\hat{Y}]$ are filled with the corresponding elements of the overall transadmittance matrix.

Note: The MNA matrix of the linear subcircuit has L nodes. Every node, that is connected to the non-linear subcircuit or/and is connected to a voltage source, is called "port" in the following text. So, there are $M + N$ ports. All these ports need to be differential ones, i.e. without ground reference. Otherwise problems may occur due to singular matrices when calculating $[\tilde{Y}]$ or $[\hat{Y}]$.

Now this should be described in more detail: By use of the MNA matrix $[A]$, the n -th column of the transimpedance matrix $[Z]$ should be calculated. The voltage source at port n is connected to node i (positive terminal) and to node j (negative terminal). This results in the following

equation. (If port n is referenced to ground, the -1 is simply omitted.)

$$[A] \cdot \begin{bmatrix} V_1 \\ \vdots \\ V_L \end{bmatrix} = \begin{bmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ -1 \\ \vdots \\ 0 \end{bmatrix} \quad \begin{array}{l} \leftarrow i\text{-th row} \\ \leftarrow j\text{-th row} \end{array} \quad (7.3)$$

After having solved it, $Z_{1,n} \dots Z_{N+M,n}$ are obtained simply by subtraction of the node voltages:

$$Z_{m,n} = V_k - V_l \quad (7.4)$$

Here the voltage source at port m is connected to node k (positive terminal) and to node l (negative terminal).

The next column of $[Z]$ is obtained by changing the right-hand side of equation (??) appropriately. As this has to be done $N + M$ times, it is strongly recommended to use LU decomposition.

As $[\tilde{Y}]$ is not square, problems encounter by trying to build its inverse matrix. Therefore, the following procedure is recommended:

- Create the transimpedance matrix for all ports (sources and interconnections).
- Compute the inverse matrix (transadmittance matrix).
- The upper left and upper right corner contains $[\tilde{Y}]$ and $[\hat{Y}]$.
- The lower left and lower right corner contains the transadmittance matrices to calculate the currents through the sources. They can be used to simplify the AC simulation at the very end.

One further thing must be mentioned: Because the non-linear components and the sources are missing in the linear MNA matrix, there are often components that are completely disconnected from the rest of the circuit. The resulting MNA matrix cannot be solved. To avoid this problem, shunt each port with a 100Ω resistor, i.e. place a resistor in parallel to each non-linear component and to each source. The effect of these resistors can be easily removed by subtracting 10mS from the first main diagonal of the transadmittance matrix.

7.2.2 Starting Values

A difficult question is how to find appropriate start values for the harmonic balance simulation. It is recommended to first perform a DC analysis and start the algorithm with this result. In many situation (perhaps always) an even better starting point can be achieved by also using the result of a linear AC simulation. However with a large signal strength and strong non-linearities, convergence may still fail. Then, the following procedure might succeed: Perform HB simulation by applying half of the desired signal levels. If convergence is reached, the result can be used as start values for the simulation with the full signal levels. Otherwise the amplitude of the signals can be further decreased in order to repeat the above-mentioned procedure.

7.2.3 Solution algorithm

To perform a HB simulation, the multi-dimensional, non-linear function ?? has to be solved. One of the best possibilities to do so is the Newton method:

$$\mathbf{V}_{n+1} = \mathbf{V}_n - \mathbf{J}_F(\mathbf{V}_n)^{-1} \cdot \mathbf{F}(\mathbf{V}_n) \quad (7.5)$$

$$\Rightarrow \quad \mathbf{J}_F(\mathbf{V}_n) \cdot \mathbf{V}_{n+1} = \mathbf{J}_F(\mathbf{V}_n) \cdot \mathbf{V}_n - \mathbf{F}(\mathbf{V}_n) \quad (7.6)$$

with $\mathbf{J}_F$ being the Jacobian matrix. DC and transient simulation also use this technique, but here a problem appears: The derivatives of the component models are not given in frequency domain. Thus, the Jacobian must be calculated starting at the HB equation ??:

$$\mathbf{J}_F(\mathbf{V}_n) = \frac{d\mathbf{F}(\mathbf{V})}{d\mathbf{V}} = \hat{\mathbf{Y}}_{N \times N} + \frac{\partial \mathbf{I}_G}{\partial \mathbf{V}} + j \cdot \Omega \frac{\partial \mathbf{Q}}{\partial \mathbf{V}} = \hat{\mathbf{Y}}_{N \times N} + \mathbf{J}_{F,G} + j \cdot \Omega \cdot \mathbf{J}_{F,Q} \quad (7.7)$$

So, two Jacobian matrices have to be built, one for the current $\mathbf{I}_G$ and one for the charge $\mathbf{Q}$. Both resulted from a Fourier Transformation. The two operations (Fourier Transformation and differentiation) are linear and thus, can be exchanged. Hence, the Jacobian matrices are built in time domain and transformed into frequency domain afterwards.

To obtain a practical algorithm of this procedure, the DFT is best written as matrix equation. By having a look at equation ?? and ??, it becomes clear how this works. The harmonic factors $\exp(j\omega_k t_n)$ build the matrix $\mathbf{\Gamma}$:

$$\text{DFT:} \quad \mathbf{U}(j\omega) = \mathbf{\Gamma} \cdot \mathbf{u}(t) \quad (7.8)$$

$$\text{IDFT:} \quad \mathbf{u}(t) = \mathbf{\Gamma}^{-1} \cdot \mathbf{U}(j\omega) \quad (7.9)$$

with $\mathbf{u}$ and $\mathbf{U}$ being the vectors of the time and frequency values, respectively. Now, it is possible to transform the desired Jacobian matrix into frequency domain:

$$\mathbf{J}_{F,G} = \frac{\partial \mathbf{I}_G}{\partial \mathbf{V}} = \frac{\partial(\mathbf{\Gamma} \cdot \mathbf{i})}{\partial(\mathbf{\Gamma} \cdot \mathbf{v})} = \mathbf{\Gamma} \cdot \frac{\partial \mathbf{i}}{\partial \mathbf{v}} \cdot \mathbf{\Gamma}^{-1} \quad (7.10)$$

Here $\mathbf{i}$ is a vector with length $K \cdot N$, i.e. first all time values of the first node are inserted, then all time values of the second node etc. The Jacobi matrix of $\mathbf{i}$ is defined as:

$$\mathbf{J}_{F,G}(\mathbf{u}) = \begin{bmatrix} \frac{\partial i_1}{\partial u_1} & \cdots & \frac{\partial i_1}{\partial u_n} \\ \vdots & \ddots & \vdots \\ \frac{\partial i_n}{\partial u_1} & \cdots & \frac{\partial i_n}{\partial u_n} \end{bmatrix} \quad (7.11)$$

Hence this matrix consists of $K \times K$ blocks (one for each node) that are diagonal matrices with time values of the derivatives in it. (Components exists that create non-diagonal blocks, but these are so special ones that they do not appear in this document.)

The formula ?? seems to be quite clear, but it has to be pointed out how this works with FFT algorithm. With $\mathbf{\Gamma}^{-1} = (\mathbf{\Gamma}^{-1})^T$ (see equation ??) and $(\mathbf{A} \cdot \mathbf{B})^T = \mathbf{B}^T \cdot \mathbf{A}^T$, it follows:

$$\mathbf{J}_{F,G} = \mathbf{\Gamma} \cdot \frac{\partial \mathbf{i}}{\partial \mathbf{v}} \cdot \mathbf{\Gamma}^{-1} = \left(\mathbf{\Gamma}^{-1} \cdot \left(\mathbf{\Gamma} \cdot \frac{\partial \mathbf{i}}{\partial \mathbf{v}} \right)^T \right)^T \quad (7.12)$$

So, there are two steps to perform an FFT-based transformation of the time domain Jacobian matrix into the frequency domain Jacobian:

1. Perform an FFT on every column of the Jacobian and build a new matrix $\mathbf{A}$ with this result, i.e. the first column of $\mathbf{A}$ is the FFTed first column of the Jacobian and so on.
2. Perform an IFFT on every row of the matrix $\mathbf{A}$ and build a new matrix $\mathbf{B}$ with this result, i.e. the first row of $\mathbf{B}$ is the IFFTed first row of $\mathbf{A}$ and so on.

As the Fourier transformation has to be applied to diagonal sub-matrices, the whole above-mentioned process can be performed by one single FFT. This is done by taking the $\partial \mathbf{i} / \partial \mathbf{v}$ values in a vector $\mathbf{J}_i$ and calculating:

$$\frac{1}{K} \cdot \text{FFT}(\mathbf{J}_i) \quad (7.13)$$

The result is the first column of $\mathbf{J}_{F,G}$. The second column equals the first one rotated down by one element. The third column is the second one rotated down by one element etc. A matrix obeying this structure is called Toeplitz matrix.

So, finally the complete HB Newton iteration step can be written down. Putting ?? and ?? into ?? leads to

$$\mathbf{J}_F(\mathbf{V}_n) \cdot \mathbf{V}_{n+1} = \mathbf{J}_{F,G} \cdot \mathbf{V}_n - \mathbf{I}_G + j \cdot \Omega \cdot (\mathbf{J}_{F,Q} \cdot \mathbf{V}_n - \mathbf{Q}) - \mathbf{I}_S \quad (7.14)$$

This is important to notice, because many non-linear components cannot be processed at every bias point (see figure ??). These components create a new voltage estimate across their nodes, whereas the new estimated absolute voltages at their nodes are not known. Thus, the term $\mathbf{J}_{F,G} \cdot \mathbf{V}_n$ can only be created in one single step, leading to the vector $\mathbf{I}_{G,J}$. Luckily, this procedure also saves computation time, as the matrix multiplication need not to be performed. The same is true for the term $\mathbf{J}_{F,Q} \cdot \mathbf{V}_n$, leading to the vector $\mathbf{Q}_J$. So it is:

$$\mathbf{J}_F \cdot \mathbf{V}_{n+1} = \mathbf{I}_{G,J} - \mathbf{I}_G + j \cdot \Omega \cdot (\mathbf{Q}_J - \mathbf{Q}) - \mathbf{I}_S \quad (7.15)$$

7.2.4 Termination Criteria

Frequency components with very different magnitude appear in harmonic balance simulation. In order to detect when the solver has found an accurate solution, an absolute as well as relative criteria must be used on all nodes and at all frequencies. The analysis is regarded as finished if one of the criteria is satisfied.

The absolute and relative criteria write as follows:

$$\left| \tilde{I}_{n,k} + \hat{I}_{n,k} \right| < \varepsilon_{abs} \quad \forall \quad n, k \quad (7.16)$$

$$2 \cdot \left| \frac{\tilde{I}_{n,k} + \hat{I}_{n,k}}{\tilde{I}_{n,k} - \hat{I}_{n,k}} \right| < \varepsilon_{rel} \quad \forall \quad n, k \quad (7.17)$$

where $\tilde{I}_{n,k}$ is the current of the linear circuit partition for node n and frequency k and $\hat{I}_{n,k}$ is the current of the non-linear circuit partition.

7.3 A Symbolic HB Algorithm

In this final section, a harmonic balance algorithm in symbolic language is presented.

Listing 7.1: symbolic HB algorithm

```

init (); // separate linear and non-linear devices
Y = calcTransMatrix (); // transadmittance matrix of linear circuit
Is = calcSourceCurrent (); // source current of linear subcircuit
(v, i, q) = calculateDC (); // DC simulation as initial HB estimate
V = FFT(v); // transform voltage into frequency domain

loop :
    I = FFT(i); // current into frequency domain
    Q = FFT(q); // charge into frequency domain
    E = Is + Y*V + j*Ω*Q + I; // HB equation
    if (abs(E) < Eterm) break; // convergence reached ?
    JG = mFFT(GJacobian(v)); // create Jacobians and transform ...
    JQ = mFFT(QJacobian(v)); // ... them into frequency domain
    J = Y + j*Ω*JQ + JG; // calculate overall Jacobian
    V = V - invert(J) * E; // Newton Raphson iteration step
    v = IFFT(V); // voltage into time domain
    i = nonlinearCurrent(v); // use component models to get ...
    q = nonlinearCharge(v); // ... values for next iteration
    goto loop;

Va = invert(Ya) * Ia; // AC simulation to get all voltages

```

7.4 Large-Signal S-Parameter Simulation

Using harmonic balance techniques, it is also possible to perform an S-parameter simulation in the large-signal regime. This is called LSSP (large-signal s-parameter). Figure ?? shows the principle. The port n excites the circuit with the simulation frequency f_0 ; meanwhile the power of all other ports is set to zero. Having voltage and current of the fundamental frequency f_0 at the ports, the S-parameters can be calculated:

$$\underline{S}_{mn} = \frac{\underline{U}_m(f_0) - \underline{I}_m(f_0) \cdot Z_m}{\underline{U}_n(f_0) + \underline{I}_n(f_0) \cdot Z_n} \cdot \sqrt{\frac{Z_n}{Z_m}} \quad (7.18)$$

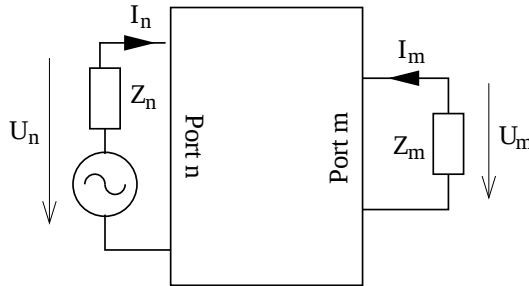


Figure 7.2: S-parameter from AC voltages and currents

An algorithm in symbolic language should describe the whole LSSP:

Listing 7.2: symbolic HB algorithm

```

for n=1 to NumberOfPorts {
  Set power of port  $n$  to  $P_n$ .
  Set power of ports  $\neq n$  to 0.
  Perform Harmonic Balance.

  for m=1 to NumberOfPorts
    Calculate  $\underline{S}_{mn}$  according to above-mentioned equation.
}

```

7.5 Autonomous Harmonic Balance

Up to here, only forced circuits were dealt with. That is, the above-mentioned methods can analyse circuits that are driven by signal sources, but do not create a signal by themselves. The typical examples are amplifiers and mixers. However, harmonic balance techniques are also capable of simulating autonomous circuits like oscillators.

This is mostly done in the following way:

1. The user enters an approximate frequency, where the oscillation is expected. (An ac simulation can be performed to get an idea about that.)
2. The user enters a frequency interval. Somewhere within this interval the oscillation must appear.
3. The user specifies a circuit node where the oscillation voltage can best be measured.
4. The simulator performs several harmonic balance analyses with different fundamental frequencies in search for the oscillation.

Chapter 8

Harmonic Balance Noise Analysis

Once a harmonic balance simulation is solved a cyclostationary noise analysis can be performed. This results in the sideband noise of each harmonic (including DC, i.e. base band noise). The method described here is based on the principle of small-signal noise. That is, the noise power is assumed small enough (compared to the signal power and its harmonics) to neglect the higher order mixing products. This procedure is the standard concept in CAE and allows for a quite simple and time-saving algorithm: Use the Jacobian to calculate a conversion matrix and then apply the noise correlation matrix to it. Two important publications for HB noise simulation exist that were used for the next subsection [?], [?].

Figure ?? shows the equivalent circuit for starting the HB noise analysis. At every connection between linear and non-linear subcircuit, there are two noise current sources: one stemming from the linear subcircuit and one stemming from the non-linear subcircuit.

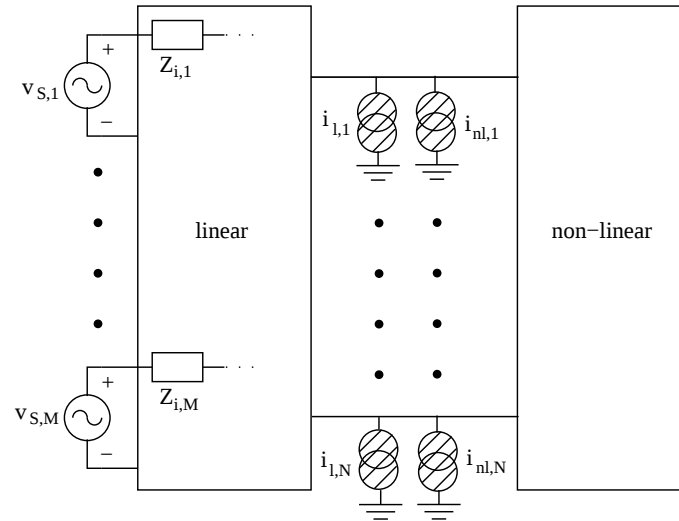


Figure 8.1: principle of harmonic balance noise model

8.1 The Linear Subcircuit

The noise stemming from the linear subcircuit is calculated in two steps:

1. An AC noise analysis (see section ??) is performed for the interconnecting nodes of linear and non-linear subcircuit. This results in the noise-voltage correlation matrix $(\underline{\mathcal{C}}_{Z,lin})_{N \times N}$.
2. The matrix $(\underline{\mathcal{C}}_{Z,lin})_{N \times N}$ is converted into a noise-current correlation matrix (see section ??):

$$(\underline{\mathcal{C}}_{Y,lin})_{N \times N} = \hat{\mathbf{Y}} \cdot \underline{\mathcal{C}}_{Z,lin} \cdot \hat{\mathbf{Y}}^+ \quad (8.1)$$

where $\hat{\mathbf{Y}}$ is taken from equation ??.

Remark: If no explicit noise sources exist in the linear subcircuit, $(\underline{\mathcal{C}}_{Z,lin})_{N \times N}$ can be computed much faster by using Bosma's theorem (equation ??).

8.2 The Non-Linear Subcircuit

The noise in the non-linear part of the circuit is calculated by using the quasi-static approach, i.e. for every moment in time the voltages and currents are regarded as a time-dependent bias point. The noise properties of these bias points are used for the noise calculation.

Remark: It is not clear whether this approach creates a valid result for noise with long-time correlation (e.g. 1/f noise), too. But up to now, no other methods were proposed and some publications reported to have achieved reasonable results with this approach and 1/f noise.

Calculating the noise-current correlation matrix $(\underline{\mathcal{C}}_{Y,nl})_{N \times N}$ needs several steps. The DC bias point taken from the result of the HB simulation is the beginning. Its values are the bias used to build the correlation matrix $(\underline{\mathcal{C}}_{Y,DC})$. Each part is a $K \times K$ diagonal submatrix. The values are the power-spectral densities (PSD) for each harmonic frequency:

$$C_{Y,DC}(\omega_R) \quad C_{Y,DC}(\omega_0 + \omega_R) \quad C_{Y,DC}(2 \cdot \omega_0 + \omega_R) \quad \dots \quad (8.2)$$

where ω_R is the desired noise frequency.

The second step creates the cyclostationary modulation that is applied to the DC correlation matrix. The modulation factor $M(t)$ originates from the current power spectral density S_i of each time step normalized to its DC bias value:

$$M(t) = \frac{S_i(u(t))}{S_i(u_{DC})} = \frac{S_i(u(t))}{C_{Y,DC}} \quad (8.3)$$

Note that this equation only holds if the frequency dependency of S_i is the same for every bias, so that $M(t)$ is frequency independent. This demand is fulfilled for all practical models. So the above-mentioned equation can be derived for an arbitrary noise frequency ω_R .

The third step transforms $M(t)$ into frequency domain. This is done by the procedure described in equation ??, resulting in a Toeplitz matrix.

The fourth and final step calculates the desired correlation matrix:

$$(\underline{\mathcal{C}}_{Y,nl}) = M \cdot (\underline{\mathcal{C}}_{Y,DC}) \cdot M^+ \quad (8.4)$$

8.3 Noise Conversion

As the noise of linear and non-linear components are uncorrelated, the noise-voltage correlation matrix at the interconnecting ports can now be calculated:

$$\mathbf{C}_Z = \mathbf{J}_F^{-1} \cdot (\mathbf{C}_{Y,lin} + \mathbf{C}_{Y,nl}) \cdot (\mathbf{J}_F^{-1})^+ \quad (8.5)$$

here $\mathbf{J}_F^{-1}$ is the inverse of the Jacobian matrix taken from the last HB iteration step (where it already was inverted). Note that it needs to be the precise Jacobian matrix. I.e. it must be taken from an iteration step very close to the solution, without any convergence helpers, and with a precise FFT algorithm (e.g. the multi-dimensional FFT).

Finally, the noise voltages from the interconnecting ports have to be used to compute all other noise voltages. This is straight forward:

1. Convert the noise-voltage correlation matrix into the noise-current correlation matrix.
2. Expand the matrix to the whole circuit, i.e. fill it up with zeros.
3. Perform an AC noise analysis for all nodes of interest.

The whole algorithm has to be performed for every noise frequency ω_R of interest.

8.4 Phase and Amplitude Noise

The harmonic balance noise analysis calculates the noise power spectral density $S_{uu,k}(\omega_R)$ at the noise frequency ω_R of the k -th harmonic. The SSB phase and amplitude noise normalized to the carrier can be obtained by using the symmetry between positive and negative harmonic numbers:

$$\langle \Phi_k \Phi_{-k}^* \rangle = \frac{S_{uu,k} + S_{uu,-k} - 2 \cdot \text{Re}(C_{Z,k,-k} \cdot \exp(-j \cdot 2 \cdot \phi_k))}{|U_k|^2} \quad (8.6)$$

$$\langle A_k A_{-k}^* \rangle = \frac{S_{uu,k} + S_{uu,-k} + 2 \cdot \text{Re}(C_{Z,k,-k} \cdot \exp(-j \cdot 2 \cdot \phi_k))}{|U_k|^2} \quad (8.7)$$

with $U_k = |U_k| \cdot \exp(j \cdot \phi_k)$ being the k -th harmonic.

Chapter 9

Linear devices

As the MNA matrix is the Y-parameter matrix of the whole circuit, components that are defined by Y-parameters can be easily inserted by adding these parameters to the MNA matrix elements (so-called 'stamping').

9.1 Extended MNA stamping

Components that cannot be defined by Y-parameters need to add additional columns and rows to the MNA matrix representation.

9.1.1 Z-parameters

Components defined by Z-parameters can be added in the following way (example for a 2-port). It is easily extendable for any number of ports.

$$\begin{bmatrix} . & . & 1 & 0 \\ . & . & 0 & 1 \\ -1 & 0 & Z_{11} & Z_{12} \\ 0 & -1 & Z_{21} & Z_{22} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ J_1 \\ J_2 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \\ 0 \end{bmatrix} \quad (9.1)$$

9.1.2 S-parameters

Components that are characterized by S-parameters (normalized to Z_0) can be put into the MNA matrix by the following scheme (example for a 3-port). It is easily extendable for any number of ports.

$$\begin{bmatrix} . & . & . & 1 & 0 & 0 \\ . & . & . & 0 & 1 & 0 \\ . & . & . & 0 & 0 & 1 \\ S_{11} - 1 & S_{12} & S_{13} & Z_0 \cdot (S_{11} + 1) & Z_0 \cdot S_{12} & Z_0 \cdot S_{13} \\ S_{21} & S_{22} - 1 & S_{23} & Z_0 \cdot S_{21} & Z_0 \cdot (S_{22} + 1) & Z_0 \cdot S_{23} \\ S_{31} & S_{32} & S_{33} - 1 & Z_0 \cdot S_{31} & Z_0 \cdot S_{32} & Z_0 \cdot (S_{33} + 1) \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ J_1 \\ J_2 \\ J_3 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.2)$$

9.1.3 H-parameters

According to eqn. (??) the MNA matrix for a 2-port H-parameter matrix can be written as:

$$\begin{bmatrix} . & . & 1 \\ . & H_{22} & H_{21} \\ -1 & H_{12} & H_{11} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ J_1 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} \quad (9.3)$$

9.1.4 G-parameters

According to eqn. (??) the MNA matrix for a 2-port G-parameter matrix can be written as:

$$\begin{bmatrix} G_{11} & . & G_{12} \\ . & . & 1 \\ G_{21} & -1 & G_{22} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ J_2 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} \quad (9.4)$$

9.1.5 A-parameters

According to eqn. (??) the MNA matrix for a 2-port A-parameter matrix can be written as:

$$\begin{bmatrix} . & A_{21} & A_{22} \\ . & . & -1 \\ -1 & A_{11} & A_{12} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ J_2 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} \quad (9.5)$$

9.1.6 T-parameters

According to eqn. (??) the MNA matrix for a 2-port T-parameter matrix can be written as:

$$\begin{bmatrix} . & . & 1 & 0 \\ . & . & 0 & 1 \\ -1 & T_{12} + T_{11} & -Z_0 & Z_0 \cdot (T_{12} - T_{11}) \\ -1 & T_{22} + T_{21} & +Z_0 & Z_0 \cdot (T_{22} - T_{21}) \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ J_1 \\ J_2 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \\ 0 \end{bmatrix} \quad (9.6)$$

9.2 Resistor

For DC and AC simulation an ideal resistor with resistance R yields:

$$Y = \frac{1}{R} \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (9.7)$$

The noise correlation matrix at temperature T yields:

$$(\underline{C}_Y) = \frac{4 \cdot k \cdot T}{R} \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (9.8)$$

The scattering parameters normalized to impedance Z_0 writes as follows.

$$S_{11} = S_{22} = \frac{R}{2 \cdot Z_0 + R} \quad (9.9)$$

$$S_{12} = S_{21} = 1 - S_{11} = \frac{2 \cdot Z_0}{2 \cdot Z_0 + R} \quad (9.10)$$

Being on temperature T , the noise wave correlation matrix writes as follows.

$$(\underline{C}) = k \cdot T \cdot \frac{4 \cdot R \cdot Z_0}{(2 \cdot Z_0 + R)^2} \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (9.11)$$

The noise wave correlation matrix of a parallel resistor with resistance R writes as follows.

$$(\underline{C}) = k \cdot T \cdot \frac{4 \cdot R \cdot Z_0}{(2 \cdot R + Z_0)^2} \cdot \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix} \quad (9.12)$$

The noise wave correlation matrix of a grounded resistor with resistance R is a matrix consisting of one element and writes as follows.

$$(\underline{C}) = k \cdot T \cdot \frac{4 \cdot R \cdot Z_0}{(R + Z_0)^2} \quad (9.13)$$

9.3 Capacitor

During DC simulation the capacitor is an open circuit. Thus, its MNA entries are all zero.

During AC simulation the y-parameter matrix of an ideal capacitor with the capacitance C writes as follows.

$$Y = \begin{pmatrix} +j\omega C & -j\omega C \\ -j\omega C & +j\omega C \end{pmatrix} \quad (9.14)$$

The scattering parameters (normalized to Z_0) of an ideal capacitor with capacitance C writes as follows.

$$S_{11} = S_{22} = \frac{1}{2 \cdot Z_0 \cdot j\omega C + 1} \quad (9.15)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (9.16)$$

An ideal capacitor is noise free. Its noise correlation matrices are, therefore, zero.

9.4 Inductor

During DC simulation an inductor is a short circuit, thus, its MNA matrix entries need an additional row and column.

$$\begin{bmatrix} \cdot & \cdot & +1 \\ \cdot & \cdot & -1 \\ +1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{br} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.17)$$

During AC simulation the Y-parameter matrix of an ideal inductor with the inductance L writes as follows.

$$Y = \begin{pmatrix} +\frac{1}{j\omega L} & -\frac{1}{j\omega L} \\ -\frac{1}{j\omega L} & +\frac{1}{j\omega L} \end{pmatrix} \quad (9.18)$$

The scattering parameters of an ideal inductor with inductance L writes as follows.

$$S_{11} = S_{22} = \frac{j\omega L}{2 \cdot Z_0 + j\omega L} \quad (9.19)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (9.20)$$

An ideal inductor is noise free.

9.5 Lossy reactances

Because of its physical implementation, both capacitors and inductors exhibit a lossy behaviour which can be assessed in terms of their quality factor. The quality factor, Q , of a reactance is defined as the ratio between the stored energy and the dissipated energy per unit time. When using a series resistance R_s to represent the losses, it is also given by

- Capacitor: $Q = \frac{1}{\omega C R_s} = \frac{1}{2\pi f \cdot C \cdot R_s}$
- Inductor: $Q = \frac{\omega L}{R_s} = \frac{2\pi f \cdot L}{R_s}$

note that in general the losses and thus the quality factor are frequency dependent. Some common models for the variation of Q with the frequency are:

- Constant: $Q(f) = Q(f_0)$
- Linear: $Q(f) = Q(f_0) \cdot (f/f_0)$
- Square root: $Q(f) = Q(f_0) \cdot \sqrt{f/f_0}$

9.5.1 Equivalent circuits and S matrices

For the lossy components models described above, the scattering parameters (normalized to Z_0) can be computed from the S-parameters of a generic series element.

- Capacitor



Figure 9.1: Equivalent circuit of a lossy capacitor

Defining $Y_s = \frac{j\omega C}{1+j\omega C R}$ we have

$$S_{11} = S_{22} = \frac{1}{1+y} \quad , \quad y = 2 \cdot Z_0 \cdot Y_s \quad (9.21)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (9.22)$$

- Inductor

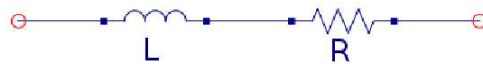


Figure 9.2: Equivalent circuit of a lossy inductor

With $Z_s = R + j\omega L$, the S-parameter matrix is given by

$$S_{11} = S_{22} = \frac{z}{z+2} \quad , \quad z = \frac{Z_s}{Z_0} \quad (9.23)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (9.24)$$

The noise parameters of these components can in general be computed using Bosma's theorem.

9.6 DC Block

A DC block is a capacitor with an infinite capacitance. During DC simulation the DC block is an open circuit. Thus, its MNA entries are all zero.

The MNA matrix entries of a DC block correspond to an ideal short circuit during AC analysis which is modeled by a voltage source with zero voltage.

$$\begin{bmatrix} \cdot & \cdot & +1 \\ \cdot & \cdot & -1 \\ +1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{br} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.25)$$

The scattering parameters writes as follows.

$$(S) = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad (9.26)$$

A DC block is noise free. A model for transient simulation does not exist. It is common practice to model it as a capacitor with finite capacitance whose value is entered by the user.

9.7 DC Feed

A DC feed is an inductor with an infinite inductance. The MNA matrix entries of a DC feed correspond to an ideal short circuit during DC analysis:

$$\begin{bmatrix} \cdot & \cdot & +1 \\ \cdot & \cdot & -1 \\ +1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{br} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.27)$$

During AC simulation the DC feed is an open circuit. Thus, its MNA entries are all zero.

The scattering parameters writes as follows.

$$(S) = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (9.28)$$

A DC feed is noise free. A model for transient simulation does not exist. It is common practice to model it as an inductor with finite inductance whose value is entered by the user.

9.8 Bias T

An ideal bias t is a combination of a DC block and a DC feed (fig. ??). During DC simulation the MNA matrix of an ideal bias t writes as follows:

$$\begin{bmatrix} \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & -1 \\ 0 & 1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.29)$$

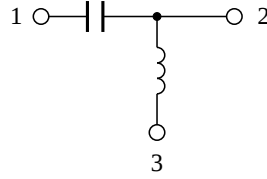


Figure 9.3: bias t

The MNA entries of the bias t during AC analysis write as follows.

$$\begin{bmatrix} \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & 0 \\ -1 & 1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.30)$$

The scattering parameters writes as follows.

$$(S) = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (9.31)$$

A bias t is noise free. A model for transient simulation does not exist. It is common practice to model it as an inductor and a capacitance with finite values which are entered by the user.

9.9 Transformer

The two winding ideal transformer, as shown in fig. ??, is determined by the following equation which introduces one more unknown in the MNA matrix.

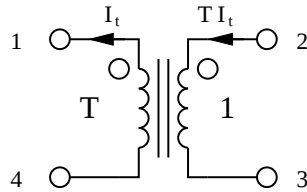


Figure 9.4: ideal two winding transformer

$$T \cdot (V_2 - V_3) = V_1 - V_4 \rightarrow V_1 - T \cdot V_2 + T \cdot V_3 - V_4 = 0 \quad (9.32)$$

The new unknown variable I_t must be considered by the four remaining simple equations.

$$I_1 = -I_t \quad I_2 = T \cdot I_t \quad I_3 = -T \cdot I_t \quad I_4 = I_t \quad (9.33)$$

And in matrix representation this is for DC and for AC simulation:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & \cdot & T \\ \cdot & \cdot & \cdot & \cdot & -T \\ \cdot & \cdot & \cdot & \cdot & 1 \\ 1 & -T & T & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_t \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.34)$$

It is noticeable that the additional row (part of the C matrix) and the corresponding column (part of the B matrix) are transposed to each other. When considering the turns ratio T being complex introducing an additional phase the transformer can be used as phase-shifting transformer. Both the vectors must be conjugated complex transposed in this case.

Using the port numbers depicted in fig. ??, the scattering parameters of an ideal transformer with voltage transformation ratio T (number of turns) writes as follows.

$$S_{14} = S_{22} = S_{33} = S_{41} = \frac{1}{T^2 + 1} \quad (9.35)$$

$$S_{12} = -S_{13} = S_{21} = -S_{24} = -S_{31} = S_{34} = -S_{42} = S_{43} = T \cdot S_{22} \quad (9.36)$$

$$S_{11} = S_{23} = S_{32} = S_{44} = T \cdot S_{12} \quad (9.37)$$

An ideal transformer is noise free.

9.10 Symmetrical transformer

The ideal symmetrical transformer, as shown in fig. ??, is determined by the following equations which introduce two more unknowns in the MNA matrix.

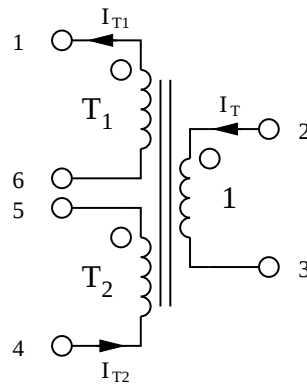


Figure 9.5: ideal three winding transformer

$$T_1 \cdot (V_2 - V_3) = V_1 - V_6 \rightarrow V_1 - T_1 \cdot V_2 + T_1 \cdot V_3 - V_6 = 0 \quad (9.38)$$

$$T_2 \cdot (V_2 - V_3) = V_5 - V_4 \rightarrow -T_2 \cdot V_2 + T_2 \cdot V_3 - V_4 + V_5 = 0 \quad (9.39)$$

The new unknown variables I_{T1} and I_{T2} must be considered by the six remaining simple equations.

$$I_2 = T_1 \cdot I_{T1} + T_2 \cdot I_{T2} \quad I_3 = -T_1 \cdot I_{T1} - T_2 \cdot I_{T2} \quad (9.40)$$

$$I_1 = -I_{T1} \quad I_4 = I_{T2} \quad I_5 = -I_{T2} \quad I_6 = I_{T1} \quad (9.41)$$

The matrix representation needs to be augmented by two more new rows and their corresponding columns. For DC and AC simulation it is:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & T_1 & T_2 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & -T_1 & -T_2 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & 0 & 1 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & 1 & 0 \\ 1 & -T_1 & T_1 & 0 & 0 & -1 & 0 & 0 \\ 0 & -T_2 & T_2 & -1 & 1 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ V_5 \\ V_6 \\ I_{T1} \\ I_{T2} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ I_5 \\ I_6 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.42)$$

Using the port numbers depicted in fig. ??, the scattering parameters of an ideal, symmetrical transformer with voltage transformation ratio (number of turns) T_1 and T_2 , respectively, writes as follows.

$$denom = 1 + T_1^2 + T_2^2 \quad (9.43)$$

$$S_{11} = S_{66} = \frac{T_1^2}{denom} \quad S_{16} = S_{61} = 1 - S_{11} \quad (9.44)$$

$$S_{44} = S_{55} = \frac{T_2^2}{denom} \quad S_{45} = S_{54} = 1 - S_{44} \quad (9.45)$$

$$S_{22} = S_{33} = \frac{1}{denom} \quad S_{23} = S_{32} = 1 - S_{22} \quad (9.46)$$

$$S_{12} = S_{21} = -S_{13} = -S_{31} = -S_{26} = -S_{62} = S_{36} = S_{63} = \frac{T_1}{denom} \quad (9.47)$$

$$-S_{24} = -S_{42} = S_{25} = S_{52} = S_{34} = S_{43} = -S_{35} = -S_{53} = \frac{T_2}{denom} \quad (9.48)$$

$$-S_{14} = -S_{41} = S_{15} = S_{51} = S_{46} = S_{64} = -S_{56} = -S_{65} = \frac{T_1 \cdot T_2}{denom} \quad (9.49)$$

An ideal symmetrical transformer is noise free.

9.11 Non-ideal transformer

Many simulators support non-ideal transformers (e.g. mutual inductor in SPICE). An often used model consists of finite inductances and an imperfect coupling (stray inductance). This model has three parameters: Inductance of the primary coil L_1 , inductance of the secondary coil L_2 and the coupling factor $k = 0 \dots 1$.

9.11.1 Mutual inductors with two or three of inductors

This model can be replaced by the equivalent circuit depicted in figure ???. The values are calculated as follows.

$$\text{turn ratio:} \quad T = \sqrt{\frac{L_1}{L_2}} \quad (9.50)$$

$$\text{mutual inductance:} \quad M = k \cdot L_1 \quad (9.51)$$

$$\text{primary inductance:} \quad L_{1,\text{new}} = L_1 - M = L_1 \cdot (1 - k) \quad (9.52)$$

$$\text{secondary inductance:} \quad L_{2,\text{new}} = L_2 - \frac{M}{T^2} = L_2 \cdot (1 - k) \quad (9.53)$$

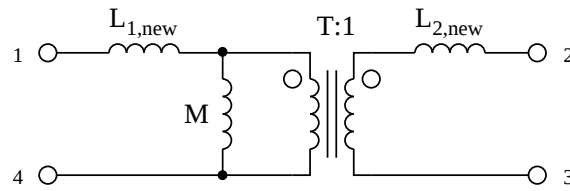


Figure 9.6: equivalent circuit of non-ideal transformer

The Y-parameters of this component are:

$$Y_{11} = Y_{44} = -Y_{41} = -Y_{14} = \frac{1}{j\omega \cdot L_1 \cdot (1 - k^2)} \quad (9.54)$$

$$Y_{22} = Y_{33} = -Y_{23} = -Y_{32} = \frac{1}{j\omega \cdot L_2 \cdot (1 - k^2)} \quad (9.55)$$

$$Y_{13} = Y_{31} = Y_{24} = Y_{42} = -Y_{12} = -Y_{21} = -Y_{34} = -Y_{43} = \frac{k}{j\omega \cdot \sqrt{L_1 \cdot L_2} \cdot (1 - k^2)} \quad (9.56)$$

Furthermore, its S-parameters are:

$$D = (k^2 - 1) \cdot \frac{\omega^2 \cdot L_1 \cdot L_2}{2 \cdot Z_0} + j\omega L_1 + j\omega L_2 + 2 \cdot Z_0 \quad (9.57)$$

$$S_{14} = S_{41} = \frac{j\omega L_2 + 2 \cdot Z_0}{D} \quad (9.58)$$

$$S_{11} = S_{44} = 1 - S_{14} \quad (9.59)$$

$$S_{23} = S_{32} = \frac{j\omega L_1 + 2 \cdot Z_0}{D} \quad (9.60)$$

$$S_{22} = S_{33} = 1 - S_{23} \quad (9.61)$$

$$S_{12} = -S_{13} = S_{21} = -S_{24} = -S_{31} = S_{34} = -S_{42} = S_{43} = \frac{j\omega \cdot k \cdot \sqrt{L_1 \cdot L_2}}{D} \quad (9.62)$$

Also including an ohmic resistance R_1 and R_2 for each coil, leads to the following Y-parameters:

$$Y_{11} = Y_{44} = -Y_{41} = -Y_{14} = \frac{1}{j\omega \cdot L_1 \cdot \left(1 - k^2 \cdot \frac{j\omega L_2}{j\omega L_2 + R_2}\right) + R_1} \quad (9.63)$$

$$Y_{22} = Y_{33} = -Y_{23} = -Y_{32} = \frac{1}{j\omega \cdot L_2 \cdot \left(1 - k^2 \cdot \frac{j\omega L_1}{j\omega L_1 + R_1}\right) + R_2} \quad (9.64)$$

$$Y_{13} = Y_{31} = Y_{24} = Y_{42} = -Y_{12} = -Y_{21} = -Y_{34} = -Y_{43} = k \cdot \frac{j\omega \sqrt{L_1 \cdot L_2}}{j\omega \cdot L_2 + R_2} \cdot Y_{11} \quad (9.65)$$

Building the S-parameters leads to too large equations. Numerically converting the Y-parameters into S-parameters is therefore recommended.

The MNA matrix entries during DC analysis and the noise correlation matrices of this transformer are:

$$(\underline{Y}) = \begin{pmatrix} 1/R_1 & 0 & 0 & -1/R_1 \\ 0 & 1/R_2 & -1/R_2 & 0 \\ 0 & -1/R_2 & 1/R_2 & 0 \\ -1/R_1 & 0 & 0 & 1/R_1 \end{pmatrix} \quad (9.66)$$

$$(\underline{C}_Y) = 4 \cdot k \cdot T \cdot \begin{pmatrix} 1/R_1 & 0 & 0 & -1/R_1 \\ 0 & 1/R_2 & -1/R_2 & 0 \\ 0 & -1/R_2 & 1/R_2 & 0 \\ -1/R_1 & 0 & 0 & 1/R_1 \end{pmatrix} \quad (9.67)$$

$$(\underline{C}_S) = 4 \cdot k \cdot T \cdot Z_0 \cdot \begin{pmatrix} \frac{R_1}{(2 \cdot Z_0 + R_1)^2} & 0 & 0 & -\frac{R_1}{(2 \cdot Z_0 + R_1)^2} \\ 0 & \frac{R_2}{(2 \cdot Z_0 + R_2)^2} & -\frac{R_2}{(2 \cdot Z_0 + R_2)^2} & 0 \\ 0 & -\frac{R_2}{(2 \cdot Z_0 + R_2)^2} & \frac{R_2}{(2 \cdot Z_0 + R_2)^2} & 0 \\ -\frac{R_1}{(2 \cdot Z_0 + R_1)^2} & 0 & 0 & \frac{R_1}{(2 \cdot Z_0 + R_1)^2} \end{pmatrix} \quad (9.68)$$

A transformer with three coupled inductors has three coupling factors k_{12} , k_{13} and k_{23} . Its Y-parameters write as follows (port numbers are according to figure ??).

$$A = j\omega \cdot (1 - k_{12}^2 - k_{13}^2 - k_{23}^2 + 2 \cdot k_{12} \cdot k_{13} \cdot k_{23}) \quad (9.69)$$

$$Y_{11} = Y_{66} = -Y_{16} = -Y_{61} = \frac{1 - k_{23}^2}{L_1 \cdot A} \quad (9.70)$$

$$Y_{22} = Y_{33} = -Y_{23} = -Y_{32} = \frac{1 - k_{12}^2}{L_3 \cdot A} \quad (9.71)$$

$$Y_{44} = Y_{55} = -Y_{45} = -Y_{54} = \frac{1 - k_{13}^2}{L_2 \cdot A} \quad (9.72)$$

$$Y_{12} = Y_{21} = Y_{36} = Y_{63} = -Y_{13} = -Y_{31} = -Y_{26} = -Y_{62} = \frac{k_{12} \cdot k_{23} - k_{13}}{\sqrt{L_1 \cdot L_3} \cdot A} \quad (9.73)$$

$$Y_{15} = Y_{51} = Y_{46} = Y_{64} = -Y_{14} = -Y_{41} = -Y_{56} = -Y_{65} = \frac{k_{13} \cdot k_{23} - k_{12}}{\sqrt{L_1 \cdot L_2} \cdot A} \quad (9.74)$$

$$Y_{25} = Y_{52} = Y_{43} = Y_{34} = -Y_{24} = -Y_{42} = -Y_{53} = -Y_{35} = \frac{k_{12} \cdot k_{13} - k_{23}}{\sqrt{L_2 \cdot L_3} \cdot A} \quad (9.75)$$

9.11.2 Mutual inductors with any number of inductors

A more general approach for coupled inductors can be obtained by using the induction law:

$$V_L = j\omega L \cdot I_L + j\omega \cdot \sum_{n=1}^N k_n \cdot \sqrt{L \cdot L_n} \cdot I_{L,n} \quad (9.76)$$

where V_L and I_L is the voltage across and the current through the inductor, respectively. L is its inductance. The inductor is coupled with N other inductances L_n . The corresponding coupling factors are k_n and $I_{L,n}$ are the currents through the inductors.

Realizing this approach with the MNA matrix is straight forward: Every inductance L needs an additional matrix row. The corresponding element in the D matrix is $j\omega L$. If two inductors are coupled the cross element in the D matrix is $j\omega k \cdot \sqrt{L_1 \cdot L_2}$. For two coupled inductors this yields:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & +1 & 0 \\ \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 & +1 \\ \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ +1 & -1 & 0 & 0 & j\omega L_1 & j\omega k \cdot \sqrt{L_1 \cdot L_2} \\ 0 & 0 & +1 & -1 & j\omega k \cdot \sqrt{L_1 \cdot L_2} & j\omega L_2 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{br1} \\ I_{br2} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.77)$$

Obviously, this approach has an advantage: It also works for zero inductances and for unity coupling factors and is extendible for any number of inductors. It has the disadvantage that it enlarges the MNA matrix.

The S-parameter matrix of this component is obtained by converting the Z-parameter matrix of the component. The Z-parameter matrix can be constructed using the following scheme: The self-inductances on the main diagonal and the mutual inductances in the off-diagonal entries.

$$(\underline{Z}') = j\omega \cdot \begin{bmatrix} L_1 & k \cdot \sqrt{L_1 \cdot L_2} \\ k \cdot \sqrt{L_1 \cdot L_2} & L_2 \end{bmatrix} \quad (9.78)$$

This matrix representation does not contain the second terminals of the inductances. That's why the Z-parameter matrix must be converted into the Y-parameter matrix representation which is then extended to contain the additional terminals.

$$(\underline{Z}') \rightarrow (\underline{Y}') = \begin{bmatrix} y_{11} & y_{12} \\ y_{21} & y_{22} \end{bmatrix} \rightarrow (\underline{Y}) = \begin{bmatrix} +y_{11} & -y_{11} & +y_{12} & -y_{12} \\ -y_{11} & +y_{11} & -y_{12} & +y_{12} \\ +y_{21} & -y_{21} & +y_{22} & -y_{22} \\ -y_{21} & +y_{21} & -y_{22} & +y_{22} \end{bmatrix} \quad (9.79)$$

The resulting Y-parameter matrix can be converted into the appropriate S-parameters numerically by eqn. (??).

9.12 Attenuator

The ideal attenuator with (power) attenuation L is frequency independent and the model is valid for DC and for AC simulation. It is determined by the following Z parameters.

$$Z_{11} = Z_{22} = Z_{ref} \cdot \frac{L+1}{L-1} \quad (9.80)$$

$$Z_{12} = Z_{21} = Z_{ref} \cdot \frac{2 \cdot \sqrt{L}}{L-1} \quad (9.81)$$

The Z parameter representation is not very practical as new lines in the MNA matrix have to be added. More useful are the Y parameters.

$$(\underline{Y}) = \frac{1}{Z_{ref} \cdot (L-1)} \cdot \begin{pmatrix} L+1 & -2 \cdot \sqrt{L} \\ -2 \cdot \sqrt{L} & L+1 \end{pmatrix} \quad (9.82)$$

Attenuator with (power) attenuation L , reference impedance Z_{ref} and temperature T :

$$(\underline{C}_Y) = 4 \cdot k \cdot T \cdot \text{Re}(\underline{Y}) = \frac{4 \cdot k \cdot T}{Z_{ref} \cdot (L-1)} \cdot \begin{pmatrix} L+1 & -2 \cdot \sqrt{L} \\ -2 \cdot \sqrt{L} & L+1 \end{pmatrix} \quad (9.83)$$

The scattering parameters and noise wave correlation matrix of an ideal attenuator with (power) attenuation L (loss) (or power gain $G = 1/L$) in reference to the impedance Z_{ref} writes as follows.

$$S_{11} = S_{22} = \frac{r \cdot (L-1)}{L-r^2} = \frac{r \cdot (1-G)}{1-r^2 \cdot G} \quad (9.84)$$

$$S_{12} = S_{21} = \frac{\sqrt{L} \cdot (1-r^2)}{L-r^2} = \frac{\sqrt{G} \cdot (1-r^2)}{1-r^2 \cdot G} \quad (9.85)$$

$$(\underline{C}) = k \cdot T \cdot \frac{(L-1) \cdot (r^2-1)}{(L-r^2)^2} \cdot \begin{pmatrix} -r^2-L & 2 \cdot r\sqrt{L} \\ 2 \cdot r\sqrt{L} & -r^2-L \end{pmatrix} \quad (9.86)$$

with

$$r = \frac{Z_{ref} - Z_0}{Z_{ref} + Z_0} \quad (9.87)$$

9.13 Amplifier

An ideal amplifier increases signal strength from input to output and blocks all signals flowing into the output. The ideal amplifier is an isolator with voltage gain G and is determined by the following Z or Y parameters (valid for DC and AC simulation).

$$Z_{11} = Z_1 \quad Z_{12} = 0 \quad (9.88)$$

$$Z_{21} = 2 \cdot \sqrt{Z_1 \cdot Z_2} \cdot G \quad Z_{22} = Z_2 \quad (9.89)$$

$$Y_{11} = \frac{1}{Z_1} \quad Y_{12} = 0 \quad (9.90)$$

$$Y_{21} = -\frac{2 \cdot G}{\sqrt{Z_1 \cdot Z_2}} \quad Y_{22} = \frac{1}{Z_2} \quad (9.91)$$

With the reference impedance of the input Z_1 and the one of the output Z_2 and the voltage amplification G , the scattering parameters of an ideal amplifier write as follows.

$$S_{11} = \frac{Z_1 - Z_0}{Z_1 + Z_0} \quad (9.92)$$

$$S_{12} = 0 \quad (9.93)$$

$$S_{22} = \frac{Z_2 - Z_0}{Z_2 + Z_0} \quad (9.94)$$

$$S_{21} = \frac{4 \cdot Z_0 \cdot \sqrt{Z_1 \cdot Z_2} \cdot G}{(Z_1 + Z_0) \cdot (Z_2 + Z_0)} \quad (9.95)$$

9.14 Isolator

An isolator is a one-way two-port, transporting incoming waves lossless from the input (port 1) to the output (port 2), but dissipating all waves flowing into the output. The ideal isolator with reference impedances Z_1 (input) and Z_2 (output) is determined by the following Z parameters (for DC and AC simulation).

$$Z_{11} = Z_1 \quad Z_{12} = 0 \quad (9.96)$$

$$Z_{21} = 2 \cdot \sqrt{Z_1 \cdot Z_2} \quad Z_{22} = Z_2 \quad (9.97)$$

A more useful MNA representation is with Y parameters.

$$(\underline{Y}) = \begin{pmatrix} \frac{1}{Z_1} & 0 \\ -\frac{2}{\sqrt{Z_1 \cdot Z_2}} & \frac{1}{Z_2} \end{pmatrix} \quad (9.98)$$

Isolator with reference impedance Z_1 (input) and Z_2 (output) and temperature T :

$$(\underline{C}_Y) = 4 \cdot k \cdot T \cdot \begin{pmatrix} \frac{1}{Z_1} & 0 \\ -\frac{2}{\sqrt{Z_1 \cdot Z_2}} & \frac{1}{Z_2} \end{pmatrix} \quad (9.99)$$

With the reference impedance of the input Z_1 and the one of the output Z_2 , the scattering parameters of an ideal isolator writes as follows.

$$S_{11} = \frac{Z_1 - Z_0}{Z_1 + Z_0} \quad (9.100)$$

$$S_{12} = 0 \quad (9.101)$$

$$S_{22} = \frac{Z_2 - Z_0}{Z_2 + Z_0} \quad (9.102)$$

$$S_{21} = \sqrt{1 - (S_{11})^2} \cdot \sqrt{1 - (S_{22})^2} \quad (9.103)$$

Being on temperature T , the noise wave correlation matrix of an ideal isolator with reference impedance Z_1 and Z_2 (input and output) writes as follows.

$$(\underline{C}) = \frac{4 \cdot k \cdot T \cdot Z_0}{(Z_1 + Z_0)^2} \cdot \begin{pmatrix} Z_1 & \sqrt{Z_1 \cdot Z_2} \cdot \frac{Z_0 - Z_1}{Z_0 + Z_2} \\ \sqrt{Z_1 \cdot Z_2} \cdot \frac{Z_0 - Z_1}{Z_0 + Z_2} & Z_2 \cdot \left(\frac{Z_1 - Z_0}{Z_2 + Z_0} \right)^2 \end{pmatrix} \quad (9.104)$$

9.15 Circulator

A circulator is a 3-port device, transporting incoming waves lossless from port 1 to port 2, from port 2 to port 3 and from port 3 to port 1. In all other directions, there is no energy flow. The ideal circulator cannot be characterized with Z or Y parameters, because their values are partly infinite. But implementing with S parameters is practical (see equation ??).

With the reference impedances Z_1 , Z_2 and Z_3 for the ports 1, 2 and 3 the scattering matrix of an ideal circulator writes as follows.

$$denom = 1 - r_1 \cdot r_2 \cdot r_3 \quad (9.105)$$

$$r_1 = \frac{Z_0 - Z_1}{Z_0 + Z_1}, \quad r_2 = \frac{Z_0 - Z_2}{Z_0 + Z_2}, \quad r_3 = \frac{Z_0 - Z_3}{Z_0 + Z_3} \quad (9.106)$$

$$S_{11} = \frac{r_2 \cdot r_3 - r_1}{denom}, \quad S_{22} = \frac{r_1 \cdot r_3 - r_2}{denom}, \quad S_{33} = \frac{r_1 \cdot r_2 - r_3}{denom} \quad (9.107)$$

$$S_{12} = \sqrt{\frac{Z_2}{Z_1}} \cdot \frac{Z_1 + Z_0}{Z_2 + Z_0} \cdot \frac{r_3 \cdot (1 - r_1^2)}{denom}, \quad S_{13} = \sqrt{\frac{Z_3}{Z_1}} \cdot \frac{Z_1 + Z_0}{Z_3 + Z_0} \cdot \frac{1 - r_1^2}{denom} \quad (9.108)$$

$$S_{21} = \sqrt{\frac{Z_1}{Z_2}} \cdot \frac{Z_2 + Z_0}{Z_1 + Z_0} \cdot \frac{1 - r_2^2}{denom}, \quad S_{23} = \sqrt{\frac{Z_3}{Z_2}} \cdot \frac{Z_2 + Z_0}{Z_3 + Z_0} \cdot \frac{r_1 \cdot (1 - r_2^2)}{denom} \quad (9.109)$$

$$S_{31} = \sqrt{\frac{Z_1}{Z_3}} \cdot \frac{Z_3 + Z_0}{Z_1 + Z_0} \cdot \frac{r_2 \cdot (1 - r_3^2)}{denom}, \quad S_{32} = \sqrt{\frac{Z_2}{Z_3}} \cdot \frac{Z_3 + Z_0}{Z_2 + Z_0} \cdot \frac{1 - r_3^2}{denom} \quad (9.110)$$

An ideal circulator is noise free.

9.16 Phase shifter

A phase shifter alters the phase of the input signal independently on the frequency. As a result the relation between input and output signal is complex. To get the DC model, some simulators use the AC formulas and create the real part or the magnitude. This procedure has no physical reason, because it uses an operation that is not defined for DC. But one can think in the following direction: As a DC quantity is constant, it doesn't change if it is phase-shifted. (An AC quantity doesn't change its magnitude, too.) Or to say it with other words, for a DC simulation the phase to shift is always zero. That leads to the result that the phase shifter is a short circuit for DC. So, this is true for all reference impedances.

For an AC simulation, the Z-parameters of a phase shifter writes as follows.

$$Z_{11} = Z_{22} = \frac{j \cdot Z_{ref}}{\tan(\phi)} \quad (9.111)$$

$$Z_{12} = Z_{21} = \frac{j \cdot Z_{ref}}{\sin(\phi)} \quad (9.112)$$

The admittance parameters required for the AC analysis result in

$$Y_{11} = Y_{22} = \frac{j}{Z_{ref} \cdot \tan(\phi)} \quad (9.113)$$

$$Y_{12} = Y_{21} = \frac{1}{j \cdot Z_{ref} \cdot \sin(\phi)} \quad (9.114)$$

where ϕ denotes the actual phase shift of the device. For a zero phase shift ($\phi = 0$) neither the Z- nor the Y-parameters are defined. That is why during AC analysis a phase shifter with zero phase shift represents an ideal short circuit regardless its reference impedance.

The MNA matrix entries of an ideal short circuit during AC and DC analysis correspond to a voltage source with zero voltage. The complete MNA matrix representation writes as follows

$$\begin{bmatrix} . & . & +1 \\ . & . & -1 \\ +1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{br} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.115)$$

whence I_{br} denote the branch current through the voltage source.

The scattering parameters of an ideal phase shifter with phase shift ϕ and reference impedance Z_{ref} writes as follows.

$$r = \frac{Z_{ref} - Z_0}{Z_{ref} + Z_0} \quad (9.116)$$

$$S_{11} = S_{22} = \frac{r \cdot (1 - \exp(j \cdot 2\phi))}{1 - r^2 \cdot \exp(j \cdot 2\phi)} \quad (9.117)$$

$$S_{12} = S_{21} = \frac{(1 - r^2) \cdot \exp(j \cdot \phi)}{1 - r^2 \cdot \exp(j \cdot 2\phi)} \quad (9.118)$$

An ideal phase shifter is noise free.

9.17 Coupler

According to the port numbers in fig. ?? the Y-parameters of a coupler write as follows.

$$Y_{11} = Y_{22} = Y_{33} = Y_{44} = \frac{A \cdot (2 - A)}{D} \quad (9.119)$$

$$Y_{12} = Y_{21} = Y_{34} = Y_{43} = \frac{-A \cdot B}{D} \quad (9.120)$$

$$Y_{13} = Y_{31} = Y_{24} = Y_{42} = \frac{C \cdot (A - 2)}{D} \quad (9.121)$$

$$Y_{14} = Y_{41} = Y_{23} = Y_{32} = \frac{B \cdot C}{D} \quad (9.122)$$

$$(9.123)$$

with

$$A = k^2 \cdot (1 + \exp(j \cdot 2\phi)) \quad (9.124)$$

$$B = 2 \cdot \sqrt{1 - k^2} \quad (9.125)$$

$$C = 2 \cdot k \cdot \exp(j \cdot \phi) \quad (9.126)$$

$$D = Z_{ref} \cdot (A^2 - C^2) \quad (9.127)$$

$$(9.128)$$

whereas $0 < k < 1$ denotes the coupling factor, ϕ the phase shift of the coupling path and Z_{ref} the reference impedance. The coupler can also be used as hybrid by setting $k = 1/\sqrt{2}$. For a 90

degree hybrid, for example, set ϕ to $\pi/2$. Note that for most couplers no real DC model exists. Taking the real part of the AC matrix often leads to non-logical results. Thus, it is better to model the coupler for DC by making a short between port 1 and port 2 and between port 3 and port 4. The rest should be an open. This leads to the following MNA matrix.

$$\begin{bmatrix} . & . & . & . & 1 & 0 \\ . & . & . & . & -1 & 0 \\ . & . & . & . & 0 & 1 \\ . & . & . & . & 0 & -1 \\ 1 & -1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{out1} \\ I_{out4} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.129)$$

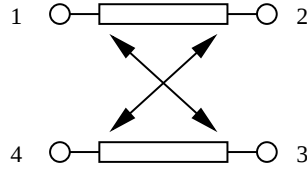


Figure 9.7: ideal coupler device

The scattering parameters of a coupler are:

$$S_{11} = S_{22} = S_{33} = S_{44} = 0 \quad (9.130)$$

$$S_{14} = S_{23} = S_{32} = S_{41} = 0 \quad (9.131)$$

$$S_{12} = S_{21} = S_{34} = S_{43} = \sqrt{1 - k^2} \quad (9.132)$$

$$S_{13} = S_{31} = S_{24} = S_{42} = k \cdot \exp(j\phi) \quad (9.133)$$

whereas $0 < k < 1$ denotes the coupling factor, ϕ the phase shift of the coupling path. Extending them for an arbitrary reference impedance Z_{ref} , they already become quite complex:

$$r = \frac{Z_0 - Z_{ref}}{Z_0 + Z_{ref}} \quad (9.134)$$

$$A = k^2 \cdot (\exp(j \cdot 2\phi) + 1) \quad (9.135)$$

$$B = r^2 \cdot (1 - A) \quad (9.136)$$

$$C = k^2 \cdot (\exp(j \cdot 2\phi) - 1) \quad (9.137)$$

$$D = 1 - 2 \cdot r^2 \cdot (1 + C) + B^2 \quad (9.138)$$

$$(9.139)$$

$$S_{11} = S_{22} = S_{33} = S_{44} = r \cdot \frac{A \cdot B + C + 2 \cdot r^2 \cdot k^2 \cdot \exp(j \cdot 2\phi)}{D} \quad (9.140)$$

$$S_{12} = S_{21} = S_{34} = S_{43} = \sqrt{1 - k^2} \cdot \frac{(1 - r^2) \cdot (1 - B)}{D} \quad (9.141)$$

$$S_{13} = S_{31} = S_{24} = S_{42} = k \cdot \exp(j\phi) \cdot \frac{(1 - r^2) \cdot (1 + B)}{D} \quad (9.142)$$

$$S_{14} = S_{23} = S_{32} = S_{41} = 2 \cdot \sqrt{1 - k^2} \cdot k \cdot \exp(j\phi) \cdot r \cdot \frac{(1 - r^2)}{D} \quad (9.143)$$

An ideal coupler is noise free.

9.18 Gyrator

A gyrator is an impedance inverter. Thus, for example, it converts a capacitance into an inductance and vice versa. The ideal gyrator, as shown in fig. ??, is determined by the following equations which introduce two more unknowns in the MNA matrix.

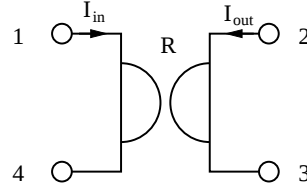


Figure 9.8: ideal gyrator

$$I_{in} = \frac{1}{R} \cdot (V_2 - V_3) \rightarrow \frac{1}{R} \cdot V_2 - \frac{1}{R} \cdot V_3 - I_{in} = 0 \quad (9.144)$$

$$I_{out} = -\frac{1}{R} \cdot (V_1 - V_4) \rightarrow -\frac{1}{R} \cdot V_1 + \frac{1}{R} \cdot V_4 - I_{out} = 0 \quad (9.145)$$

The new unknown variables I_{out} and I_{in} must be considered by the four remaining simple equations.

$$I_1 = I_{in} \quad I_2 = I_{out} \quad I_3 = -I_{out} \quad I_4 = -I_{in} \quad (9.146)$$

As can be seen, a gyrator consists of two cross-connected VCCS (section ??). Hence, its y-parameter matrix is:

$$(\underline{Y}) = \begin{bmatrix} 0 & \frac{1}{R} & -\frac{1}{R} & 0 \\ -\frac{1}{R} & 0 & 0 & \frac{1}{R} \\ \frac{1}{R} & 0 & 0 & -\frac{1}{R} \\ 0 & -\frac{1}{R} & \frac{1}{R} & 0 \end{bmatrix} \quad (9.147)$$

The scattering matrix of an ideal gyrator with the ratio R writes as follows.

$$r = \frac{R}{Z_{ref}} = \frac{1}{G \cdot Z_{ref}} \quad (9.148)$$

$$S_{11} = S_{22} = S_{33} = S_{44} = \frac{R^2}{4 \cdot Z_{ref}^2 + R^2} = \frac{r^2}{r^2 + 4} \quad (9.149)$$

$$S_{14} = S_{23} = S_{32} = S_{41} = 1 - S_{11} \quad (9.150)$$

$$S_{12} = -S_{13} = -S_{21} = S_{24} = S_{31} = -S_{34} = -S_{42} = S_{43} = \frac{2 \cdot r}{r^2 + 4} \quad (9.151)$$

9.19 Voltage and current sources

For an AC analysis, DC sources are short circuit (voltage source) or open circuit (current source), respectively. Accordingly, for a DC analysis, AC sources are short circuit (voltage source) or open circuit (current source), respectively. As these sources have no internal resistance, they are noise-free.

The MNA matrix of a current source is (with short circuit current I_0 flowing into node 1 and out of node 2):

$$\begin{bmatrix} \cdot & \cdot \\ \cdot & \cdot \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} I_0 \\ -I_0 \end{bmatrix} \quad (9.152)$$

The MNA matrix of a voltage source is (with open circuit voltage U_0 across node 1 to node 2):

$$\begin{bmatrix} \cdot & \cdot & 1 \\ \cdot & \cdot & -1 \\ 1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{in} \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ U_0 \end{bmatrix} \quad (9.153)$$

The MNA matrix of a power source is (with internal resistance R and available power P):

$$\begin{bmatrix} \frac{1}{R} & -\frac{1}{R} \\ -\frac{1}{R} & \frac{1}{R} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} \sqrt{\frac{8 \cdot P}{R}} \\ -\sqrt{\frac{8 \cdot P}{R}} \end{bmatrix} \quad (9.154)$$

The factor "8" is because of:

- transforming peak current value into effective value (factor two)
- at power matching the internal resistance dissipates the same power as the load (gives factor four).

The noise current correlation matrix of a power source equals the one of a resistor with resistance R .

All voltage sources (AC and DC) are short circuits and therefore their S-parameter matrix equals the one of the DC block. All current sources are open circuits and therefore their S-parameter matrix equals the one of the DC feed.

9.20 Noise sources

To implement the frequency dependencies of all common noise PSDs the following equation can be used.

$$PSD = \frac{PSD_0}{a + b \cdot f^c} \quad (9.155)$$

Where f is frequency and a , b , c are the parameters. The following PSDs appear in electric devices.

white noise (thermal noise, shot noise):	$a = 0, b = 1, c = 0$
pink noise (flicker noise):	$a = 0, b = 1, c = 1$
Lorentzian PSD (generation-recombination noise):	$a = 1, b = 1/f_c^2, c = 2$

9.20.1 Noise current source

Noise current source with a current power spectral density of $cPSD$:

$$(\underline{C}_Y) = cPSD \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (9.156)$$

The MNA matrix entries for DC and AC analysis are all zero.

The noise wave correlation matrix of a noise current source with current power spectral density $cPSD$ and its S parameter matrix write as follows.

$$(\underline{C}) = cPSD \cdot Z_0 \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (\underline{S}) = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (9.157)$$

9.20.2 Noise voltage source

A noise voltage source (voltage power spectral density $vPSD$) cannot be modeled with the noise current matrix. That is why one has to use a noise current source (current power spectral density $cPSD$) connected to a gyrator (transimpedance R) satisfying the equation

$$vPSD = cPSD \cdot R^2 \quad (9.158)$$

Figure ?? shows an example.

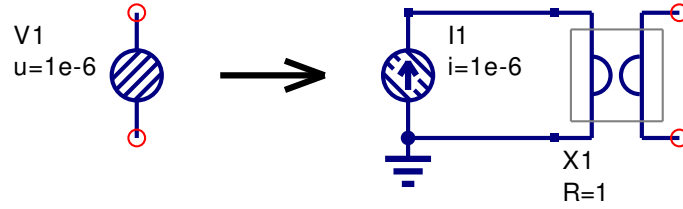


Figure 9.9: noise voltage source (left-hand side) and its equivalent circuit (right-hand side)

The MNA matrix entries of the above construct (gyrator ratio $R = 1$) is similar to a voltage source with zero voltage.

$$\begin{bmatrix} . & . & -1 \\ . & . & 1 \\ 1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_x \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.159)$$

The appropriate noise current correlation matrix yields:

$$(\underline{C}_Y) = cPSD \cdot \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (9.160)$$

The noise wave correlation matrix of a noise voltage source with voltage power spectral density $vPSD$ and its S parameter matrix write as follows.

$$(\underline{C}) = \frac{vPSD}{4 \cdot Z_0} \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (\underline{S}) = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad (9.161)$$

9.20.3 Correlated noise sources

For two correlated noise current sources the (normalized) correlation coefficient K must be known (with $|K| = 0 \dots 1$). If the first noise source has the current power spectral density S_{i1} and is

connected to node 1 and 2, and if furthermore the second noise source has the spectral density S_{i2} and is connected to node 3 and 4, then the correlation matrix writes:

$$(\underline{C}_Y) = \begin{pmatrix} S_{i1} & -S_{i1} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} \\ -S_{i1} & S_{i1} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} \\ K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & S_{i2} & -S_{i2} \\ -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -S_{i2} & S_{i2} \end{pmatrix} \quad (9.162)$$

The MNA matrix entries for DC and AC analysis are all zero.

The noise wave correlation matrix of two correlated noise current sources with current power spectral densities S_{i1} and S_{i2} and correlation coefficient K writes as follows.

$$(\underline{C}) = Z_0 \cdot \begin{pmatrix} S_{i1} & -S_{i1} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} \\ -S_{i1} & S_{i1} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} \\ K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & S_{i2} & -S_{i2} \\ -K \cdot \sqrt{S_{i1} \cdot S_{i2}} & K \cdot \sqrt{S_{i1} \cdot S_{i2}} & -S_{i2} & S_{i2} \end{pmatrix} \quad (9.163)$$

$$(\underline{S}) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \quad (9.164)$$

For two correlated noise voltage sources two extra rows and columns are needed in the MNA matrix:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ \cdot & \cdot & \cdot & \cdot & 0 & 1 \\ 1 & -1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{x1} \\ I_{x2} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.165)$$

The appropriate noise current correlation matrix (with the noise voltage power spectral densities S_{v1} and S_{v2} and the correlation coefficient K) yields:

$$(\underline{C}_Y) = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & S_{v1} & K \cdot \sqrt{S_{v1} \cdot S_{v2}} \\ 0 & 0 & 0 & 0 & K \cdot \sqrt{S_{v1} \cdot S_{v2}} & S_{v2} \end{pmatrix} \quad (9.166)$$

The noise wave correlation matrix of two correlated noise voltage sources with voltage power spectral densities S_{v1} and S_{v2} and correlation coefficient K and its S parameter matrix write as follows.

$$(\underline{C}) = \frac{1}{4 \cdot Z_0} \cdot \begin{pmatrix} S_{v1} & -S_{v1} & K \cdot \sqrt{S_{v1} \cdot S_{v2}} & -K \cdot \sqrt{S_{v1} \cdot S_{v2}} \\ -S_{v1} & S_{v1} & -K \cdot \sqrt{S_{v1} \cdot S_{v2}} & K \cdot \sqrt{S_{v1} \cdot S_{v2}} \\ K \cdot \sqrt{S_{v1} \cdot S_{v2}} & -K \cdot \sqrt{S_{v1} \cdot S_{v2}} & S_{v2} & -S_{v2} \\ -K \cdot \sqrt{S_{v1} \cdot S_{v2}} & K \cdot \sqrt{S_{v1} \cdot S_{v2}} & -S_{v2} & S_{v2} \end{pmatrix} \quad (9.167)$$

$$(\underline{S}) = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} \quad (9.168)$$

If a noise current source (ports 1 and 2) and a noise voltage source (ports 3 and 4) are correlated, the MNA matrix entries are as follows.

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & \cdot & 1 \\ 0 & 0 & 1 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_x \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.169)$$

The appropriate noise current correlation matrix (with the noise power spectral densities S_{i1} and S_{v2} and the correlation coefficient K) yields:

$$(\underline{C}_Y) = \begin{pmatrix} S_{i1} & -S_{i1} & 0 & 0 & K \cdot \sqrt{S_{i1} \cdot S_{v2}} \\ -S_{i1} & S_{i1} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ K \cdot \sqrt{S_{i1} \cdot S_{v2}} & 0 & 0 & 0 & S_{v2} \end{pmatrix} \quad (9.170)$$

Note: Because the gyrator factor (It is unity.) has been omitted in the above matrix the units are not correct. This can be ignored.

The noise wave correlation matrix of one correlated noise current source S_{i1} and one noise voltage source S_{v2} with correlation coefficient K writes as follows.

$$(\underline{C}) = \begin{pmatrix} Z_0 \cdot S_{i1} & -Z_0 \cdot S_{i1} & K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & -K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} \\ -Z_0 \cdot S_{i1} & Z_0 \cdot S_{i1} & -K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} \\ K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & -K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & S_{v2}/4/Z_0 & -S_{v2}/4/Z_0 \\ -K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & K/2 \cdot \sqrt{S_{i1} \cdot S_{v2}} & -S_{v2}/4/Z_0 & S_{v2}/4/Z_0 \end{pmatrix} \quad (9.171)$$

$$(\underline{S}) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} \quad (9.172)$$

9.21 Controlled sources

The models of the controlled sources contain the transfer factor G . It is complex because of the delay time T and frequency f .

$$\underline{G} = G \cdot e^{j\omega T} = G \cdot e^{j \cdot 2\pi f \cdot T} \quad (9.173)$$

During a DC analysis (frequency zero) it becomes real because the exponent factor is unity.

9.21.1 Voltage controlled current source

The voltage-dependent current source (VCCS), as shown in fig. ??, is determined by the following equation which introduces one more unknown in the MNA matrix.

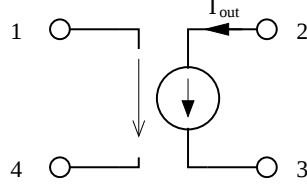


Figure 9.10: voltage controlled current source

$$I_{out} = G \cdot (V_1 - V_4) \quad \rightarrow \quad V_1 - V_4 - \frac{1}{G} \cdot I_{out} = 0 \quad (9.174)$$

The new unknown variable I_{out} must be considered by the four remaining simple equations.

$$I_1 = 0 \quad I_2 = I_{out} \quad I_3 = -I_{out} \quad I_4 = 0 \quad (9.175)$$

And in matrix representation this is:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & \cdot & 0 \\ 1 & 0 & 0 & -1 & -\frac{1}{G} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.176)$$

As you can see the last row which has been added by the VCCS represents the determining equation (??). The additional right hand column in the matrix keeps the system consistent.

When pivoting the above MNA stamp (??) the additional row and column can be saved ensuring G keeps finite (the pivot element must be non-zero). Both representations are equivalent. If G is zero the below representation must be used.

$$\begin{bmatrix} 0 & 0 & 0 & 0 \\ G & 0 & 0 & -G \\ -G & 0 & 0 & G \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.177)$$

The scattering matrix of the voltage controlled current source writes as follows (τ is time delay).

$$S_{11} = S_{22} = S_{33} = S_{44} = 1 \quad (9.178)$$

$$S_{12} = S_{13} = S_{14} = S_{23} = S_{32} = S_{41} = S_{42} = S_{43} = 0 \quad (9.179)$$

$$S_{21} = S_{34} = -2 \cdot G \cdot \exp(-j\omega\tau) \quad (9.180)$$

$$S_{24} = S_{31} = 2 \cdot G \cdot \exp(-j\omega\tau) \quad (9.181)$$

9.21.2 Current controlled current source

The current-dependent current source (CCCS), as shown in fig. ??, is determined by the following equation which introduces one more unknown in the MNA matrix.

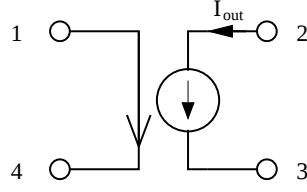


Figure 9.11: current controlled current source

$$V_1 - V_4 = 0 \quad (9.182)$$

The new unknown variable I_{out} must be considered by the four remaining simple equations.

$$I_1 = +\frac{1}{G} \cdot I_{out} \quad I_2 = I_{out} \quad I_3 = -I_{out} \quad I_4 = -\frac{1}{G} \cdot I_{out} \quad (9.183)$$

And in matrix representation this is:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & \frac{1}{G} \\ \cdot & \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & \cdot & -\frac{1}{G} \\ 1 & 0 & 0 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.184)$$

The scattering matrix of the current controlled current source writes as follows (τ is time delay).

$$S_{14} = S_{22} = S_{33} = S_{41} = 1 \quad (9.185)$$

$$S_{11} = S_{12} = S_{13} = S_{23} = S_{32} = S_{42} = S_{43} = S_{44} = 0 \quad (9.186)$$

$$S_{21} = S_{34} = -G \cdot \exp(-j\omega\tau) \quad (9.187)$$

$$S_{24} = S_{31} = G \cdot \exp(-j\omega\tau) \quad (9.188)$$

9.21.3 Voltage controlled voltage source

The voltage-dependent voltage source (VCVS), as shown in fig. ??, is determined by the following equation which introduces one more unknown in the MNA matrix.

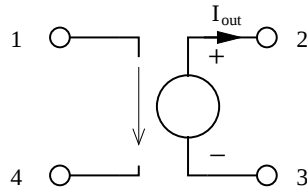


Figure 9.12: voltage controlled voltage source

$$V_2 - V_3 = G \cdot (V_1 - V_4) \quad \rightarrow \quad V_1 \cdot G - V_2 + V_3 - V_4 \cdot G = 0 \quad (9.189)$$

The new unknown variable I_{out} must be considered by the four remaining simple equations.

$$I_1 = 0 \quad I_2 = -I_{out} \quad I_3 = I_{out} \quad I_4 = 0 \quad (9.190)$$

And in matrix representation this is:

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & -1 \\ \cdot & \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & \cdot & 0 \\ G & -1 & 1 & -G & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.191)$$

The scattering matrix of the voltage controlled voltage source writes as follows (τ is time delay).

$$S_{11} = S_{23} = S_{32} = S_{44} = 1 \quad (9.192)$$

$$S_{12} = S_{13} = S_{14} = S_{22} = S_{33} = S_{41} = S_{42} = S_{43} = 0 \quad (9.193)$$

$$S_{21} = S_{34} = G \cdot \exp(-j\omega\tau) \quad (9.194)$$

$$S_{24} = S_{31} = -G \cdot \exp(-j\omega\tau) \quad (9.195)$$

9.21.4 Current controlled voltage source

The current-dependent voltage source (CCVS), as shown in fig. ??, is determined by the following equations which introduce two more unknowns in the MNA matrix.

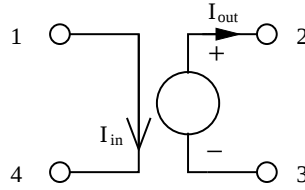


Figure 9.13: current controlled voltage source

$$V_1 - V_4 = 0 \quad (9.196)$$

$$V_2 - V_3 = G \cdot I_{in} \rightarrow V_2 - V_3 - I_{in} \cdot G = 0 \quad (9.197)$$

The new unknown variables I_{out} and I_{in} must be considered by the four remaining simple equations.

$$I_1 = I_{in} \quad I_2 = -I_{out} \quad I_3 = I_{out} \quad I_4 = -I_{in} \quad (9.198)$$

The matrix representation needs to be augmented by two more new rows (for the new unknown variables) and their corresponding columns.

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & 1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ \cdot & \cdot & \cdot & \cdot & 0 & 1 \\ \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ 0 & 1 & -1 & 0 & -G & 0 \\ 1 & 0 & 0 & -1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \\ I_{in} \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ I_4 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (9.199)$$

The scattering matrix of the current controlled voltage source writes as follows (τ is time delay).

$$S_{14} = S_{23} = S_{32} = S_{41} = 1 \quad (9.200)$$

$$S_{11} = S_{12} = S_{13} = S_{22} = S_{33} = S_{42} = S_{43} = S_{44} = 0 \quad (9.201)$$

$$S_{21} = S_{34} = \frac{G}{2} \cdot \exp(-j\omega\tau) \quad (9.202)$$

$$S_{24} = S_{31} = -\frac{G}{2} \cdot \exp(-j\omega\tau) \quad (9.203)$$

9.22 Transmission Line

A transmission line is usually described by its ABCD-matrix. (Note that in ABCD-matrices, i.e. the chain matrix representation, the current i_2 is defined to flow out of the output port.)

$$(\underline{A}) = \begin{pmatrix} \cosh(\gamma \cdot l) & Z_L \cdot \sinh(\gamma \cdot l) \\ \sinh(\gamma \cdot l)/Z_L & \cosh(\gamma \cdot l) \end{pmatrix} \quad (9.204)$$

These can easily be recalculated into impedance parameters.

$$Z_{11} = Z_{22} = \frac{Z_L}{\tanh(\gamma \cdot l)} \quad (9.205)$$

$$Z_{12} = Z_{21} = \frac{Z_L}{\sinh(\gamma \cdot l)} \quad (9.206)$$

Or in admittance parameter representation it yields

$$\begin{aligned} Y_{11} = Y_{22} &= \frac{1}{Z_L \cdot \tanh(\gamma \cdot l)} \\ Y_{12} = Y_{21} &= \frac{-1}{Z_L \cdot \sinh(\gamma \cdot l)} \end{aligned} \quad (9.207)$$

whence γ denotes the propagation constant $\alpha + j\beta$ and l is the length of the transmission line. Z_L represents the characteristic impedance of the transmission line. The Y-parameters as defined by eq. (??) can be used for the microstrip line. For an ideal, i.e. lossless, transmission lines they write accordingly.

$$Z_{11} = Z_{22} = \frac{Z_L}{j \cdot \tan(\beta \cdot l)} \quad (9.208)$$

$$Z_{12} = Z_{21} = \frac{Z_L}{j \cdot \sin(\beta \cdot l)} \quad (9.209)$$

$$Y_{11} = Y_{22} = \frac{1}{j \cdot Z_L \cdot \tan(\beta \cdot l)} \quad (9.210)$$

$$Y_{12} = Y_{21} = \frac{j}{Z_L \cdot \sin(\beta \cdot l)} \quad (9.211)$$

The scattering matrix of an ideal, lossless transmission line with impedance Z and electrical length l writes as follows.

$$r = \frac{Z - Z_0}{Z + Z_0} \quad (9.212)$$

$$p = \exp\left(-j\omega \frac{l}{c_0}\right) \quad (9.213)$$

$$S_{11} = S_{22} = \frac{r \cdot (1 - p^2)}{1 - r^2 \cdot p^2} \quad , \quad S_{12} = S_{21} = \frac{p \cdot (1 - r^2)}{1 - r^2 \cdot p^2} \quad (9.214)$$

With $c_0 = 299\,792\,458$ m/s being the vacuum light velocity. Adding attenuation to the transmission line, the quantity p extends to:

$$p = \exp\left(-j\omega \frac{l}{c_0} - \alpha \cdot l\right) \quad (9.215)$$

Another equivalent equation set for the calculation of the scattering parameters is the following: With the physical length l of the component, its impedance Z_L and propagation constant γ , the complex propagation constant γ is given by

$$\gamma = \alpha + j\beta \quad (9.216)$$

where α is the attenuation factor and β is the (real) propagation constant given by

$$\beta = \sqrt{\varepsilon_{eff}(\omega)} \cdot k_0 \quad (9.217)$$

where $\varepsilon_{eff}(\omega)$ is the effective dielectric constant and k_0 is the TEM propagation constant k_0 for the equivalent transmission line with an air dielectric.

$$k_0 = \omega \sqrt{\varepsilon_0 \mu_0} \quad (9.218)$$

The expressions used to calculate the scattering parameters are given by

$$S_{11} = S_{22} = \frac{(z - y) \sinh \gamma l}{2 \cosh \gamma l + (z + y) \sinh \gamma l} \quad (9.219)$$

$$S_{12} = S_{21} = \frac{2}{2 \cosh \gamma l + (z + y) \sinh \gamma l} \quad (9.220)$$

with z being the normalized impedance and y is the normalized admittance.

9.23 Differential Transmission Line

A differential (4-port) transmission line is not referenced to ground potential, i.e. the wave from the input (port 1 and 4) is distributed to the output (port 2 and 3). Its admittance parameters are:

$$Y_{11} = Y_{22} = Y_{33} = Y_{44} = -Y_{14} = -Y_{41} = -Y_{23} = -Y_{32} = \frac{1}{Z_L \cdot \tanh(\gamma \cdot l)} \quad (9.221)$$

$$Y_{13} = Y_{31} = Y_{24} = Y_{42} = -Y_{12} = -Y_{21} = -Y_{34} = -Y_{43} = \frac{1}{Z_L \cdot \sinh(\gamma \cdot l)} \quad (9.222)$$

The scattering parameters writes:

$$S_{11} = S_{22} = S_{33} = S_{44} = Z_L \cdot \frac{(2 \cdot Z_0 + Z_L) \cdot \exp(2 \cdot \gamma \cdot l) + (2 \cdot Z_0 - Z_L)}{(2 \cdot Z_0 + Z_L)^2 \cdot \exp(2 \cdot \gamma \cdot l) - (2 \cdot Z_0 - Z_L)^2} \quad (9.223)$$

$$S_{14} = S_{41} = S_{23} = S_{32} = 1 - S_{11} \quad (9.224)$$

$$S_{12} = S_{21} = S_{34} = S_{43} = -S_{13} = -S_{31} = -S_{24} = -S_{42} \quad (9.225)$$

$$= \frac{4 \cdot Z_L \cdot Z_0 \cdot \exp(\gamma \cdot l)}{(2 \cdot Z_0 + Z_L)^2 \cdot \exp(2 \cdot \gamma \cdot l) - (2 \cdot Z_0 - Z_L)^2} \quad (9.226)$$

Note: As already stated, this is a pure differential transmission line without ground reference. It is *not* a three-wire system. I.e. there is only one mode. The next section describes a differential line with ground reference.

9.24 Coupled transmission line

A coupled transmission line is described by two identical transmission line ABCD-matrices, one for the even mode (or common mode) and one for the odd mode (or differential mode). Because the coupled lines are a symmetrical 3-line system, the matrices are completely independent of each other. Therefore, its Y-parameters write as follows.

$$Y_{11} = Y_{22} = Y_{33} = Y_{44} = \frac{1}{2 \cdot Z_{L,e} \cdot \tanh(\gamma_e \cdot l)} + \frac{1}{2 \cdot Z_{L,o} \cdot \tanh(\gamma_o \cdot l)} \quad (9.227)$$

$$Y_{12} = Y_{21} = Y_{34} = Y_{43} = \frac{-1}{2 \cdot Z_{L,e} \cdot \sinh(\gamma_e \cdot l)} + \frac{-1}{2 \cdot Z_{L,o} \cdot \sinh(\gamma_o \cdot l)} \quad (9.228)$$

$$Y_{13} = Y_{31} = Y_{24} = Y_{42} = \frac{-1}{2 \cdot Z_{L,e} \cdot \sinh(\gamma_e \cdot l)} + \frac{1}{2 \cdot Z_{L,o} \cdot \sinh(\gamma_o \cdot l)} \quad (9.229)$$

$$Y_{14} = Y_{41} = Y_{23} = Y_{32} = \frac{1}{2 \cdot Z_{L,e} \cdot \tanh(\gamma_e \cdot l)} + \frac{-1}{2 \cdot Z_{L,o} \cdot \tanh(\gamma_o \cdot l)} \quad (9.230)$$

The S-parameters (according to the port numbering in fig. ??) are as followed [?].

reflection coefficients

$$S_{11} = S_{22} = S_{33} = S_{44} = X_e + X_o \quad (9.231)$$

through paths

$$S_{12} = S_{21} = S_{34} = S_{43} = Y_e + Y_o \quad (9.232)$$

coupled paths

$$S_{14} = S_{41} = S_{23} = S_{32} = X_e - X_o \quad (9.233)$$

isolated paths

$$S_{13} = S_{31} = S_{24} = S_{42} = Y_e - Y_o \quad (9.234)$$

with the denominator

$$D_{e,o} = 2 \cdot Z_{L,e,o} \cdot Z_0 \cdot \cosh(\gamma_{e,o} \cdot l) + (Z_{L,e,o}^2 + Z_0^2) \cdot \sinh(\gamma_{e,o} \cdot l) \quad (9.235)$$

and

$$X_{e,o} = \frac{(Z_{L,e,o}^2 - Z_0^2) \cdot \sinh(\gamma_{e,o} \cdot l)}{2 \cdot D_{e,o}} \quad (9.236)$$

$$Y_{e,o} = \frac{Z_{L,e,o} \cdot Z_0}{D_{e,o}} \quad (9.237)$$

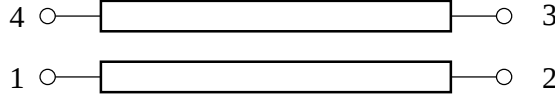


Figure 9.14: coupled transmission line

9.25 Tapered transmission line

The tapered transmission line is a component used mainly for achieving broadband impedance matching between a load, Z_L , and a source impedance, Z_S (both real). It differs from the conventional ideal transmission line in the sense that its characteristic impedance, $Z_0(l)$, depends on the distance from the origin.

The reflection coefficient at a given distance from the origin, $\Gamma(l)$, is given by:

$$\Gamma(l) = \frac{(Z_0(l) + \frac{\partial Z_0(l)}{\partial l}) - Z_0(l)}{(Z_0(l) + \frac{\partial Z_0(l)}{\partial l}) + Z_0(l)} = \frac{\frac{\partial Z_0(l)}{\partial l}}{2Z_0(l) + \frac{\partial Z_0(l)}{\partial l}} \quad (9.238)$$

Under the assumption that $2Z_0 \gg \frac{\partial Z_0(l)}{\partial l}$:

$$\Gamma(l) \sim \frac{\frac{\partial Z_0(l)}{\partial l}}{2Z_0(l)} \quad (9.239)$$

Then, the derivative of the reflection coefficient is given by:

$$\frac{\partial \Gamma(l)}{\partial l} = \frac{\partial \ln(Z_0(l))}{2 \cdot \partial l} \quad (9.240)$$

The *Theory of the small reflections* states that the reflection coefficient of a cascade of N lines of length l each is given by the sum of the reflection coefficients at each junction ([?], Eq. 5.44):

$$\Gamma(\beta L) = \sum_{k=0}^N \Gamma_i e^{-2j\beta l \cdot k} \quad , \quad L = Nl \quad (9.241)$$

In the case of a tapered line, the length of the portions of line where the characteristic impedance is constant tend to 0, so the summation in ?? becomes an integral expression as follows:

$$\Gamma(\beta L) = \int_0^L \frac{\partial \Gamma(l)}{\partial l} e^{-2j\beta l} dl = \frac{1}{2} \int_0^L \frac{\partial}{\partial l} \ln(Z_0(l)) e^{-2j\beta l} dl \quad (9.242)$$

The most common types of tapered transmission lines are the following:

Exponential taper The characteristic impedance and the reflection coefficient are given by the following expressions:

$$Z(l) = Z_0 \cdot e^{\frac{1}{L} \ln\left(\frac{Z_L}{Z_S}\right)l} \quad \forall l \in [0, L] \quad (9.243)$$

$$\Gamma(\beta L) = \frac{1}{2} \ln\left(\frac{Z_L}{Z_S}\right) e^{-j\beta L} \frac{\sin(\beta L)}{\beta L} \quad (9.244)$$

Triangular taper The characteristic impedance and the reflection coefficient are given by the following expressions:

$$Z(l) = \begin{cases} Z_S e^{2(\frac{l}{L})^2 \ln(\frac{Z_L}{Z_S})} & \text{for } 0 \leq l \leq L/2 \\ Z_S e^{(4\frac{l}{L} - 2)(\frac{l}{L})^2 \ln(\frac{Z_L}{Z_S})} & \text{for } L/2 \leq l \leq L \end{cases} \quad (9.245)$$

Note that the derivative of the reflection coefficient, computed according to ??, has a *triangular* shape.

$$\Gamma(\beta L) = \frac{1}{2} e^{-j\beta L} \ln\left(\frac{Z_L}{Z_S}\right) \left(\frac{\sin(\beta L/2)}{\beta L/2}\right)^2 \quad (9.246)$$

Klopfenstein taper The characteristic impedance of the Klopfenstein taper satisfies the following condition:

$$\ln(Z(l)) = \frac{1}{2} \ln(Z_S Z_L) + \frac{\Gamma_0}{\cosh(A)} A^2 \phi\left(2\frac{l}{L} - 1, A\right) \quad \forall l \in [0, L] \quad (9.247)$$

where

$$\phi(x, A) = -\phi(-x, A) = \int_0^x \frac{I_1\left(A\sqrt{1-y^2}\right)}{A\sqrt{1-y^2}} dy \quad \forall |x| \leq 1 \quad (9.248)$$

$I_1(x)$ is the modified Bessel function of the first kind and A can be determined from the specified maximum ripple of the Klopfenstein taper, Γ_m and the reflection coefficient at zero frequency, Γ_0 :

$$\Gamma_0 = \frac{Z_L - Z_S}{Z_L + Z_S} \sim \frac{1}{2} \ln\left(\frac{Z_L}{Z_S}\right) \quad (9.249)$$

$$A = \cosh^{-1}\left(\frac{\Gamma_0}{\Gamma_m}\right) \quad (9.250)$$

The reflection coefficient, $\Gamma(\beta L)$, of the Klopfenstein taper is given by the following expression:

$$\Gamma(\beta L) = \Gamma_0 e^{-j\beta L} \frac{\cos\left(\sqrt{(\beta L)^2 - A^2}\right)}{\cosh(A)} \quad \forall \beta L > A \quad (9.251)$$

The complexity of ?? makes it unsuitable for simulation purposes. In this sense, the following procedure reduces significantly the computational cost of evaluating $\phi(x, A)$ [?]:

Expanding $I_1(x)$ in power series, $\phi(x, A)$ can be rewritten as:

$$\phi(x, A) = \int_0^x \frac{1}{2} \sum_{k=0}^{\infty} \frac{\left(\frac{A^2}{4}\right)^k (1-y^2)^k}{k!(k+1)!} dy, \quad y \in [-1, 1] \quad (9.252)$$

The series can be integrated term by term and the function above expressed as

$$\phi(x, A) = \sum_{k=0}^{\infty} a_k \cdot b_k \quad (9.253)$$

where

$$a_k = \frac{A^{2k}}{4^k k! (k+1)!} \quad (9.254)$$

$$b_k = \frac{1}{2} \int_0^x (1-y^2)^k dy \quad (9.255)$$

Finally, the a_k and b_k coefficients can be expressed in terms of the following recursion:

$$a_0 = 1, \quad a_k = a_{k-1} \cdot \frac{A^2}{4k(k+1)} \quad (9.256)$$

$$b_0 = \frac{x}{2}, \quad b_k = \frac{\frac{x}{2}(1-x^2)^k + 2 \cdot k \cdot b_{k-1}}{2 \cdot k + 1} \quad (9.257)$$

Linear The characteristic impedance profile of the linear taper is given by:

$$Z_0(l) = Z_S + \left(l \cdot \frac{Z_L - Z_S}{L} \right) \quad (9.258)$$

The reflection coefficient is calculated using ??,

$$\Gamma(\beta L) = \frac{1}{2} \int_0^L e^{-j2\beta l} \frac{\partial}{\partial l} \ln(Z_0(l)) dl = \frac{1}{2} \int_0^L e^{-j2\beta l} \frac{\partial}{\partial l} \ln \left(Z_S + l \cdot \left(\frac{Z_L - Z_S}{L} \right) \right) dl \quad (9.259)$$

$$\Gamma(\beta L) = \frac{Z_L - Z_S}{2L} \int_0^L \frac{e^{-j2\beta l}}{Z_S + l \left(\frac{Z_L - Z_S}{L} \right)} dl \quad (9.260)$$

The following equivalence was found using a symbolic mathematical software (like [?]):

$$\int \frac{e^{-jAx}}{B + xC} = \frac{e^{jAB} \cdot \mathbb{E}_i \left(-\frac{jA(B+Cx)}{C} \right)}{C} \quad (9.261)$$

where $\mathbb{E}_i(x) = \int_x^\infty \frac{e^{-t}}{t} dt$ is the exponential integral. Finally, using ?? in ?? it is found that:

$$\Gamma(\beta L) = \frac{Z_L - Z_S}{2L} \cdot e^{j2\beta Z_S} \cdot \left(\mathbb{E}_i \left(\frac{-j2\beta Z_L L}{Z_L - Z_S} \right) - \mathbb{E}_i \left(\frac{-j2\beta Z_S L}{Z_L - Z_S} \right) \right) \quad (9.262)$$

The following figure shows a comparison of the normalized reflection coefficient between the tapers described above

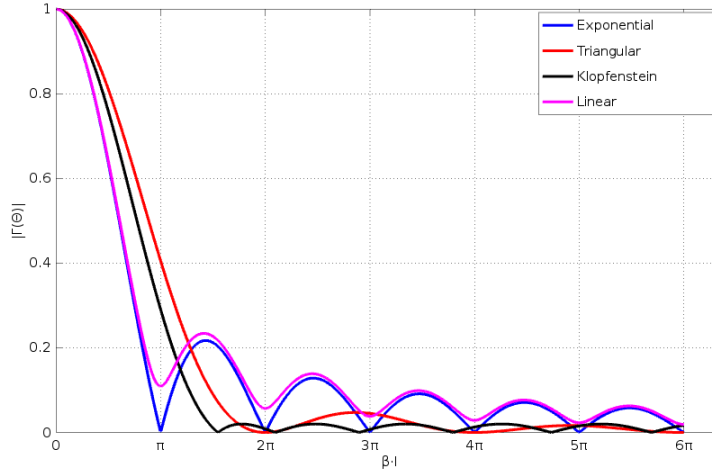


Figure 9.15: Comparison of the reflection coefficient between different tapers

In order to calculate the quadripole parameters, the line is discretized into a finite number of slots where the characteristic impedance is supposed to be uniform. The ABCD parameters of each slot is calculated according to ?? and, finally, the ABCD parameters of the tapered line are obtained as the cascade of these slots:

Let be

$$Z_k = Z(l)|_{l=\frac{k \cdot L}{N}}, \quad \forall k \in [0, N-1] \quad (9.263)$$

the characteristic impedance of the k^{th} slot, where L is the length of the tapered line and N is the number of slots. The ABCD parameters of the k^{th} slot are:

$$ABCD_k = \begin{pmatrix} \cosh(\gamma \cdot dl) & Z_k \cdot \sinh(\gamma \cdot dl) \\ \sinh(\gamma \cdot dl)/Z_k & \cosh(\gamma \cdot dl) \end{pmatrix} \quad (9.264)$$

where dl is the length of the slots and γ is the propagation constant. Then, the ABCD parameters of the whole tapered line are:

$$ABCD = \prod_{k=0}^{N-1} ABCD_k \quad (9.265)$$

Reference: [?], pages 261-271

9.26 S-parameter container with additional reference port

The Y-parameters of a multi-port component defined by its S-parameters required for a small signal AC analysis can be obtained by converting the S-parameters into Y-parameters.

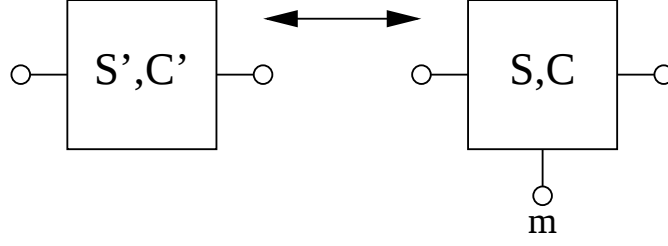


Figure 9.16: S-parameter container with noise wave correlation matrix

In order to extend a $m - 1$ -port to have a S-parameter device with m ports assuming that the original reference port had a reflection coefficient Γ_m the new S-parameters are according to T. O. Grosch and L. A. Carpenter [?]:

$$S_{mm} = \frac{2 - \Gamma_m - m + \sum_{i=1}^{m-1} \sum_{j=1}^{m-1} S'_{ij}}{1 - m \cdot \Gamma_m - \sum_{i=1}^{m-1} \sum_{j=1}^{m-1} S'_{ij}} \quad (9.266)$$

$$S_{im} = \left(\frac{1 - \Gamma_m \cdot S_{mm}}{1 - \Gamma_m} \right) \cdot \left(1 - \sum_{j=1}^{m-1} S'_{ij} \right) \quad \text{for } i = 1, 2 \dots m - 1 \quad (9.267)$$

$$S_{mj} = \left(\frac{1 - \Gamma_m \cdot S_{mm}}{1 - \Gamma_m} \right) \cdot \left(1 - \sum_{i=1}^{m-1} S'_{ij} \right) \quad \text{for } j = 1, 2 \dots m - 1 \quad (9.268)$$

$$S_{ij} = S'_{ij} - \left(\frac{\Gamma_m \cdot S_{im} \cdot S_{mj}}{1 - \Gamma_m \cdot S_{mm}} \right) \quad \text{for } i, j = 1, 2 \dots m - 1 \quad (9.269)$$

If the reference port has been ground potential, then Γ_m simply folds to -1. The reverse transformation by connecting a termination with a reflection coefficient of Γ_m to the m th port writes as follows.

$$S'_{ij} = S_{ij} + \left(\frac{\Gamma_m \cdot S_{im} \cdot S_{mj}}{1 - \Gamma_m \cdot S_{mm}} \right) \quad \text{for } i, j = 1, 2 \dots m - 1 \quad (9.270)$$

With the S-parameter transformation done the m -port noise wave correlation matrix is

$$C_m = \left| \frac{1}{1 - \Gamma_m} \right|^2 \cdot \left(K \cdot C_{m-1} \cdot K^+ - T_s \cdot k_B \cdot \left| 1 - |\Gamma_m|^2 \right| \cdot D \cdot D^+ \right) \quad (9.271)$$

with

$$K = \begin{bmatrix} 1 + \Gamma_m (S_{1m} - 1) & \Gamma_m S_{1m} & \cdots & \Gamma_m S_{1m} \\ \Gamma_m S_{2m} & 1 + \Gamma_m (S_{2m} - 1) & \cdots & \Gamma_m S_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ \Gamma_m S_{(m-1)m} & \Gamma_m S_{(m-1)m} & \cdots & 1 + \Gamma_m (S_{(m-1)m} - 1) \\ \Gamma_m S_{mm} - 1 & \Gamma_m S_{mm} - 1 & \cdots & \Gamma_m S_{mm} - 1 \end{bmatrix} \quad (9.272)$$

$$D = \begin{bmatrix} S_{1m} \\ S_{2m} \\ \vdots \\ S_{(m-1)m} \\ S_{mm} - 1 \end{bmatrix} \quad (9.273)$$

whence T_s denotes the equivalent noise temperature of the original reference port and the $^+$ operator indicates the transposed conjugate matrix (also called adjoint or adjugate).

The reverse transformation can be written as

$$C_{m-1} = K' \cdot C_m \cdot K'^+ + T_s \cdot k_B \cdot \frac{|1 - |\Gamma_m|^2|}{|1 - \Gamma_m S_{mm}|^2} \cdot D' \cdot D'^+ \quad (9.274)$$

with

$$K' = \begin{bmatrix} 1 & 0 & \cdots & 0 & \frac{\Gamma_m S_{1m}}{1 - \Gamma_m S_{mm}} \\ 0 & 1 & \cdots & 0 & \frac{\Gamma_m S_{2m}}{1 - \Gamma_m S_{mm}} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \cdots & 1 & \frac{\Gamma_m S_{(m-1)m}}{1 - \Gamma_m S_{mm}} \end{bmatrix} \quad (9.275)$$

$$D' = \begin{bmatrix} S_{1m} \\ S_{2m} \\ \vdots \\ S_{(m-1)m} \end{bmatrix} \quad (9.276)$$

9.27 Real-Life Models

Non-ideal electronic components exhibit parasitic effects. Depending on the usage, they may show a very different behaviour than the ideal ones. More precise models can sometimes be obtained from their manufacturers or vendors. However, first order approximations exists that can give satisfactory result in many cases. A few of these simple models are presented in this chapter.

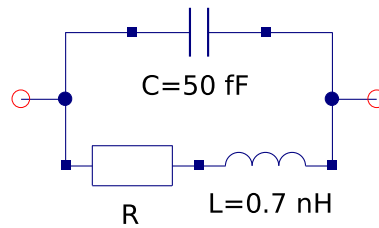


Figure 9.17: simple equivalent circuit of a 0603 resistor

A model for a resistor (case 0603) is depicted in figure ?? . Conclusion:

- useful up to 1GHz
- values around 150Ω are useful up to 10GHz

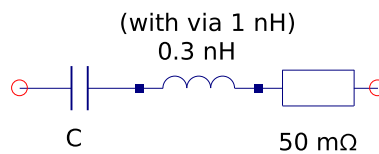


Figure 9.18: simple equivalent circuit of a 0603 ceramic capacitor

A model for a (ceramic) capacitor (case 0603) is depicted in figure ?? . Conclusion:

- as coupling capacitor useful wide into GHz band
- as blocking capacitor a via is necessary, i.e. 10nF has resonance at about 50MHz

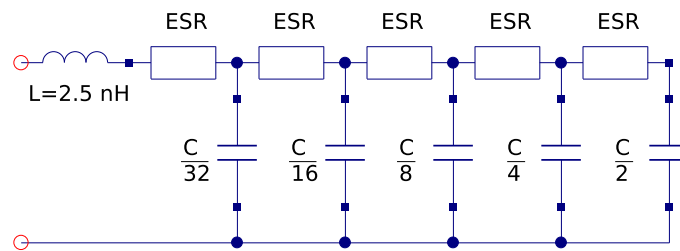


Figure 9.19: simple equivalent circuit of an electrolyte capacitor

Electrolyte capacitors are quite complicate to model. They also show the biggest differences from sample to sample. Nonetheless, figure ?? gives an idea how a model may look like. Conclusion:

- very broad resonance
- useful up to about 10MHz (strongly depending on capacitance)

Chapter 10

Non-linear devices

10.1 Operational amplifier

The ideal operational amplifier, as shown in fig. ??, is determined by the following equation which introduces one more unknown in the MNA matrix.

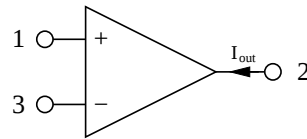


Figure 10.1: ideal operational amplifier

$$V_1 - V_3 = 0 \quad (10.1)$$

The new unknown variable I_{out} must be considered by the three remaining simple equations.

$$I_1 = 0 \quad I_2 = I_{out} \quad I_3 = 0 \quad (10.2)$$

And in matrix representation this is (for DC and AC simulation):

$$\begin{bmatrix} \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & 0 \\ 1 & 0 & -1 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \end{bmatrix} \quad (10.3)$$

The operational amplifier could be considered as a special case of a voltage controlled current source with infinite forward transconductance G . Please note that the presented matrix form is only valid in cases where there is a finite feedback impedance between the output and the inverting input port.

To allow a feedback circuit to the non-inverting input (e.g. for a Schmitt trigger), one needs a limited output voltage swing. The following equations are often used to model the transmission characteristic of operational amplifiers.

$$I_1 = 0 \quad I_3 = 0 \quad (10.4)$$

$$V_2 = V_{max} \cdot \frac{2}{\pi} \arctan \left(\frac{\pi}{2 \cdot V_{max}} \cdot G \cdot (V_1 - V_3) \right) \quad (10.5)$$

with V_{max} being the maximum output voltage swing and G the voltage amplification. To model the small-signal behaviour (AC analysis), it is necessary to differentiate:

$$g = \frac{\partial V_2}{\partial (V_1 - V_3)} = \frac{G}{1 + \left(\frac{\pi}{2 \cdot V_{max}} \cdot G \cdot (V_1 - V_3) \right)^2} \quad (10.6)$$

This leads to the following matrix representation being a specialised three node voltage controlled voltage source (see section ?? on page ??).

$$\begin{bmatrix} . & . & . & 0 \\ . & . & . & 1 \\ . & . & . & 0 \\ g & -1 & -g & 0 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \end{bmatrix} \quad (10.7)$$

The above MNA matrix entries are also used during the non-linear DC analysis with the 0 in the right hand side vector replaced by an equivalent voltage

$$V_{eq} = g \cdot (V_1 - V_3) - V_{out} \quad (10.8)$$

with V_{out} computed using eq. (??).

With the given small-signal matrix representation, building the S-parameters is easy.

$$(\underline{S}) = \begin{bmatrix} 1 & 0 & 0 \\ 4g & -1 & -4g \\ 0 & 0 & 1 \end{bmatrix} \quad (10.9)$$

10.2 PN-Junction Diode

The following table contains the model parameters for the pn-junction diode model.

		Name	Symbol	Description	Unit	Default	
Is	I_S	saturation current			A	10^{-14}	
N	N	emission coefficient				1.0	
Isr	I_{SR}	recombination current parameter			A	0.0	
Nr	N_R	emission coefficient for Isr				2.0	
Rs	R_S	ohmic resistance			Ω	0.0	
Cj0	C_{j0}	zero-bias junction capacitance			F	0.0	
M	M	grading coefficient				0.5	
Vj	V_j	junction potential			V	0.7	
Fc	F_c	forward-bias depletion capacitance coefficient				0.5	
Cp	C_p	linear capacitance			F	0.0	
Tt	τ	transit time			s	0.0	
Bv	B_v	reverse breakdown voltage			V	∞	
Ibv	I_{Bv}	current at reverse breakdown voltage			A	0.001	
Kf	K_F	flicker noise coefficient				0.0	

		Name	Symbol	Description	Unit	Default
Af	A_F	flicker noise exponent				1.0
Ffe	F_{FE}	flicker noise frequency exponent				1.0
Temp	T	device temperature			°C	26.85
Xti	X_{TI}	saturation current exponent				3.0
Eg	E_G	energy bandgap			eV	1.11
Tbv	T_{Bv}	Bv linear temperature coefficient			1/°C	0.0
Trs	T_{RS}	Rs linear temperature coefficient			1/°C	0.0
Ttt1	$T_{\tau 1}$	Tt linear temperature coefficient			1/°C	0.0
Ttt2	$T_{\tau 2}$	Tt quadratic temperature coefficient			1/°C ²	0.0
Tm1	T_{M1}	M linear temperature coefficient			1/°C	0.0
Tm2	T_{M2}	M quadratic temperature coefficient			1/°C ²	0.0
Tnom	T_{NOM}	temperature at which parameters were extracted			°C	26.85
Area	A	default area for diode				1.0

10.2.1 Large signal model

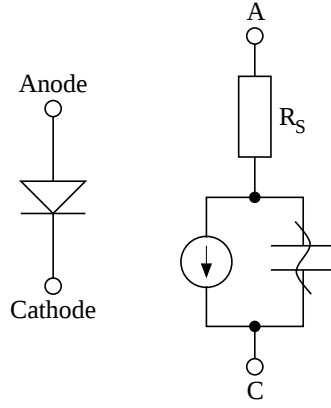


Figure 10.2: pn-junction diode symbol and large signal model

The current equation of the diode and its derivative writes as follows:

$$I_d = I_S \cdot \left(e^{\frac{V_d}{N \cdot V_T}} - 1 \right) + I_{SR} \cdot \left(e^{\frac{V_d}{N_R \cdot V_T}} - 1 \right) \quad (10.10)$$

$$g_d = \frac{\partial I_d}{\partial V_d} = \frac{I_S}{N \cdot V_T} \cdot e^{\frac{V_d}{N \cdot V_T}} + \frac{I_{SR}}{N_R \cdot V_T} \cdot e^{\frac{V_d}{N_R \cdot V_T}} \quad (10.11)$$

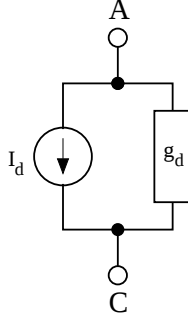


Figure 10.3: accompanied DC model of intrinsic diode

The complete MNA matrix entries are:

$$\begin{bmatrix} g_d & -g_d \\ -g_d & g_d \end{bmatrix} \cdot \begin{bmatrix} V_C \\ V_A \end{bmatrix} = \begin{bmatrix} +I_d - g_d \cdot V_d \\ -I_d + g_d \cdot V_d \end{bmatrix} \quad (10.12)$$

10.2.2 Small signal model

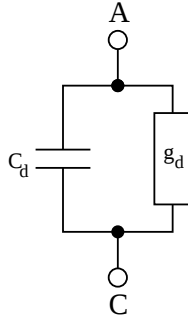


Figure 10.4: small signal model of intrinsic diode

The voltage dependent capacitance consists of a diffusion capacitance, a junction capacitance and an additional linear capacitance which is usually modeled by the following equations.

$$C_d = C_p + \tau \cdot g_d + \begin{cases} C_{j0} \cdot \left(1 - \frac{V_d}{V_j}\right)^{-M} & \text{for } V_d \leq F_c \cdot V_j \\ \frac{C_{j0}}{(1 - F_c)^M} \cdot \left(1 + \frac{M \cdot (V_d - F_c \cdot V_j)}{V_j \cdot (1 - F_c)}\right) & \text{for } V_d > F_c \cdot V_j \end{cases} \quad (10.13)$$

The S-parameters of the passive circuit shown in fig. ?? can be written as

$$S_{11} = S_{22} = \frac{1}{1 + 2 \cdot y} \quad (10.14)$$

$$S_{12} = S_{21} = 1 - S_{11} = \frac{2 \cdot y}{1 + 2 \cdot y} \quad (10.15)$$

with

$$y = Z_0 \cdot (g_d + j\omega C_d) \quad (10.16)$$

10.2.3 Noise model

The thermal noise generated by the series resistor is characterized by the following spectral density.

$$\frac{\overline{i_{RS}^2}}{\Delta f} = \frac{4k_B T}{R_S} \quad (10.17)$$

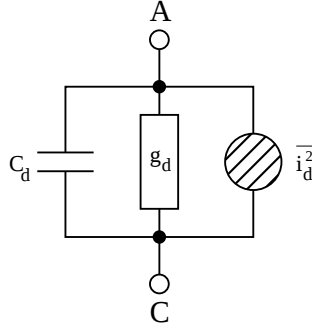


Figure 10.5: noise model of intrinsic diode

The shot noise and flicker noise generated by the DC current flow through the diode is characterized by the following spectral density.

$$\frac{\overline{i_d^2}}{\Delta f} = 2eI_d + K_F \frac{I_d^{A_F}}{f^{F_{FE}}} \quad (10.18)$$

Thus the noise current correlation matrix can be formed. This matrix can be easily converted to the noise wave correlation matrix representation using the formulas given in section ?? on page ??.

$$\underline{C}_Y = \Delta f \begin{bmatrix} +\overline{i_d^2} & -\overline{i_d^2} \\ -\overline{i_d^2} & +\overline{i_d^2} \end{bmatrix} \quad (10.19)$$

An ideal diode (pn- or schottky-diode) generates shot noise. Both types of current (field and diffusion) contribute independently to it. That is, even though the two currents flow in different directions ("minus" in dc current equation), they have to be added in the noise equation (current is proportional to noise power spectral density). Taking into account the dynamic conductance g_d in parallel to the noise current source, the noise wave correlation matrix writes as follows.

$$\begin{aligned} (\underline{C}) &= \left| \frac{0.5 \cdot Y_0}{g_d + j\omega C_d + 0.5 \cdot Y_0} \right|^2 \cdot 2 \cdot e \cdot I_S \cdot \left(\exp \left(\frac{V_d}{N \cdot V_T} \right) + 1 \right) \cdot Z_0 \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \\ &= 2 \cdot e \cdot Z_0 \cdot (I_d + 2 \cdot I_S) \cdot \left| \frac{1}{2 \cdot Z_0 \cdot (g_d + j\omega C_d) + 1} \right|^2 \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \end{aligned} \quad (10.20)$$

Where e is charge of an electron, V_T the temperature voltage, g_d the (dynamic) conductance of the diode and C_d its junction capacitance.

To be very precise, the equation above only holds for diodes whose field and diffusion current dominate absolutely (diffusion limited diode), i.e. $N = 1$. Many diodes also generate a generation/recombination current ($N \approx 2$), which produces shot noise, too. But depending on where and how the charge carriers generate or recombine, their effective charge is somewhat smaller than e . To take this into account, one needs a further factor K . Several opinions exist according to the value of K . Some say 1 and $2/3$ are common values, others say $K = 1/N$ with K and N being bias dependent. Altogether it is:

$$(\underline{C}) = 2 \cdot e \cdot Z_0 \cdot K \cdot (I_d + 2 \cdot I_S) \cdot \left| \frac{1}{2 \cdot Z_0 \cdot (g_d + j\omega C_d) + 1} \right|^2 \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (10.21)$$

with $\frac{1}{2} \leq K \leq 1$

Remark: Believing the diode equation $I_D = I_S \cdot (\exp(V/(N \cdot V_T)) - 1)$ is the whole truth, it is logical to define $K = 1/N$, because at $V = 0$ the conductance g_d of the diode must create thermal noise.

Some special diodes have additional current or noise components (tunnel diodes, avalanche diodes etc.). All these mechanisms are not taken into account in equation (??).

The parasitic ohmic resistance in a non-ideal diode, of course, creates thermal noise. Noise current correlation matrix (for details on the parameters see above):

$$(\underline{C}_Y) = 2 \cdot e \cdot K \cdot (I_d + 2 \cdot I_S) \cdot \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad (10.22)$$

10.2.4 Temperature model

This section mathematically describes the dependencies of the diode characteristics on temperature. For a junction diode a typical value for X_{TI} is 3.0, for a Schottky barrier diode it is 2.0. The energy band gap at zero temperature E_G is by default 1.11eV. For other materials than Si, 0.69eV (for a Schottky barrier diode), 0.67eV (for Ge) and 1.43eV (for GaAs) should be used.

$$n_i^2(T) = B \cdot T^3 \cdot e^{-E_G(T)/k_B T} \quad (10.23)$$

$$n_i(T) = 1.45 \cdot 10^{10} \cdot \left(\frac{T}{300K} \right)^{1.5} \cdot \exp \left(\frac{e \cdot E_G(300K)}{2 \cdot k_B \cdot 300K} - \frac{e \cdot E_G(T)}{2 \cdot k_B \cdot T} \right) \quad (10.24)$$

$$E_G(T) = E_G - \frac{\alpha \cdot T^2}{\beta + T} \quad (10.25)$$

with experimental values for Si given by

$$\begin{aligned} \alpha &= 7.02 \cdot 10^{-4} \\ \beta &= 1108 \\ E_G &= 1.16eV \end{aligned}$$

The following equations show the temperature dependencies of the diode parameters. The reference temperature T_1 in these equations denotes the nominal temperature T_{NOM} specified by

the diode model.

$$I_S(T_2) = I_S(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TI}/N} \cdot \exp\left[-\frac{e \cdot E_G(300K)}{N \cdot k_B \cdot T_2} \cdot \left(1 - \frac{T_2}{T_1}\right)\right] \quad (10.26)$$

$$V_j(T_2) = \frac{T_2}{T_1} \cdot V_j(T_1) + \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{n_i(T_1)}{n_i(T_2)}\right) \quad (10.27)$$

$$= \frac{T_2}{T_1} \cdot V_j(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.28)$$

$$C_{j0}(T_2) = C_{j0}(T_1) \cdot \left(1 + M \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{V_j(T_2) - V_j(T_1)}{V_j(T_1)}\right)\right) \quad (10.29)$$

Some additional temperature coefficients determine the temperature dependence of even more model parameters.

$$B_v(T_2) = B_v(T_1) - T_{Bv} \cdot (T_2 - T_1) \quad (10.30)$$

$$\tau(T_2) = \tau(T_1) \cdot \left(1 + T_{\tau 1} \cdot (T_2 - T_1) + T_{\tau 2} \cdot (T_2 - T_1)^2\right) \quad (10.31)$$

$$M(T_2) = M(T_1) \cdot \left(1 + T_{M1} \cdot (T_2 - T_1) + T_{M2} \cdot (T_2 - T_1)^2\right) \quad (10.32)$$

$$R_S(T_2) = R_S(T_1) \cdot (1 + T_{RS} \cdot (T_2 - T_1)) \quad (10.33)$$

10.2.5 Area dependence of the model

The area factor A used in the diode model determines the number of equivalent parallel devices of the specified model. The diode model parameters affected by the A factor are:

$$I_S(A) = I_S \cdot A \quad (10.34)$$

$$C_{j0}(A) = C_{j0} \cdot A \quad (10.35)$$

$$R_S(A) = \frac{R_S}{A} \quad (10.36)$$

10.3 Junction FET

The following table contains the model parameters for the JFET model.

		Name	Symbol	Description	Unit	Default
Vt0	V_{Th}	zero -bias threshold voltage			V	-2.0
Beta	β	transconductance parameter			A/V ²	10 ⁻⁴
Lambda	λ	channel-length modulation parameter			1/V	0.0
Rd	R_D	drain ohmic resistance			Ω	0.0
Rs	R_S	source ohmic resistance			Ω	0.0
Is	I_S	gate-junction saturation current			A	10 ⁻¹⁴
N	N	gate P-N emission coefficient				1.0
Isr	I_{SR}	gate-junction recombination current parameter			A	0.0
Nr	N_R	Isr emission coefficient				2.0
Cgs	C_{gs}	zero-bias gate-source junction capacitance			F	0.0
Cgd	C_{gd}	zero-bias gate-drain junction capacitance			F	0.0
Pb	P_b	gate-junction potential			V	1.0

		Name	Symbol	Description	Unit	Default
Fc	F_c	forward-bias junction capacitance coefficient				0.5
M	M	gate P-N grading coefficient				0.5
Kf	K_F	flicker noise coefficient				0.0
Af	A_F	flicker noise exponent				1.0
Ffe	F_{FE}	flicker noise frequency exponent				1.0
Temp	T	device temperature			°C	26.85
Xti	X_{TI}	saturation current exponent				3.0
Vt0tc	V_{ThTC}	Vt0 temperature coefficient			V/°C	0.0
Betatce	β_{TCE}	Beta exponential temperature coefficient			%/°C	0.0
Tnom	T_{NOM}	temperature at which parameters were extracted			°C	26.85
Area	A	default area for JFET				1.0

10.3.1 Large signal model

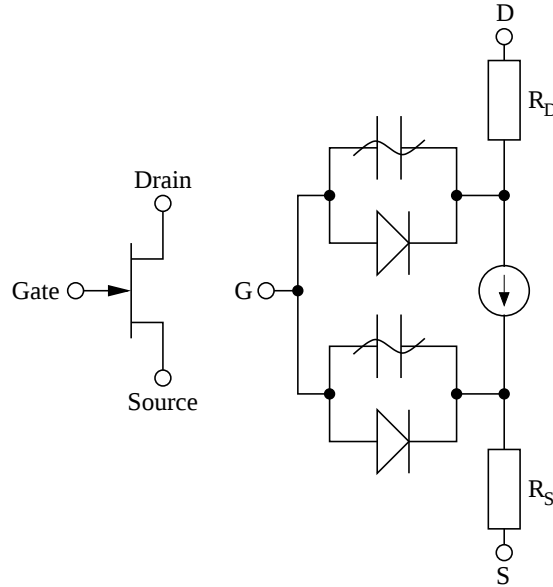


Figure 10.6: junction FET symbol and large signal model

The current equation of the gate source diode and its derivative writes as follows:

$$I_{GS} = I_S \cdot \left(e^{\frac{V_{GS}}{N \cdot V_T}} - 1 \right) + I_{SR} \cdot \left(e^{\frac{V_{GS}}{N_R \cdot V_T}} - 1 \right) \quad (10.37)$$

$$g_{gs} = \frac{\partial I_{GS}}{\partial V_{GS}} = \frac{I_S}{N \cdot V_T} \cdot e^{\frac{V_{GS}}{N \cdot V_T}} + \frac{I_{SR}}{N_R \cdot V_T} \cdot e^{\frac{V_{GS}}{N_R \cdot V_T}} \quad (10.38)$$

The current equation of the gate drain diode and its derivative writes as follows:

$$I_{GD} = I_S \cdot \left(e^{\frac{V_{GD}}{N \cdot V_T}} - 1 \right) + I_{SR} \cdot \left(e^{\frac{V_{GD}}{N_R \cdot V_T}} - 1 \right) \quad (10.39)$$

$$g_{gd} = \frac{\partial I_{GD}}{\partial V_{GD}} = \frac{I_S}{N \cdot V_T} \cdot e^{\frac{V_{GD}}{N \cdot V_T}} + \frac{I_{SR}}{N_R \cdot V_T} \cdot e^{\frac{V_{GD}}{N_R \cdot V_T}} \quad (10.40)$$

Both equations contain the gate-junction saturation current I_S , the gate P-N emission coefficient N and the temperature voltage V_T with the Boltzmann's constant k_B and the electron charge q . The operating temperature T must be specified in Kelvin.

$$V_T = \frac{k_B \cdot T}{q} \quad (10.41)$$

The controlled drain currents have been defined by Shichman and Hodges [?] for different modes of operations.

$$g_m = \frac{\partial I_d}{\partial V_{GS}} \quad \text{and} \quad g_{ds} = \frac{\partial I_d}{\partial V_{DS}} \quad \text{with} \quad V_{GD} = V_{GS} - V_{DS} \quad (10.42)$$

- normal mode: $V_{DS} > 0$

- normal mode, cutoff region: $V_{GS} - V_{Th} < 0$

$$I_d = 0 \quad (10.43)$$

$$g_m = 0 \quad (10.44)$$

$$g_{ds} = 0 \quad (10.45)$$

- normal mode, saturation region: $0 < V_{GS} - V_{Th} < V_{DS}$

$$I_d = \beta \cdot (1 + \lambda V_{DS}) \cdot (V_{GS} - V_{Th})^2 \quad (10.46)$$

$$g_m = \beta \cdot (1 + \lambda V_{DS}) \cdot 2(V_{GS} - V_{Th}) \quad (10.47)$$

$$g_{ds} = \beta \cdot \lambda (V_{GS} - V_{Th})^2 \quad (10.48)$$

- normal mode, linear region: $V_{DS} < V_{GS} - V_{Th}$

$$I_d = \beta \cdot (1 + \lambda V_{DS}) \cdot (2(V_{GS} - V_{Th}) - V_{DS}) \cdot V_{DS} \quad (10.49)$$

$$g_m = \beta \cdot (1 + \lambda V_{DS}) \cdot 2 \cdot V_{DS} \quad (10.50)$$

$$g_{ds} = \beta \cdot (1 + \lambda V_{DS}) \cdot 2(V_{GS} - V_{Th} - V_{DS}) + \beta \cdot \lambda V_{DS} \cdot (2(V_{GS} - V_{Th}) - V_{DS}) \quad (10.51)$$

- inverse mode: $V_{DS} < 0$

- inverse mode, cutoff region: $V_{GD} - V_{Th} < 0$

$$I_d = 0 \quad (10.52)$$

$$g_m = 0 \quad (10.53)$$

$$g_{ds} = 0 \quad (10.54)$$

– inverse mode, saturation region: $0 < V_{GD} - V_{Th} < -V_{DS}$

$$I_d = -\beta \cdot (1 - \lambda V_{DS}) \cdot (V_{GD} - V_{Th})^2 \quad (10.55)$$

$$g_m = -\beta \cdot (1 - \lambda V_{DS}) \cdot 2(V_{GD} - V_{Th}) \quad (10.56)$$

$$g_{ds} = \beta \cdot \lambda (V_{GD} - V_{Th})^2 + \beta \cdot (1 - \lambda V_{DS}) \cdot 2(V_{GD} - V_{Th}) \quad (10.57)$$

– inverse mode, linear region: $-V_{DS} < V_{GD} - V_{Th}$

$$I_d = \beta \cdot (1 - \lambda V_{DS}) \cdot (2(V_{GD} - V_{Th}) + V_{DS}) \cdot V_{DS} \quad (10.58)$$

$$g_m = \beta \cdot (1 - \lambda V_{DS}) \cdot 2 \cdot V_{DS} \quad (10.59)$$

$$g_{ds} = \beta \cdot (1 - \lambda V_{DS}) \cdot 2(V_{GD} - V_{Th}) - \beta \cdot \lambda V_{DS} \cdot (2(V_{GD} - V_{Th}) + V_{DS}) \quad (10.60)$$

The MNA matrix entries for the voltage controlled drain current source can be written as:

	G	S	controlling nodes
D	$+g_m$	$-g_m$	
S	$-g_m$	$+g_m$	
controlled nodes			

With the accompanied DC model shown in fig. ?? using the same principles as explained in section ?? on page ?? it is possible to build the complete MNA matrix of the intrinsic JFET.

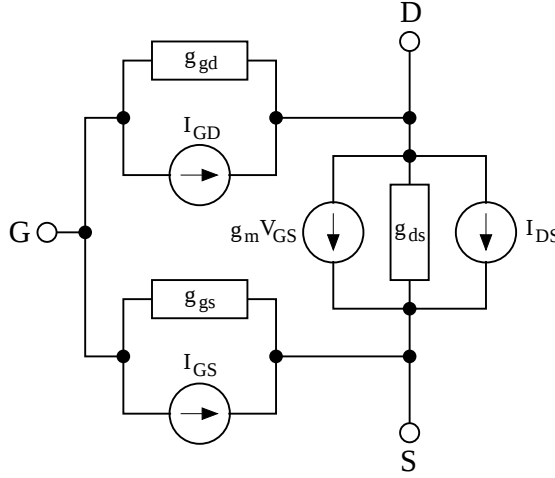


Figure 10.7: accompanied DC model of intrinsic JFET

Applying the rules for creating the MNA matrix of an arbitrary network the complete MNA matrix entries (admittance matrix and current vector) for the intrinsic junction FET are:

$$\begin{bmatrix} g_{gd} + g_{gs} & -g_{gd} & -g_{gs} \\ -g_{gd} + g_m & g_{ds} + g_{gd} & -g_{ds} - g_m \\ -g_{gs} - g_m & -g_{ds} & g_{gs} + g_{ds} + g_m \end{bmatrix} \cdot \begin{bmatrix} V_G \\ V_D \\ V_S \end{bmatrix} = \begin{bmatrix} -I_{GD_{eq}} - I_{GS_{eq}} \\ +I_{GD_{eq}} - I_{DS_{eq}} \\ +I_{GS_{eq}} + I_{DS_{eq}} \end{bmatrix} \quad (10.61)$$

with

$$I_{GS_{eq}} = I_{GS} - g_{gs} \cdot V_{GS} \quad (10.62)$$

$$I_{GD_{eq}} = I_{GD} - g_{gd} \cdot V_{GD} \quad (10.63)$$

$$I_{DS_{eq}} = I_d - g_m \cdot V_{GS} - g_{ds} \cdot V_{DS} \quad (10.64)$$

10.3.2 Small signal model

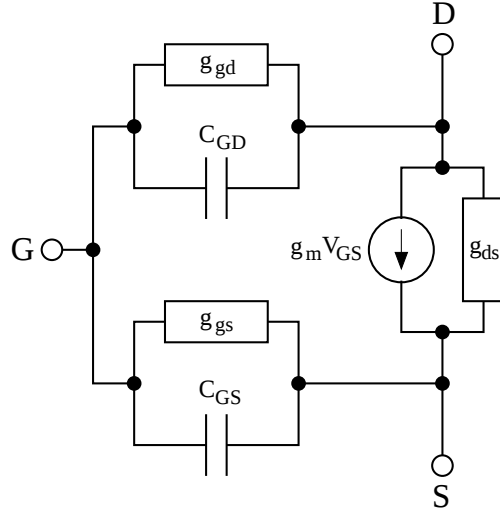


Figure 10.8: small signal model of intrinsic junction FET

The small signal Y-parameter matrix of the intrinsic junction FET writes as follows. It can be converted to S-parameters.

$$Y = \begin{bmatrix} Y_{GD} + Y_{GS} & -Y_{GD} & -Y_{GS} \\ g_m - Y_{GD} & Y_{GD} + Y_{DS} & -Y_{DS} - g_m \\ -g_m - Y_{GS} & -Y_{DS} & Y_{GS} + Y_{DS} + g_m \end{bmatrix} \quad (10.65)$$

with

$$Y_{GD} = g_{gd} + j\omega C_{GD} \quad (10.66)$$

$$Y_{GS} = g_{gs} + j\omega C_{GS} \quad (10.67)$$

$$Y_{DS} = g_{ds} \quad (10.68)$$

The junction capacitances are modeled with the following equations.

$$C_{GD} = \begin{cases} C_{gd} \cdot \left(1 - \frac{V_{GD}}{P_b}\right)^{-M} & \text{for } V_{GD} \leq F_c \cdot P_b \\ \frac{C_{gd}}{(1 - F_c)^M} \cdot \left(1 + \frac{M \cdot (V_{GD} - F_c \cdot P_b)}{P_b \cdot (1 - F_c)}\right) & \text{for } V_{GD} > F_c \cdot P_b \end{cases} \quad (10.69)$$

$$C_{GS} = \begin{cases} C_{gs} \cdot \left(1 - \frac{V_{GS}}{P_b}\right)^{-M} & \text{for } V_{GS} \leq F_c \cdot P_b \\ \frac{C_{gs}}{(1 - F_c)^M} \cdot \left(1 + \frac{M \cdot (V_{GS} - F_c \cdot P_b)}{P_b \cdot (1 - F_c)}\right) & \text{for } V_{GS} > F_c \cdot P_b \end{cases} \quad (10.70)$$

10.3.3 Noise model

Both the drain and source resistance R_D and R_S generate thermal noise characterized by the following spectral density.

$$\frac{\overline{i_{R_D}^2}}{\Delta f} = \frac{4k_B T}{R_D} \quad \text{and} \quad \frac{\overline{i_{R_S}^2}}{\Delta f} = \frac{4k_B T}{R_S} \quad (10.71)$$

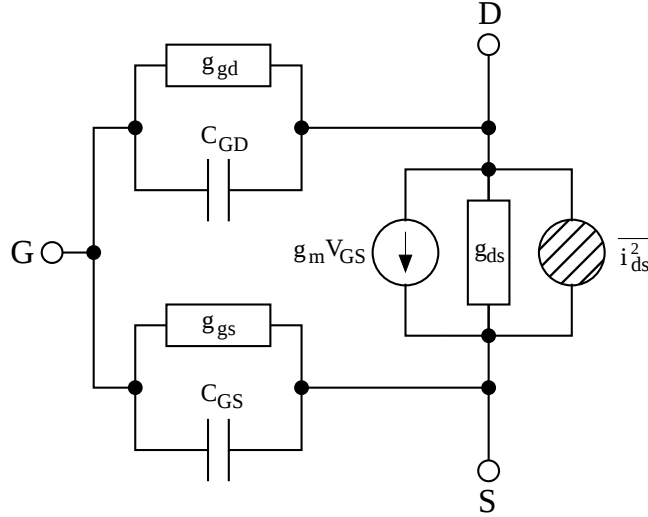


Figure 10.9: noise model of intrinsic junction FET

Channel noise and flicker noise generated by the DC transconductance g_m and current flow from drain to source is characterized by the following spectral density.

$$\frac{\overline{i_{ds}^2}}{\Delta f} = \frac{8k_B T g_m}{3} + K_F \frac{I_{DS}^{A_F}}{f^{F_{FE}}} \quad (10.72)$$

The noise current correlation matrix (admittance representation) of the intrinsic junction FET can be expressed by

$$\underline{C}_Y = \Delta f \begin{bmatrix} 0 & 0 & 0 \\ 0 & +\overline{i_{ds}^2} & -\overline{i_{ds}^2} \\ 0 & -\overline{i_{ds}^2} & +\overline{i_{ds}^2} \end{bmatrix} \quad (10.73)$$

This matrix representation can be easily converted to the noise-wave representation $\underline{C}_S$ if the small signal S-parameter matrix is known.

10.3.4 Temperature model

Temperature appears explicitly in the exponential terms of the JFET model equations. In addition, saturation current, gate-junction potential and zero-bias junction capacitances have built-in temperature dependence.

$$I_S(T_2) = I_S(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TI}/N} \cdot \exp\left[-\frac{e \cdot E_G(300K)}{N \cdot k_B \cdot T_2} \cdot \left(1 - \frac{T_2}{T_1}\right)\right] \quad (10.74)$$

$$I_{SR}(T_2) = I_{SR}(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TI}/N_R} \cdot \exp\left[-\frac{e \cdot E_G(300K)}{N_R \cdot k_B \cdot T_2} \cdot \left(1 - \frac{T_2}{T_1}\right)\right] \quad (10.75)$$

$$P_b(T_2) = \frac{T_2}{T_1} \cdot P_b(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.76)$$

$$C_{gs}(T_2) = C_{gs}(T_1) \cdot \left(1 + M \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{P_b(T_2) - P_b(T_1)}{P_b(T_1)}\right)\right) \quad (10.77)$$

$$C_{gd}(T_2) = C_{gd}(T_1) \cdot \left(1 + M \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{P_b(T_2) - P_b(T_1)}{P_b(T_1)}\right)\right) \quad (10.78)$$

where the $E_G(T)$ dependency has already been described in section ?? on page ?. Also the threshold voltage as well as the transconductance parameter have a temperature dependence determined by

$$V_{Th}(T_2) = V_{Th}(T_1) + V_{ThTC} \cdot (T_2 - T_1) \quad (10.79)$$

$$\beta(T_2) = \beta(T_1) \cdot 1.01^{\beta_{TCE} \cdot (T_2 - T_1)} \quad (10.80)$$

10.3.5 Area dependence of the model

The area factor A used for the JFET model determines the number of equivalent parallel devices of a specified model. The following parameters are affected by the area factor.

$$\beta(A) = \beta \cdot A \quad I_S(A) = I_S \cdot A \quad (10.81)$$

$$R_D(A) = \frac{R_D}{A} \quad R_S(A) = \frac{R_S}{A} \quad (10.82)$$

$$C_{gs}(A) = C_{gs} \cdot A \quad C_{gd}(A) = C_{gd} \cdot A \quad (10.83)$$

10.4 Homo-Junction Bipolar Transistor

The following table contains the model parameters for the BJT (Spice Gummel-Poon) model.

	Name	Symbol	Description	Unit	Default
Is	I_S		saturation current	A	10^{-16}
Nf	N_F		forward emission coefficient		1.0
Nr	N_R		reverse emission coefficient		1.0
Ikf	I_{KF}		high current corner for forward beta	A	∞

		Name	Symbol	Description	Unit	Default
Ikr	I_{KR}	high current corner for reverse beta			A	∞
Vaf	V_{AF}	forward early voltage			V	∞
Var	V_{AR}	reverse early voltage			V	∞
Ise	I_{SE}	base-emitter leakage saturation current			A	0
Ne	N_E	base-emitter leakage emission coefficient				1.5
Isc	I_{SC}	base-collector leakage saturation current			A	0
Nc	N_C	base-collector leakage emission coefficient				2.0
Bf	B_F	forward beta				100
Br	B_R	reverse beta				1
Rbm	R_{Bm}	minimum base resistance for high currents			Ω	0.0
Irb	I_{RB}	current for base resistance midpoint			A	∞
Rc	R_C	collector ohmic resistance			Ω	0.0
Re	R_E	emitter ohmic resistance			Ω	0.0
Rb	R_B	zero-bias base resistance (may be high-current dependent)			Ω	0.0
Cje	C_{JE}	base-emitter zero-bias depletion capacitance			F	0.0
Vje	V_{JE}	base-emitter junction built-in potential			V	0.75
Mje	M_{JE}	base-emitter junction exponential factor				0.33
Cjc	C_{JC}	base-collector zero-bias depletion capacitance			F	0.0
Vjc	V_{JC}	base-collector junction built-in potential			V	0.75
Mjc	M_{JC}	base-collector junction exponential factor				0.33
Xcjc	X_{CJC}	fraction of Cjc that goes to internal base pin				1.0
Cjs	C_{JS}	zero-bias collector-substrate capacitance			F	0.0
Vjs	V_{JS}	substrate junction built-in potential			V	0.75
Mjs	M_{JS}	substrate junction exponential factor				0.0
Fc	F_C	forward-bias depletion capacitance coefficient				0.5
Tf	T_F	ideal forward transit time			s	0.0
Xtf	X_{TF}	coefficient of bias-dependence for Tf				0.0
Vtf	V_{TF}	voltage dependence of Tf on base-collector voltage			V	∞
Itf	I_{TF}	high-current effect on Tf			A	0.0
Ptf	φ_{TF}	excess phase at the frequency $1/(2\pi T_F)$			$^\circ$	0.0
Tr	T_R	ideal reverse transit time			s	0.0
Kf	K_F	flicker noise coefficient				0.0
Af	A_F	flicker noise exponent				1.0
Ffe	F_{FE}	flicker noise frequency exponent				1.0
Kb	K_B	burst noise coefficient				0.0
Ab	A_B	burst noise exponent				1.0
Fb	F_B	burst noise corner frequency			Hz	1.0
Temp	T	device temperature			$^\circ\text{C}$	26.85
Xti	X_{TI}	saturation current exponent				3.0
Xtb	X_{TB}	temperature exponent for forward- and reverse-beta				0.0
Eg	E_G	energy bandgap			eV	1.11
Tnom	T_{NOM}	temperature at which parameters were extracted			$^\circ\text{C}$	26.85
Area	A	default area for bipolar transistor				1.0

10.4.1 Large signal model

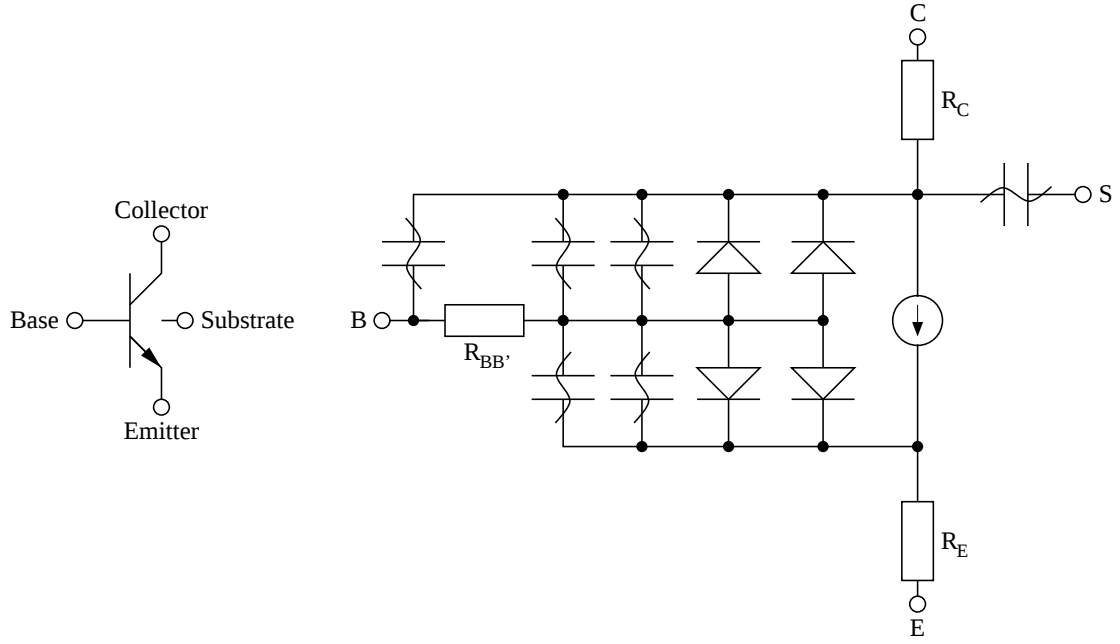


Figure 10.10: bipolar transistor symbol and large signal model for vertical device

The SGP (**SPICE Gummel-Poon**) model is basically a transport model, i.e. the voltage dependent ideal transfer currents (forward I_F and backward I_R) are reference currents in the model. The ideal base current parts are defined dependent on the ideal transfer currents. The ideal forward transfer current starts flowing when applying a positive control voltage at the base-emitter junction. It is defined by:

$$I_F = I_S \cdot \left(e^{\frac{V_{BE}}{N_F \cdot V_T}} - 1 \right) \quad (10.84)$$

The ideal base current components are defined by the ideal transfer currents. The non-ideal components are independently defined by dedicated saturation currents and emission coefficients.

$$I_{BEI} = \frac{I_F}{B_F} \quad g_{BEI} = \frac{\partial I_{BEI}}{\partial V_{BE}} = \frac{I_S}{N_F \cdot V_T \cdot B_F} \cdot e^{\frac{V_{BE}}{N_F \cdot V_T}} \quad (10.85)$$

$$I_{BEN} = I_{SE} \cdot \left(e^{\frac{V_{BE}}{N_E \cdot V_T}} - 1 \right) \quad g_{BEN} = \frac{\partial I_{BEN}}{\partial V_{BE}} = \frac{I_{SE}}{N_E \cdot V_T} \cdot e^{\frac{V_{BE}}{N_E \cdot V_T}} \quad (10.86)$$

$$I_{BE} = I_{BEI} + I_{BEN} \quad (10.87)$$

$$g_{\pi} = g_{BE} = g_{BEI} + g_{BEN} \quad (10.88)$$

The ideal backward transfer current arises when applying a positive control voltage at the base-collector junction (e.g. in the active inverse mode). It is defined by:

$$I_R = I_S \cdot \left(e^{\frac{V_{BC}}{N_R \cdot V_T}} - 1 \right) \quad (10.89)$$

Again, the ideal base current component through the base-collector junction is defined in reference to the ideal backward transfer current and the non-ideal component is defined by a dedicated saturation current and emission coefficient.

$$I_{BCI} = \frac{I_R}{B_R} \quad g_{BCI} = \frac{\partial I_{BCI}}{\partial V_{BC}} = \frac{I_S}{N_R \cdot V_T \cdot B_R} \cdot e^{\frac{V_{BC}}{N_R \cdot V_T}} \quad (10.90)$$

$$I_{BCN} = I_{SC} \cdot \left(e^{\frac{V_{BC}}{N_C \cdot V_T}} - 1 \right) \quad g_{BCN} = \frac{\partial I_{BCN}}{\partial V_{BC}} = \frac{I_{SC}}{N_C \cdot V_T} \cdot e^{\frac{V_{BC}}{N_C \cdot V_T}} \quad (10.91)$$

$$I_{BC} = I_{BCI} + I_{BCN} \quad (10.92)$$

$$g_{\mu} = g_{BC} = g_{BCI} + g_{BCN} \quad (10.93)$$

With these definitions it is possible to calculate the overall base current flowing into the device using all the base current components.

$$I_B = I_{BE} + I_{BC} = I_{BEI} + I_{BEN} + I_{BCI} + I_{BCN} \quad (10.94)$$

The overall transfer current I_T can be calculated using the normalized base charge Q_B and the ideal forward and backward transfer currents.

$$I_T = I_{TF} - I_{TR} = \frac{I_F - I_R}{Q_B} \quad (10.95)$$

The normalized base charge Q_B has no dimension and has the value 1 for $V_{BE} = V_{BC} = 0$. It is used to model two effects: the influence of the base width modulation on the transfer current (Early effect) and the ideal transfer currents deviation at high currents, i.e. the decreasing current gain at high currents.

$$Q_B = \frac{Q_1}{2} \cdot \left(1 + \sqrt{1 + 4 \cdot Q_2} \right) \quad (10.96)$$

The Q_1 term is used to describe the Early effect and Q_2 is responsible for the high current effects.

$$Q_1 = \frac{1}{1 - \frac{V_{BC}}{V_{AF}} - \frac{V_{BE}}{V_{AR}}} \quad \text{and} \quad Q_2 = \frac{I_F}{I_{KF}} + \frac{I_R}{I_{KR}} \quad (10.97)$$

The transfer current I_T depends on V_{BE} and V_{BC} by the normalized base charge Q_B and the forward transfer current I_F and the backward transfer current I_R . That is why both of the partial derivatives are required.

The forward transconductance g_{mf} of the transfer current I_T is obtained by differentiating it with respect to V_{BE} . The reverse transconductance g_{mr} can be calculated by differentiating the transfer current with respect to V_{BC} .

$$g_{mf} = \frac{\partial I_T}{\partial V_{BE}} = \frac{\partial I_{TF}}{\partial V_{BE}} - \frac{\partial I_{TR}}{\partial V_{BE}} = \frac{1}{Q_B} \cdot \left(+g_{IF} - I_T \cdot \frac{\partial Q_B}{\partial V_{BE}} \right) \quad (10.98)$$

$$g_{mr} = \frac{\partial I_T}{\partial V_{BC}} = \frac{\partial I_{TF}}{\partial V_{BC}} - \frac{\partial I_{TR}}{\partial V_{BC}} = \frac{1}{Q_B} \cdot \left(-g_{IR} - I_T \cdot \frac{\partial Q_B}{\partial V_{BC}} \right) \quad (10.99)$$

With g_{IF} being the forward conductance of the ideal forward transfer current and g_{IR} being the reverse conductance of the ideal backward transfer current.

$$g_{IF} = \frac{\partial I_F}{\partial V_{BE}} = g_{BEI} \cdot B_F \quad (10.100)$$

$$g_{IR} = \frac{\partial I_R}{\partial V_{BC}} = g_{BCI} \cdot B_R \quad (10.101)$$

The remaining derivatives in eq. (??), (??), (??) and (??) can be written as

$$\frac{\partial Q_B}{\partial V_{BE}} = Q_1 \cdot \left(\frac{Q_B}{V_{AR}} + \frac{g_{IF}}{I_{KF} \cdot \sqrt{1 + 4 \cdot Q_2}} \right) \quad (10.102)$$

$$\frac{\partial Q_B}{\partial V_{BC}} = Q_1 \cdot \left(\frac{Q_B}{V_{AF}} + \frac{g_{IR}}{I_{KR} \cdot \sqrt{1 + 4 \cdot Q_2}} \right) \quad (10.103)$$

For the calculation of the bias dependent base resistance $R_{BB'}$ there are two different ways within the SGP model. If the model parameter I_{RB} is not given it is determined by the normalized base charge Q_B . Otherwise I_{RB} specifies the base current at which the base resistance drops half way to the minimum (i.e. the constant component) base resistance R_{Bm} .

$$R_{BB'} = \begin{cases} R_{Bm} + \frac{R_B - R_{Bm}}{Q_B} & \text{for } I_{RB} = \infty \\ R_{Bm} + 3 \cdot (R_B - R_{Bm}) \cdot \frac{\tan z - z}{z \cdot \tan^2 z} & \text{for } I_{RB} \neq \infty \end{cases} \quad (10.104)$$

$$\text{with } z = \frac{\sqrt{1 + \frac{144}{\pi^2} \cdot \frac{I_B}{I_{RB}}} - 1}{\frac{24}{\pi^2} \cdot \sqrt{\frac{I_B}{I_{RB}}}} \quad (10.105)$$

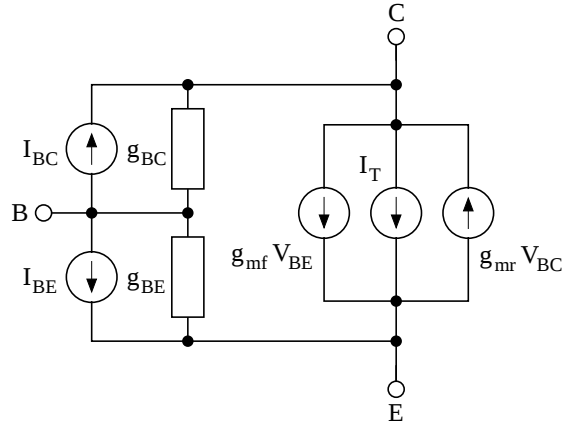


Figure 10.11: accompanied DC model of intrinsic BJT

With the accompanied DC model shown in fig. ?? the MNA matrix entries as well as the current

vector entries differ.

$$\begin{bmatrix} g_\mu + g_\pi & -g_\mu & -g_\pi & 0 \\ -g_\mu + g_{mf} - g_{mr} & g_\mu + g_{mr} & -g_{mf} & 0 \\ -g_\pi - g_{mf} + g_{mr} & -g_{mr} & g_\pi + g_{mf} & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_B \\ V_C \\ V_E \\ V_S \end{bmatrix} = \begin{bmatrix} -I_{BE_{eq}} - I_{BC_{eq}} \\ +I_{BC_{eq}} - I_{CE_{eq}} \\ +I_{BE_{eq}} + I_{CE_{eq}} \\ 0 \end{bmatrix} \quad (10.106)$$

$$I_{BE_{eq}} = I_{BE} - g_\pi \cdot V_{BE} \quad (10.107)$$

$$I_{BC_{eq}} = I_{BC} - g_\mu \cdot V_{BC} \quad (10.108)$$

$$I_{CE_{eq}} = I_T - g_{mf} \cdot V_{BE} + g_{mr} \cdot V_{BC} \quad (10.109)$$

In order to implement the influence of the excess phase parameter φ_{TF} – denoting the phase shift of the current gain at the transit frequency – the method developed by P.B. Weil and L.P. McNamee [?] can be used. They propose to use a second-order Bessel polynomial to modify the forward transfer current:

$$I_{Tx} = I_T \cdot \Phi(s) = I_T \cdot \frac{3 \cdot \omega_0^2}{s^2 + 3 \cdot \omega_0 \cdot s + 3 \cdot \omega_0^2} \quad (10.110)$$

This polynomial is formulated to closely resemble a time domain delay for a Gaussian curve which is similar to the physical phenomenon exhibited by bipolar transistor action.

Applying the inverse Laplace transformation to eq. (??) and using finite difference methods the transfer current can be written as

$$I_{Tx}^{n+1} = C_1 \cdot I_T^{n+1} + C_2 \cdot I_{Tx}^n - C_3 \cdot I_{Tx}^{n-1} \quad (10.111)$$

with

$$C_1 = \frac{3 \cdot \omega_0^2 \cdot \Delta t^2}{1 + 3 \cdot \omega_0 \cdot \Delta t + 3 \cdot \omega_0^2 \cdot \Delta t^2} \quad (10.112)$$

$$C_2 = \frac{2 + 3 \cdot \omega_0 \cdot \Delta t}{1 + 3 \cdot \omega_0 \cdot \Delta t + 3 \cdot \omega_0^2 \cdot \Delta t^2} \quad (10.113)$$

$$C_3 = \frac{1}{1 + 3 \cdot \omega_0 \cdot \Delta t + 3 \cdot \omega_0^2 \cdot \Delta t^2} \quad (10.114)$$

and

$$\omega_0 = \frac{\pi}{180} \cdot \frac{1}{\varphi_{TF} \cdot T_F} \quad (10.115)$$

The appropriate modified derivative writes as

$$g_{mx}^{n+1} = C_1 \cdot g_m^{n+1} \quad (10.116)$$

It should be noted that the excess phase implementation during the transient analysis (and thus in the AC analysis as well) holds for the forward part of the transfer current only.

With non-equidistant integration time steps during transient analysis present eqs. (??) and (??) yield

$$C_2 = \frac{1 + \Delta t / \Delta t_1 + 3 \cdot \omega_0 \cdot \Delta t}{1 + 3 \cdot \omega_0 \cdot \Delta t + 3 \cdot \omega_0^2 \cdot \Delta t^2} \quad (10.117)$$

$$C_3 = \frac{\Delta t / \Delta t_1}{1 + 3 \cdot \omega_0 \cdot \Delta t + 3 \cdot \omega_0^2 \cdot \Delta t^2} \quad (10.118)$$

whereas Δt denotes the current time step and Δt_1 the previous one.

Original SPICE model

The original SGP model implementation defines the output conductance g_0 and the transconductance value g_m . Thus the SPICE simulator is able to compute the BJT circuit using a single voltage controlled current source. These definitions are given here.

$$g_0 = \left. \frac{\partial I_T}{\partial V_{CE}} \right|_{V_{BE}=\text{const}} = -\frac{\partial I_T}{\partial V_{BC}} = -g_{mr} = \frac{1}{Q_B} \cdot \left(g_{IR} + I_T \cdot \frac{\partial Q_B}{\partial V_{BC}} \right) \quad (10.119)$$

$$g_m = \left. \frac{\partial I_T}{\partial V_{BE}} \right|_{V_{CE}=\text{const}} = \frac{\partial I_T}{\partial V_{BE}} + \frac{\partial I_T}{\partial V_{BC}} = g_{mf} + g_{mr} = \frac{1}{Q_B} \cdot \left(g_{IF} - I_T \cdot \frac{\partial Q_B}{\partial V_{BE}} \right) - g_0 \quad (10.120)$$

There are two possible ways to compute the MNA matrix of the SGP model. One using a single voltage controlled current source with an accompanied output conductance and the other using two independent voltage controlled current sources (see fig.??). Both possibilities are equivalent.

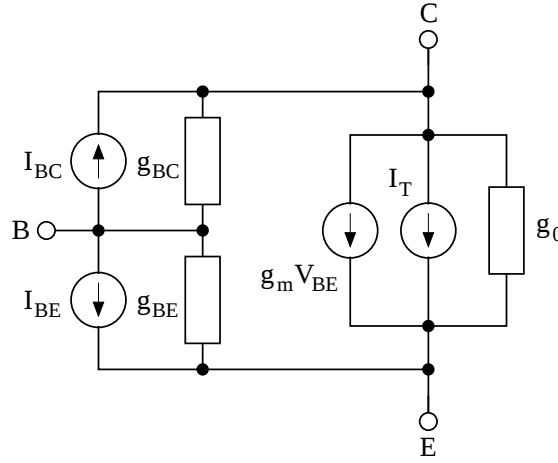


Figure 10.12: accompanied DC model of intrinsic BJT in SPICE

With the accompanied DC model shown in fig. ?? it is possible to build the complete MNA matrix of the intrinsic BJT and the current vector.

$$\begin{bmatrix} g_\mu + g_\pi & -g_\mu & -g_\pi & 0 \\ -g_\mu + g_m & g_0 + g_\mu & -g_0 - g_m & 0 \\ -g_\pi - g_m & -g_0 & g_\pi + g_0 + g_m & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} V_B \\ V_C \\ V_E \\ V_S \end{bmatrix} = \begin{bmatrix} -I_{BE_{eq}} - I_{BC_{eq}} \\ +I_{BC_{eq}} - I_{CE_{eq}} \\ +I_{BE_{eq}} + I_{CE_{eq}} \\ 0 \end{bmatrix} \quad (10.121)$$

$$I_{BE_{eq}} = I_{BE} - g_\pi \cdot V_{BE} \quad (10.122)$$

$$I_{BC_{eq}} = I_{BC} - g_\mu \cdot V_{BC} \quad (10.123)$$

$$I_{CE_{eq}} = I_T - g_m \cdot V_{BE} - g_0 \cdot V_{CE} \quad (10.124)$$

10.4.2 Small signal model

Equations for the real valued conductances in both equivalent circuits for the intrinsic BJT have already been given.

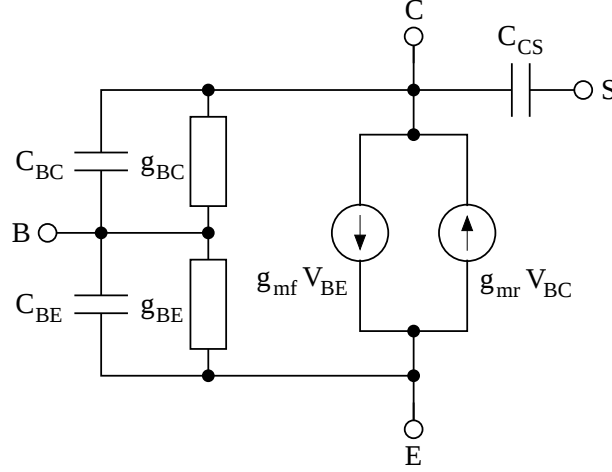


Figure 10.13: small signal model of intrinsic BJT

The junctions depletion capacitances in the SGP model write as follows:

$$C_{BE_{dep}} = \begin{cases} C_{JE} \cdot \left(1 - \frac{V_{BE}}{V_{JE}}\right)^{-M_{JE}} & \text{for } V_{BE} \leq F_C \cdot V_{JE} \\ \frac{C_{JE}}{(1 - F_C)^{M_{JE}}} \cdot \left(1 + \frac{M_{JE} \cdot (V_{BE} - F_C \cdot V_{JE})}{V_{JE} \cdot (1 - F_C)}\right) & \text{for } V_{BE} > F_C \cdot V_{JE} \end{cases} \quad (10.125)$$

$$C_{BC_{dep}} = \begin{cases} C_{JC} \cdot \left(1 - \frac{V_{BC}}{V_{JC}}\right)^{-M_{JC}} & \text{for } V_{BC} \leq F_C \cdot V_{JC} \\ \frac{C_{JC}}{(1 - F_C)^{M_{JC}}} \cdot \left(1 + \frac{M_{JC} \cdot (V_{BC} - F_C \cdot V_{JC})}{V_{JC} \cdot (1 - F_C)}\right) & \text{for } V_{BC} > F_C \cdot V_{JC} \end{cases} \quad (10.126)$$

$$C_{CS_{dep}} = \begin{cases} C_{JS} \cdot \left(1 - \frac{V_{CS}}{V_{JS}}\right)^{-M_{JS}} & \text{for } V_{CS} \leq 0 \\ C_{JS} \cdot \left(1 + M_{JS} \cdot \frac{V_{CS}}{V_{JS}}\right) & \text{for } V_{CS} > 0 \end{cases} \quad (10.127)$$

The base-collector depletion capacitance is split into two components: an external and an internal.

$$C_{BCI_{dep}} = X_{CJC} \cdot C_{BC_{dep}} \quad (10.128)$$

$$C_{BCX_{dep}} = (1 - X_{CJC}) \cdot C_{BC_{dep}} \quad (10.129)$$

The base-emitter diffusion capacitance can be obtained using the following equation.

$$C_{BE_{diff}} = \frac{\partial Q_{BE}}{\partial V_{BE}} \quad \text{with} \quad Q_{BE} = \frac{I_F}{Q_B} \cdot T_{FF} \quad (10.130)$$

Thus the diffusion capacitance depends on the bias-dependent effective forward transit time T_{FF} which is defined as:

$$T_{FF} = T_F \cdot \left(1 + X_{TF} \cdot \left(\frac{I_F}{I_F + I_{TF}} \right)^2 \cdot \exp \left(\frac{V_{BC}}{1.44 \cdot V_{TF}} \right) \right) \quad (10.131)$$

With

$$\frac{\partial T_{FF}}{\partial V_{BE}} = \frac{T_F \cdot X_{TF} \cdot 2 \cdot g_{IF} \cdot I_F \cdot I_{TF}}{(I_F + I_{TF})^3} \cdot \exp \left(\frac{V_{BC}}{1.44 \cdot V_{TF}} \right) \quad (10.132)$$

the base-emitter diffusion capacitance can finally be written as:

$$C_{BE_{diff}} = \frac{\partial Q_{BE}}{\partial V_{BE}} = \frac{1}{Q_B} \cdot \left(I_F \cdot \frac{\partial T_{FF}}{\partial V_{BE}} + T_{FF} \cdot \left(g_{IF} - \frac{I_F}{Q_B} \cdot \frac{\partial Q_B}{\partial V_{BE}} \right) \right) \quad (10.133)$$

Because the base-emitter charge Q_{BE} in eq. (??) also depends on the voltage across the base-collector junction, it is necessary to find the appropriate derivative as well:

$$C_{BE_{BC}} = \frac{\partial Q_{BE}}{\partial V_{BC}} = \frac{I_F}{Q_B} \cdot \left(\frac{\partial T_{FF}}{\partial V_{BC}} - \frac{T_{FF}}{Q_B} \cdot \frac{\partial Q_B}{\partial V_{BC}} \right) \quad (10.134)$$

which turns out to be a so called transcapacitance. It additionally requires:

$$\frac{\partial T_{FF}}{\partial V_{BC}} = \frac{T_F \cdot X_{TF}}{1.44 \cdot V_{TF}} \cdot \left(\frac{I_F}{I_F + I_{TF}} \right)^2 \cdot \exp \left(\frac{V_{BC}}{1.44 \cdot V_{TF}} \right) \quad (10.135)$$

The base-collector diffusion capacitance writes as follows:

$$C_{BC_{diff}} = \frac{\partial Q_{BC}}{\partial V_{BC}} = T_R \cdot g_{IR} \quad (10.136)$$

To take the excess phase parameter φ_{TF} into account the forward transconductance is going to be a complex quantity.

$$g_{mf} = g_{mf} \cdot e^{-j\varphi_{ex}} \quad \text{with} \quad \varphi_{ex} = \left(\frac{\pi}{180} \cdot \varphi_{TF} \right) \cdot T_F \cdot 2\pi f \quad (10.137)$$

With these calculations made it is now possible to define the small signal Y-parameters of the intrinsic BJT. The Y-parameter matrix can be converted to S-parameters.

$$Y = \begin{bmatrix} Y_{BC} + Y_{BE} + Y_{BE_{BC}} & -Y_{BC} - Y_{BE_{BC}} & -Y_{BE} & 0 \\ g_{mf} - Y_{BC} - g_{mr} & Y_{CS} + Y_{BC} + g_{mr} & -g_{mf} & -Y_{CS} \\ g_{mr} - g_{mf} - Y_{BE} - Y_{BE_{BC}} & -g_{mr} + Y_{BE_{BC}} & Y_{BE} + g_{mf} & 0 \\ 0 & -Y_{CS} & 0 & Y_{CS} \end{bmatrix} \quad (10.138)$$

with

$$Y_{BC} = g_{\mu} + j\omega (C_{BCI_{dep}} + C_{BC_{diff}}) \quad (10.139)$$

$$Y_{BE} = g_{\pi} + j\omega (C_{BE_{dep}} + C_{BE_{diff}}) \quad (10.140)$$

$$Y_{CS} = j\omega \cdot C_{CS_{dep}} \quad (10.141)$$

$$Y_{BE_{BC}} = j\omega \cdot C_{BE_{BC}} \quad (10.142)$$

The external capacitance C_{BCX} connected between the internal collector node and the external base node is separately modeled if it is non-zero and if there is a non-zero base resistance.

Original SPICE model

The original SPICE variant of the above small signal equivalent circuit with the transconductance g_m and the output conductance g_o is depicted in fig. ??.

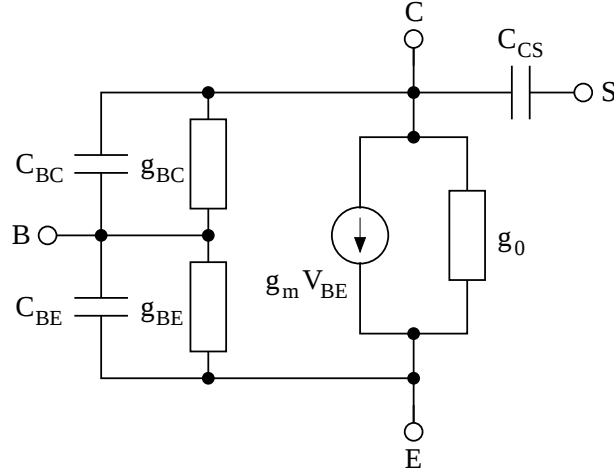


Figure 10.14: small signal model of intrinsic BJT in SPICE

The appropriate MNA matrix (Y-parameters) during the small signal analysis can be written as

$$Y = \begin{bmatrix} Y_{BC} + Y_{BE} + Y_{BE_{BC}} & -Y_{BC} - Y_{BE_{BC}} & -Y_{BE} & 0 \\ g_m - Y_{BC} & Y_{CS} + Y_{BC} + g_0 & -g_m - g_0 & -Y_{CS} \\ -g_m - Y_{BE} - Y_{BE_{BC}} & -g_0 + Y_{BE_{BC}} & Y_{BE} + g_m + g_0 & 0 \\ 0 & -Y_{CS} & 0 & Y_{CS} \end{bmatrix} \quad (10.143)$$

10.4.3 Noise model

The ohmic resistances $R_{BB'}$, R_C and R_E generate thermal noise characterized by the following spectral densities.

$$\frac{\overline{i_{R_{BB'}}^2}}{\Delta f} = \frac{4k_B T}{R_{BB'}} \quad \text{and} \quad \frac{\overline{i_{R_C}^2}}{\Delta f} = \frac{4k_B T}{R_C} \quad \text{and} \quad \frac{\overline{i_{R_E}^2}}{\Delta f} = \frac{4k_B T}{R_E} \quad (10.144)$$

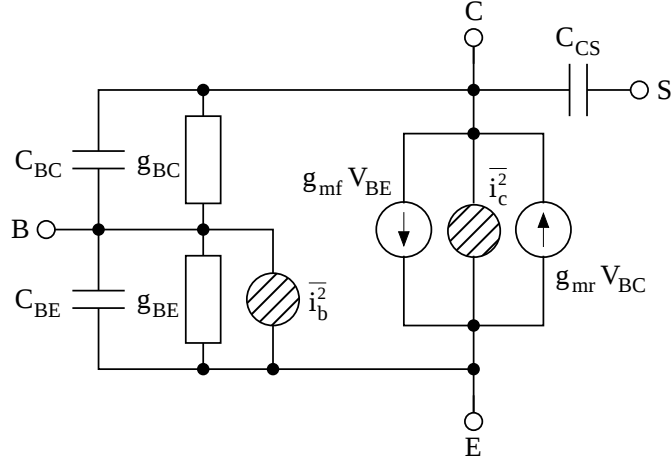


Figure 10.15: noise model of intrinsic BJT

Shot noise, flicker noise and burst noise generated by the DC base current is characterized by the spectral density

$$\frac{\overline{i_b^2}}{\Delta f} = 2eI_{BE} + K_F \frac{I_{BE}^{A_F}}{f^{F_{FE}}} + K_B \frac{I_{BE}^{A_B}}{1 + \left(\frac{f}{F_B}\right)^2} \quad (10.145)$$

The shot noise generated by the DC collector to emitter current flow is characterized by the spectral density

$$\frac{\overline{i_c^2}}{\Delta f} = 2eI_T \quad (10.146)$$

The noise current correlation matrix of the four port intrinsic bipolar transistor can then be written as

$$\underline{C}_Y = \Delta f \begin{bmatrix} +\overline{i_b^2} & 0 & -\overline{i_b^2} & 0 \\ 0 & +\overline{i_c^2} & -\overline{i_c^2} & 0 \\ -\overline{i_b^2} & -\overline{i_c^2} & +\overline{i_c^2} + \overline{i_b^2} & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \quad (10.147)$$

This matrix representation can be converted to the noise wave correlation matrix representation $\underline{C}_S$ using the formulas given in section ?? on page ??.

10.4.4 Temperature model

Temperature appears explicitly in the exponential term of the bipolar transistor model equations. In addition, the model parameters are modified to reflect changes in the temperature. The reference temperature T_1 in these equations denotes the nominal temperature T_{NOM} specified by the bipolar transistor model.

$$I_S(T_2) = I_S(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TI}} \cdot \exp\left[-\frac{e \cdot E_G(300K)}{k_B \cdot T_2} \cdot \left(1 - \frac{T_2}{T_1}\right)\right] \quad (10.148)$$

$$V_{JE}(T_2) = \frac{T_2}{T_1} \cdot V_{JE}(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.149)$$

$$V_{JC}(T_2) = \frac{T_2}{T_1} \cdot V_{JC}(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.150)$$

$$V_{JS}(T_2) = \frac{T_2}{T_1} \cdot V_{JS}(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.151)$$

where the $E_G(T)$ dependency has already been described in section ?? on page ?. The temperature dependence of B_F and B_R is determined by

$$B_F(T_2) = B_F(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TB}} \quad (10.152)$$

$$B_R(T_2) = B_R(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{X_{TB}} \quad (10.153)$$

Through the parameters I_{SE} and I_{SC} , respectively, the temperature dependence of the non-ideal saturation currents is determined by

$$I_{SE}(T_2) = I_{SE}(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{-X_{TB}} \cdot \left[\frac{I_S(T_2)}{I_S(T_1)}\right]^{1/N_E} \quad (10.154)$$

$$I_{SC}(T_2) = I_{SC}(T_1) \cdot \left(\frac{T_2}{T_1}\right)^{-X_{TB}} \cdot \left[\frac{I_S(T_2)}{I_S(T_1)}\right]^{1/N_C} \quad (10.155)$$

The temperature dependence of the zero-bias depletion capacitances C_{JE} , C_{JC} and C_{JS} are determined by

$$C_{JE}(T_2) = C_{JE}(T_1) \cdot \left(1 + M_{JE} \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{V_{JE}(T_2) - V_{JE}(T_1)}{V_{JE}(T_1)}\right)\right) \quad (10.156)$$

$$C_{JC}(T_2) = C_{JC}(T_1) \cdot \left(1 + M_{JC} \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{V_{JC}(T_2) - V_{JC}(T_1)}{V_{JC}(T_1)}\right)\right) \quad (10.157)$$

$$C_{JS}(T_2) = C_{JS}(T_1) \cdot \left(1 + M_{JS} \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{V_{JS}(T_2) - V_{JS}(T_1)}{V_{JS}(T_1)}\right)\right) \quad (10.158)$$

10.4.5 Area dependence of the model

The area factor A used in the bipolar transistor model determines the number of equivalent parallel devices of a specified model. The bipolar transistor model parameters affected by the A factor are:

$$I_S(A) = I_S \cdot A \quad (10.159)$$

$$I_{SE}(A) = I_{SE} \cdot A \quad I_{SC}(A) = I_{SC} \cdot A \quad (10.160)$$

$$I_{KF}(A) = I_{KF} \cdot A \quad I_{KR}(A) = I_{KR} \cdot A \quad (10.161)$$

$$I_{RB}(A) = I_{RB} \cdot A \quad I_{TF}(A) = I_{TF} \cdot A \quad (10.162)$$

$$C_{JE}(A) = C_{JE} \cdot A \quad C_{JC}(A) = C_{JC} \cdot A \quad (10.163)$$

$$C_{JS}(A) = C_{JS} \cdot A \quad (10.164)$$

$$R_B(A) = \frac{R_B}{A} \quad R_{Bm}(A) = \frac{R_{Bm}}{A} \quad (10.165)$$

$$R_E(A) = \frac{R_E}{A} \quad R_C(A) = \frac{R_C}{A} \quad (10.166)$$

10.5 MOS Field-Effect Transistor

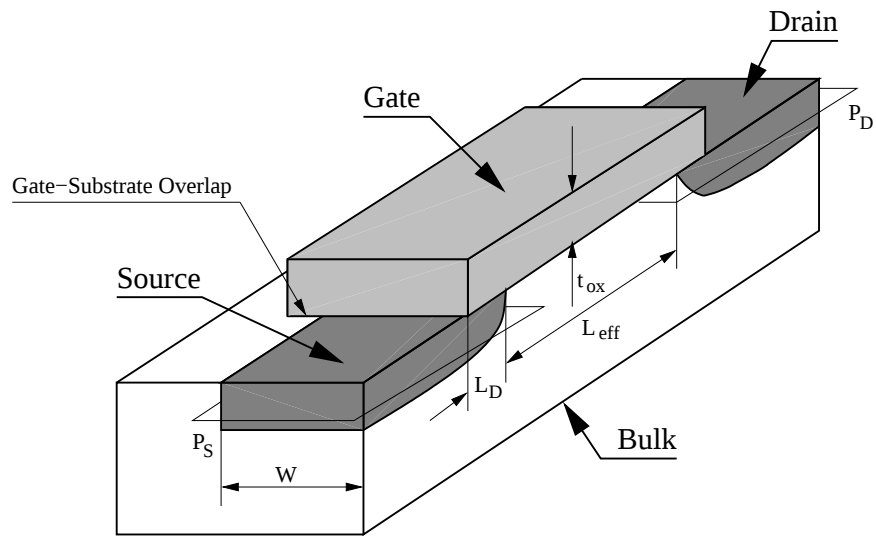


Figure 10.16: vertical section of integrated MOSFET

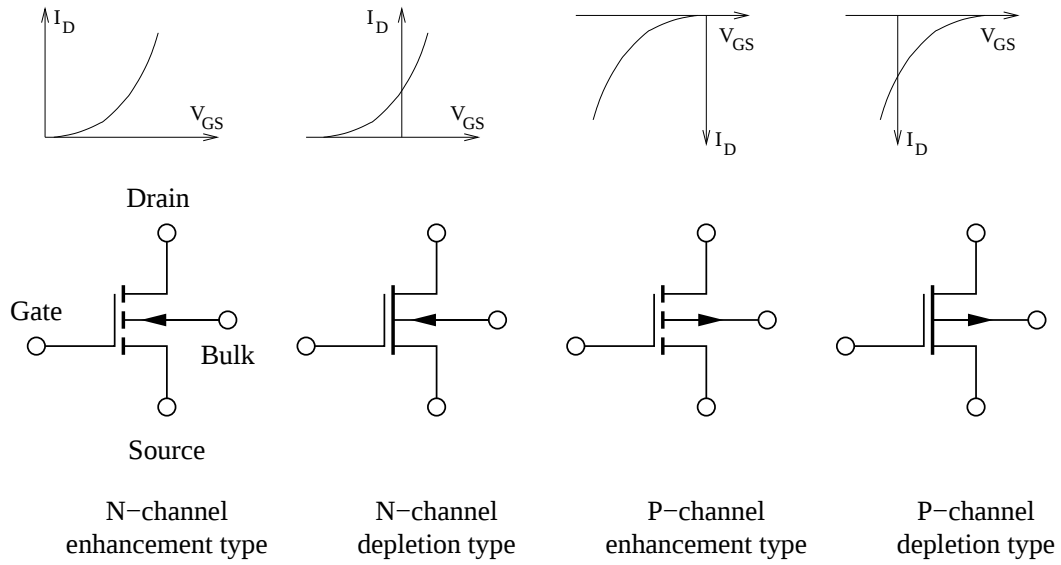


Figure 10.17: four types of MOS field effect transistors and their symbols

There are four different types of MOS field effect transistors as shown in fig. ?? all covered by the model going to be explained here. The “First Order Model” is a physical model with the drain current equations according to Harold Shichman and David A. Hodges [?].

The following table contains the model and device parameters for the MOSFET level 1.

	Name	Symbol	Description	Unit	Default	Typical
Is	I_S		bulk junction saturation current	A	10^{-14}	10^{-15}
N	N		bulk junction emission coefficient		1.0	
Vt0	V_{T0}		zero-bias threshold voltage	V	0.0	0.7
Lambda	λ		channel-length modulation parameter	1/V	0.0	0.02
Kp	K_P		transconductance coefficient	A/V ²	$2 \cdot 10^{-5}$	$6 \cdot 10^{-5}$
Gamma	γ		bulk threshold	$\sqrt{V}$	0.0	0.37
Phi	Φ		surface potential	V	0.6	0.65
Rd	R_D		drain ohmic resistance	Ω	0.0	1.0
Rs	R_S		source ohmic resistance	Ω	0.0	1.0
Rg	R_G		gate ohmic resistance	Ω	0.0	
L	L		channel length	m	100 μ	
Ld	L_D		lateral diffusion length	m	0.0	10^{-7}
W	W		channel width	m	100 μ	
Tox	T_{OX}		oxide thickness	m	0.1 μ	$2 \cdot 10^{-8}$
Cgso	C_{GSO}		gate-source overlap capacitance per meter of channel width	F/m	0.0	$4 \cdot 10^{-11}$
Cgdo	C_{GDO}		gate-drain overlap capacitance per meter of channel width	F/m	0.0	$4 \cdot 10^{-11}$
Cgbo	C_{GBO}		gate-bulk overlap capacitance per meter of channel length	F/m	0.0	$2 \cdot 10^{-10}$

		Name	Symbol	Description	Unit	Default	Typical	
Cbd	C_{BD}	zero-bias		bulk-drain junction capaci-	F	0.0		$6 \cdot 10^{-17}$
Cbs	C_{BS}	tance		zero-bias bulk-source junction capaci-	F	0.0		$6 \cdot 10^{-17}$
Pb	Φ_B	tance		bulk junction potential	V	0.8		0.87
Mj	M_J	bulk junction bottom grading coeffi-				0.5		0.5
Fc	F_C	cient		bulk junction forward-bias depletion ca-		0.5		
Cjsw	C_{JSW}	pacitance coefficient		zero-bias bulk junction periphery ca-	F/m	0.0		
Mjsw	M_{JSW}	pacitance per meter of junction perime-						
Tt	T_T	ter		bulk junction periphery grading coeffi-		0.33		0.33
Kf	K_F	cient		bulk transit time	s	0.0		
Af	A_F	flicker noise coefficient				0.0		
Ffe	F_{FE}	flicker noise exponent				1.0		
Nsub	N_{SUB}	flicker noise frequency exponent				1.0		
Nss	N_{SS}	substrate (bulk) doping density			$1/\text{cm}^3$	0.0		$4 \cdot 10^{15}$
Tpg	T_{PG}	surface state density			$1/\text{cm}^2$	0.0		10^{10}
Uo	μ_0	gate material type (0 = alumina, -1 =				1		
Rsh	R_{SH}	same as bulk, 1 = opposite to bulk)						
Nrd	N_{RD}	surface mobility			cm^2/Vs	600.0		400.0
Nrs	N_{RS}	drain and source diffusion sheet resis-			Ω/square	0.0		10.0
Cj	C_J	tance		number of equivalent drain squares		1		
Js	J_S	number of equivalent source squares				1		
Ad	A_D	zero-bias bulk junction bottom capaci-			F/m^2	0.0		$2 \cdot 10^{-4}$
As	A_S	tance per square meter of junction area						
Pd	P_D	bulk junction saturation current per			A/m^2	0.0		10^{-8}
Ps	P_S	square meter of junction area						
Temp	T	drain diffusion area			m^2	0.0		
Tnom	T_{NOM}	source diffusion area			m^2	0.0		
		drain junction perimeter			m	0.0		
		source junction perimeter			m	0.0		
		device temperature			$^{\circ}\text{C}$	26.85		
		parameter measurement temperature			$^{\circ}\text{C}$	26.85		

10.5.1 Large signal model

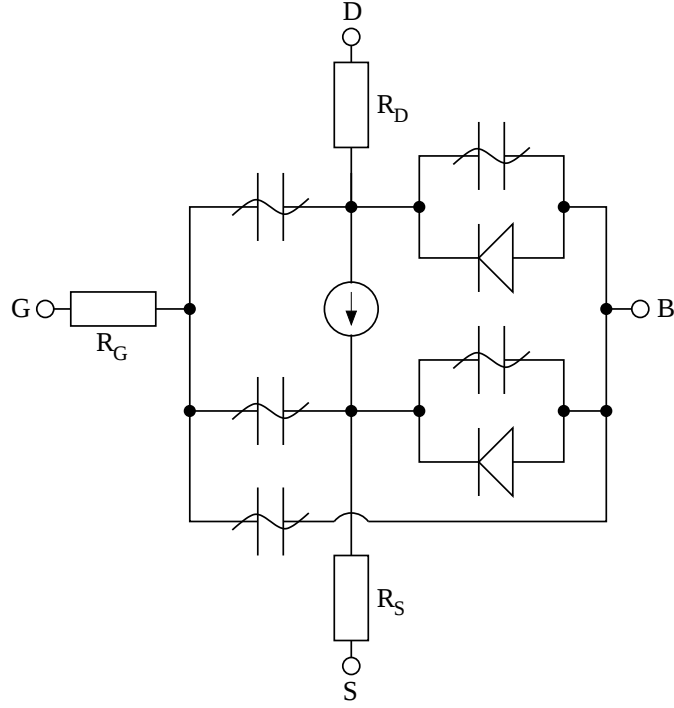


Figure 10.18: n-channel MOSFET large signal model

Beforehand some useful abbreviation are made to simplify the DC current equations.

$$L_{eff} = L - 2 \cdot L_D \quad (10.167)$$

$$\beta = K_P \cdot \frac{W}{L_{eff}} \quad (10.168)$$

The bias-dependent threshold voltage depends on the bulk-source voltage V_{BS} or the bulk-drain voltage V_{BD} depending on the mode of operation.

$$V_{Th} = V_{T0} + \begin{cases} \gamma \cdot \left(\sqrt{\Phi - V_{BS}} - \sqrt{\Phi} \right) & \text{for } V_{DS} \geq 0, \text{ i.e. } V_{BS} \geq V_{BD} \\ \gamma \cdot \left(\sqrt{\Phi - V_{BD}} - \sqrt{\Phi} \right) & \text{for } V_{DS} < 0, \text{ i.e. } V_{BD} > V_{BS} \end{cases} \quad (10.169)$$

The following equations describe the DC current behaviour of a N-channel MOSFET in normal mode, i.e. $V_{DS} > 0$, according to Shichman and Hodges.

- cutoff region: $V_{GS} - V_{Th} < 0$

$$I_d = 0 \quad (10.170)$$

$$g_{ds} = 0 \quad (10.171)$$

$$g_m = 0 \quad (10.172)$$

$$g_{mb} = 0 \quad (10.173)$$

- saturation region: $0 < V_{GS} - V_{Th} < V_{DS}$

$$I_d = \beta/2 \cdot (1 + \lambda V_{DS}) \cdot (V_{GS} - V_{Th})^2 \quad (10.174)$$

$$g_{ds} = \beta/2 \cdot \lambda (V_{GS} - V_{Th})^2 \quad (10.175)$$

$$g_m = \beta \cdot (1 + \lambda V_{DS}) (V_{GS} - V_{Th}) \quad (10.176)$$

$$g_{mb} = g_m \cdot \frac{\gamma}{2\sqrt{\Phi - V_{BS}}} \quad (10.177)$$

- linear region: $V_{DS} < V_{GS} - V_{Th}$

$$I_d = \beta \cdot (1 + \lambda V_{DS}) \cdot (V_{GS} - V_{Th} - V_{DS}/2) \cdot V_{DS} \quad (10.178)$$

$$g_{ds} = \beta \cdot (1 + \lambda V_{DS}) \cdot (V_{GS} - V_{Th} - V_{DS}) + \beta \cdot \lambda V_{DS} \cdot (V_{GS} - V_{Th} - V_{DS}/2) \quad (10.179)$$

$$g_m = \beta \cdot (1 + \lambda V_{DS}) \cdot V_{DS} \quad (10.180)$$

$$g_{mb} = g_m \cdot \frac{\gamma}{2\sqrt{\Phi - V_{BS}}} \quad (10.181)$$

with

$$g_{ds} = \frac{\partial I_d}{\partial V_{DS}} \quad \text{and} \quad g_m = \frac{\partial I_d}{\partial V_{GS}} \quad \text{and} \quad g_{mb} = \frac{\partial I_d}{\partial V_{BS}} \quad (10.182)$$

In the inverse mode of operation, i.e. $V_{DS} < 0$, the same equations can be applied with the following modifications. Replace V_{BS} with V_{BD} , V_{GS} with V_{GD} and V_{DS} with $-V_{DS}$. The drain current I_d gets reversed. Furthermore the transconductances alter their controlling nodes, i.e.

$$g_m = \frac{\partial I_d}{\partial V_{GD}} \quad \text{and} \quad g_{mb} = \frac{\partial I_d}{\partial V_{BD}} \quad (10.183)$$

The current equations of the two parasitic diodes at the bulk node and their derivatives write as follows.

$$I_{BD} = I_{SD} \cdot \left(e^{\frac{V_{BD}}{N \cdot V_T}} - 1 \right) \quad g_{bd} = \frac{\partial I_{BD}}{\partial V_{BD}} = \frac{I_{SD}}{N \cdot V_T} \cdot e^{\frac{V_{BD}}{N \cdot V_T}} \quad (10.184)$$

$$I_{BS} = I_{SS} \cdot \left(e^{\frac{V_{BS}}{N \cdot V_T}} - 1 \right) \quad g_{bs} = \frac{\partial I_{BS}}{\partial V_{BS}} = \frac{I_{SS}}{N \cdot V_T} \cdot e^{\frac{V_{BS}}{N \cdot V_T}} \quad (10.185)$$

with

$$I_{SD} = I_S \quad \text{and} \quad I_{SS} = I_S \quad (10.186)$$

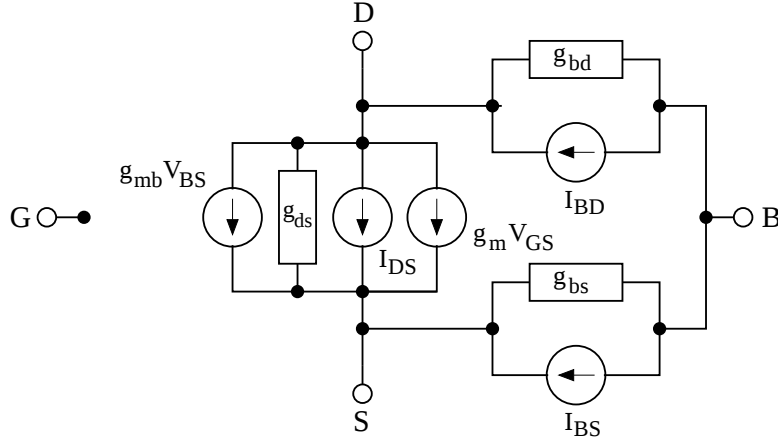


Figure 10.19: accompanied DC model of intrinsic MOSFET

With the accompanied DC model shown in fig. ?? it is possible to form the MNA matrix and the current vector of the intrinsic MOSFET device.

$$\begin{bmatrix} 0 & 0 & 0 & 0 \\ g_m & g_{ds} + g_{bd} & -g_{ds} - g_m - g_{mb} & g_{mb} - g_{bd} \\ -g_m & -g_{ds} & g_{bs} + g_{ds} + g_m + g_{mb} & -g_{bs} - g_{mb} \\ 0 & -g_{bd} & -g_{bs} & g_{bs} + g_{bd} \end{bmatrix} \cdot \begin{bmatrix} V_G \\ V_D \\ V_S \\ V_B \end{bmatrix} = \begin{bmatrix} 0 \\ +I_{BD_{eq}} - I_{DS_{eq}} \\ +I_{BS_{eq}} + I_{DS_{eq}} \\ -I_{BD_{eq}} - I_{BS_{eq}} \end{bmatrix} \quad (10.187)$$

$$I_{BD_{eq}} = I_{BD} - g_{bd} \cdot V_{BD} \quad (10.188)$$

$$I_{BS_{eq}} = I_{BS} - g_{bs} \cdot V_{BS} \quad (10.189)$$

$$I_{DS_{eq}} = I_d - g_m \cdot V_{GS} - g_{mb} \cdot V_{BS} - g_{ds} \cdot V_{DS} \quad (10.190)$$

10.5.2 Physical model

There are electrical parameters as well as physical and geometry parameters in the set of model parameters for the MOSFETs “First Order Model”. Some of the electrical parameters can be derived from the geometry and physical parameters.

The oxide capacitance per square meter of the channel area can be computed as

$$C'_{ox} = \varepsilon_0 \cdot \frac{\varepsilon_{ox}}{T_{ox}} \quad \text{with} \quad \varepsilon_{ox} = \varepsilon_{SiO_2} = 3.9 \quad (10.191)$$

Then the overall oxide capacitance can be written as

$$C_{ox} = C'_{ox} \cdot W \cdot L_{eff} \quad (10.192)$$

The transconductance coefficient K_P can be calculated using

$$K_P = \mu_0 \cdot C'_{ox} \quad (10.193)$$

The surface potential Φ is given by (with temperature voltage V_T)

$$\Phi = 2 \cdot V_T \cdot \ln \left(\frac{N_{SUB}}{n_i} \right) \quad \text{with the intrinsic density } n_i = 1.45 \cdot 10^{16} \text{1/m}^3 \quad (10.194)$$

Equation (??) holds for acceptor concentrations N_A (N_{SUB}) essentially greater than the donor concentration N_D . The bulk threshold γ (also sometimes called the body effect coefficient) is

$$\gamma = \frac{\sqrt{2 \cdot e \cdot \varepsilon_{Si} \cdot \varepsilon_0 \cdot N_{SUB}}}{C'_{ox}} \quad \text{with } \varepsilon_{Si} = 11.7 \quad (10.195)$$

And finally the zero-bias threshold voltage V_{T0} writes as follows.

$$V_{T0} = V_{FB} + \Phi + \gamma \cdot \sqrt{\Phi} \quad (10.196)$$

Whereas V_{FB} denotes the flat band voltage consisting of the work function difference Φ_{MS} between the gate and substrate material and an additional potential due to the oxide surface charge.

$$V_{FB} = \Phi_{MS} - \frac{e \cdot N_{SS}}{C'_{ox}} \quad (10.197)$$

The temperature dependent bandgap potential E_G of silicon (substrate material Si) writes as follows. With $T = 290\text{K}$ the bandgap is approximately 1.12eV .

$$E_G(T) = 1.16 - \frac{7.02 \cdot 10^{-4} \cdot T^2}{T + 1108} \quad (10.198)$$

The work function difference Φ_{MS} gets computed dependent on the gate conductor material. This can be either alumina ($\Phi_M = 4.1\text{eV}$), n-polysilicon ($\Phi_M \approx 4.15\text{eV}$) or p-polysilicon ($\Phi_M \approx 5.27\text{eV}$). The work function of a semiconductor, which is the energy difference between the vacuum level and the Fermi level (see fig. ??), varies with the doping concentration.

$$\Phi_{MS} = \Phi_M - \Phi_S = \Phi_M - \left(4.15 + \frac{1}{2}E_G + \frac{1}{2}\Phi \right) \quad (10.199)$$

$$\Phi_M = \begin{cases} 4.1 & \text{for } T_{PG} = +0, \text{ i.e. alumina} \\ 4.15 & \text{for } T_{PG} = +1, \text{ i.e. opposite to bulk} \\ 4.15 + E_G & \text{for } T_{PG} = -1, \text{ i.e. same as bulk} \end{cases} \quad (10.200)$$

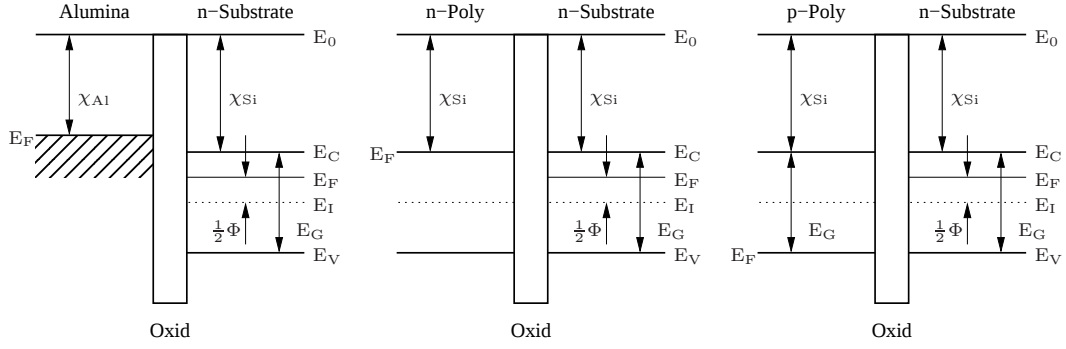


Figure 10.20: energy band diagrams of isolated (flat band) MOS materials

The expression in eq. (??) is visualized in fig. ???. The abbreviations denote

- χ_{Al} electron affinity of alumina = $4.1eV$
- χ_{Si} electron affinity of silicon = $4.15eV$
- E_0 vacuum level
- E_C conduction band
- E_V valence band
- E_F Fermi level
- E_I intrinsic Fermi level
- E_G bandgap of silicon $\approx 1.12eV$ at room temperature

Please note that the potential $1/2 \cdot \Phi$ is positive in p-MOS and negative in n-MOS as the following equation reveals.

$$\Phi_F = \frac{E_F - E_I}{e} \quad (10.201)$$

When the gate conductor material is a heavily doped polycrystalline silicon (also called polysilicon) then the model assumes that the Fermi level of this semiconductor is the same as the conduction band (for n-poly) or the valence band (for p-poly). In alumina the Fermi level, valence and conduction band all equal the electron affinity.

If the zero-bias bulk junction bottom capacitance per square meter of junction area C_J is not given it can be computed as follows.

$$C_J = \sqrt{\frac{\varepsilon_{Si} \cdot \varepsilon_0 \cdot e \cdot N_{SUB}}{2 \cdot \Phi_B}} \quad (10.202)$$

That's it for the physical parameters. The geometry parameters account for the electrical parameters per length, area or volume. Thus the MOS model is scalable.

The diffusion resistances at drain and gate are computed as follows. The sheet resistance R_{SH} refers to the thickness of the diffusion area.

$$R_D = N_{RD} \cdot R_{SH} \quad \text{and} \quad R_S = N_{RS} \cdot R_{SH} \quad (10.203)$$

If the bulk junction saturation current per square meter of the junction area J_S and the drain and source areas are given the according saturation currents are calculated with the following

equations.

$$I_{SD} = A_D \cdot J_S \quad \text{and} \quad I_{SS} = A_S \cdot J_S \quad (10.204)$$

If the parameters C_{BD} and C_{BS} are not given the zero-bias depletion capacitances for the bottom and sidewall capacitances are computed as follows.

$$C_{BD} = C_J \cdot A_D \quad (10.205)$$

$$C_{BS} = C_J \cdot A_S \quad (10.206)$$

$$C_{BDS} = C_{JSW} \cdot P_D \quad (10.207)$$

$$C_{BSS} = C_{JSW} \cdot P_S \quad (10.208)$$

10.5.3 Small signal model

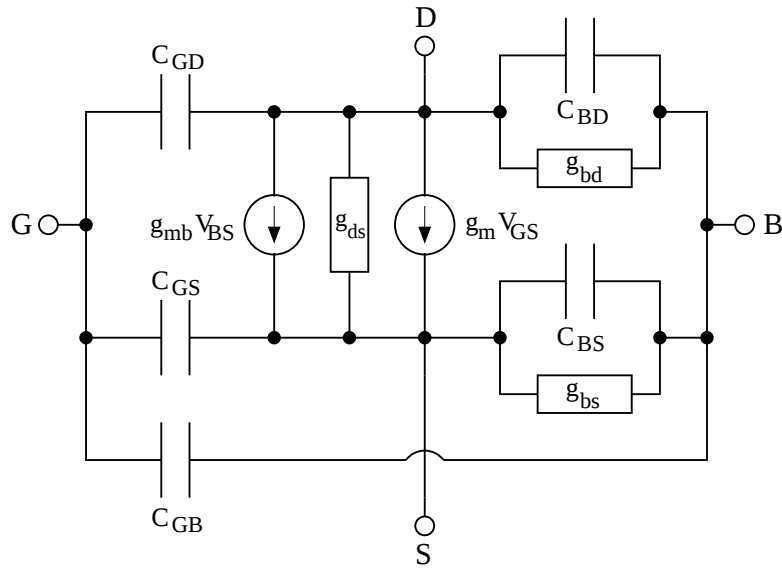


Figure 10.21: small signal model of intrinsic MOSFET

The bulk-drain and bulk-source capacitances in the MOSFET model split into three parts: the junctions depletion capacitance which consists of an area and a sidewall part and the diffusion capacitance.

$$C_{BD_{dep}} = \begin{cases} C_{BD} \cdot \left(1 - \frac{V_{BD}}{\Phi_B}\right)^{-M_J} & \text{for } V_{BD} \leq F_C \cdot \Phi_B \\ \frac{C_{BD}}{(1 - F_C)^{M_J}} \cdot \left(1 + \frac{M_J \cdot (V_{BD} - F_C \cdot \Phi_B)}{\Phi_B \cdot (1 - F_C)}\right) & \text{for } V_{BD} > F_C \cdot \Phi_B \end{cases} \quad (10.209)$$

$$C_{BDS_{dep}} = \begin{cases} C_{BDS} \cdot \left(1 - \frac{V_{BD}}{\Phi_B}\right)^{-M_{JSW}} & \text{for } V_{BD} \leq F_C \cdot \Phi_B \\ \frac{C_{BDS}}{(1 - F_C)^{M_{JSW}}} \cdot \left(1 + \frac{M_{JSW} \cdot (V_{BD} - F_C \cdot \Phi_B)}{\Phi_B \cdot (1 - F_C)}\right) & \text{for } V_{BD} > F_C \cdot \Phi_B \end{cases} \quad (10.210)$$

$$C_{BS_{dep}} = \begin{cases} C_{BS} \cdot \left(1 - \frac{V_{BS}}{\Phi_B}\right)^{-M_J} & \text{for } V_{BS} \leq F_C \cdot \Phi_B \\ \frac{C_{BS}}{(1 - F_C)^{M_J}} \cdot \left(1 + \frac{M_J \cdot (V_{BS} - F_C \cdot \Phi_B)}{\Phi_B \cdot (1 - F_C)}\right) & \text{for } V_{BS} > F_C \cdot \Phi_B \end{cases} \quad (10.211)$$

$$C_{BSS_{dep}} = \begin{cases} C_{BSS} \cdot \left(1 - \frac{V_{BS}}{\Phi_B}\right)^{-M_{JSW}} & \text{for } V_{BS} \leq F_C \cdot \Phi_B \\ \frac{C_{BSS}}{(1 - F_C)^{M_{JSW}}} \cdot \left(1 + \frac{M_{JSW} \cdot (V_{BS} - F_C \cdot \Phi_B)}{\Phi_B \cdot (1 - F_C)}\right) & \text{for } V_{BS} > F_C \cdot \Phi_B \end{cases} \quad (10.212)$$

The diffusion capacitances of the bulk-drain and bulk-source junctions are determined by the transit time of the minority charges through the junction.

$$C_{BD_{diff}} = g_{bd} \cdot T_T \quad (10.213)$$

$$C_{BS_{diff}} = g_{bs} \cdot T_T \quad (10.214)$$

Charge storage in the MOSFET consists of capacitances associated with parasitics and the intrinsic device. Parasitic capacitances consist of three constant overlap capacitances. The intrinsic capacitances consist of the nonlinear thin-oxide capacitance, which is distributed among the gate, drain, source and bulk regions. The MOS gate capacitances, as a nonlinear function of the terminal voltages, are modeled by J.E. Meyer's piece-wise linear model [?].

The bias-dependent gate-oxide capacitances distribute according to the Meyer model [?] as follows.

- cutoff regions: $V_{GS} - V_{Th} < 0$

$$- V_{GS} - V_{Th} \leq -\Phi$$

$$C_{GS} = 0 \quad (10.215)$$

$$C_{GD} = 0 \quad (10.216)$$

$$C_{GB} = C_{ox} \quad (10.217)$$

$$- \Phi < V_{GS} - V_{Th} \leq -\Phi/2$$

$$C_{GS} = 0 \quad (10.218)$$

$$C_{GD} = 0 \quad (10.219)$$

$$C_{GB} = -C_{ox} \cdot \frac{V_{GS} - V_{Th}}{\Phi} \quad (10.220)$$

$$- \Phi/2 < V_{GS} - V_{Th} \leq 0$$

$$C_{GS} = \frac{2}{3} \cdot C_{ox} + \frac{4}{3} \cdot C_{ox} \cdot \frac{V_{GS} - V_{Th}}{\Phi} \quad (10.221)$$

$$C_{GD} = 0 \quad (10.222)$$

$$C_{GB} = -C_{ox} \cdot \frac{V_{GS} - V_{Th}}{\Phi} \quad (10.223)$$

- saturation region: $0 < V_{GS} - V_{Th} < V_{DS}$

$$C_{GS} = \frac{2}{3} \cdot C_{ox} \quad (10.224)$$

$$C_{GD} = 0 \quad (10.225)$$

$$C_{GB} = 0 \quad (10.226)$$

- linear region: $V_{DS} < V_{GS} - V_{Th}$

$$C_{GS} = \frac{2}{3} \cdot C_{ox} \cdot \left(1 - \frac{(V_{Dsat} - V_{DS})^2}{(2 \cdot V_{Dsat} - V_{DS})^2} \right) \quad (10.227)$$

$$C_{GD} = \frac{2}{3} \cdot C_{ox} \cdot \left(1 - \frac{V_{Dsat}^2}{(2 \cdot V_{Dsat} - V_{DS})^2} \right) \quad (10.228)$$

$$C_{GB} = 0 \quad (10.229)$$

with

$$V_{Dsat} = \begin{cases} V_{GS} - V_{Th} & \text{for } V_{GS} - V_{Th} > 0 \\ 0 & \text{otherwise} \end{cases} \quad (10.230)$$

In the inverse mode of operation V_{GS} and V_{GD} need to be exchanged, V_{DS} changes its sign, then the above formulas can be applied as well.

The constance overlap capacitances compute as follows.

$$C_{GS_{OVL}} = C_{GSO} \cdot W \quad (10.231)$$

$$C_{GD_{OVL}} = C_{GDO} \cdot W \quad (10.232)$$

$$C_{GB_{OVL}} = C_{GBO} \cdot L_{eff} \quad (10.233)$$

With these definitions it is possible to form the small signal Y-parameter matrix of the intrinsic MOSFET device in an operating point which can be converted into S-parameters.

$$Y = \begin{bmatrix} Y_{GS} + & -Y_{GD} & -Y_{GS} & -Y_{GB} \\ Y_{GD} + Y_{GB} & Y_{GD} + Y_{BD} + Y_{DS} & -Y_{DS} - g_m - g_{mb} & -Y_{BD} + g_{mb} \\ g_m - Y_{GD} & -Y_{DS} & Y_{GS} + Y_{DS} + & -Y_{BS} - g_{mb} \\ -g_m - Y_{GS} & & Y_{BS} + g_m + g_{mb} & \\ -Y_{GB} & -Y_{BD} & -Y_{BS} & Y_{BD} + Y_{BS} + Y_{GB} \end{bmatrix} \quad (10.234)$$

with

$$Y_{GS} = j\omega (C_{GS} + C_{GS_{OVL}}) \quad (10.235)$$

$$Y_{GD} = j\omega (C_{GD} + C_{GD_{OVL}}) \quad (10.236)$$

$$Y_{GB} = j\omega (C_{GB} + C_{GB_{OVL}}) \quad (10.237)$$

$$Y_{BD} = g_{bd} + j\omega (C_{BD_{dep}} + C_{BDS_{dep}} + C_{BD_{diff}}) \quad (10.238)$$

$$Y_{BS} = g_{bs} + j\omega (C_{BS_{dep}} + C_{BSS_{dep}} + C_{BS_{diff}}) \quad (10.239)$$

$$Y_{DS} = g_{ds} \quad (10.240)$$

10.5.4 Noise model

The thermal noise generated by the external resistors R_G , R_S and R_D is characterized by the following spectral density.

$$\frac{\overline{i_{R_G}^2}}{\Delta f} = \frac{4k_B T}{R_G} \quad \text{and} \quad \frac{\overline{i_{R_D}^2}}{\Delta f} = \frac{4k_B T}{R_D} \quad \text{and} \quad \frac{\overline{i_{R_S}^2}}{\Delta f} = \frac{4k_B T}{R_S} \quad (10.241)$$

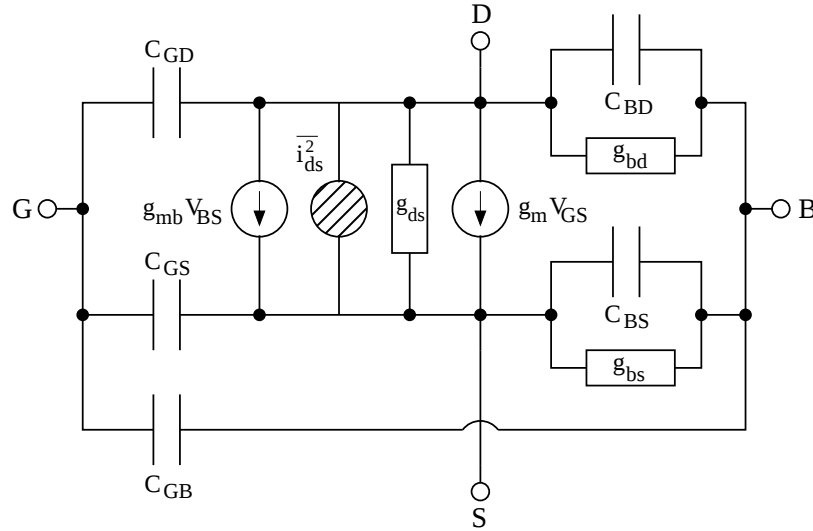


Figure 10.22: noise model of intrinsic MOSFET

Channel and flicker noise generated by the DC transconductance g_m and current flow from drain to source is characterized by the spectral density

$$\frac{\overline{i_{ds}^2}}{\Delta f} = \frac{8k_B T g_m}{3} + K_F \frac{I_{DS}^{A_F}}{f^{F_{FE}}} \quad (10.242)$$

The noise current correlation matrix (admittance representation) of the intrinsic MOSFET can be expressed as

$$\underline{C}_Y = \Delta f \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & +\overline{i_{ds}^2} & -\overline{i_{ds}^2} & 0 \\ 0 & -\overline{i_{ds}^2} & +\overline{i_{ds}^2} & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \quad (10.243)$$

This matrix representation can be easily converted to the noise-wave representation $\underline{C}_S$ if the small signal S-parameter matrix is known.

10.5.5 Temperature model

Temperature affects some MOS model parameters which are updated according to the new temperature. The reference temperature T_1 in the following equations denotes the nominal temperature T_{NOM} specified by the MOS transistor model. The temperature dependence of K_P and μ_0 is determined by

$$K_P(T_2) = K_P(T_1) \cdot \left(\frac{T_1}{T_2}\right)^{1.5} \quad (10.244)$$

$$\mu_0(T_2) = \mu_0(T_1) \cdot \left(\frac{T_1}{T_2}\right)^{1.5} \quad (10.245)$$

The effect of temperature on Φ_B and Φ is modeled by

$$\Phi(T_2) = \frac{T_2}{T_1} \cdot \Phi(T_1) - \frac{2 \cdot k_B \cdot T_2}{e} \cdot \ln\left(\frac{T_2}{T_1}\right)^{1.5} - \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right) \quad (10.246)$$

where the $E_G(T)$ dependency has already been described in section ?? on page ?. The temperature dependence of C_{BD} , C_{BS} , C_J and C_{JSW} is described by the following relations

$$C_{BD}(T_2) = C_{BD}(T_1) \cdot \left(1 + M_J \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{\Phi_B(T_2) - \Phi_B(T_1)}{\Phi_B(T_1)}\right)\right) \quad (10.247)$$

$$C_{BS}(T_2) = C_{BS}(T_1) \cdot \left(1 + M_J \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{\Phi_B(T_2) - \Phi_B(T_1)}{\Phi_B(T_1)}\right)\right) \quad (10.248)$$

$$C_J(T_2) = C_J(T_1) \cdot \left(1 + M_J \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{\Phi_B(T_2) - \Phi_B(T_1)}{\Phi_B(T_1)}\right)\right) \quad (10.249)$$

$$C_{JSW}(T_2) = C_{JSW}(T_1) \cdot \left(1 + M_{JSW} \cdot \left(400 \cdot 10^{-6} \cdot (T_2 - T_1) - \frac{\Phi_B(T_2) - \Phi_B(T_1)}{\Phi_B(T_1)}\right)\right) \quad (10.250)$$

The temperature dependence of I_S is given by the relation

$$I_S(T_2) = I_S(T_1) \cdot \exp\left[-\frac{e}{k_B \cdot T_2} \cdot \left(\frac{T_2}{T_1} \cdot E_G(T_1) - E_G(T_2)\right)\right] \quad (10.251)$$

An analogue dependence holds for J_S .

10.6 Models for boolean devices

Logical (boolean) functions (OR, AND, XOR etc.) can be modeled using macro models. Here, each input gets the transfer characteristic and its derivative described as follows:

$$u_i = \tanh(10 \cdot (u_{in} - 0.5)) \quad (10.252)$$

$$u'_i = 10 \cdot (1 - \tanh^2(10 \cdot (u_{in} - 0.5))) \quad (10.253)$$

The resulting voltages u_i for each input are combined to create the wanted function for a device with N inputs:

$$\text{Inverter: } u_{out} = 0.5 \cdot (1 - u_i) \quad (10.254)$$

$$\text{NOR: } u_{out} = \frac{N}{\sum_m \frac{2}{1 - u_{i,m}}} \quad (10.255)$$

$$\text{OR: } u_{out} = 1 - u_{out,NOR} \quad (10.256)$$

$$\text{AND: } u_{out} = \frac{N}{\sum_m \frac{2}{1 + u_{i,m}}} \quad (10.257)$$

$$\text{NAND: } u_{out} = 1 - u_{out,AND} \quad (10.258)$$

$$\text{XOR: } u_{out} = 0.5 \cdot \left(1 - \prod_m -u_{i,m} \right) \quad (10.259)$$

$$\text{XNOR: } u_{out} = 0.5 \cdot \left(1 + \prod_m u_{i,m} \right) \quad (10.260)$$

The above-mentioned functions model devices with 0V as logical low-level and 1V as logical high-level. Of course, they can be easily transformed into higher voltage levels by multiplying the desired high-level voltage to the output voltage u_{out} and dividing the input voltages u_{in} by the desired high-level voltage. Note: The derivatives also get u_{in} divided by the desired high-level voltage, but they are not multiplied by the desired high-level voltage.

To perform a simulation on these devices, the first derivatives are also needed:

$$\text{Inverter: } \frac{\partial u_{out}}{\partial u_{in}} = -0.5 \cdot u'_i \quad (10.261)$$

$$\text{OR: } \frac{\partial u_{out}}{\partial u_{in,n}} = \frac{2 \cdot N \cdot u'_{i,n}}{\left((1 - u_{i,n}) \cdot \sum_m \frac{2}{1 - u_{i,m}} \right)^2} \quad (10.262)$$

$$\text{NOR: } \frac{\partial u_{out}}{\partial u_{in,n}} = - \frac{\partial u_{out}}{\partial u_{in,n}} \Big|_{\text{OR}} \quad (10.263)$$

$$\text{AND: } \frac{\partial u_{out}}{\partial u_{in,n}} = \frac{2 \cdot N \cdot u'_{i,n}}{\left((1 + u_{i,n}) \cdot \sum_m \frac{2}{1 + u_{i,m}} \right)^2} \quad (10.264)$$

$$\text{NAND: } \frac{\partial u_{out}}{\partial u_{in,n}} = - \frac{\partial u_{out}}{\partial u_{in,n}} \Big|_{\text{AND}} \quad (10.265)$$

$$\text{XOR: } \frac{\partial u_{out}}{\partial u_{in,n}} = 0.5 \cdot u'_{i,n} \cdot \prod_{m \neq n} -u_{i,m} \quad (10.266)$$

$$\text{XNOR: } \frac{\partial u_{out}}{\partial u_{in,n}} = 0.5 \cdot u'_{i,n} \cdot \prod_{m \neq n} u_{i,m} \quad (10.267)$$

A problem of these macro models are the numbers of input ports. The output voltage levels worsen with increasing number of ports. The practical limit lies around eight input ports.

With that knowledge it is now easy to create the MNA matrix. The first port is the output port of the device. So, for a 2-input port device, it is:

$$\begin{bmatrix} \cdot & \cdot & \cdot & 1 \\ \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & 0 \\ -1 & \partial u_{out}/\partial u_{in,1} & \partial u_{out}/\partial u_{in,2} & 0 \end{bmatrix} \cdot \begin{bmatrix} V_{out} \\ V_{in,1} \\ V_{in,2} \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_0 \\ I_1 \\ I_2 \\ 0 \end{bmatrix} \quad (10.268)$$

The above MNA matrix entries are also used during the non-linear DC and transient analysis with the 0 in the right hand side vector replaced by an equivalent voltage

$$V_{eq} = \frac{\partial u_{out}}{\partial u_{in,1}} \cdot V_{in,1} + \frac{\partial u_{out}}{\partial u_{in,2}} \cdot V_{in,2} - u_{out} \quad (10.269)$$

with u_{out} computed using equations (??) to (??).

With the given small-signal matrix representation, building the S-parameters is easy.

$$(\underline{S}) = \begin{bmatrix} -1 & 4 \cdot \partial u_{out}/\partial u_{in,1} & 4 \cdot \partial u_{out}/\partial u_{in,2} \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad (10.270)$$

These matrices can easily be extended to any number of input ports.

10.7 Equation defined models

Often it will happen that a user needs to implement his own model. Therefore, it is useful to supply devices that are defined by arbitrary equations.

10.7.1 Models with Explicit Equations

For example the user must enter an equation $i(V)$ describing how the port current I depends on the port voltage $V = V_1 - V_2$ and an equation $q(V)$ describing how much charge Q is held due to the voltage V . These are time domain equations. The most simple way then is a device with two nodes. Defining

$$I = i(V) \quad \text{and} \quad g = \frac{\partial I}{\partial V} = \lim_{h \rightarrow 0} \frac{I(V+h) - I(V)}{h} \quad (10.271)$$

as well as

$$Q = q(V) \quad \text{and} \quad c = \frac{\partial Q}{\partial V} = \lim_{h \rightarrow 0} \frac{Q(V+h) - Q(V)}{h} \quad (10.272)$$

the MNA matrix for a (non-linear) DC analysis writes:

$$\begin{aligned} \begin{bmatrix} +g^{(m)} & -g^{(m)} \\ -g^{(m)} & +g^{(m)} \end{bmatrix} \cdot \begin{bmatrix} V_1^{(m+1)} \\ V_2^{(m+1)} \end{bmatrix} &= \begin{bmatrix} -I^{(m)} + g^{(m)} \cdot V^{(m)} \\ +I^{(m)} - g^{(m)} \cdot V^{(m)} \end{bmatrix} \\ &= \begin{bmatrix} -I^{(m)} + g^{(m)} \cdot (V_1^{(m)} - V_2^{(m)}) \\ +I^{(m)} - g^{(m)} \cdot (V_1^{(m)} - V_2^{(m)}) \end{bmatrix} \end{aligned} \quad (10.273)$$

For a transient simulation, equation (??) on page ?? has to be used with Q and c .

For an AC analysis the MNA matrix writes:

$$(\underline{Y}) = (g + j\omega \cdot c) \cdot \begin{bmatrix} +1 & -1 \\ -1 & +1 \end{bmatrix} \quad (10.274)$$

And the S-parameter matrix writes:

$$S_{11} = S_{22} = \frac{1}{2 \cdot Z_0 \cdot Y + 1} \quad (10.275)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (10.276)$$

$$Y = g + j\omega \cdot c \quad (10.277)$$

The simulator needs to create the derivatives g and c by its own. This can be done numerically or symbolically. One might ask why the non-linear capacitance is modeled as charge, not as capacitance. Indeed this may be changed, but with a computer algorithm, creating the derivative is easier than to integrate.

The component described above can be expanded to one with two ports (two pairs of terminals: terminal 1 and 2 and terminal 3 and 4). That is, the currents and charges of both ports depend on both port voltages $V_{12} = V_1 - V_2$ and $V_{34} = V_3 - V_4$. Thus, the defining equations are:

$$I_1 = i_1(V_{12}, V_{34}) \quad \text{and} \quad g_{11} = \frac{\partial I_1}{\partial V_{12}} \quad \text{and} \quad g_{12} = \frac{\partial I_1}{\partial V_{34}} \quad (10.278)$$

$$I_2 = i_2(V_{12}, V_{34}) \quad \text{and} \quad g_{21} = \frac{\partial I_2}{\partial V_{12}} \quad \text{and} \quad g_{22} = \frac{\partial I_2}{\partial V_{34}} \quad (10.279)$$

as well as

$$Q_1 = q_1(V_{12}, V_{34}) \quad \text{and} \quad c_{11} = \frac{\partial Q_1}{\partial V_{12}} \quad \text{and} \quad c_{12} = \frac{\partial Q_1}{\partial V_{34}} \quad (10.280)$$

$$Q_2 = q_2(V_{12}, V_{34}) \quad \text{and} \quad c_{21} = \frac{\partial Q_2}{\partial V_{12}} \quad \text{and} \quad c_{22} = \frac{\partial Q_2}{\partial V_{34}} \quad (10.281)$$

The MNA matrix for the DC analysis writes:

$$\begin{bmatrix} +g_{11}^{(m)} & -g_{11}^{(m)} & +g_{12}^{(m)} & -g_{12}^{(m)} \\ -g_{11}^{(m)} & +g_{11}^{(m)} & -g_{12}^{(m)} & +g_{12}^{(m)} \\ +g_{21}^{(m)} & -g_{21}^{(m)} & +g_{22}^{(m)} & -g_{22}^{(m)} \\ -g_{21}^{(m)} & +g_{21}^{(m)} & -g_{22}^{(m)} & +g_{22}^{(m)} \end{bmatrix} \cdot \begin{bmatrix} V_1^{(m+1)} \\ V_2^{(m+1)} \\ V_3^{(m+1)} \\ V_4^{(m+1)} \end{bmatrix} = \begin{bmatrix} -I_1^{(m)} + g_{11}^{(m)} \cdot V_{12}^{(m)} + g_{12}^{(m)} \cdot V_{34}^{(m)} \\ +I_1^{(m)} - g_{11}^{(m)} \cdot V_{12}^{(m)} - g_{12}^{(m)} \cdot V_{34}^{(m)} \\ -I_2^{(m)} + g_{21}^{(m)} \cdot V_{12}^{(m)} + g_{22}^{(m)} \cdot V_{34}^{(m)} \\ +I_2^{(m)} - g_{21}^{(m)} \cdot V_{12}^{(m)} - g_{22}^{(m)} \cdot V_{34}^{(m)} \end{bmatrix} \quad (10.282)$$

For a transient simulation, the DC equations have to be extended by the non-linear (trans-)

capacitances, e.g. for backward Euler:

$$I_{C11}^{n+1,m} = \underbrace{\frac{c_{11}(V_{12}^{n+1,m})}{h^n}}_{g_{eq,11}} \cdot V_{12}^{n+1} - \underbrace{\frac{c_{11}(V_{12}^n)}{h^n}}_{I_{eq,11}} \cdot V_{12}^n \quad (10.283)$$

$$I_{C12}^{n+1,m} = \underbrace{\frac{c_{12}(V_{12}^{n+1,m})}{h^n}}_{g_{eq,12}} \cdot V_{12}^{n+1} - \underbrace{\frac{c_{12}(V_{12}^n)}{h^n}}_{I_{eq,12}} \cdot V_{12}^n \quad (10.284)$$

$$I_{C21}^{n+1,m} = \underbrace{\frac{c_{21}(V_{34}^{n+1,m})}{h^n}}_{g_{eq,21}} \cdot V_{34}^{n+1} - \underbrace{\frac{c_{21}(V_{34}^n)}{h^n}}_{I_{eq,21}} \cdot V_{34}^n \quad (10.285)$$

$$I_{C22}^{n+1,m} = \underbrace{\frac{c_{22}(V_{34}^{n+1,m})}{h^n}}_{g_{eq,22}} \cdot V_{34}^{n+1} - \underbrace{\frac{c_{22}(V_{34}^n)}{h^n}}_{I_{eq,22}} \cdot V_{34}^n \quad (10.286)$$

So with $g_{tr} = g + g_{eq}$ it is:

$$\begin{aligned} & \begin{bmatrix} +g_{tr,11}^{(m)} & -g_{tr,11}^{(m)} & +g_{tr,12}^{(m)} & -g_{tr,12}^{(m)} \\ -g_{tr,11}^{(m)} & +g_{tr,11}^{(m)} & -g_{tr,12}^{(m)} & +g_{tr,12}^{(m)} \\ +g_{tr,21}^{(m)} & -g_{tr,21}^{(m)} & +g_{tr,22}^{(m)} & -g_{tr,22}^{(m)} \\ -g_{tr,21}^{(m)} & +g_{tr,21}^{(m)} & -g_{tr,22}^{(m)} & +g_{tr,22}^{(m)} \end{bmatrix} \cdot \begin{bmatrix} V_1^{(n+1,m+1)} \\ V_2^{(n+1,m+1)} \\ V_3^{(n+1,m+1)} \\ V_4^{(n+1,m+1)} \end{bmatrix} \\ &= \begin{bmatrix} -I_1^{(m)} + g_{11}^{(m)} \cdot V_{12}^{(m)} + g_{12}^{(m)} \cdot V_{34}^{(m)} - I_{eq,11}^{(n)} - I_{eq,12}^{(n)} \\ +I_1^{(m)} - g_{11}^{(m)} \cdot V_{12}^{(m)} - g_{12}^{(m)} \cdot V_{34}^{(m)} + I_{eq,12}^{(n)} + I_{eq,11}^{(n)} \\ -I_2^{(m)} + g_{21}^{(m)} \cdot V_{12}^{(m)} + g_{22}^{(m)} \cdot V_{34}^{(m)} - I_{eq,21}^{(n)} - I_{eq,22}^{(n)} \\ +I_2^{(m)} - g_{21}^{(m)} \cdot V_{12}^{(m)} - g_{22}^{(m)} \cdot V_{34}^{(m)} + I_{eq,22}^{(n)} + I_{eq,21}^{(n)} \end{bmatrix} \end{aligned} \quad (10.287)$$

For an AC analysis the MNA matrix writes:

$$(\underline{Y}) = \begin{bmatrix} +g_{11} + j\omega \cdot c_{11} & -g_{11} - j\omega \cdot c_{11} & +g_{12} + j\omega \cdot c_{12} & -g_{12} - j\omega \cdot c_{12} \\ -g_{11} - j\omega \cdot c_{11} & +g_{11} + j\omega \cdot c_{11} & -g_{12} - j\omega \cdot c_{12} & +g_{12} + j\omega \cdot c_{12} \\ +g_{21} + j\omega \cdot c_{21} & -g_{21} - j\omega \cdot c_{21} & +g_{22} + j\omega \cdot c_{22} & -g_{22} - j\omega \cdot c_{22} \\ -g_{21} - j\omega \cdot c_{21} & +g_{21} + j\omega \cdot c_{21} & -g_{22} - j\omega \cdot c_{22} & +g_{22} + j\omega \cdot c_{22} \end{bmatrix} \quad (10.288)$$

As can be seen, this scheme can be expanded to any number of ports. The matrices soon become quite complex, but fortunately modern computers are able to cope with this complexity. S-parameters must be obtained numerical by setting equation ?? into equation ??.

10.7.2 Models with Implicit Equations

The above-mentioned explicit models are not useable for all components. If the Y-parameters do not exist or if the equations cannot be analytically transformed into the explicit form, then an implicit representation must be taken. That is, for a one-port (two-terminal) component the following formulas are defined by the user:

$$0 = f(V, I) \quad \text{and} \quad g_V = \frac{\partial f(V, I)}{\partial V} = \lim_{h \rightarrow 0} \frac{f(V + h, I) - f(V, I)}{h} \quad (10.289)$$

$$\text{and} \quad g_I = \frac{\partial f(V, I)}{\partial I} = \lim_{h \rightarrow 0} \frac{f(V, I + h) - f(V, I)}{h} \quad (10.290)$$

The MNA matrix for the AC analysis writes as follows:

$$\begin{bmatrix} \cdot & \cdot & +1 \\ \cdot & \cdot & -1 \\ +g_V & -g_V & g_I \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{out} \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad (10.291)$$

As usual, for the DC analysis the last zero on the right hand side has to be replaced by the iteration formula:

$$g_V \cdot (V_1 - V_2) + g_I \cdot I_{out} - f(V_1 - V_2, I_{out}) \quad (10.292)$$

The S-parameters are:

$$S_{11} = S_{22} = \frac{g_I}{g_I - 2 \cdot Z_0 \cdot g_V} \quad (10.293)$$

$$S_{12} = S_{21} = 1 - S_{11} \quad (10.294)$$

Consequently, for a two-port device two equation are necessary: One for first port and one for second port:

$$0 = f_1(V_{12}, V_{34}, I_1, I_2) \quad (10.295)$$

$$0 = f_2(V_{12}, V_{34}, I_1, I_2) \quad (10.296)$$

Building the MNA matrix is again straight forward:

$$\begin{aligned} & \begin{bmatrix} \cdot & \cdot & \cdot & \cdot & +1 & 0 \\ \cdot & \cdot & \cdot & \cdot & -1 & 0 \\ \cdot & \cdot & \cdot & \cdot & 0 & +1 \\ \cdot & \cdot & \cdot & \cdot & 0 & -1 \\ +g_{f1,V12} & -g_{f1,V12} & +g_{f1,V34} & -g_{f1,V34} & g_{f1,I1} & g_{f1,I2} \\ +g_{f2,V12} & -g_{f2,V12} & +g_{f2,V34} & -g_{f2,V34} & g_{f2,I1} & g_{f2,I2} \end{bmatrix} \cdot \begin{bmatrix} V_1^{(m+1)} \\ V_2^{(m+1)} \\ V_3^{(m+1)} \\ V_4^{(m+1)} \\ I_{out1}^{(m+1)} \\ I_{out2}^{(m+1)} \end{bmatrix} \\ &= \begin{bmatrix} g_{f1,V12} \cdot V_{12} + g_{f1,V34} \cdot V_{34} + g_{f1,I1} \cdot I_{out1} + g_{f1,I2} \cdot I_{out2} - f_1(V_{12}, V_{34}, I_{out1}, I_{out2}) \\ g_{f2,V12} \cdot V_{12} + g_{f2,V34} \cdot V_{34} + g_{f2,I1} \cdot I_{out1} + g_{f2,I2} \cdot I_{out2} - f_2(V_{12}, V_{34}, I_{out1}, I_{out2}) \end{bmatrix} \end{aligned} \quad (10.297)$$

Once more, this concept can easily expanded to any number of ports. It is also possible mix implicit and explicit definitions, i.e. some ports of the device may be defined by explicit equations whereas the others are defined by implicit equations.

The calculation of the S-parameters is not that trivial. The Y-parameters as well as the Z-parameters might be infinite. A small trick can avoid this problem, as will be shown in the following 2-port example. First, the small-signal Y-parameters should be derived by using the law about implicit functions:

$$(\underline{J}) = \begin{pmatrix} y_{11} & y_{12} \\ y_{21} & y_{22} \end{pmatrix} = - \underbrace{\begin{pmatrix} \frac{\partial f_1}{\partial I_1} & \frac{\partial f_1}{\partial I_2} \\ \frac{\partial f_2}{\partial I_1} & \frac{\partial f_2}{\partial I_2} \end{pmatrix}}_{J_i}^{-1} \cdot \underbrace{\begin{pmatrix} \frac{\partial f_1}{\partial V_1} & \frac{\partial f_1}{\partial V_2} \\ \frac{\partial f_2}{\partial V_1} & \frac{\partial f_2}{\partial V_2} \end{pmatrix}}_{J_v} \quad (10.298)$$

The equation reveals immediately the difficulty: The inverse of the current Jacobi matrix J_i may not exist. But this problem can be outsourced to one single scalar number by using Cramer's rule for matrix inversion:

$$J_i^{-1} = \frac{1}{\Delta J_i} \cdot A_{J_i} \quad (10.299)$$

The matrix A_{J_i} is built of the sub-determinantes of J_i in the way that $a_{(n,m)}$ is the determinante of J_i without row m and without column n but multiplied with $(-1)^{n+m}$. It therefore always exists, whereas dividing by the determinante of J_i may become infinity. Now parameters can be defined as follows:

$$(\underline{J}') = \begin{pmatrix} y'_{11} & y'_{12} \\ y'_{21} & y'_{22} \end{pmatrix} = -A_{J_i} \cdot J_v \quad (10.300)$$

Before converting to S-parameters the matrix must be expanded to a 4-port matrix, because the 2-ports are not referenced to ground:

$$(\underline{J}'') = \begin{pmatrix} +y'_{11} & -y'_{11} & +y'_{12} & -y'_{12} \\ -y'_{11} & +y'_{11} & -y'_{12} & +y'_{12} \\ +y'_{21} & -y'_{21} & +y'_{22} & -y'_{22} \\ -y'_{21} & +y'_{21} & -y'_{22} & +y'_{22} \end{pmatrix} = -A'_{J_i} \cdot J'_v \quad (10.301)$$

Finally, equation (??) converts the parameters to S-parameters:

$$(\underline{S}) = ((E) - Z_0 \cdot (\underline{Y})) \cdot ((E) + Z_0 \cdot (\underline{Y}))^{-1} \quad (10.302)$$

$$= \left((E) + Z_0 \cdot \frac{1}{\Delta J_i} \cdot A'_{J_i} \cdot J'_v \right) \cdot \left((E) - Z_0 \cdot \frac{1}{\Delta J_i} \cdot A'_{J_i} \cdot J'_v \right)^{-1} \quad (10.303)$$

$$= (\Delta J_i \cdot (E) + Z_0 \cdot A'_{J_i} \cdot J'_v) \cdot (\Delta J_i \cdot (E) - Z_0 \cdot A'_{J_i} \cdot J'_v)^{-1} \quad (10.304)$$

The calculations proofs that the critical factor $1/\Delta J_i$ disappears and a solution exists if and only if the S-parameters of this device exist.

Chapter 11

Microstrip components

11.1 Single microstrip line

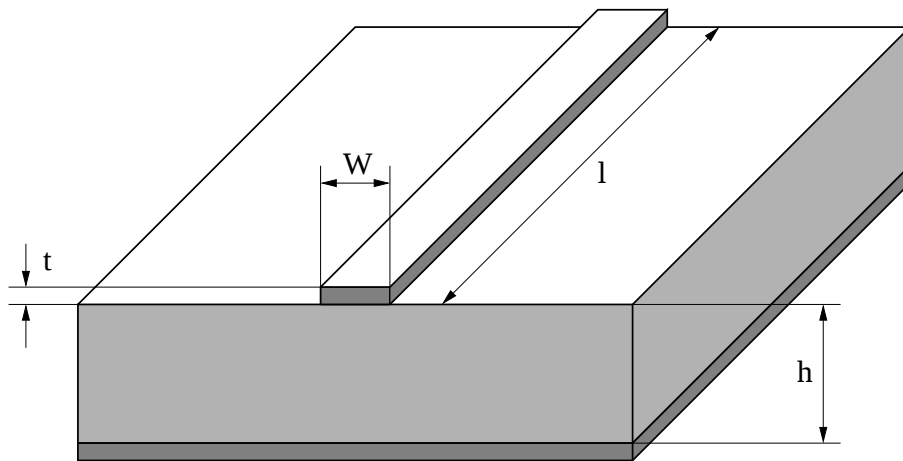


Figure 11.1: single microstrip line

The electrical parameters of microstrip lines which are required for circuit design are impedance, attenuation, wavelength and propagation constant. These parameters are interrelated for all microstrips assuming that the propagation mode is a transverse electromagnetic mode, or it can be approximated by a transverse electromagnetic mode. The Y and S parameters can be found in section ??.

11.1.1 Quasi-static characteristic impedance

Wheeler

Harold A. Wheeler [?] formulated his synthesis and analysis equations based upon a conformal mapping's approximation of the dielectric boundary with parallel conductor strips separated by a dielectric sheet.

For wide strips ($W/h > 3.3$) he obtains the approximation

$$Z_L(W, h, \varepsilon_r) = \frac{Z_{F0}}{2\sqrt{\varepsilon_r}} \cdot \frac{1}{\frac{W}{2h} + \frac{1}{\pi} \ln 4 + \frac{\varepsilon_r + 1}{2\pi\varepsilon_r} \ln \left(\frac{\pi e}{2} \left(\frac{W}{2h} + 0.94 \right) \right) + \frac{\varepsilon_r - 1}{2\pi\varepsilon_r^2} \cdot \ln \frac{e\pi^2}{16}} \quad (11.1)$$

For narrow strips ($W/h \leq 3.3$) he obtains the approximation

$$Z_L(W, h, \varepsilon_r) = \frac{Z_{F0}}{\pi\sqrt{2(\varepsilon_r + 1)}} \cdot \left(\ln \left(\frac{4h}{W} + \sqrt{\left(\frac{4h}{W} \right)^2 + 2} \right) - \frac{1}{2} \cdot \frac{\varepsilon_r - 1}{\varepsilon_r + 1} \left(\ln \frac{\pi}{2} + \frac{1}{\varepsilon_r} \ln \frac{4}{\pi} \right) \right) \quad (11.2)$$

The formulae are applicable to alumina-type substrates ($8 \leq \varepsilon_r \leq 12$) and have an estimated relative error less than 1 per cent.

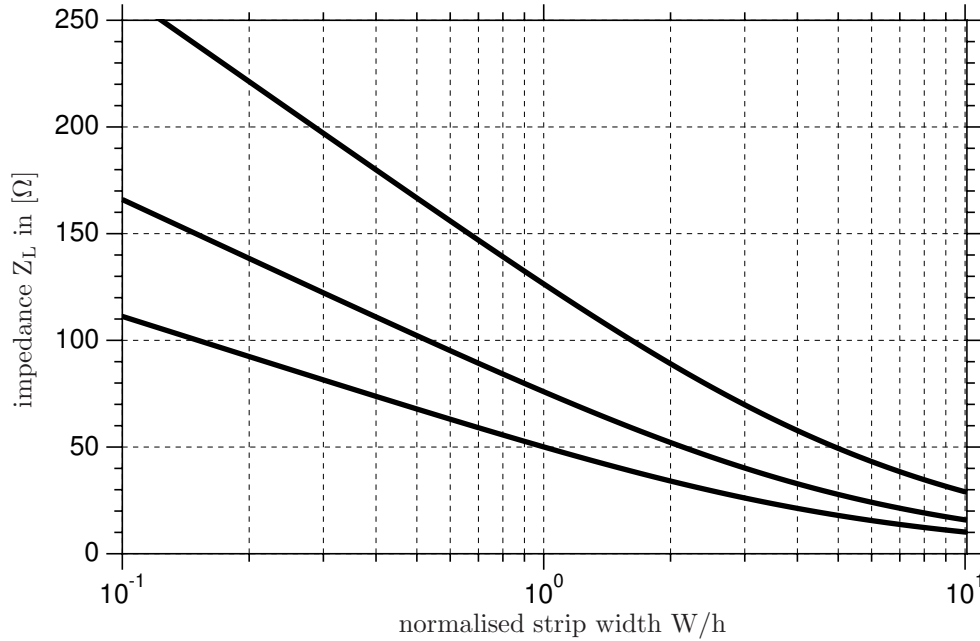


Figure 11.2: characteristic impedance as approximated by Hammerstad for $\varepsilon_r = 1.0$ (air), 3.78 (quartz) and 9.5 (alumina)

Schneider

The following formulas obtained by rational function approximation give accuracy of $\pm 0.25\%$ for $0 \leq W/h \leq 10$ which is the range of importance for most engineering applications. M.V. Schneider [?] found these approximations for the complete elliptic integrals of the first kind as accurate as $\pm 1\%$ for $W/h > 10$.

$$Z_L = \frac{Z_{F0}}{\sqrt{\varepsilon_{eff}}} \cdot \begin{cases} \frac{1}{2\pi} \cdot \ln \left(\frac{8 \cdot h}{W} + \frac{W}{4 \cdot h} \right) & \text{for } \frac{W}{h} \leq 1 \\ \frac{1}{\frac{W}{h} + 2.42 - 0.44 \cdot \frac{h}{W} + \left(1 - \frac{h}{W}\right)^6} & \text{for } \frac{W}{h} > 1 \end{cases} \quad (11.3)$$

Hammerstad and Jensen

The equations for the single microstrip line presented by E. Hammerstad and Ø. Jensen [?] are based upon an equation for the impedance of microstrip in an homogeneous medium and an equation for the microstrip effective dielectric constant. The obtained accuracy gives errors at least less than those caused by physical tolerances and is better than 0.01% for $W/h \leq 1$ and 0.03% for $W/h \leq 1000$.

$$Z_{L1}(W, h) = \frac{Z_{F0}}{2\pi} \cdot \ln \left(f_u \frac{h}{W} + \sqrt{1 + \left(\frac{2h}{W} \right)^2} \right) \quad (11.4)$$

$$Z_L(W, h, \varepsilon_r) = \frac{Z_{L1}(W, h)}{\sqrt{\varepsilon_r}} = \frac{Z_{F0}}{2\pi \cdot \sqrt{\varepsilon_r}} \cdot \ln \left(f_u \frac{h}{W} + \sqrt{1 + \left(\frac{2h}{W} \right)^2} \right) \quad (11.5)$$

with

$$f_u = 6 + (2\pi - 6) \cdot \exp \left(- \left(30.666 \cdot \frac{h}{W} \right)^{0.7528} \right) \quad (11.6)$$

The comparison of the expression given for the quasi-static impedance as shown in fig. ?? has been done with respect to E. Hammerstad and Ø. Jensen. It reveals the advantage of closed-form expressions. The impedance step for Wheelers formulae at $W/h = 3.3$ is approximately 0.1Ω .

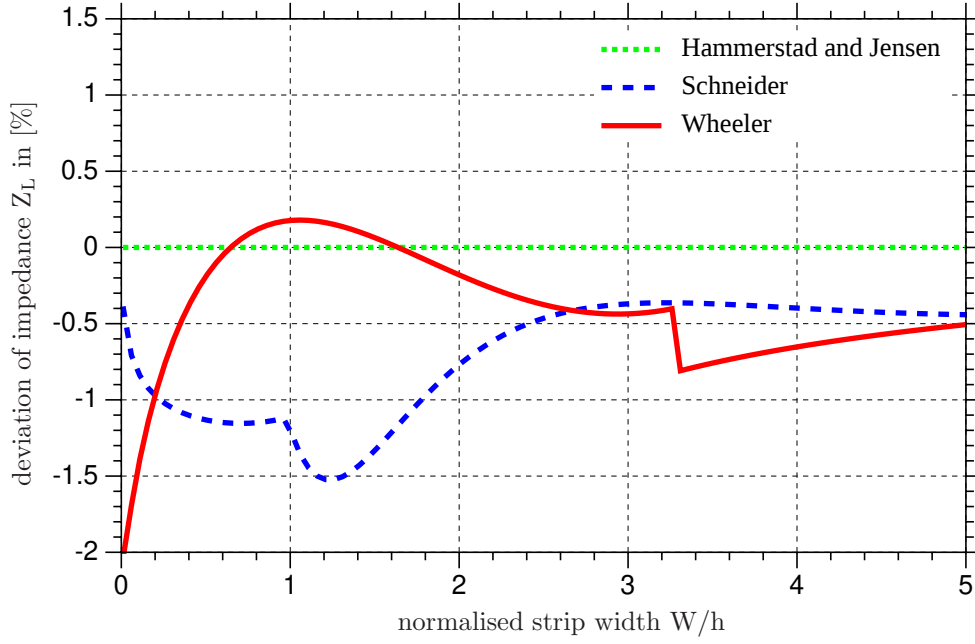


Figure 11.3: characteristic impedance in comparison for $\varepsilon_r = 9.8$

11.1.2 Quasi-static effective dielectric constant

Wheeler

Harold A. Wheeler [?] gives the following approximation for narrow strips ($W/h < 3$) based upon the characteristic impedance Z_L . The estimated relative error is less than 1%.

$$\varepsilon_{r_{eff}} = \frac{\varepsilon_r + 1}{2} + \frac{Z_{F0}}{2\pi Z_L} \cdot \frac{\varepsilon_r - 1}{2} \cdot \left(\ln \frac{\pi}{2} + \frac{1}{\varepsilon_r} \ln \frac{4}{\pi} \right) \quad (11.7)$$

For narrow strips ($W/h \leq 1.3$):

$$\varepsilon_{r_{eff}} = \frac{1 + \varepsilon_r}{2} \cdot \left(\frac{A}{A - B} \right)^2 \quad (11.8)$$

with

$$A = \ln \left(8 \frac{h}{W} \right) + \frac{1}{32} \cdot \left(\frac{W}{h} \right)^2 \quad (11.9)$$

$$B = \frac{1}{2} \cdot \frac{\varepsilon_r - 1}{\varepsilon_r + 1} \cdot \left(\ln \frac{\pi}{2} + \frac{1}{\varepsilon_r} \ln \frac{4}{\pi} \right) \quad (11.10)$$

For wide strips ($W/h > 1.3$):

$$\varepsilon_{r_{eff}} = \varepsilon_r \cdot \left(\frac{E - D}{E} \right)^2 \quad (11.11)$$

with

$$D = \frac{\varepsilon_r - 1}{2\pi\varepsilon_r} \cdot \left(\ln \left(\frac{\pi e}{2} \left(\frac{W}{2h} + 0.94 \right) \right) - \frac{1}{\varepsilon_r} \ln \frac{e\pi^2}{16} \right) \quad (11.12)$$

$$E = \frac{1}{2} \cdot \frac{W}{h} + \frac{1}{\pi} \cdot \ln \left(\pi e \frac{W}{h} + 16.0547 \right) \quad (11.13)$$

Schneider

The approximate function found by M.V. Schneider [?] is meant to have an accuracy of $\pm 2\%$ for $\varepsilon_{r_{eff}}$ and an accuracy of $\pm 1\%$ for $\sqrt{\varepsilon_{r_{eff}}}$.

$$\varepsilon_{r_{eff}} = \frac{\varepsilon_r + 1}{2} + \frac{\varepsilon_r - 1}{2} \cdot \frac{1}{\sqrt{1 + 10 \frac{h}{W}}} \quad (11.14)$$

Hammerstad and Jensen

The accuracy of the E. Hammerstad and Ø. Jensen [?] model is better than 0.2% at least for $\varepsilon_r < 128$ and $0.01 \leq W/h \leq 100$.

$$\varepsilon_{r_{eff}}(W, h, \varepsilon_r) = \frac{\varepsilon_r + 1}{2} + \frac{\varepsilon_r - 1}{2} \cdot \left(1 + 10 \frac{h}{W} \right)^{-ab} \quad (11.15)$$

with

$$a(u) = 1 + \frac{1}{49} \cdot \ln \left(\frac{u^4 + (u/52)^2}{u^4 + 0.432} \right) + \frac{1}{18.7} \cdot \ln \left(1 + \left(\frac{u}{18.1} \right)^3 \right) \quad (11.16)$$

$$b(\varepsilon_r) = 0.564 \cdot \left(\frac{\varepsilon_r - 0.9}{\varepsilon_r + 3} \right)^{0.053} \quad (11.17)$$

$$u = \frac{W}{h} \quad (11.18)$$

11.1.3 Strip thickness correction

The formulas given for the quasi-static characteristic impedance and effective dielectric constant in the previous sections are based upon an infinite thin microstrip line thickness $t = 0$. A finite thickness t can be compensated by a reduction of width. That means a strip with the width W and the finite thickness t appears to be a wider strip.

Wheeler

Harold A. Wheeler [?] proposes the following equation to account for the strip thickness effect based on free space without dielectric.

$$\Delta W_1 = \frac{t}{\pi} \ln \frac{4e}{\sqrt{\left(\frac{t}{h} \right)^2 + \left(\frac{1/\pi}{W/t + 1.10} \right)}} \quad (11.19)$$

For the mixed media case with dielectric he obtains the approximation

$$\Delta W_r = \frac{1}{2} \Delta W_1 \left(1 + \frac{1}{\varepsilon_r} \right) \quad (11.20)$$

Schneider

M.V. Schneider [?] derived the following approximate expressions.

$$\Delta W = \begin{cases} \frac{t}{\pi} \cdot \left(1 + \ln \frac{4 \cdot \pi \cdot W}{t} \right) & \text{for } \frac{W}{h} \leq \frac{1}{2\pi} \\ \frac{t}{\pi} \cdot \left(1 + \ln \frac{2 \cdot h}{t} \right) & \text{for } \frac{W}{h} > \frac{1}{2\pi} \end{cases} \quad (11.21)$$

Additional restrictions for applying these expressions are $t \ll h$, $t < W/2$ and $t/\Delta W < 0.75$. Notice also that the ratio $\Delta W/t$ is divergent for $t \rightarrow 0$.

Hammerstad and Jensen

E. Hammerstad and Ø. Jensen are using the method described by Wheeler [?] to account for a non-zero strip thickness. However, some modifications in his equations have been made, which give better accuracy for narrow strips and for substrates with low dielectric constant. For the homogeneous media the correction is

$$\Delta W_1 = \frac{t}{h \cdot \pi} \ln \left(1 + \frac{4e}{\frac{t}{h} \cdot \coth^2 \sqrt{6.517W}} \right) \quad (11.22)$$

and for the mixed media the correction is

$$\Delta W_r = \frac{1}{2} \Delta W_1 (1 + \operatorname{sech} \sqrt{\varepsilon_r - 1}) \quad (11.23)$$

By defining corrected strip widths, $W_1 = W + \Delta W_1$ and $W_r = W + \Delta W_r$, the effect of strip thickness may be included in the equations (??) and (??).

$$Z_L(W, h, t, \varepsilon_r) = \frac{Z_{L1}(W_r, h)}{\sqrt{\varepsilon_{r_{eff}}(W_r, h, \varepsilon_r)}} \quad (11.24)$$

$$\varepsilon_{r_{eff}}(W, h, t, \varepsilon_r) = \varepsilon_{r_{eff}}(W_r, h, \varepsilon_r) \cdot \left(\frac{Z_{L1}(W_1, h)}{Z_{L1}(W_r, h)} \right)^2 \quad (11.25)$$

11.1.4 Dispersion

Dispersion can be a strong effect in microstrip transmission lines due to their inhomogeneity. Typically, as frequency is increased, $\varepsilon_{r_{eff}}$ increases in a non-linear manner, approaching an asymptotic value. Dispersion affects characteristic impedance in a similar way.

Kirschning and Jansen

The dispersion formulae given by Kirschning and Jansen [?] is meant to have an accuracy better than 0.6% in the range $0.1 \leq W/h \leq 100$, $1 \leq \varepsilon_r \leq 20$ and $0 \leq h/\lambda_0 \leq 0.13$, i.e. up to about 60GHz for 25mm substrates.

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{r_{eff}}}{1 + P(f)} \quad (11.26)$$

with

$$P(f) = P_1 P_2 \cdot ((0.1844 + P_3 P_4) \cdot f_n)^{1.5763} \quad (11.27)$$

$$P_1 = 0.27488 + \left(0.6315 + \frac{0.525}{(1 + 0.0157 \cdot f_n)^{20}} \right) \cdot \frac{W}{h} - 0.065683 \cdot \exp \left(-8.7513 \frac{W}{h} \right) \quad (11.28)$$

$$P_2 = 0.33622 \cdot (1 - \exp(-0.03442 \cdot \varepsilon_r)) \quad (11.29)$$

$$P_3 = 0.0363 \cdot \exp \left(-4.6 \frac{W}{h} \right) \cdot \left(1 - \exp \left(- \left(\frac{f_n}{38.7} \right)^{4.97} \right) \right) \quad (11.30)$$

$$P_4 = 1 + 2.751 \cdot \left(1 - \exp \left(- \left(\frac{\varepsilon_r}{15.916} \right)^8 \right) \right) \quad (11.31)$$

$$f_n = f \cdot h = \text{normalised frequency in } [\text{GHz} \cdot \text{mm}] \quad (11.32)$$

Dispersion of the characteristic impedance according to [?] can be applied for the range $0 \leq h/\lambda_0 \leq 0.1$, $0.1 \leq W/h \leq 10$ and for substrates with $1 \leq \varepsilon_r \leq 18$ and is given by the following set of equations.

$$R_1 = 0.03891 \cdot \varepsilon_r^{1.4} \quad (11.33)$$

$$R_2 = 0.267 \cdot u^{7.0} \quad (11.34)$$

$$R_3 = 4.766 \cdot \exp(-3.228 \cdot u^{0.641}) \quad (11.35)$$

$$R_4 = 0.016 + (0.0514 \cdot \varepsilon_r)^{4.524} \quad (11.36)$$

$$R_5 = (f_n/28.843)^{12.0} \quad (11.37)$$

$$R_6 = 22.20 \cdot u^{1.92} \quad (11.38)$$

and

$$R_7 = 1.206 - 0.3144 \cdot \exp(-R_1) \cdot (1 - \exp(-R_2)) \quad (11.39)$$

$$R_8 = 1 + 1.275 \cdot \left(1 - \exp(-0.004625 \cdot R_3 \cdot \varepsilon_r^{1.674}) \cdot (f_n/18.365)^{2.745} \right) \quad (11.40)$$

$$R_9 = 5.086 \cdot \frac{R_4 \cdot R_5}{0.3838 + 0.386 \cdot R_4} \cdot \frac{\exp(-R_6)}{1 + 1.2992 \cdot R_5} \cdot \frac{(\varepsilon_r - 1)^6}{1 + 10 \cdot (\varepsilon_r - 1)^6} \quad (11.41)$$

and

$$R_{10} = 0.00044 \cdot \varepsilon_r^{2.136} + 0.0184 \quad (11.42)$$

$$R_{11} = \frac{(f_n/19.47)^6}{1 + 0.0962 \cdot (f_n/19.47)^6} \quad (11.43)$$

$$R_{12} = \frac{1}{1 + 0.00245 \cdot u^2} \quad (11.44)$$

$$R_{13} = 0.9408 \cdot \varepsilon_r(f)^{R_8} - 0.9603 \quad (11.45)$$

$$R_{14} = (0.9408 - R_9) \cdot \varepsilon_{r_{eff}}^{R_8} - 0.9603 \quad (11.46)$$

$$R_{15} = 0.707 \cdot R_{10} \cdot (f_n/12.3)^{1.097} \quad (11.47)$$

$$R_{16} = 1 + 0.0503 \cdot \varepsilon_r^2 \cdot R_{11} \cdot \left(1 - \exp \left(- (u/15)^6 \right) \right) \quad (11.48)$$

$$R_{17} = R_7 \cdot \left(1 - 1.1241 \cdot \frac{R_{12}}{R_{16}} \cdot \exp(-0.026 \cdot f_n^{1.15656} - R_{15}) \right) \quad (11.49)$$

Finally the frequency-dependent characteristic impedance can be written as

$$Z_L(f_n) = Z_L(0) \cdot \left(\frac{R_{13}}{R_{14}} \right)^{R_{17}} \quad (11.50)$$

The abbreviations used in these expressions are f_n for the normalized frequency as denoted in eq. (??) and $u = W/h$ for the microstrip width normalised with respect to the substrate height. The terms $Z_L(0)$ and $\varepsilon_{r_{eff}}$ denote the static values of microstrip characteristic impedance and effective dielectric constant. The value $\varepsilon_r(f)$ is the frequency dependent effective dielectric constant computed according to [?].

R.H. Jansen and M. Kirschning remark in [?] for the implementation of the expressions on a computer, R_1 , R_2 and R_6 should be restricted to numerical values less or equal 20 in order to prevent overflow.

Yamashita

The values obtained by the approximate dispersion formula as given by E. Yamashita [?] deviate within 1% in a wide frequency range compared to the integral equation method used to derive the functional approximation. The formula assumes the knowledge of the quasi-static effective dielectric constant. The applicable ranges of the formula are $2 < \varepsilon_r < 16$, $0.06 < W/h < 16$ and $0.1\text{GHz} < f < 100\text{GHz}$. Though the lowest usable frequency is limited by 0.1GHz, the propagation constant for frequencies less than 0.1GHz has been given as the quasi-static one.

$$\varepsilon_r(f) = \varepsilon_{r_{eff}} \cdot \left(\frac{1 + \frac{1}{4} \cdot k \cdot F^{1.5}}{1 + \frac{1}{4} \cdot F^{1.5}} \right)^2 \quad (11.51)$$

with

$$k = \sqrt{\frac{\varepsilon_r}{\varepsilon_{r_{eff}}}} \quad (11.52)$$

$$F = \frac{4 \cdot h \cdot f \cdot \sqrt{\varepsilon_r - 1}}{c_0} \cdot \left(0.5 + \left(1 + 2 \cdot \log \left(1 + \frac{W}{h} \right) \right)^2 \right) \quad (11.53)$$

Kobayashi

The dispersion formula presented by M. Kobayashi [?], derived by comparison to a numerical model, has a high degree of accuracy, better than 0.6% in the range $0.1 \leq W/h \leq 10$, $1 < \varepsilon_r \leq 128$ and any h/λ_0 (no frequency limits).

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{r_{eff}}}{1 + \left(\frac{f}{f_{50}} \right)^m} \quad (11.54)$$

with

$$f_{50} = \frac{c_0}{2\pi \cdot h \cdot \left(0.75 + \left(0.75 - \frac{0.332}{\varepsilon_r^{1.73}}\right) \frac{W}{h}\right)} \cdot \frac{\arctan\left(\varepsilon_r \cdot \sqrt{\frac{\varepsilon_{reff} - 1}{\varepsilon_r - \varepsilon_{reff}}}\right)}{\sqrt{\varepsilon_r - \varepsilon_{reff}}} \quad (11.55)$$

$$m = m_0 \cdot m_c \quad (\leq 2.32) \quad (11.56)$$

$$m_0 = 1 + \frac{1}{1 + \sqrt{\frac{W}{h}}} + 0.32 \cdot \left(\frac{1}{1 + \sqrt{\frac{W}{h}}}\right)^3 \quad (11.57)$$

$$m_c = \begin{cases} 1 + \frac{1.4}{1 + \frac{W}{h}} \cdot \left(0.15 - 0.235 \cdot \exp\left(-0.45 \frac{f}{f_{50}}\right)\right) & \text{for } W/h \leq 0.7 \\ 1 & \text{for } W/h \geq 0.7 \end{cases} \quad (11.58)$$

Getsinger

Based upon measurements of dispersion curves for microstrip lines on alumina substrates 0.025 and 0.050 inch thick W. J. Getsinger [?] developed a very simple , closed-form expression that allow slide-rule prediction of microstrip dispersion.

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{reff}}{1 + G \cdot \left(\frac{f}{f_p}\right)^2} \quad (11.59)$$

with

$$f_p = \frac{Z_L}{2\mu_0 h} \quad (11.60)$$

$$G = 0.6 + 0.009 \cdot Z_L \quad (11.61)$$

Also based upon measurements of microstrip lines 0.1, 0.25 and 0.5 inch in width on a 0.250 inch thick alumina substrate Getsinger [?] developed two different dispersion models for the characteristic impedance.

- wave impedance model published in [?]

$$Z_L(f) = Z_L \cdot \sqrt{\frac{\varepsilon_{reff}}{\varepsilon_r(f)}} \quad (11.62)$$

- group-delay model published in [?]

$$Z_L(f) = Z_L \cdot \sqrt{\frac{\varepsilon_r(f)}{\varepsilon_{reff}}} \cdot \frac{1}{1 + D(f)} \quad (11.63)$$

with

$$D(f) = \frac{(\varepsilon_r - \varepsilon_r(f)) \cdot (\varepsilon_r(f) - \varepsilon_{reff})}{\varepsilon_r(f) \cdot (\varepsilon_r - \varepsilon_{reff})} \quad (11.64)$$

Hammerstad and Jensen

The dispersion formulae of E. Hammerstad and Ø. Jensen [?] give good results for all types of substrates (not as limited as Getsinger's formulae). The impedance dispersion model is based upon a parallel-plate model using the theory of dielectrics.

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{reff}}{1 + G \cdot \left(\frac{f}{f_p}\right)^2} \quad (11.65)$$

with

$$f_p = \frac{Z_L}{2\mu_0 h} \quad (11.66)$$

$$G = \frac{\pi^2}{12} \cdot \frac{\varepsilon_r - 1}{\varepsilon_{reff}} \cdot \sqrt{\frac{2\pi \cdot Z_L}{Z_{F0}}} \quad (11.67)$$

$$Z_L(f) = Z_L \cdot \sqrt{\frac{\varepsilon_{reff}}{\varepsilon_r(f)}} \cdot \frac{\varepsilon_r(f) - 1}{\varepsilon_{reff} - 1} \quad (11.68)$$

Edwards and Owens

The authors T. C. Edwards and R. P. Owens [?] developed a dispersion formula based upon measurements of microstrip lines on sapphire in the range $10\Omega \leq Z_L \leq 100\Omega$ and up to 18GHz. The procedure was repeated for several microstrip width-to-substrate-height ratios (W/h) between 0.1 and 10.

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{reff}}{1 + P} \quad (11.69)$$

with

$$P = \left(\frac{h}{Z_L}\right)^{1.33} \cdot (0.43f^2 - 0.009f^3) \quad (11.70)$$

where h is in millimeters and f is in gigahertz. Their new dispersion equation involving the polynomial, which was developed to predict the fine detail of the experimental $\varepsilon_r(f)$ versus frequency curves, includes two empirical parameters. However, it seems the formula is not too sensitive to changes in substrate parameters.

Pramanick and Bhartia

P. Bhartia and P. Pramanick [?] developed dispersion equations without any empirical quantity. Their work expresses dispersion of the dielectric constant and characteristic impedance in terms of a single inflection frequency.

For the frequency-dependent relative dielectric constant they propose

$$\varepsilon_r(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{reff}}{1 + \left(\frac{f}{f_T}\right)^2} \quad (11.71)$$

where

$$f_T = \sqrt{\frac{\varepsilon_r}{\varepsilon_{reff}}} \cdot \frac{Z_L}{2\mu_0 h} \quad (11.72)$$

Dispersion of the characteristic impedance is accounted by

$$Z_L(f) = \frac{Z_{F0} \cdot h}{W_e(f) \cdot \sqrt{\varepsilon_r(f)}} \quad (11.73)$$

whence

$$W_e(f) = W + \frac{W_{eff} - W}{1 + \left(\frac{f}{f_T}\right)^2} \quad \text{and} \quad W_{eff} = \frac{Z_{F0} \cdot h}{Z_L \cdot \sqrt{\varepsilon_{reff}}} \quad (11.74)$$

Schneider

Martin V. Schneider [?] proposed the following equation for the dispersion of the effective dielectric constant of a single microstrip line. The estimated error is less than 3%.

$$\varepsilon_r(f) = \varepsilon_{reff} \cdot \left(\frac{1 + f_n^2}{1 + k \cdot f_n^2} \right)^2 \quad (11.75)$$

with

$$f_n = \frac{4h \cdot f \cdot \sqrt{\varepsilon_r - 1}}{c_0} \quad \text{and} \quad k = \sqrt{\frac{\varepsilon_{reff}}{\varepsilon_r}} \quad (11.76)$$

For the dispersion of the characteristic impedance he uses the same wave guide impedance model as Getsinger in his first approach to the problem.

$$Z_L(f) = Z_L \cdot \sqrt{\frac{\varepsilon_{reff}}{\varepsilon_r(f)}} \quad (11.77)$$

11.1.5 Transmission losses

The attenuation of a microstrip line consists of conductor (ohmic) losses, dielectric (substrate) losses, losses due to radiation and propagation of surface waves and higher order modes.

$$\alpha = \alpha_c + \alpha_d + \alpha_r + \alpha_s \quad (11.78)$$

Dielectric losses

Dielectric loss is due to the effects of finite loss tangent $\tan \delta_d$. Basically the losses rise proportional over the operating frequency. For common microwave substrate materials like Al_2O_3 ceramics with a loss tangent δ_d less than 10^{-3} the dielectric losses can be neglected compared to the conductor losses.

For the inhomogeneous line, an effective dielectric filling fraction give that proportion of the transmission line's cross section not filled by air. For microstrip lines, the result is

$$\alpha_d = \frac{\varepsilon_r}{\sqrt{\varepsilon_{reff}}} \cdot \frac{\varepsilon_{reff} - 1}{\varepsilon_r - 1} \cdot \frac{\pi}{\lambda_0} \cdot \tan \delta_d \quad (11.79)$$

whereas

δ_d dielectric loss tangent

Conductor losses

E. Hammerstad and Ø. Jensen [?] proposed the following equation for the conductor losses. The surface roughness of the substrate is necessary to account for an asymptotic increase seen in the apparent surface resistance with decreasing skin depth. This effect is considered by the correction factor K_r . The current distribution factor K_i is a very good approximation provided that the strip thickness exceeds three skin depths ($t > 3\delta$).

$$\alpha_c = \frac{R^\square}{Z_L \cdot W} \cdot K_r \cdot K_i \quad (11.80)$$

with

$$R^\square = \frac{\rho}{\delta} = \sqrt{\rho \cdot \frac{\omega \cdot \mu}{2}} = \sqrt{\rho \cdot \pi \cdot f \cdot \mu} \quad (11.81)$$

$$K_i = \exp \left(-1.2 \left(\frac{Z_L}{Z_{F0}} \right)^{0.7} \right) \quad (11.82)$$

$$K_r = 1 + \frac{2}{\pi} \arctan \left(1.4 \left(\frac{\Delta}{\delta} \right)^2 \right) \quad (11.83)$$

whereas

- $R^\square$ sheet resistance of conductor material (skin resistance)
- ρ specific resistance of conductor
- δ skin depth
- K_i current distribution factor
- K_r correction term due to surface roughness
- Δ effective (rms) surface roughness of substrate
- Z_{F0} wave impedance in vacuum

11.2 Microstrip corner

The equivalent circuit of a microstrip corner is shown in fig. ???. The values of the components are as follows [?].

$$C \text{ [pF]} = W \cdot \left((10.35 \cdot \varepsilon_r + 2.5) \cdot \frac{W}{h} + (2.6 \cdot \varepsilon_r + 5.64) \right) \quad (11.84)$$

$$L \text{ [nH]} = 220 \cdot h \cdot \left(1 - 1.35 \cdot \exp \left(-0.18 \cdot \left(\frac{W}{h} \right)^{1.39} \right) \right) \quad (11.85)$$

The values for a 50% mitered bend are [?].

$$C \text{ [pF]} = W \cdot \left((3.93 \cdot \varepsilon_r + 0.62) \cdot \frac{W}{h} + (7.6 \cdot \varepsilon_r + 3.80) \right) \quad (11.86)$$

$$L \text{ [nH]} = 440 \cdot h \cdot \left(1 - 1.062 \cdot \exp \left(-0.177 \cdot \left(\frac{W}{h} \right)^{0.947} \right) \right) \quad (11.87)$$

With W being width of the microstrip line and h height of the substrate. These formulas are valid for $W/h = 0.2$ to 6.0 and for $\varepsilon_r = 2.36$ to 10.4 and up to 14 GHz. The precision is approximately 0.3% .

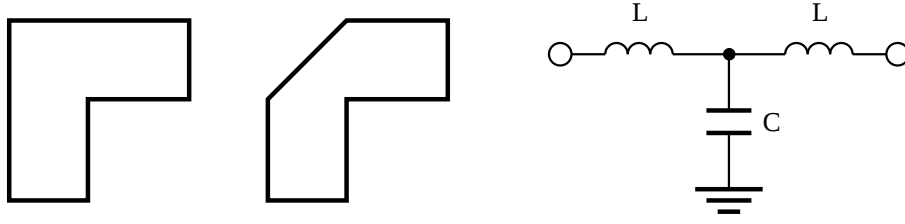


Figure 11.4: microstrip corner (left), mitered corner (middle) and equivalent circuit (right)

The Z-parameters for the given equivalent small signal circuit can be written as stated in eq. (11.88) and are easy to convert to scattering parameters.

$$Z = \begin{bmatrix} j\omega L + \frac{1}{j\omega C} & \frac{1}{j\omega C} \\ \frac{1}{j\omega C} & j\omega L + \frac{1}{j\omega C} \end{bmatrix} \quad (11.88)$$

11.3 Parallel coupled microstrip lines

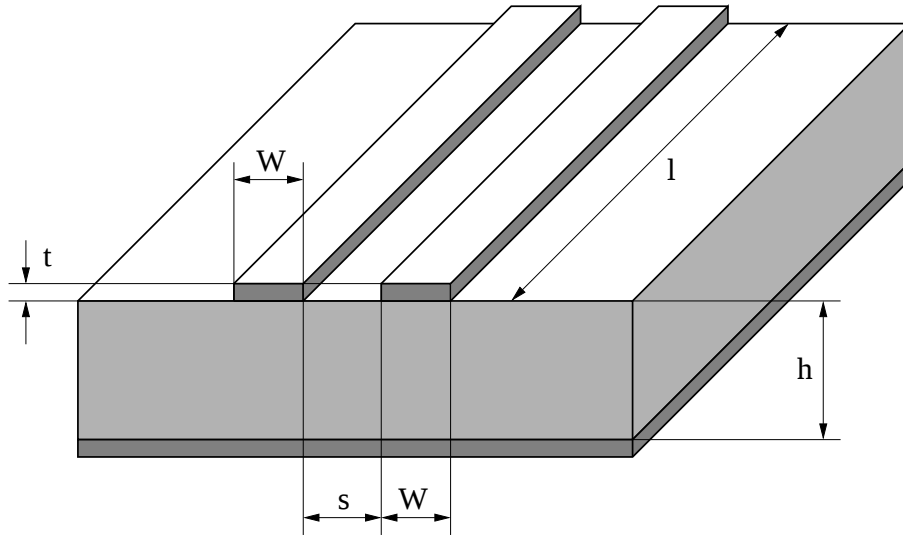


Figure 11.5: parallel coupled microstrip lines

11.3.1 Characteristic impedance and effective dielectric constant

Parallel coupled microstrip lines are defined by the characteristic impedance and the effective permittivity of the even and the odd mode. The y- and S-parameters are depicted in section ??.

Kirschning and Jansen

These quantities can very precisely be modeled by the following equations [?], [?].

Beforehand some normalised quantities (with microstrip line width W , spacing s between the lines and substrate height h) are introduced:

$$u = \frac{W}{h} \quad , \quad g = \frac{s}{h} \quad , \quad f_n = \frac{f}{\text{GHz}} \cdot \frac{h}{\text{mm}} = \frac{f}{\text{MHz}} \cdot h \quad (11.89)$$

The applicability of the described model is

$$0.1 \leq u \leq 10 \quad , \quad 0.1 \leq g \leq 10 \quad , \quad 1 \leq \varepsilon_r \leq 18 \quad (11.90)$$

The accuracies of the formulas holds for these ranges.

Static effective permittivity of even mode:

$$\varepsilon_{eff,e}(0) = 0.5 \cdot (\varepsilon_r + 1) + 0.5 \cdot (\varepsilon_r - 1) \cdot \left(1 + \frac{10}{v}\right)^{-a_e(v) \cdot b_e(\varepsilon_r)} \quad (11.91)$$

with

$$v = u \cdot \frac{20 + g^2}{10 + g^2} + g \cdot \exp(-g) \quad (11.92)$$

$$a_e(v) = 1 + \frac{1}{49} \cdot \ln \left(\frac{v^4 + (v/52)^2}{v^4 + 0.432} \right) + \frac{1}{18.7} \cdot \ln \left(1 + \left(\frac{v}{18.1} \right)^3 \right) \quad (11.93)$$

$$b_e(\varepsilon_r) = 0.564 \cdot \left(\frac{\varepsilon_r - 0.9}{\varepsilon_r + 3} \right)^{0.053} \quad (11.94)$$

Static effective permittivity of odd mode:

$$\varepsilon_{eff,o}(0) = (0.5 \cdot (\varepsilon_r + 1) + a_o(u, \varepsilon_r) - \varepsilon_{eff}(0)) \cdot \exp(-c_o \cdot g^{d_o}) + \varepsilon_{eff}(0) \quad (11.95)$$

with

$$a_o(u, \varepsilon_r) = 0.7287 \cdot (\varepsilon_{eff}(0) - 0.5 \cdot (\varepsilon_r + 1)) \cdot (1 - \exp(-0.179 \cdot u)) \quad (11.96)$$

$$b_o(\varepsilon_r) = 0.747 \cdot \frac{\varepsilon_r}{0.15 + \varepsilon_r} \quad (11.97)$$

$$c_o = b_o(\varepsilon_r) - (b_o(\varepsilon_r) - 0.207) \cdot \exp(-0.414 \cdot u) \quad (11.98)$$

$$d_o = 0.593 + 0.694 \cdot \exp(-0.562 \cdot u) \quad (11.99)$$

whence $\varepsilon_{eff}(0)$ refers to the zero-thickness single microstrip line of width W according to [?] (see also eq. (??)).

The dispersion formulae for the odd and even mode write as follows.

$$\varepsilon_{eff,e,o}(f_n) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{eff,e,o}(0)}{1 + F_{e,o}(f_n)} \quad (11.100)$$

The frequency dependence for the even mode is

$$F_e(f_n) = P_1 \cdot P_2 \cdot ((P_3 \cdot P_4 + 0.1844 \cdot P_7) \cdot f_n)^{1.5763} \quad (11.101)$$

with

$$P_1 = 0.27488 + \left(0.6315 + \frac{0.525}{(1 + 0.0157 \cdot f_n)^{20}} \right) \cdot u - 0.065683 \cdot \exp(-8.7513 \cdot u) \quad (11.102)$$

$$P_2 = 0.33622 \cdot (1 - \exp(-0.03442 \cdot \varepsilon_r)) \quad (11.103)$$

$$P_3 = 0.0363 \cdot \exp(-4.6 \cdot u) \cdot \left(1 - \exp\left(- (f_n/38.7)^{4.97}\right) \right) \quad (11.104)$$

$$P_4 = 1 + 2.751 \cdot \left(1 - \exp\left(- (\varepsilon_r/15.916)^8\right) \right) \quad (11.105)$$

$$P_5 = 0.334 \cdot \exp\left(-3.3 \cdot (\varepsilon_r/15)^3\right) + 0.746 \quad (11.106)$$

$$P_6 = P_5 \cdot \exp\left(- (f_n/18)^{0.368}\right) \quad (11.107)$$

$$P_7 = 1 + 4.069 \cdot P_6 \cdot g^{0.479} \cdot \exp\left(-1.347 \cdot g^{0.595} - 0.17 \cdot g^{2.5}\right) \quad (11.108)$$

The frequency dependence for the odd mode is

$$F_o(f_n) = P_1 \cdot P_2 \cdot ((P_3 \cdot P_4 + 0.1844) \cdot f_n \cdot P_{15})^{1.5763} \quad (11.109)$$

with

$$P_8 = 0.7168 \cdot \left(1 + \frac{1.076}{1 + 0.0576 \cdot (\varepsilon_r - 1)} \right) \quad (11.110)$$

$$P_9 = P_8 - 0.7913 \cdot \left(1 - \exp\left(- (f_n/20)^{1.424}\right) \right) \cdot \arctan\left(2.481 \cdot (\varepsilon_r/8)^{0.946}\right) \quad (11.111)$$

$$P_{10} = 0.242 \cdot (\varepsilon_r - 1)^{0.55} \quad (11.112)$$

$$P_{11} = 0.6366 \cdot (\exp(-0.3401 \cdot f_n) - 1) \cdot \arctan\left(1.263 \cdot (u/3)^{1.629}\right) \quad (11.113)$$

$$P_{12} = P_9 + \frac{1 - P_9}{1 + 1.183 \cdot u^{1.376}} \quad (11.114)$$

$$P_{13} = 1.695 \cdot \frac{P_{10}}{0.414 + 1.605 \cdot P_{10}} \quad (11.115)$$

$$P_{14} = 0.8928 + 0.1072 \cdot \left(1 - \exp\left(-0.42 \cdot (f_n/20)^{3.215}\right) \right) \quad (11.116)$$

$$P_{15} = |1 - 0.8928 \cdot (1 + P_{11}) \cdot \exp(-P_{13} \cdot g^{1.092}) \cdot P_{12}/P_{14}| \quad (11.117)$$

Up to $f_n = 25$ the maximum error of these equations is 1.4%.

The static characteristic impedance for the even mode writes as follows.

$$Z_{L,e}(0) = \sqrt{\frac{\varepsilon_{eff}(0)}{\varepsilon_{eff,e}(0)}} \cdot \frac{Z_L(0)}{1 - \frac{Z_L(0)}{377\Omega} \cdot \sqrt{\varepsilon_{eff}(0)} \cdot Q_4} \quad (11.118)$$

with

$$Q_1 = 0.8695 \cdot u^{0.194} \quad (11.119)$$

$$Q_2 = 1 + 0.7519 \cdot g + 0.189 \cdot g^{2.31} \quad (11.120)$$

$$Q_3 = 0.1975 + \left(16.6 + (8.4/g)^6 \right)^{-0.387} + \frac{1}{241} \cdot \ln\left(\frac{g^{10}}{1 + (g/3.4)^{10}}\right) \quad (11.121)$$

$$Q_4 = \frac{Q_1}{Q_2} \cdot \frac{2}{\exp(-g) \cdot u^{Q_3} + (2 - \exp(-g)) \cdot u^{-Q_3}} \quad (11.122)$$

with $Z_L(0)$ and $\varepsilon_{eff}(0)$ being again quantities for a zero-thickness single microstrip line of width W according to [?] (see also eq. (??) and (??)).

The static characteristic impedance for the odd mode writes as follows.

$$Z_{L,o}(0) = \sqrt{\frac{\varepsilon_{eff}(0)}{\varepsilon_{eff,o}(0)}} \cdot \frac{Z_L(0)}{1 - \frac{Z_L(0)}{377\Omega} \cdot \sqrt{\varepsilon_{eff}(0)} \cdot Q_{10}} \quad (11.123)$$

with

$$Q_5 = 1.794 + 1.14 \cdot \ln \left(1 + \frac{0.638}{g + 0.517 \cdot g^{2.43}} \right) \quad (11.124)$$

$$Q_6 = 0.2305 + \frac{1}{281.3} \cdot \ln \left(\frac{g^{10}}{1 + (g/5.8)^{10}} \right) + \frac{1}{5.1} \cdot \ln (1 + 0.598 \cdot g^{1.154}) \quad (11.125)$$

$$Q_7 = \frac{10 + 190 \cdot g^2}{1 + 82.3 \cdot g^3} \quad (11.126)$$

$$Q_8 = \exp \left(-6.5 - 0.95 \cdot \ln(g) - (g/0.15)^5 \right) \quad (11.127)$$

$$Q_9 = \ln(Q_7) \cdot (Q_8 + 1/16.5) \quad (11.128)$$

$$Q_{10} = \frac{Q_2 \cdot Q_4 - Q_5 \cdot \exp(\ln(u) \cdot Q_6 \cdot u^{-Q_9})}{Q_2} = Q_4 - \frac{Q_5}{Q_2} \cdot u^{Q_6 \cdot u^{-Q_9}} \quad (11.129)$$

The accuracy of the static impedances is better than 0.6%.

Dispersion of the characteristic impedance for the even mode can be modeled by the following equations.

$$Z_{L,e}(f_n) = Z_{L,e}(0) \cdot \left(\frac{0.9408 \cdot (\varepsilon_{eff}(f_n))^{C_e} - 0.9603}{(0.9408 - d_e) \cdot (\varepsilon_{eff}(0))^{C_e} - 0.9603} \right)^{Q_0} \quad (11.130)$$

with

$$C_e = 1 + 1.275 \cdot \left(1 - \exp \left(-0.004625 \cdot p_e \cdot \varepsilon_r^{1.674} \cdot (f_n/18.365)^{2.745} \right) \right) - Q_{12} + Q_{16} - Q_{17} + Q_{18} + Q_{20} \quad (11.131)$$

$$d_e = 5.086 \cdot q_e \cdot \frac{r_e}{0.3838 + 0.386 \cdot q_e} \cdot \frac{\exp(-22.2 \cdot u^{1.92})}{1 + 1.2992 \cdot r_e} \cdot \frac{(\varepsilon_r - 1)^6}{1 + 10 \cdot (\varepsilon_r - 1)^6} \quad (11.132)$$

$$p_e = 4.766 \cdot \exp(-3.228 \cdot u^{0.641}) \quad (11.133)$$

$$q_e = 0.016 + (0.0514 \cdot \varepsilon_r \cdot Q_{21})^{4.524} \quad (11.134)$$

$$r_e = (f_n/28.843)^{12} \quad (11.135)$$

and

$$Q_{11} = 0.893 \cdot \left(1 - \frac{0.3}{1 + 0.7 \cdot (\varepsilon_r - 1)} \right) \quad (11.136)$$

$$Q_{12} = 2.121 \cdot \frac{(f_n/20)^{4.91}}{1 + Q_{11} \cdot (f_n/20)^{4.91}} \cdot \exp(-2.87 \cdot g) \cdot g^{0.902} \quad (11.137)$$

$$Q_{13} = 1 + 0.038 \cdot (\varepsilon_r/8)^{5.1} \quad (11.138)$$

$$Q_{14} = 1 + 1.203 \cdot \frac{(\varepsilon_r/15)^4}{1 + (\varepsilon_r/15)^4} \quad (11.139)$$

$$Q_{15} = \frac{1.887 \cdot \exp(-1.5 \cdot g^{0.84}) \cdot g^{Q_{14}}}{1 + 0.41 \cdot (f_n/15)^3 \cdot \frac{u^{2/Q_{13}}}{0.125 + u^{1.626/Q_{13}}}} \quad (11.140)$$

$$Q_{16} = Q_{15} \cdot \left(1 + \frac{9}{1 + 0.403 \cdot (\varepsilon_r - 1)^2} \right) \quad (11.141)$$

$$Q_{17} = 0.394 \cdot \left(1 - \exp(-1.47 \cdot (u/7)^{0.672}) \right) \cdot \left(1 - \exp(-4.25 \cdot (f_n/20)^{1.87}) \right) \quad (11.142)$$

$$Q_{18} = 0.61 \cdot \frac{1 - \exp(-2.13 \cdot (u/8)^{1.593})}{1 + 6.544 \cdot g^{4.17}} \quad (11.143)$$

$$Q_{19} = \frac{0.21 \cdot g^4}{(1 + 0.18 \cdot g^{4.9}) \cdot (1 + 0.1 \cdot u^2) \cdot (1 + (f_n/24)^3)} \quad (11.144)$$

$$Q_{20} = Q_{19} \cdot \left(0.09 + \frac{1}{1 + 0.1 \cdot (\varepsilon_r - 1)^{2.7}} \right) \quad (11.145)$$

$$Q_{21} = \left| 1 - 42.54 \cdot g^{0.133} \cdot \exp(-0.812 \cdot g) \cdot \frac{u^{2.5}}{1 + 0.033 \cdot u^{2.5}} \right| \quad (11.146)$$

With $\varepsilon_{eff}(f_n)$ being the single microstrip effective dielectric constant according to [?] (see eq. (??)) and Q_0 single microstrip impedance dispersion according to [?] (there denoted as R_{17} , see eq. (??)).

Dispersion of the characteristic impedance for the odd mode can be modeled by the following equations.

$$Z_{L,o}(f_n) = Z_L(f_n) + \frac{Z_{L,o}(0) \cdot \left(\frac{\varepsilon_{eff,o}(f_n)}{\varepsilon_{eff,o}(0)} \right)^{Q_{22}} - Z_L(f_n) \cdot Q_{23}}{1 + Q_{24} + (0.46 \cdot g)^{2.2} \cdot Q_{25}} \quad (11.147)$$

with

$$Q_{22} = 0.925 \cdot \frac{(f_n/Q_{26})^{1.536}}{1 + 0.3 \cdot (f_n/30)^{1.536}} \quad (11.148)$$

$$Q_{23} = 1 + \frac{0.005 \cdot f_n \cdot Q_{27}}{\left(1 + 0.812 \cdot (f_n/15)^{1.9}\right) \cdot (1 + 0.025 \cdot u^2)} \quad (11.149)$$

$$Q_{24} = \frac{2.506 \cdot Q_{28} \cdot u^{0.894}}{3.575 + u^{0.894}} \cdot \left(\frac{(1 + 1.3 \cdot u) \cdot f_n}{99.25}\right)^{4.29} \quad (11.150)$$

$$Q_{25} = \frac{0.3 \cdot f_n^2}{10 + f_n^2} \cdot \left(1 + \frac{2.333 \cdot (\varepsilon_r - 1)^2}{5 + (\varepsilon_r - 1)^2}\right) \quad (11.151)$$

$$Q_{26} = 30 - \frac{22.2 \cdot \left(\frac{\varepsilon_r - 1}{13}\right)^{12}}{1 + 3 \cdot \left(\frac{\varepsilon_r - 1}{13}\right)^{12}} - Q_{29} \quad (11.152)$$

$$Q_{27} = 0.4 \cdot g^{0.84} \cdot \left(1 + \frac{2.5 \cdot (\varepsilon_r - 1)^{1.5}}{5 + (\varepsilon_r - 1)^{1.5}}\right) \quad (11.153)$$

$$Q_{28} = 0.149 \cdot \frac{(\varepsilon_r - 1)^3}{94.5 + 0.038 \cdot (\varepsilon_r - 1)^3} \quad (11.154)$$

$$Q_{29} = \frac{15.16}{1 + 0.196 \cdot (\varepsilon_r - 1)^2} \quad (11.155)$$

with $Z_L(f_n)$ being the frequency-dependent power-current characteristic impedance formulation of a single microstrip with width W according to [?] (see eq. (??)). Up to $f_n = 20$, the numerical error of $Z_{L,o}(f_n)$ and $Z_{L,e}(f_n)$ is less than 2.5%.

Hammerstad and Jensen

The equations given by E. Hammerstad and Ø. Jensen [?] represent the first generally valid model of coupled microstrips with an acceptable accuracy. The model equations have been validated in the range $0.1 \leq u \leq 10$ and $g \geq 0.01$, a range which should cover that used in practice.

The homogeneous mode impedances are

$$Z_{L,e,o}(u, g) = \frac{Z_L(u)}{1 - Z_L(u) \cdot \Phi_{e,o}(u, g) / Z_{F0}} \quad (11.156)$$

The effective dielectric constants are

$$\varepsilon_{eff,e,o}(u, g, \varepsilon_r) = \frac{\varepsilon_r + 1}{2} + \frac{\varepsilon_r - 1}{2} \cdot F_{e,o}(u, g, \varepsilon_r) \quad (11.157)$$

with

$$F_e(u, g, \varepsilon_r) = \left(1 + \frac{10}{\mu(u, g)}\right)^{-a(\mu) \cdot b(\varepsilon_r)} \quad (11.158)$$

$$F_o(u, g, \varepsilon_r) = f_o(u, g, \varepsilon_r) \cdot \left(1 + \frac{10}{u}\right)^{-a(u) \cdot b(\varepsilon_r)} \quad (11.159)$$

whence $a(u)$ and $b(\varepsilon_r)$ denote eqs. (??) and (??) of the single microstrip line. The characteristic impedance of the single microstrip line $Z_L(u)$ also defined in [?] is given by eq. (??). The modifying equations for the even mode are as follows

$$\Phi_e(u, g) = \frac{\varphi(u)}{\Psi(g) \cdot (\alpha(g) \cdot u^{m(g)} + (1 - \alpha(g)) \cdot u^{-m(g)})} \quad (11.160)$$

$$\varphi(u) = 0.8645 \cdot u^{0.172} \quad (11.161)$$

$$\Psi(g) = 1 + \frac{g}{1.45} + \frac{g^{2.09}}{3.95} \quad (11.162)$$

$$\alpha(g) = 0.5 \cdot e^{-g} \quad (11.163)$$

$$m(g) = 0.2175 + \left(4.113 + (20.36/g)^6\right)^{-0.251} + \frac{1}{323} \cdot \ln \left(\frac{g^{10}}{1 + (g/13.8)^{10}} \right) \quad (11.164)$$

The modifying equations for the odd mode are as follows

$$\Phi_o(u, g) = \Phi_e(u, g) - \frac{\theta(g)}{\Psi(g)} \cdot \exp \left(\beta(g) \cdot u^{-n(g)} \cdot \ln u \right) \quad (11.165)$$

$$\theta(g) = 1.729 + 1.175 \cdot \ln \left(1 + \frac{0.627}{g + 0.327 \cdot g^{2.17}} \right) \quad (11.166)$$

$$\beta(g) = 0.2306 + \frac{1}{301.8} \cdot \ln \left(\frac{g^{10}}{1 + (g/3.73)^{10}} \right) + \frac{1}{5.3} \cdot \ln \left(1 + 0.646 \cdot g^{1.175} \right) \quad (11.167)$$

$$n(g) = \left(\frac{1}{17.7} + \exp \left(-6.424 - 0.76 \cdot \ln g - (g/0.23)^5 \right) \right) \cdot \ln \left(\frac{10 + 68.3 \cdot g^2}{1 + 32.5 \cdot g^{3.093}} \right) \quad (11.168)$$

Furthermore

$$\mu(u, g) = g \cdot e^{-g} + u \cdot \frac{20 + g^2}{10 + g^2} \quad (11.169)$$

$$f_o(u, g, \varepsilon_r) = f_{o1}(g, \varepsilon_r) \cdot \exp(p(g) \cdot \ln u + q(g) \cdot \sin(\pi \cdot \log u)) \quad (11.170)$$

$$p(g) = \frac{\exp(-0.745 \cdot g^{0.295})}{\cosh(g^{0.68})} \quad (11.171)$$

$$q(g) = \exp(-1.366 - g) \quad (11.172)$$

$$f_{o1}(g, \varepsilon_r) = 1 - \exp \left(-0.179 \cdot g^{0.15} - \frac{0.328 \cdot g^{r(g, \varepsilon_r)}}{\ln(e + (g/7)^{2.8})} \right) \quad (11.173)$$

$$r(g, \varepsilon_r) = 1 + 0.15 \cdot \left(1 - \frac{\exp \left(1 - (\varepsilon_r - 1)^2 / 8.2 \right)}{1 + g^{-6}} \right) \quad (11.174)$$

The quasi-static characteristic impedance $Z_L(u)$ of a zero-thickness single microstrip line denoted in eq. (??) can either be calculated using the below equations with ε_{reff} being the quasi-static effective dielectric constant defined by eq. (??) or using eqs. (??) and (??).

$$Z_{L1}(u) = \frac{Z_{F0}}{u + 1.98 \cdot u^{0.172}} \quad (11.175)$$

$$Z_L(u) = \frac{Z_{L1}(u)}{\sqrt{\varepsilon_{reff}}} \quad (11.176)$$

The errors in the even and odd mode impedances $Z_{L,e}$ and $Z_{L,o}$ were found to be less than 0.8% and less than 0.3% for the wavelengths.

The model does not include the effect of non-zero strip thickness or asymmetry. Dispersion is also not included. W. J. Getsinger [?] has proposed modifications to his single strip dispersion model, but unfortunately it is easily shown that the results are asymptotically wrong for extreme values of gap width.

In fact he correctly assumes that in the even mode the two strips are at the same potential, and the total current is twice that on a single strip, and dispersion for even-mode propagation is computed by substituting $Z_{L,e}/2$ for Z_L in eqs. (??) and (??). In the odd mode the two strips are at opposite potentials, and the voltage between strips is twice that of a single strip to ground. Thus the total mode impedance is twice that of a single strip, and the dispersion for odd-mode propagation is computed substituting $2Z_{L,o}$ for Z_L in eqs. (??) and (??).

$$\varepsilon_{r,e,o}(f) = \varepsilon_r - \frac{\varepsilon_r - \varepsilon_{reff,e,o}}{1 + G \cdot \left(\frac{f}{f_p}\right)^2} \quad (11.177)$$

with

$$f_p = \begin{cases} \frac{Z_{L,e}}{4\mu_0 h} & \text{even mode} \\ \frac{Z_{L,o}}{\mu_0 h} & \text{odd mode} \end{cases} \quad (11.178)$$

$$G = \begin{cases} 0.6 + Z_{L,e} \cdot 0.0045 & \text{even mode} \\ 0.6 + Z_{L,o} \cdot 0.018 & \text{odd mode} \end{cases} \quad (11.179)$$

11.3.2 Strip thickness correction

According to R.H. Jansen [?] corrected strip width values have been found in the range of technologically meaningful geometries to be

$$W_{t,e} = W + \Delta W \cdot \left(1 - 0.5 \cdot \exp\left(-0.69 \cdot \frac{\Delta W}{\Delta t}\right)\right) \quad (11.180)$$

$$W_{t,o} = W_{t,e} + \Delta t \quad (11.181)$$

with

$$\Delta t = \frac{2 \cdot t \cdot h}{s \cdot \varepsilon_r} \quad \text{for } s \gg 2t \quad (11.182)$$

The author refers to the modifications of the strip width of a single microstrip line ΔW given by Hammerstad and Bekkadal. See also eq. (??) on page ??.

$$\Delta W = \begin{cases} \frac{t}{\pi} \cdot \left(1 + \ln\left(\frac{2h}{t}\right)\right) & \text{for } W > \frac{h}{2\pi} > 2t \\ \frac{t}{\pi} \cdot \left(1 + \ln\left(\frac{4\pi W}{t}\right)\right) & \text{for } \frac{h}{2\pi} \geq W > 2t \end{cases} \quad (11.183)$$

For large spacings s the single line formulae (??) applies.

11.3.3 Transmission losses

The loss equations given by E. Hammerstad and Ø. Jensen [?] for the single microstrip line are also valid for coupled microstrips, provided that the dielectric filling factor, homogeneous impedance, and current distribution factor of the actual mode are used. The following approximation gives good results for odd and even current distribution factors (modification of eq. (??)).

$$K_{i,e} = K_{i,o} = \exp \left(-1.2 \cdot \left(\frac{Z_{L,e} + Z_{L,o}}{2 \cdot Z_{F0}} \right)^{0.7} \right) \quad (11.184)$$

11.4 Microstrip open

A microstrip open end can be modeled by a longer effective microstrip line length Δl as described by M. Kirschning, R.H. Jansen and N.H.L. Koster [?].

$$\frac{\Delta l}{h} = \frac{Q_1 \cdot Q_3 \cdot Q_5}{Q_4} \quad (11.185)$$

with

$$Q_1 = 0.434907 \cdot \frac{\varepsilon_{r,eff}^{0.81} + 0.26}{\varepsilon_{r,eff}^{0.81} - 0.189} \cdot \frac{(W/h)^{0.8544} + 0.236}{(W/h)^{0.8544} + 0.87} \quad (11.186)$$

$$Q_2 = 1 + \frac{(W/h)^{0.371}}{2.358 \cdot \varepsilon_r + 1} \quad (11.187)$$

$$Q_3 = 1 + \frac{0.5274}{\varepsilon_{r,eff}^{0.9236}} \cdot \arctan \left(0.084 \cdot (W/h)^{\frac{1.9413}{Q_2}} \right) \quad (11.188)$$

$$Q_4 = 1 + 0.0377 \cdot (6 - 5 \cdot \exp(0.036 \cdot (1 - \varepsilon_r))) \cdot \arctan \left(0.067 \cdot (W/h)^{1.456} \right) \quad (11.189)$$

$$Q_5 = 1 - 0.218 \cdot \exp(-7.5 \cdot W/h) \quad (11.190)$$

The numerical error is less than 2.5% for $0.01 \leq W/h \leq 100$ and $1 \leq \varepsilon_r \leq 50$.

Another microstrip open end model was published by E. Hammerstad [?]:

$$\frac{\Delta l}{h} = 0.102 \cdot \frac{W/h + 0.106}{W/h + 0.264} \cdot \left(1.166 + \frac{\varepsilon_r + 1}{\varepsilon_r} \cdot (0.9 + \ln(W/h + 2.475)) \right) \quad (11.191)$$

Here the numerical error is less than 1.7% for $W/h < 20$.

In order to simplify calculations, the equivalent additional line length Δl can be transformed into an equivalent open end capacitance C_{end} :

$$C_{end} = C' \cdot \Delta l = \frac{\sqrt{\varepsilon_{r,eff}}}{c_0 \cdot Z_L} \Delta l \quad (11.192)$$

With C' being the capacitance per length and $c_0 = 299\,792\,458$ m/s being the vacuum light velocity.

11.5 Microstrip gap

A symmetrical microstrip gap can be modeled by two open ends with a capacitive series coupling between the two ends. The physical layout is shown in fig. ??.

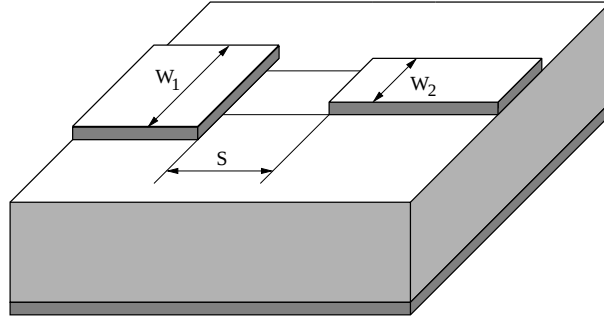


Figure 11.6: symmetrical microstrip gap layout

The equivalent π -network of a microstrip gap is shown in figure ??. The values of the components are according to [?] and [?].

$$C_S \text{ [pF]} = 500 \cdot h \cdot \exp\left(-1.86 \cdot \frac{s}{h}\right) \cdot Q_1 \cdot \left(1 + 4.19 \left(1 - \exp\left(-0.785 \cdot \sqrt{\frac{h}{W_1}} \cdot \frac{W_2}{W_1}\right)\right)\right) \quad (11.193)$$

$$C_{P1} = C_1 \cdot \frac{Q_2 + Q_3}{Q_2 + 1} \quad (11.194)$$

$$C_{P2} = C_2 \cdot \frac{Q_2 + Q_4}{Q_2 + 1} \quad (11.195)$$

with

$$Q_1 = 0.04598 \cdot \left(0.03 + \left(\frac{W_1}{h}\right)^{Q_5}\right) \cdot (0.272 + 0.07 \cdot \varepsilon_r) \quad (11.196)$$

$$Q_2 = 0.107 \cdot \left(\frac{W_1}{h} + 9\right) \cdot \left(\frac{s}{h}\right)^{3.23} + 2.09 \cdot \left(\frac{s}{h}\right)^{1.05} \cdot \frac{1.5 + 0.3 \cdot W_1/h}{1 + 0.6 \cdot W_1/h} \quad (11.197)$$

$$Q_3 = \exp\left(-0.5978 \cdot \left(\frac{W_2}{W_1}\right)^{1.35}\right) - 0.55 \quad (11.198)$$

$$Q_4 = \exp\left(-0.5978 \cdot \left(\frac{W_1}{W_2}\right)^{1.35}\right) - 0.55 \quad (11.199)$$

$$Q_5 = \frac{1.23}{1 + 0.12 \cdot (W_2/W_1 - 1)^{0.9}} \quad (11.200)$$

with C_1 and C_2 being the open end capacitances of a microstrip line (see eq. (??)). The numerical

error of the capacitive admittances is less than 0.1mS for

$$\begin{aligned} 0.1 &\leq W_1/h \leq 3 \\ 0.1 &\leq W_2/h \leq 3 \\ 1 &\leq W_2/W_1 \leq 3 \\ 6 &\leq \varepsilon_r \leq 13 \\ 0.2 &\leq s/h \leq \infty \\ 0.2\text{GHz} &\leq f \leq 18\text{GHz} \end{aligned}$$

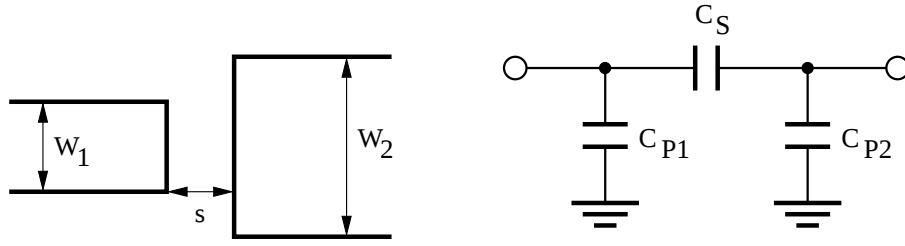


Figure 11.7: microstrip gap and its equivalent circuit

The Y-parameters for the given equivalent small signal circuit can be written as stated in eq. (??) and are easy to convert to scattering parameters.

$$Y = \begin{bmatrix} j\omega \cdot (C_{P1} + C_S) & -j\omega C_S \\ -j\omega C_S & j\omega \cdot (C_{P2} + C_S) \end{bmatrix} \quad (11.201)$$

11.6 Microstrip impedance step

The equivalent circuit of a microstrip impedance step is the same as for the microstrip corner (figure ??). The values are according to [?]:

$$C_S [\text{pF}] = \sqrt{W_1 \cdot W_2} \cdot \left((10.1 \cdot \log \varepsilon_r + 2.33) \cdot \frac{W_1}{W_2} - 12.6 \cdot \log \varepsilon_r - 3.17 \right) \quad (11.202)$$

for $\varepsilon_r \leq 10$ and $1.5 \leq W_1/W_2 \leq 3.5$ the error is $< 10\%$.

$$L_1 = \frac{L_{W1}}{L_{W1} + L_{W2}} \cdot L_S \quad (11.203)$$

$$L_2 = \frac{L_{W2}}{L_{W1} + L_{W2}} \cdot L_S \quad (11.204)$$

with

$$L_{W1,2} = \frac{Z_{L,1,2} \cdot \sqrt{\varepsilon_{r,eff,1,2}}}{c_0} \quad (11.205)$$

$$\frac{L_S}{h} [\text{nH/m}] = 40.5 \cdot \left(\frac{W_1}{W_2} - 1 \right) - 75 \cdot \log \frac{W_1}{W_2} + 0.2 \cdot \left(\frac{W_1}{W_2} - 1 \right)^2 \quad (11.206)$$

With $c_0 = 299\,792\,458$ m/s being the vacuum light velocity. The error is less than 5% for $W_1/W_2 \leq 5$ and $W_2/h = 1$.

11.7 Microstrip tee junction

A model of a microstrip tee junction is published in [?]. Figure ?? shows a unsymmetrical microstrip tee with the main arms consisting of port a and b and with the side arm consisting of port 2. The following model describes the gray area. The equivalent circuit is depicted in figure ?? . It consists of a shunt reactance B_T , one transformer in each main arm (ratios T_a and T_b) and a microstrip line in each arm (width W_a , W_b and W_2).

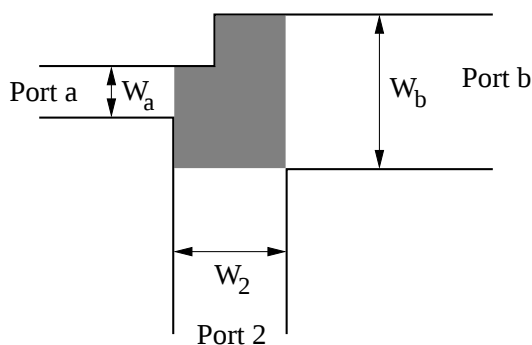


Figure 11.8: unsymmetrical microstrip tee (see text)

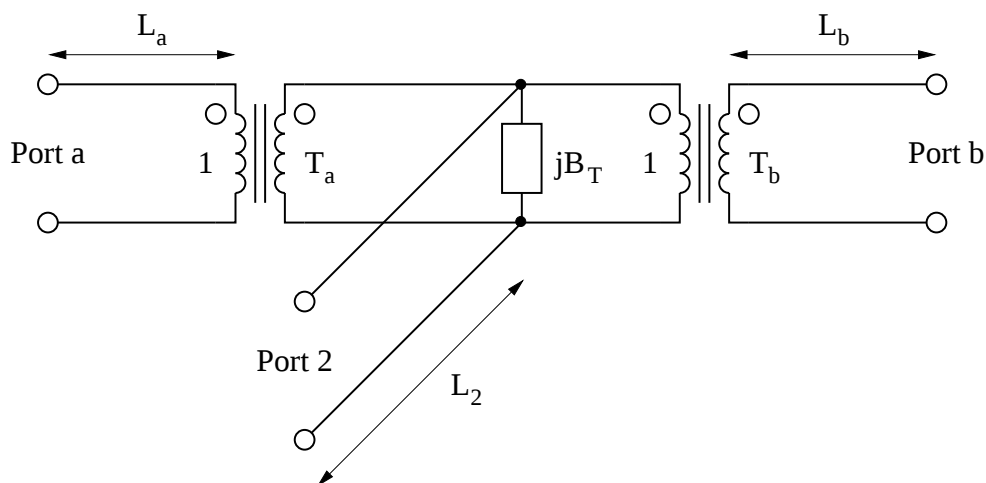


Figure 11.9: equivalent circuit of unsymmetrical microstrip tee

First, let us define some quantities. Each of them is used in the equations below with an index of the arm they belong to (a , b or 2).

$$\text{equivalent parallel plate line width:} \quad D = \frac{Z_{F0}}{\sqrt{\epsilon_{r,eff}}} \cdot \frac{h}{Z_L} \quad (11.207)$$

where Z_{F0} is vacuum field impedance, h height of substrate, $\epsilon_{r,eff}$ effective, relative dielectric

constant, Z_L microstrip line impedance.

$$\text{first higher order mode cut-off frequency: } f_p = 4 \cdot 10^5 \cdot \frac{Z_L}{h} \quad (11.208)$$

$$\text{effective wave length of the microstrip quasi-TEM mode: } \lambda = \frac{c_0}{\sqrt{\varepsilon_{r,eff}} \cdot f} \quad (11.209)$$

The main arm displacements of the reference planes from the center lines are (index x stand for a or b):

$$d_x = 0.055 \cdot D_2 \cdot \frac{Z_{L,x}}{Z_{L,2}} \cdot \left(1 - 2 \cdot \frac{Z_{L,x}}{Z_{L,2}} \cdot \left(\frac{f}{f_{p,x}} \right)^2 \right) \quad (11.210)$$

The length of the line in the main arms is:

$$L_x = 0.5 \cdot W_2 - d_x \quad (11.211)$$

where f is frequency.

The side arm displacement of the reference planes from the center lines is:

$$d_2 = \sqrt{D_a \cdot D_b} \cdot (0.5 - R \cdot (0.05 + 0.7 \cdot \exp(-1.6 \cdot R) + 0.25 \cdot R \cdot Q - 0.17 \cdot \ln R)) \quad (11.212)$$

The length of the line in the side arm is:

$$L_2 = 0.5 \cdot \max(W_a, W_b) - d_2 \quad (11.213)$$

where $\max(x, y)$ is the larger of the both quantities, R and Q are:

$$R = \frac{\sqrt{Z_{L,a} \cdot Z_{L,b}}}{Z_{L,2}} \quad Q = \frac{f^2}{f_{p,a} \cdot f_{p,b}} \quad (11.214)$$

Turn ratio of transformers in the side arms:

$$T_x^2 = 1 - \pi \cdot \left(\frac{f}{f_{p,x}} \right)^2 \cdot \left(\frac{1}{12} \cdot \left(\frac{Z_{L,x}}{Z_{L,2}} \right)^2 + \left(0.5 - \frac{d_2}{D_x} \right)^2 \right) \quad (11.215)$$

Shunt susceptance:

$$B_T = 5.5 \cdot \sqrt{\frac{D_a \cdot D_b}{\lambda_a \cdot \lambda_b}} \cdot \frac{\varepsilon_r + 2}{\varepsilon_r} \cdot \frac{1}{Z_{L,2} \cdot T_a \cdot T_b} \cdot \frac{\sqrt{d_a \cdot d_b}}{D_2} \cdot \left(1 + 0.9 \cdot \ln R + 4.5 \cdot R \cdot Q - 4.4 \cdot \exp(-1.3 \cdot R) - 20 \cdot \left(\frac{Z_{L,2}}{Z_{F0}} \right)^2 \right) \quad (11.216)$$

For better implementation of the microstrip tee (figure ??) the device parameter of the internal equivalent circuit (two transformers and the shunt susceptance) are given below. The port numbering for them is port a = 1, port b = 2 and port 2 = 3.

$$(\underline{Y}) = \text{infinity} \quad (11.217)$$

$$(\underline{Z}) = \frac{1}{j \cdot B_T} \cdot \begin{bmatrix} \frac{1}{n_a^2} & \frac{1}{n_a \cdot n_b} & \frac{1}{n_a} \\ \frac{1}{n_a \cdot n_b} & \frac{1}{n_b^2} & \frac{1}{n_b} \\ \frac{1}{n_a} & \frac{1}{n_b} & 1 \end{bmatrix} \quad (11.218)$$

$$\underline{S}_{11} = \frac{1 - n_a^2 \cdot (j \cdot B_T \cdot Z_0 + \frac{1}{n_b^2} + 1)}{1 + n_a^2 \cdot (j \cdot B_T \cdot Z_0 + \frac{1}{n_b^2} + 1)} \quad (11.219)$$

$$\underline{S}_{22} = \frac{1 - n_b^2 \cdot (j \cdot B_T \cdot Z_0 + \frac{1}{n_a^2} + 1)}{1 + n_b^2 \cdot (j \cdot B_T \cdot Z_0 + \frac{1}{n_a^2} + 1)} \quad (11.220)$$

$$\underline{S}_{33} = \frac{1 - \left(\frac{1}{n_a^2} + \frac{1}{n_b^2} + j \cdot B_T \cdot Z_0 \right)}{1 + \left(\frac{1}{n_a^2} + \frac{1}{n_b^2} + j \cdot B_T \cdot Z_0 \right)} \quad (11.221)$$

$$\underline{S}_{13} = \underline{S}_{31} = \frac{2 \cdot n_a}{n_a^2 \cdot \left(\frac{1}{n_b^2} + j \cdot B_T \cdot Z_0 + 1 \right) + 1} \quad (11.222)$$

$$\underline{S}_{23} = \underline{S}_{32} = \frac{2 \cdot n_b}{n_b^2 \cdot \left(\frac{1}{n_a^2} + j \cdot B_T \cdot Z_0 + 1 \right) + 1} \quad (11.223)$$

$$\underline{S}_{12} = \underline{S}_{21} = \frac{2}{n_a \cdot n_b \cdot (j \cdot B_T \cdot Z_0 + 1) + \frac{n_a}{n_b} + \frac{n_b}{n_a}} \quad (11.224)$$

The MNA matrix representation can be derived from the Z parameters in the following way.

$$\begin{bmatrix} \cdot & \cdot & \cdot & 1 & 0 & 0 \\ \cdot & \cdot & \cdot & 0 & 1 & 0 \\ \cdot & \cdot & \cdot & 0 & 0 & 1 \\ -1 & 0 & 0 & Z_{11} & Z_{12} & Z_{13} \\ 0 & -1 & 0 & Z_{21} & Z_{22} & Z_{23} \\ 0 & 0 & -1 & Z_{31} & Z_{32} & Z_{33} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ I_{1,in} \\ I_{2,in} \\ I_{3,in} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ I_3 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (11.225)$$

Please note that the main arm displacements in eq. (??) yield two small microstrip lines at each main arm and the side arm displacement of eq. (??) results in a small microstrip strip line as well, but with negative length, i.e. kind of phaseshifter here.

The transformer ratios defined in eq. (??) are going to be negative with increasing frequency which produces complex values in the Z-parameter matrix as well as in the S-parameter matrix. That is why the ratios are delimited to a minimum value.

11.8 Microstrip cross

The most useful model of a microstrip cross have been published in [?, ?]. Fig. ?? shows the equivalent circuit (right-hand side) and the scheme with dimensions (left-hand side). The hatched area in the scheme marks the area modeled by the equivalent circuit. As can be seen the model require the microstrip width of line 1 and 3, as well as the one of line 2 and 4 to equal each other. Furthermore the permittivity of the substrat must be $\epsilon_r = 9.9$. The component values are calculated as follows:

$$X = \log_{10} \left(\frac{W_1}{h} \right) \cdot \left(86.6 \cdot \frac{W_2}{h} - 30.9 \cdot \sqrt{\frac{W_2}{h}} + 367 \right) + \left(\frac{W_2}{h} \right)^3 + 74 \cdot \frac{W_2}{h} + 130 \quad (11.226)$$

$$C_1 = C_2 = C_3 = C_4$$

$$= 10^{-12} \cdot W_1 \cdot \left(0.25 \cdot X \cdot \left(\frac{h}{W_1} \right)^{1/3} - 60 + \frac{h}{2 \cdot W_2} - 0.375 \cdot \frac{W_1}{h} \cdot \left(1 - \frac{W_2}{h} \right) \right) \quad (11.227)$$

$$Y = 165.6 \cdot \frac{W_2}{h} + 31.2 \sqrt{\frac{W_2}{h}} - 11.8 \cdot \left(\frac{W_2}{h} \right)^2 \quad (11.228)$$

$$L_1 = 10^{-9} \cdot h \cdot \left(Y \cdot \frac{W_1}{h} - 32 \cdot \frac{W_2}{h} + 3 \right) \cdot \left(\frac{h}{W_1} \right)^{1.5} \quad (11.229)$$

$$L_3 = 10^{-9} \cdot h \cdot \left(5 \cdot \frac{W_2}{h} \cdot \cos \left(\frac{\pi}{2} \cdot \left(1.5 - \frac{W_1}{h} \right) \right) - \left(1 + \frac{7 \cdot h}{W_1} \right) \cdot \frac{h}{W_2} - 337.5 \right) \quad (11.230)$$

The equation of L_2 is obtained from the one of L_1 by exchanging the indices (W_1 and W_2). Note that L_3 is negative, so the model is unphysical without external microstrip lines. The above-mentioned equations are accurate to within 5% for $0.3 \leq W_1/h \leq 3$ and $0.1 \leq W_2/h \leq 3$ (value of $C_1 \dots C_4$) or for $0.5 \leq W_{1,2}/h \leq 2$ (value of $L_1 \dots L_3$), respectively.

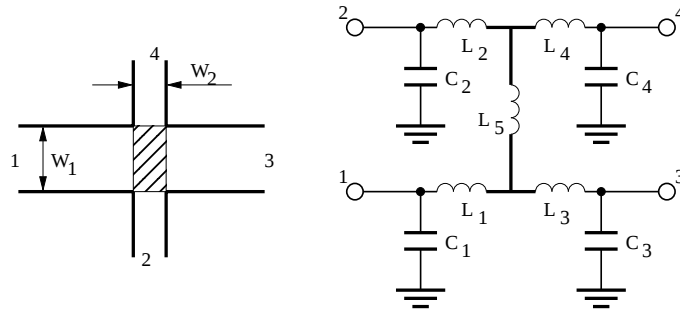


Figure 11.10: single-symmetrical microstrip cross and its model

Some improvement should be added to the original model:

1. Comparisons with real life show that the value of L_3 is too large. Multiplying it by 0.8 leads to much better results.
2. The model can be expanded for substrates with $\epsilon_r \neq 9.9$ by modifying the values of the capacitances:

$$C_x = C_x(\epsilon_r = 9.9) \cdot \frac{Z_0(\epsilon_r = 9.9, W = W_x)}{Z_0(\epsilon_r = \epsilon_{r,sub}, W = W_x)} \cdot \sqrt{\frac{\epsilon_{eff}(\epsilon_r = \epsilon_{r,sub}, W = W_x)}{\epsilon_{eff}(\epsilon_r = 9.9, W = W_x)}} \quad (11.231)$$

The equations of Z_0 and ϵ_{eff} are the ones from the microstrip lines.

A useful model for an unsymmetrical cross junction has never been published. Nonetheless, as long as the lines that lie opposite are not too different in width, the model described here can be used as a first order approximation. This is performed by replacing W_1 and W_2 by the arithmetic mean of the line widths that lie opposite. This is done:

- In equation (??) and (??) for W_2 only, whereas W_1 is replaced by the width of the line.
- In equation (??) and (??) for W_2 only, whereas W_1 is replaced by the width of the line.
- In equation (??) for W_1 and W_2 .

Another closed-form expression describing the non-ideal behaviour of a microstrip cross junction was published by [?]. Additionally there have been published papers [?, ?, ?] giving analytic (but not closed-form) expressions or just simple equivalent circuits with only a few expressions for certain topologies and dielectric constants which are actually of no practical use.

11.9 Microstrip via hole

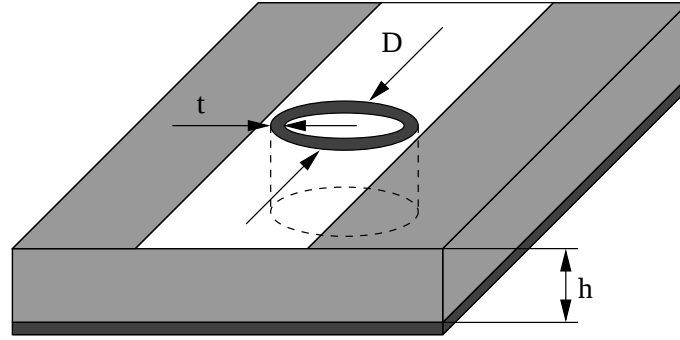


Figure 11.11: microstrip via hole to ground

According to Marc E. Goldfarb and Robert A. Pucel [?] a via hole ground in microstrip is a series of a resistance and an inductance. The given model for a cylindrical via hole has been verified numerically and experimentally for a range of $h < 0.03 \cdot \lambda_0$.

$$L = \frac{\mu_0}{2\pi} \cdot \left(h \cdot \ln \left(\frac{h + \sqrt{r^2 + h^2}}{r} \right) + \frac{3}{2} \cdot \left(r - \sqrt{r^2 + h^2} \right) \right) \quad (11.232)$$

whence h is the via length (substrate height) and $r = D/2$ the via's radius.

$$R = R(f = 0) \cdot \sqrt{1 + \frac{f}{f_\delta}} \quad (11.233)$$

with

$$f_\delta = \frac{\rho}{\pi \cdot \mu_0 \cdot t^2} \quad (11.234)$$

The relationship for the via resistance can be used as a close approximation and is valid independent of the ratio of the metalization thickness t to the skin depth. In the formula ρ denotes the specific resistance of the conductor material.

11.10 Bondwire

Wire inductors, so called bond wire connections, are used to connect active and passive circuit components as well as micro devices to the real world.

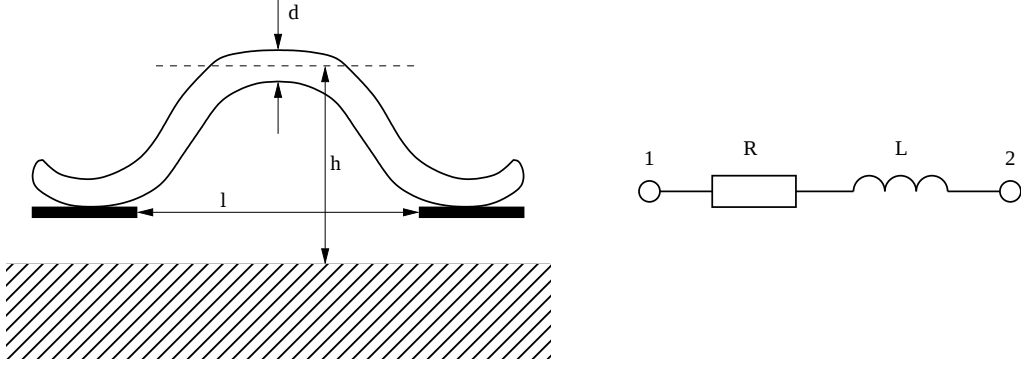


Figure 11.12: bond wire and its equivalent circuit

11.10.1 Freespace model

The freespace inductance L of a wire of diameter d and length l is given [?, ?] by

$$L = \frac{\mu_0}{2\pi} \cdot l \left[\ln \left\{ \frac{2l}{d} + \sqrt{1 + \left(\frac{2l}{d} \right)^2} \right\} + \frac{d}{2l} - \sqrt{1 + \left(\frac{d}{2l} \right)^2} + C \right] \quad (11.235)$$

where the frequency-dependent correction factor C is a function of bond wire diameter and its material skin depth δ is expressed as

$$C = \frac{\mu_r}{4} \cdot \tanh \left(\frac{4\delta}{d} \right) \quad (11.236)$$

$$\delta = \frac{1}{\sqrt{\pi \cdot \sigma \cdot f \cdot \mu_0 \cdot \mu_r}} \quad (11.237)$$

where σ is the conductivity of the wire material. When δ/d is small, $C = \delta/d$. The wire resistance R is given by

$$R = \frac{\rho \cdot l}{\pi \cdot r^2} \quad (11.238)$$

with $\rho = 1/\sigma$ and $r = d/2$.

11.10.2 Mirror model

The effect of the ground plane on the inductance value of a wire has also been considered. If the wire is at a distance h above the ground plane, it sees its image at $2h$ from it. The wire and its image result in a mutual inductance. Since the image wire carries a current opposite to the current flow in the bond wire, the effective inductance of the bond wire becomes

$$L = \frac{\mu_0}{2\pi} \cdot l \left[\ln \left(\frac{4h}{d} \right) + \ln \left(\frac{l + \sqrt{l^2 + d^2/4}}{l + \sqrt{l^2 + 4h^2}} \right) + \sqrt{1 + \frac{4h^2}{l^2}} - \sqrt{1 + \frac{d^2}{4l^2}} - 2\frac{h}{l} + \frac{d}{2l} \right] \quad (11.239)$$

Mirror is a strange model that is frequency independent. Whereas computations are valid, hypothesis are arguable. Indeed, they did the assumption that the ground plane is perfect that is really a zero order model in the high frequency domain.

11.11 Planar inductors

The planar inductors are widely used both in MMIC and PCB designs as an alternative to coil inductors. The model implemented in Qucs is based on following equivalent circuit [?]:

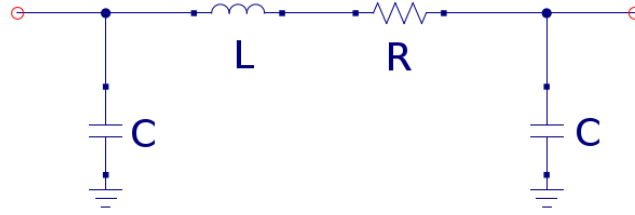


Figure 11.13: Equivalent circuit of a printed inductor

11.11.1 Circular loop

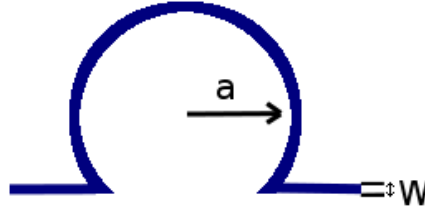


Figure 11.14: Design parameters of the circular loop inductor

The parameters for the equivalent circuit model ?? are given by the following equations:

$$L(nH) = 1.257 \cdot 10^{-3} a \left(\ln \left(\frac{a}{W+t} \right) + 0.078 \right) \cdot K_g \quad (11.240)$$

$$R(\Omega) = \frac{\pi \cdot a \cdot K \cdot R_s}{W+t} \quad (11.241)$$

$$C(pF) = 33.33 \cdot 10^{-4} \cdot \frac{\pi \cdot a \cdot \sqrt{\epsilon_{re}}}{Z_0} \quad (11.242)$$

where

$$K_g = 0.57 - 0.145 \cdot \ln\left(\frac{W}{h}\right) \quad \frac{W}{h} > 0.05 \quad (11.243)$$

$$K = 1.4 + 0.217 \cdot \ln\left(\frac{W}{5 \cdot t}\right) \quad 5 < \frac{W}{t} < 100 \quad (11.244)$$

K_g accounts for the ground plane effect and K is a correction factor that accounts for the crowding effect at the corners of the transmission line [?].

11.11.2 Spiral inductor

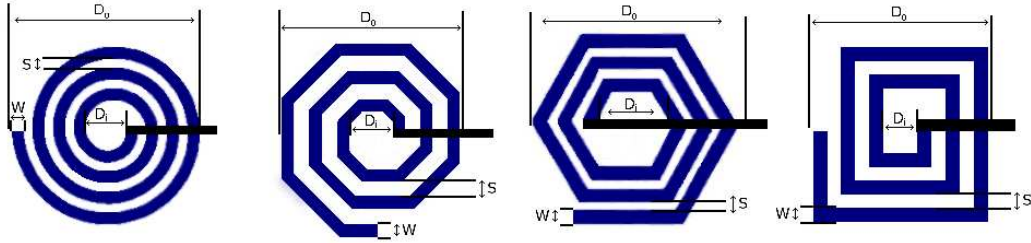


Figure 11.15: Circular, octagonal, hexagonal and square spiral planar inductors

The parameters for the equivalent circuit model ?? are given by the following equations:

$$L = \frac{\mu_0 \cdot n^2 \cdot D_{av} c_1}{2} \left(\ln\left(\frac{c_2}{\rho}\right) + c_3 \cdot \rho + c_4 \cdot \rho^2 \right) \quad (11.245)$$

$$R = \frac{K \cdot \pi \cdot a \cdot n \cdot R_s}{W} \quad (11.246)$$

$$C(pF) = 3.5 \cdot 10^{-5} D_0 + 0.06 \quad (11.247)$$

where n is the number of turns, K is the same as in ?? and a and D_{av} are defined as:

$$a = \frac{D_o + D_i}{4} \quad (11.248)$$

$$D_{av} = \frac{D_o + D_i}{2} \quad (11.249)$$

The coefficients of the inductance formula are the following:

Geometry	c_1	c_2	c_3	c_4
Square	1.27	2.07	0.18	0.13
Hexagonal	1.09	2.23	0	0.17
Octogonal	1.07	2.29	0	0.19
Circular	1	2.46	0	0.2

Table 11.1: Coefficients of the inductance formula depending on the geometry

Chapter 12

Coplanar components

12.1 Coplanar waveguides (CPW)

12.1.1 Definition

A *coplanar line* is a structure in which all the conductors supporting wave propagation are located on the same plane, i.e. generally the top of a dielectric substrate. There exist two main types of coplanar lines: the first, called *coplanar waveguide (CPW)*, that we will study here, is composed of a median metallic strip separated by two narrow slits from a infinite ground plane, as may be seen on the figure below.

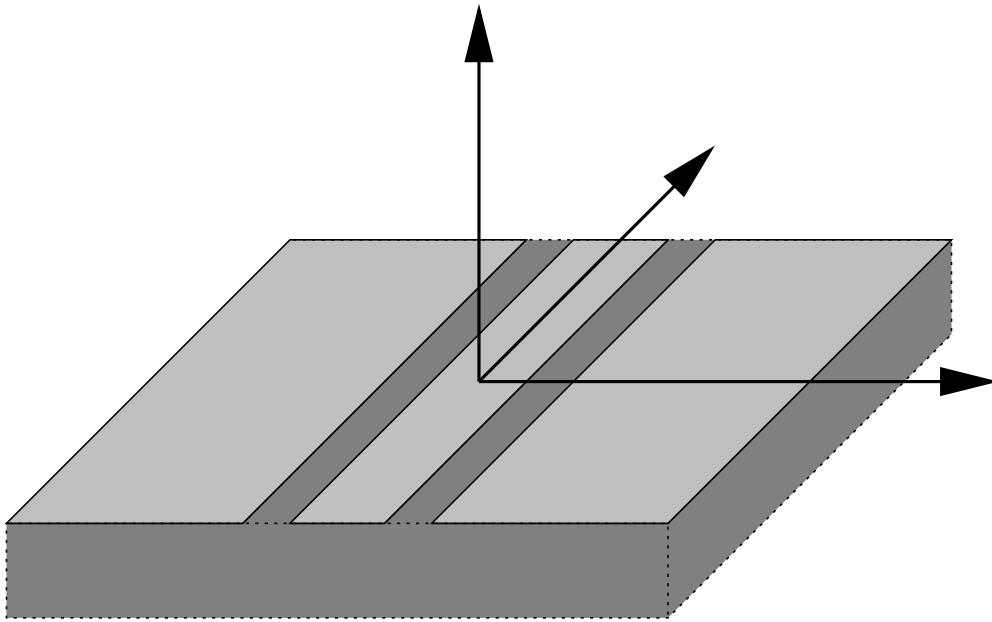
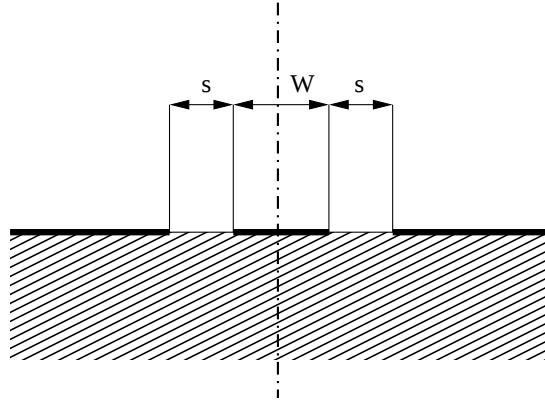


Figure 12.1: coplanar waveguide line

The characteristic dimensions of a CPW are the central strip width W and the width of the slots

s . The structure is obviously symmetrical along a vertical plane running in the middle of the central strip.



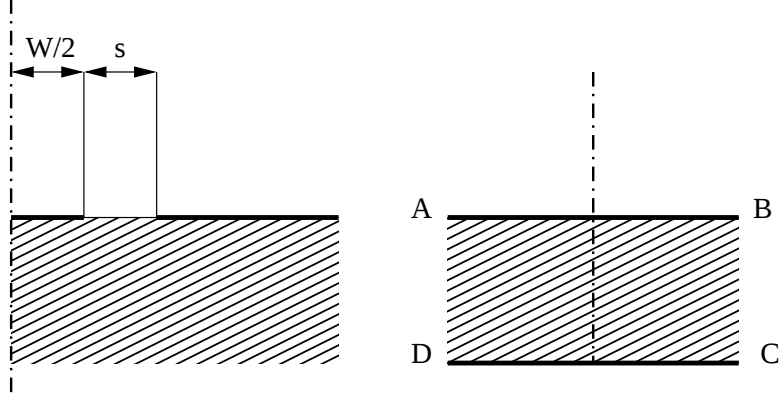
The other coplanar line, called a *coplanar slot (CPS)* is the complementary of that topology, consisting of two strips running side by side.

12.1.2 Quasi-static analysis by conformal mappings

A CPW can be quasi-statically analysed by the use of *conformal mappings*. Briefly speaking, it consists in transforming the geometry of the PCB into another conformation, whose properties make the computations straightforward. The interested reader can consult the pp. 886 - 910 of [?] which has a correct coverage of both the theoretical and applied methods. The French reader interested in the mathematical arcans involved is referred to the second chapter of [?] (which may be out of print nowadays), for an extensive review of all the theoretical framework. The following analysis is mainly borrowed from [?], pp. 375 *et seq.* with additions from [?].

The CPW of negligible thickness located on top of an infinitely deep substrate, as shown on the left of the figure below, can be mapped into a parallel plate capacitor filled with dielectric $ABCD$ using the conformal function:

$$w = \int_{z_0}^z \frac{dz}{\sqrt{(z - W/2)(z - W/2 - s)}}. \quad (12.1)$$



To further simplify the analysis, the original dielectric boundary is assumed to constitute a magnetic wall, so that BC and AD become magnetic walls too and there is no resulting fringing field in the resulting capacitor. With that assumption, the capacitance per unit length is merely the sum of the top (air filled) and bottom (dielectric filled) partial capacitances. The latter is given by:

$$C_d = 2 \cdot \varepsilon_0 \cdot \varepsilon_r \cdot \frac{K(k_1)}{K'(k_1)} \quad (12.2)$$

while the former is:

$$C_a = 2 \cdot \varepsilon_0 \cdot \frac{K(k_1)}{K'(k_1)} \quad (12.3)$$

In both formulae $K(k)$ and $K'(k)$ represent the complete elliptic integral of the first kind and its complement, and $k_1 = \frac{W}{W+2s}$.

$K(k)$ (and $K'(k)$) can be computed easily by using the *arithmetic-geometric mean* method [?]: starting with

$$a_0 = 1, \quad b_0 = \sqrt{1 - k^2} \quad (12.4)$$

the following recursion formulae are iterated

$$a_n = \frac{1}{2}(a_{n-1} + b_{n-1}), \quad b_n = \sqrt{a_{n-1}b_{n-1}}, \quad c_n = a_{n-1} - b_{n-1} \quad (12.5)$$

until $c_N = 0$ within a predetermined accuracy.

$K(k)$ is then found from

$$K(k) = \frac{\pi}{2a_N} \quad (12.6)$$

When only the K/K' ratio is needed it can be approximated efficiently through the following formulae:

$$\frac{K(k)}{K'(k)} = \frac{\pi}{\ln \left(2 \frac{1+\sqrt{k'}}{1-\sqrt{k'}} \right)} \quad \text{for} \quad 0 \leq k \leq \frac{1}{\sqrt{2}} \quad (12.7)$$

$$\frac{K(k)}{K'(k)} = \frac{\ln \left(2 \frac{1+\sqrt{k}}{1-\sqrt{k}} \right)}{\pi} \quad \text{for} \quad \frac{1}{\sqrt{2}} \leq k \leq 1 \quad (12.8)$$

with k' being the complementary modulus: $k' = \sqrt{1 - k^2}$. This formula was first derived in [?]. According to [?] the accuracy of the above formulae is close to 10^{-5} , [?] claims it to be $3 \cdot 10^{-6}$ while comparison with GNU Octave using the *ellipke()* function shows a maximum relative error of about $2 \cdot 10^{-6}$. More accurate approximations can be found in [?] and [?] but the above formulae can already be considered as exact for any practical purpose.

The total line capacitance is thus the sum of C_d and C_a . The effective permittivity is therefore:

$$\varepsilon_{re} = \frac{\varepsilon_r + 1}{2} \quad (12.9)$$

and the impedance:

$$Z = \frac{30\pi}{\sqrt{\varepsilon_{re}}} \cdot \frac{K'(k_1)}{K(k_1)} \quad (12.10)$$

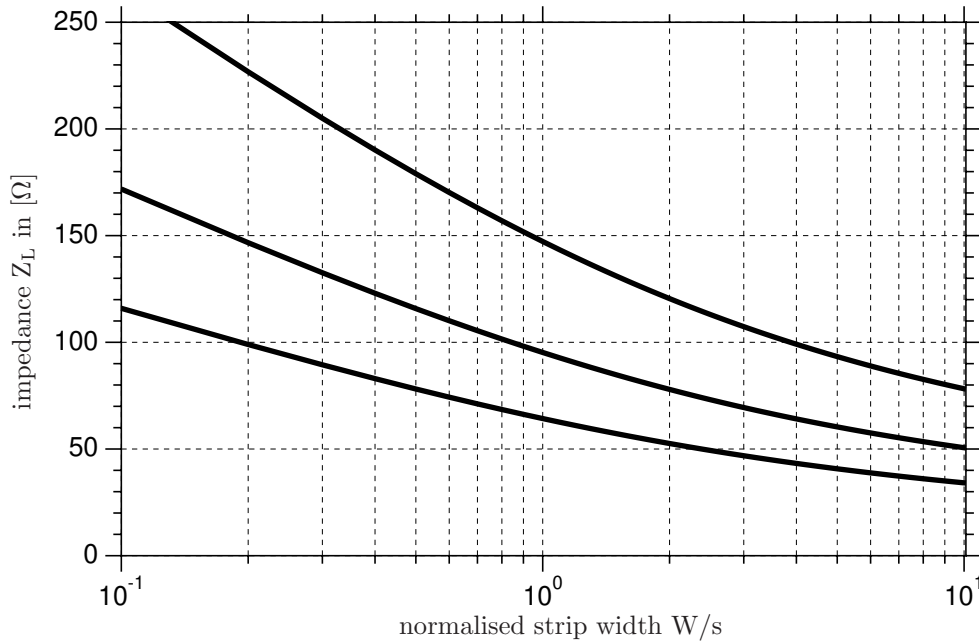


Figure 12.2: characteristic impedance as approximated by eq. (??) for $\varepsilon_r = 1.0$ (air), 3.78 (quartz) and 9.5 (alumina)

In practical cases, the substrate has a finite thickness h . To carry out the analysis of this conformation, a preliminary conformal mapping transforms the finite thickness dielectric into an infinite thickness one. Only the effective permittivity is altered; it becomes:

$$\varepsilon_{re} = 1 + \frac{\varepsilon_r - 1}{2} \cdot \frac{K(k_2)}{K'(k_2)} \cdot \frac{K'(k_1)}{K(k_1)} \quad (12.11)$$

where k_1 is given above and

$$k_2 = \frac{\sinh\left(\frac{\pi W}{4h}\right)}{\sinh\left(\frac{\pi \cdot (W + 2s)}{4h}\right)}. \quad (12.12)$$

Finally, let us consider a CPW over a finite thickness dielectric backed by an infinite ground plane. In this case, the quasi-TEM wave is an hybrid between microstrip and true CPW mode. The equations then become:

$$\varepsilon_{re} = 1 + q \cdot (\varepsilon_r - 1) \quad (12.13)$$

where q , called *filling factor* is given by:

$$q = \frac{\frac{K(k_3)}{K'(k_3)}}{\frac{K(k_1)}{K'(k_1)} + \frac{K(k_3)}{K'(k_3)}} \quad (12.14)$$

and

$$k_3 = \frac{\tanh\left(\frac{\pi W}{4h}\right)}{\tanh\left(\frac{\pi \cdot (W + 2s)}{4h}\right)} \quad (12.15)$$

The impedance of this line amounts to:

$$Z = \frac{60\pi}{\sqrt{\varepsilon_{re}}} \cdot \frac{1}{\frac{K(k_1)}{K'(k_1)} + \frac{K(k_3)}{K'(k_3)}} \quad (12.16)$$

12.1.3 Effects of metalization thickness

In most practical cases, the strips are very thin, yet their thickness cannot be entirely neglected. A first order correction to take into account the non-zero thickness of the conductor is given by [?]:

$$s_e = s - \Delta \quad (12.17)$$

and

$$W_e = W + \Delta \quad (12.18)$$

where

$$\Delta = \frac{1.25t}{\pi} \cdot \left(1 + \ln\left(\frac{4\pi W}{t}\right)\right) \quad (12.19)$$

In the computation of the impedance, both the k_1 and the effective dielectric constant are affected, wherefore k_1 must be substituted by an “effective” modulus k_e , with:

$$k_e = \frac{W_e}{W_e + 2s_e} \approx k_1 + (1 - k_1^2) \cdot \frac{\Delta}{2s} \quad (12.20)$$

and

$$\varepsilon_{re}^t = \varepsilon_{re} - \frac{0.7 \cdot (\varepsilon_{re} - 1) \cdot \frac{t}{s}}{\frac{K(k_1)}{K'(k_1)} + 0.7 \cdot \frac{t}{s}} \quad (12.21)$$

12.1.4 Effects of dispersion

The effects of dispersion in CPW are similar to those encountered in the microstrip lines, though the net effect on impedance is somewhat different. [?] gives a closed form expression to compute $\varepsilon_{re}(f)$ from its quasi-static value:

$$\sqrt{\varepsilon_{re}(f)} = \sqrt{\varepsilon_{re}(0)} + \frac{\sqrt{\varepsilon_r} - \sqrt{\varepsilon_{re}(0)}}{1 + G \cdot \left(\frac{f}{f_{TE}}\right)^{-1.8}} \quad (12.22)$$

where:

$$G = e^{u \cdot \ln\left(\frac{W}{s}\right) + v} \quad (12.23)$$

$$u = 0.54 - 0.64p + 0.015p^2 \quad (12.24)$$

$$v = 0.43 - 0.86p + 0.54p^2 \quad (12.25)$$

$$p = \ln\left(\frac{W}{h}\right) \quad (12.26)$$

and f_{TE} is the cut-off frequency of the TE_0 mode, defined by:

$$f_{TE} = \frac{c}{4h \cdot \sqrt{\varepsilon_r - 1}}. \quad (12.27)$$

This dispersion expression was first reported by [?] and has been reused and extended in [?]. The accuracy of this expression is claimed to be better than 5% for $0.1 \leq W/h \leq 5$, $0.1 \leq W/s \leq 5$, $1.5 \leq \varepsilon_r \leq 50$ and $0 \leq f/f_{TE} \leq 10$.

12.1.5 Evaluation of losses

As for microstrip lines, the losses in CPW results of at least two factors: a dielectric loss α_d and conductor losses α_c^{CW} . The dielectric loss α_d is identical to the microstrip case, see eq. (??) on page ??.

The α_c^{CW} part of the losses is more complex to evaluate. As a general rule, it might be written:

$$\alpha_c^{CW} = 0.023 \cdot \frac{R_s}{Z_{0cp}} \left[\frac{\partial Z_{0cp}^a}{\partial s} - \frac{\partial Z_{0cp}^a}{\partial W} - \frac{\partial Z_{0cp}^a}{\partial t} \right] \text{ in dB/unit length} \quad (12.28)$$

where Z_{0cp}^a stands for the impedance of the coplanar waveguide with air as dielectric and R_s is the surface resistivity of the conductors (see eq. (??) on page ??).

Through a direct approach evaluating the losses by conformal mapping of the current density, one obtains [?], first reported in [?] and finally applied to coplanar lines by [?]:

$$\alpha_c^{CW} = \frac{R_s \cdot \sqrt{\varepsilon_{re}}}{480\pi \cdot K(k_1) \cdot K'(k_1) \cdot (1 - k_1^2)} \cdot \left(\frac{1}{a} \left[\pi + \ln \frac{8\pi a \cdot (1 - k_1)}{t \cdot (1 + k_1)} \right] + \frac{1}{b} \left[\pi + \ln \frac{8\pi b \cdot (1 - k_1)}{t \cdot (1 + k_1)} \right] \right) \quad (12.29)$$

In the formula above, $a = W/2$, $b = s + W/2$ and it is assumed that $t > 3\delta$, $t \ll W$ and $t \ll s$.

12.1.6 S- and Y-parameters of the single coplanar line

The computation of the coplanar waveguide lines S- and Y-parameters is equal to all transmission lines (see section ?? on page ??).

12.2 Coplanar waveguide open

The behaviour of an open circuit as shown in fig. ?? is very similar to that in a microstrip line; that is, the open circuit is capacitive in nature.

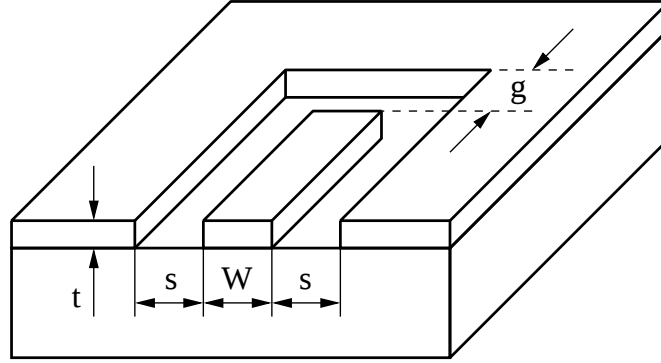


Figure 12.3: coplanar waveguide open-circuit

A very simple approximation for the equivalent length extension Δl associated with the fringing fields has been given by K.Beilenhoff [?].

$$\Delta l_{open} = \frac{C_{open}}{C'} \approx \frac{W + 2s}{4} \quad (12.30)$$

For the open end, the value of Δl is not influenced significantly by the metalization thickness and the gap width g when $g > W + 2s$. Also, the effect of frequency and aspect ration $W/(W + 2s)$ is relatively weak. The above approximation is valid for $0.2 \leq W/(W + 2s) \leq 0.8$.

The open end capacitance C_{open} can be written in terms of the capacitance per unit length and the wave resistance.

$$C_{open} = C' \cdot \Delta l_{open} = \frac{\sqrt{\epsilon_{r,eff}}}{c_0 \cdot Z_L} \cdot \Delta l_{open} \quad (12.31)$$

12.3 Coplanar waveguide short

There is a similar simple approximation for a coplanar waveguide short-circuit, also given in [?]. The short circuit is inductive in nature.

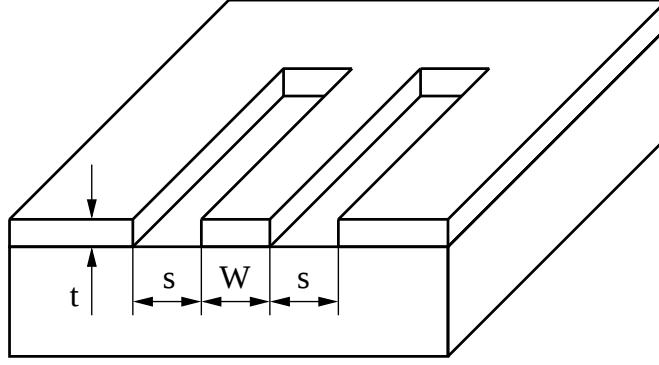


Figure 12.4: coplanar waveguide short-circuit

The equivalent length extension Δl associated with the fringing fields is

$$\Delta l_{short} = \frac{L_{short}}{L'} \approx \frac{W + 2s}{8} \quad (12.32)$$

Equation (??) is valid when the metalization thickness t does not become too large ($t < s/3$).

The short end inductance L_{short} can be written in terms of the inductance per unit length and the wave resistance.

$$L_{short} = L' \cdot \Delta l_{short} = \frac{\sqrt{\varepsilon_{r,eff}} \cdot Z_L}{c_0} \cdot \Delta l_{short} \quad (12.33)$$

According to W.J.Getsinger [?] the CPW short-circuit inductance per unit length can also be modeled by

$$L_{short} = \frac{2}{\pi} \cdot \varepsilon_0 \cdot \varepsilon_{r,eff} \cdot (W + s) \cdot Z_L^2 \cdot \left(1 - \operatorname{sech} \left(\frac{\pi \cdot Z_{F0}}{2 \cdot Z_L \cdot \sqrt{\varepsilon_{r,eff}}} \right) \right) \quad (12.34)$$

based on his duality [?] theory.

12.4 Coplanar waveguide gap

According to W.J.Getsinger [?] a coplanar series gap (see fig. ??) is supposed to be the dual problem of the inductance of a connecting strip between twin strip lines.

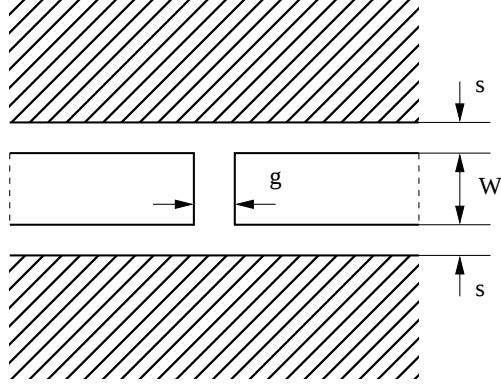


Figure 12.5: coplanar waveguide series gap

The inductance of such a thin strip with a width g and the length W is given to a good approximation by

$$L = \frac{\mu_0 \cdot W}{2\pi} \cdot \left(p - \sqrt{1 + p^2} + \ln \left(\frac{1 + \sqrt{1 + p^2}}{p} \right) \right) \quad (12.35)$$

where $p = g/4W$ and $g, W \ll \lambda$. Substituting this inductance by its equivalent capacitance of the gap in CPW yields

$$\begin{aligned} C &= L \cdot \frac{4 \cdot \varepsilon_{r,eff}}{Z_{F0}^2} \\ &= \frac{2 \cdot \varepsilon_0 \cdot \varepsilon_{r,eff} \cdot W}{\pi} \cdot \left(p - \sqrt{1 + p^2} + \ln \left(\frac{1 + \sqrt{1 + p^2}}{p} \right) \right) \end{aligned} \quad (12.36)$$

12.5 Coplanar waveguide step

The coplanar step discontinuity shown in figure ?? has been analysed by C. Sinclair [?].

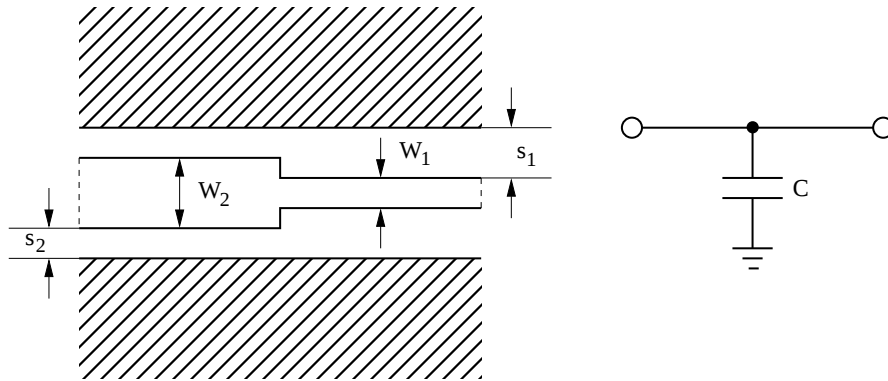


Figure 12.6: coplanar waveguide impedance step and equivalent circuit

The symmetric step change in width of the centre conductor is considered to have a similar equivalent circuit as a step of a parallel plate guide - this is a reasonable approximation to the CPW step as in the CPW the majority of the field is between the inner and outer conductors with some fringing.

The actual CPW capacitance can be expressed as

$$C = \bar{x} \cdot \frac{\varepsilon_0}{\pi} \cdot \left(\frac{\alpha^2 + 1}{\alpha} \cdot \ln \left(\frac{1 + \alpha}{1 - \alpha} \right) - 2 \cdot \ln \left(\frac{4 \cdot \alpha}{1 - \alpha^2} \right) \right) \quad (12.37)$$

where

$$\alpha = \frac{s_1}{s_2}, \alpha < 1 \quad \text{and} \quad \bar{x} = \frac{x_1 + x_2}{2} \quad (12.38)$$

The capacitance per unit length equivalence yields

$$x_1 = \frac{C'(W_1, s_1) \cdot s_1}{\varepsilon_0} \quad \text{and} \quad x_2 = \frac{C'(W_2, s_2) \cdot s_2}{\varepsilon_0} \quad (12.39)$$

with

$$C' = \frac{\sqrt{\varepsilon_{r,eff}}}{c_0 \cdot Z_L} \quad (12.40)$$

The average equivalent width $\bar{x}$ of the parallel plate guide can be adjusted with an expression that uses weighted average of the gaps s_1 and s_2 . The final expression has not been discussed in [?]. The given equations are validated over the following ranges: $2 < \varepsilon_r < 14$, $h > W + 2s$ and $f < 40\text{GHz}$.

The Z-parameters of the equivalent circuit depicted in fig. ?? are

$$Z_{11} = Z_{21} = Z_{12} = Z_{22} = \frac{1}{j\omega C} \quad (12.41)$$

The MNA matrix representation for the AC analysis can be derived from the Z-parameters in the following way.

$$\begin{bmatrix} . & . & 1 & 0 \\ . & . & 0 & 1 \\ -1 & 0 & Z_{11} & Z_{12} \\ 0 & -1 & Z_{21} & Z_{22} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ I_{in} \\ I_{out} \end{bmatrix} = \begin{bmatrix} I_1 \\ I_2 \\ 0 \\ 0 \end{bmatrix} \quad (12.42)$$

The above expanded representation using the Z-parameters is necessary because the Y-parameters are infinite. During DC analysis the equivalent circuit is a voltage source between both terminals with zero voltage.

The S-parameters of the topology are

$$S_{11} = S_{22} = -\frac{Z_0}{2Z + Z_0} \quad (12.43)$$

$$S_{12} = S_{21} = 1 + S_{11} = \frac{2Z}{2Z + Z_0} \quad (12.44)$$

Chapter 13

Other types of transmission lines

13.1 Coaxial cable

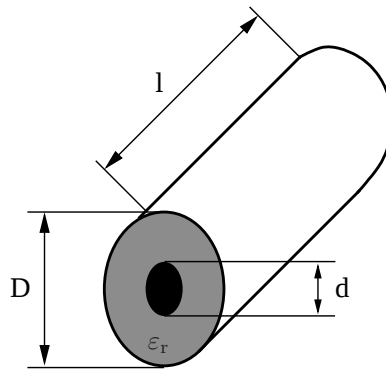


Figure 13.1: coaxial line

13.1.1 Characteristic impedance

The characteristic impedance of a coaxial line can be calculated as follows:

$$Z_L = \frac{Z_{F0}}{2\pi \cdot \sqrt{\epsilon_r}} \cdot \ln \left(\frac{D}{d} \right) \quad (13.1)$$

13.1.2 Losses

Overall losses in a coaxial cable consist of dielectric and conductor losses. The dielectric losses compute as follows:

$$\alpha_d = \frac{\pi}{c_0} \cdot f \cdot \sqrt{\epsilon_r} \cdot \tan \delta \quad (13.2)$$

The conductor (i.e. ohmic) losses are specified by

$$\alpha_c = \sqrt{\varepsilon_r} \cdot \left(\frac{\frac{1}{D} + \frac{1}{d}}{\ln\left(\frac{D}{d}\right)} \right) \cdot \frac{R_S}{Z_{F0}} \quad (13.3)$$

with R_S denoting the sheet resistance of the conductor material, i.e. the skin resistance

$$R_S = \sqrt{\pi \cdot f \cdot \mu_r \cdot \mu_o \cdot \rho} \quad (13.4)$$

13.1.3 Cutoff frequencies

In normal operation a signal wave passes through the coaxial line as a TEM wave with no electrical or magnetic field component in the direction of propagation. Beyond a certain cutoff frequency additional (unwanted) higher order modes are excited.

$$f_{TE} \approx \frac{2 \cdot c_0}{\pi \cdot (D + d) \cdot \sqrt{\varepsilon_r}} \rightarrow \text{TE}(1,1) \text{ mode} \quad (13.5)$$

$$f_{TM} \approx \frac{c_0}{(D - d) \cdot \sqrt{\varepsilon_r}} \rightarrow \text{TM}(n,1) \text{ mode} \quad (13.6)$$

13.2 Rectangular waveguide

In most of practical applications, the rectangular waveguides are designed to propagate the dominant mode only. This allows the designer to use conventional transmission line equations and avoid power leaking to higher order modes. In this sense, the Qucs rectangular waveguide component was implemented to extract the transmission line parameters from the TE_{10} mode (dominant mode) and generate the quadripole parameters. For this reason, the simulation data is valid only for those frequencies between the TE_{10} and TM_{01} cutoff, where only TE_{10} propagation is guaranteed.

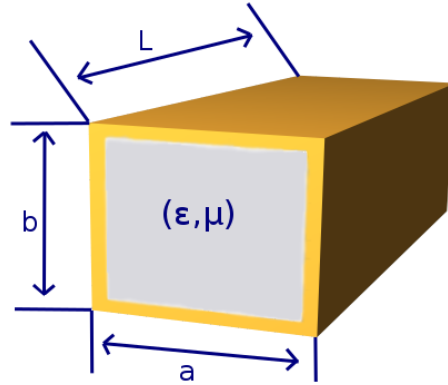


Figure 13.2: Rectangular waveguide of width a , height b and length L

As detailed in ??, the quadripole parameters of an ideal transmission line depends on the propagation constant (γ), the characteristic impedance of the line (Z_0) and its length (L).

The propagation constant is defined as $\gamma = \alpha + j \cdot \beta$, where α is the attenuation constant and β is the phase constant. The attenuation sources are the conductive loss due to the finite conductivity of the wall material (α_c) and the loss caused by dielectric material (α_d). In this sense, the attenuation constant is defined as $\alpha = \alpha_c + \alpha_d$

The conductor loss depends on the wall material and the geometry of the waveguide (width and height) whereas the dielectric loss mainly depends on the $\tan(\delta)$ parameter of the dielectric material:

$$\text{Conductive loss: } \alpha_c = \frac{R_s}{a^3 b \beta k \eta} (2b\pi^2 + a^3 k^2) \quad (13.7)$$

$$\text{Dielectric loss: } \alpha_d = \frac{k^2 \tan(\delta)}{2\beta} \quad (13.8)$$

where $k = \omega \sqrt{\mu\epsilon}$ is the wave number and $\eta = \sqrt{\frac{\mu}{\epsilon}}$ is the characteristic impedance of the dielectric material. The wall surface resistance (R_s) can be calculated as follows:

$$R_s = \sqrt{\frac{\omega \mu_0}{2\sigma}} \quad (13.9)$$

where σ is the conductivity of the walls.

The phase constant of the TE_{10} mode is given by:

$$\beta = \sqrt{k^2 - \left(\frac{\pi}{a}\right)^2} \quad (13.10)$$

and the wave impedance of the TE modes can be calculated as:

$$Z_0 = \frac{k \cdot \eta}{\beta} \quad (13.11)$$

Finally, the cutoff frequencies of the TE_{10} and TM_{11} modes are given by:

$$f_c^{TE_{10}} = \frac{1}{2\pi\sqrt{\mu\epsilon}} \sqrt{\left(\frac{\pi}{a}\right)^2} \quad (13.12)$$

$$f_c^{TM_{11}} = \frac{1}{2\pi\sqrt{\mu\epsilon}} \sqrt{\left(\frac{\pi}{a}\right)^2 + \left(\frac{\pi}{b}\right)^2} \quad (13.13)$$

Reference: [?], pages 110-121

13.3 Circular waveguide

As seen in ??, most of practical applications involving waveguides are designed to propagate the dominant mode only. For this reason, the Qucs circular waveguide component only models the

dominant mode (TE_{11}). Again, please notice that the simulation data is valid only for those frequencies between the TE_{11} and the TM_{01} cutoff, where only TE_{11} propagation is guaranteed.

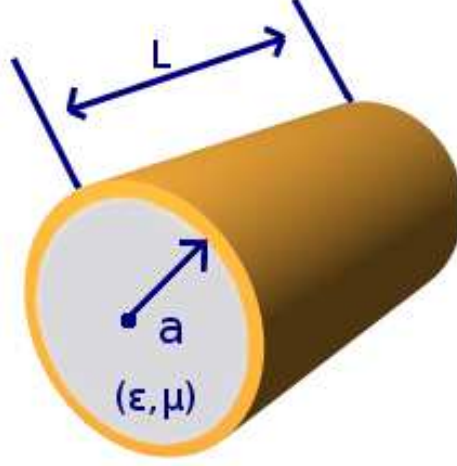


Figure 13.3: Circular waveguide of radius a and length L

As introduced in ?? the Qucs model extracts the quadripole parameters using the ideal transmission line model. The model parameters are obtained as follows:

The attenuation constant $\alpha = \alpha_c + \alpha_d$ involves both the loss in the conductive walls (α_c) and the dielectric loss (α_d):

$$\text{Conductive loss: } \alpha_c = \frac{R_s}{ak\eta\beta} \left(k_c^2 + \frac{k^2}{2.389} \right) \quad (13.14)$$

$$\text{Dielectric loss: } \alpha_d = \frac{k^2 \cdot \tan(\delta)}{2\beta} \quad (13.15)$$

where $k = \omega\sqrt{\mu\epsilon}$ is the wave number, $k_c = \frac{1.841}{a}$, $\tan(\delta)$ is a parameter that accounts for the loss of the dielectric material and surface resistance of the wall (R_s) is calculated according to ??.

The phase constant of the TE_{11} mode is given by:

$$\beta = \sqrt{k^2 - \left(\frac{1.841}{a} \right)^2} \quad (13.16)$$

As in the rectangular waveguide case, the characteristic impedance of a circular waveguide (TE modes) is given by:

$$Z_0 = \frac{k\eta}{\beta} \quad (13.17)$$

Finally, the cutoff frequencies of the TE_{11} and TM_{01} modes are:

$$f_c^{TE_{11}} = \frac{1.841}{2\pi a \sqrt{\mu\epsilon}} \quad (13.18)$$

$$f_c^{TM_{01}} = \frac{2.405}{2\pi a \sqrt{\mu\epsilon}} \quad (13.19)$$

Reference: [?], pages 121-130

13.4 Twisted pair

The twisted pair configurations as shown in fig. ?? provides good low frequency shielding. Undesired signals tend to be coupled equally into each line of the pair. A differential receiver will therefore completely cancel the interference.

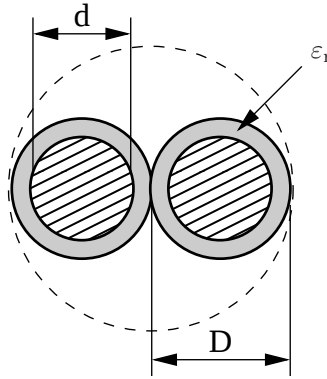


Figure 13.4: twisted pair configuration

13.4.1 Quasi-static model

According to P. Lefferson [?] the characteristic impedance and effective dielectric constant of a twisted pair can be calculated as follows.

$$Z_L = \frac{Z_{F0}}{\pi \cdot \sqrt{\epsilon_{r,eff}}} \cdot \text{acosh} \left(\frac{D}{d} \right) \quad (13.20)$$

$$\epsilon_{r,eff} = \epsilon_{r,1} + q \cdot (\epsilon_r - \epsilon_{r,1}) \quad (13.21)$$

with

$$q = 0.25 + 0.0004 \cdot \theta^2 \quad \text{and} \quad \theta = \text{atan}(T \cdot \pi \cdot D) \quad (13.22)$$

whereas θ is the pitch angle of the twist; the angle between the twisted pair's center line and the twist. It was found to be optimal for θ to be between 20° and 45° . T denotes the twists per length. Eq. (??) is valid for film insulations, for the softer PTFE material it should be modified as follows.

$$q = 0.25 + 0.001 \cdot \theta^2 \quad (13.23)$$

Assuming air as dielectric around the wires yields 1's replacing $\varepsilon_{r,1}$ in eq. (??). The wire's total length before twisting in terms of the number of turns N is

$$l = N \cdot \pi \cdot D \cdot \sqrt{1 + \frac{1}{\tan^2 \theta}} \quad (13.24)$$

13.4.2 Transmission losses

The propagation constant γ of a general transmission line is given by

$$\gamma = \sqrt{(R' + j\omega L') \cdot (G' + j\omega C')} \quad (13.25)$$

Using some transformations of the formula gives an expression with and without the angular frequency.

$$\begin{aligned} \gamma &= \sqrt{(R' + j\omega L') \cdot (G' + j\omega C')} \\ &= \sqrt{L'C'} \cdot \sqrt{\frac{R'G'}{L'C'} + j\omega \left(\frac{R'}{L'} + \frac{G'}{C'} \right) - \omega^2} \\ &= \sqrt{L'C'} \cdot \sqrt{\left(\frac{1}{2} \cdot \left(\frac{R'}{L'} + \frac{G'}{C'} \right) + j\omega \right)^2 - \frac{1}{4} \cdot \left(\frac{R'}{L'} + \frac{G'}{C'} \right)^2 + \frac{R'G'}{L'C'}} \end{aligned} \quad (13.26)$$

For high frequencies eq.(??) can be approximated to

$$\gamma \approx \sqrt{L'C'} \cdot \left(\frac{1}{2} \cdot \left(\frac{R'}{L'} + \frac{G'}{C'} \right) + j\omega \right) \quad (13.27)$$

Thus the real part of the propagation constant γ yields

$$\alpha = \text{Re} \{ \gamma \} = \sqrt{L'C'} \cdot \frac{1}{2} \cdot \left(\frac{R'}{L'} + \frac{G'}{C'} \right) \quad (13.28)$$

With

$$Z_L = \sqrt{\frac{L'}{C'}} \quad (13.29)$$

the expression in eq.(??) can be written as

$$\alpha = \alpha_c + \alpha_d = \frac{1}{2} \cdot \left(\frac{R'}{Z_L} + G' Z_L \right) \quad (13.30)$$

whereas α_c denotes the conductor losses and α_d the dielectric losses.

Conductor losses

The sheet resistance R' of a transmission line conductor is given by

$$R' = \frac{\rho}{A_{eff}} \quad (13.31)$$

whereas ρ is the specific resistance of the conductor material and A_{eff} the effective area of the conductor perpendicular to the propagation direction. At higher frequencies the area of the conductor is reduced by the skin effect. The skin depth is given by

$$\delta_s = \sqrt{\frac{\rho}{\pi \cdot f \cdot \mu}} \quad (13.32)$$

Thus the effective area of a single round wire yields

$$A_{eff} = \pi \cdot (r^2 - (r - \delta_s)^2) \quad (13.33)$$

whereas r denotes the radius of the wire. This means the overall conductor attenuation constant α_c for a single wire gives

$$\alpha_c = \frac{R'}{2 \cdot Z_L} = \frac{\rho}{2 \cdot Z_L \cdot \pi \cdot (r^2 - (r - \delta_s)^2)} \quad (13.34)$$

Dielectric losses

The dielectric losses are determined by the dielectric loss tangent.

$$\tan \delta_d = \frac{G'}{\omega C'} \quad \rightarrow \quad G' = \omega C' \cdot \tan \delta_d \quad (13.35)$$

With

$$C' = \frac{1}{\omega} \cdot \text{Im} \left\{ \frac{\gamma}{Z_L} \right\} \quad (13.36)$$

the equation (??) can be rewritten to

$$\begin{aligned} G' &= \frac{\beta}{Z_L} \cdot \tan \delta_d = \frac{\omega}{v_{ph} \cdot Z_L} \cdot \tan \delta_d \\ &= \frac{2\pi \cdot f \cdot \sqrt{\varepsilon_{r,eff}}}{c_0 \cdot Z_L} \cdot \tan \delta_d = \frac{2\pi \cdot \sqrt{\varepsilon_{r,eff}}}{\lambda_0 \cdot Z_L} \cdot \tan \delta_d \end{aligned} \quad (13.37)$$

whereas v_{ph} denotes the phase velocity, c_0 the speed of light, $\varepsilon_{r,eff}$ the effective dielectric constant and λ_0 the freespace wavelength. With these expressions at hand it is possible to find a formula for the dielectric losses of the transmission line.

$$\alpha_d = \frac{1}{2} \cdot G' Z_L = \frac{\pi \cdot \sqrt{\varepsilon_{r,eff}}}{\lambda_0} \cdot \tan \delta_d \quad (13.38)$$

Overall losses of the twisted pair configuration

Transmission losses consist of conductor losses, dielectric losses as well as radiation losses. The above expressions for the conductor and dielectric losses are considered to be first order approximations. The conductor losses have been derived for a single round wire. The overall conductor losses due to the twin wires must be doubled. The dielectric losses can be used as is. Radiation losses are neglected.

Chapter 14

Synthesizing circuits

14.1 Attenuators

Attenuators are used to damp a signal. Using pure ohmic resistors the circuit can be realized for a very high bandwidth, i.e. from DC to many GHz. The power attenuation $0 < L \leq 1$ is defined as:

$$L = \frac{P_{in}}{P_{out}} = \frac{V_{in}^2}{Z_{in}} \cdot \frac{Z_{out}}{V_{out}^2} = \left(\frac{V_{in}}{V_{out}} \right)^2 \cdot \frac{Z_{out}}{Z_{in}} \quad (14.1)$$

where P_{in} and P_{out} are the input and output power and V_{in} and V_{out} are the input and output voltages.

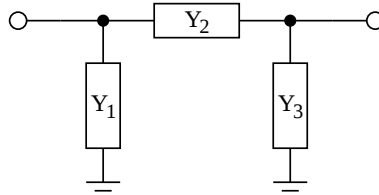


Figure 14.1: π -topology of an attenuator

Fig. ?? shows an attenuator using the π -topology. The conductances can be calculated as follows.

$$Y_2 = \frac{L - 1}{2 \cdot \sqrt{L \cdot Z_{in} \cdot Z_{out}}} \quad (14.2)$$

$$Y_1 = Y_2 \cdot \left(\sqrt{\frac{Z_{out}}{Z_{in}}} \cdot L - 1 \right) \quad (14.3)$$

$$Y_3 = Y_2 \cdot \left(\sqrt{\frac{Z_{in}}{Z_{out}}} \cdot L - 1 \right) \quad (14.4)$$

where Z_{in} and Z_{out} are the input and output reference impedances, respectively. The π -attenuator can be used for an impedance ratio of:

$$\frac{1}{L} \leq \frac{Z_{out}}{Z_{in}} \leq L \quad (14.5)$$

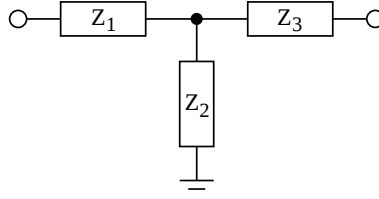


Figure 14.2: T-topology of an attenuator

Fig. ?? shows an attenuator using the T-topology. The resistances can be calculated as follows.

$$Z_2 = \frac{2 \cdot \sqrt{L \cdot Z_{in} \cdot Z_{out}}}{L - 1} \quad (14.6)$$

$$Z_1 = Z_{in} \cdot A - Z_2 \quad (14.7)$$

$$Z_3 = Z_{out} \cdot A - Z_2 \quad (14.8)$$

$$\text{with } A = \frac{L + 1}{L - 1} \quad (14.9)$$

where L is the attenuation ($0 < L \leq 1$) according to equation ?? and Z_{in} and Z_{out} are the input and output reference impedance, respectively. The T-attenuator can be used for an impedance ratio of:

$$\frac{Z_{out}}{Z_{in}} \leq \frac{(L + 1)^2}{4 \cdot L} \quad (14.10)$$

14.2 Filters

One of the most common tasks in microwave technologies is to extract a frequency band from others. Optimized filters exist in order to easily create a filter with an appropriate characteristic. The most popular ones are:

Name	Property
Bessel filter (Thomson filter)	as constant group delay as possible
Butterworth filter (power-term filter)	as constant amplitude transfer function as possible
Chebyshev filter type I	constant ripple in pass band
Chebyshev filter type II	constant ripple in stop band
Cauer filter (elliptical filter)	constant ripple in pass and stop band

From top to bottom the following properties increase:

- ringing of step response
- phase distortion
- steepness of amplitude transfer function at the beginning of the pass band

The order n of a filter denotes the number of poles of its (voltage) transfer function. It is:

$$\text{slope of asymptote} = \pm n \cdot 20\text{dB/decade} \quad (14.11)$$

Note that this equation holds for all filter characteristics, but there are big differences concerning the attenuation near the pass band.

14.2.1 LC ladder filters

The best possibility to realize a filters in VHF and UHF bands are LC ladder filters. The usual way to synthesize them is to first calculate a low-pass (LP) filter and afterwards transform it into a high-pass (HP), band-pass (BP) or band-stop (BS) filter. To do so, each component must be transformed into another.

In a low-pass filter, there are parallel capacitors C_{LP} and series inductors L_{LP} in alternating order. The other filter classes can be derived from it:

In a high-pass filter:

$$C_{LP} \rightarrow L_{HP} = \frac{1}{\omega_B^2 \cdot C_{LP}} \quad (14.12)$$

$$L_{LP} \rightarrow C_{HP} = \frac{1}{\omega_B^2 \cdot L_{LP}} \quad (14.13)$$

In a band-pass filter:

$$C_{LP} \rightarrow \text{parallel resonance circuit with} \quad (14.14)$$

$$C_{BP} = \frac{C_{LP}}{\Delta\Omega} \quad (14.15)$$

$$L_{BP} = \frac{\Delta\Omega}{\omega_1 \cdot \omega_2 \cdot C_{LP}} \quad (14.16)$$

$$L_{LP} \rightarrow \text{series resonance circuit with} \quad (14.17)$$

$$C_{BP} = \frac{\Delta\Omega}{\omega_1 \cdot \omega_2 \cdot L_{LP}} \quad (14.18)$$

$$L_{BP} = \frac{L_{LP}}{\Delta\Omega} \quad (14.19)$$

In a band-stop filter:

$$C_{LP} \rightarrow \text{series resonance circuit with} \quad (14.20)$$

$$C_{BP} = \frac{C_{LP}}{2 \cdot \left| \frac{\omega_2}{\omega_1} - \frac{\omega_1}{\omega_2} \right|} \quad (14.21)$$

$$L_{BP} = \frac{1}{\omega^2 \cdot \Delta\Omega \cdot C_{LP}} \quad (14.22)$$

$$L_{LP} \rightarrow \text{parallel resonance circuit with} \quad (14.23)$$

$$C_{BP} = \frac{1}{\omega^2 \cdot \Delta\Omega \cdot L_{LP}} \quad (14.24)$$

$$L_{BP} = \frac{L_{LP}}{2 \cdot \left| \frac{\omega_2}{\omega_1} - \frac{\omega_1}{\omega_2} \right|} \quad (14.25)$$

Where

$$\omega_1 \rightarrow \text{lower corner frequency of frequency band} \quad (14.26)$$

$$\omega_2 \rightarrow \text{upper corner frequency of frequency band} \quad (14.27)$$

$$\omega \rightarrow \text{center frequency of frequency band} \quad \omega = 0.5 \cdot (\omega_1 + \omega_2) \quad (14.28)$$

$$\Delta\Omega \rightarrow \Delta\Omega = \frac{|\omega_2 - \omega_1|}{\omega} \quad (14.29)$$

Butterworth

The k -th element of an n order Butterworth low-pass ladder filter is:

$$\text{capacitance:} \quad C_k = \frac{X_k}{Z_0} \quad (14.30)$$

$$\text{inductance:} \quad L_k = X_k \cdot Z_0 \quad (14.31)$$

$$\text{with} \quad X_k = \frac{2}{\omega_B} \cdot \sin \frac{(2 \cdot k + 1) \cdot \pi}{2 \cdot n} \quad (14.32)$$

The order of the Butterworth filter is dependent on the specifications provided by the user. These specifications include the edge frequencies and gains.

$$n = \frac{\log \left(\frac{10^{-0.1 \cdot \alpha_{stop}} - 1}{10^{-0.1 \cdot \alpha_{pass}} - 1} \right)}{2 \cdot \log \left(\frac{\omega_{stop}}{\omega_{pass}} \right)} \quad (14.33)$$

Chebyshev I

The equations for a Chebyshev type I filter are defined recursively. With R_{dB} being the passband ripple in decibel, the k -th element of an n order low-pass ladder filter is:

$$\text{capacitance:} \quad C_k = \frac{X_k}{Z_0} \quad (14.34)$$

$$\text{inductance:} \quad L_k = X_k \cdot Z_0 \quad (14.35)$$

$$\text{with} \quad X_k = \frac{2}{\omega_B} \cdot g_k \quad (14.36)$$

$$r = \sinh \left(\frac{1}{n} \cdot \operatorname{arsinh} \frac{1}{\sqrt{10^{R_{dB}/10} - 1}} \right) \quad (14.37)$$

$$a_k = \sin \frac{(2 \cdot k + 1) \cdot \pi}{2 \cdot n} \quad (14.38)$$

$$g_k = \begin{cases} \frac{a_k}{r} & \text{for } k = 0 \\ \frac{a_{k-1} \cdot a_k}{g_{k-1} \cdot \left(r^2 + \sin^2 \frac{k \cdot \pi}{n} \right)} & \text{for } k \geq 1 \end{cases} \quad (14.39)$$

$$X_k = \frac{2}{\omega_B} \cdot g_k \quad (14.40)$$

The order of the Chebyshev filter is dependent on the specifications provided by the user. The general form of the calculation for the order is the same as for the Butterworth, except that the

inverse hyperbolic cosine function is used in place of the common logarithm function.

$$n = \frac{\operatorname{sech}\left(\frac{10^{-0.1 \cdot \alpha_{stop}} - 1}{10^{-0.1 \cdot \alpha_{pass}} - 1}\right)}{2 \cdot \operatorname{sech}\left(\frac{\omega_{stop}}{\omega_{pass}}\right)} \quad (14.41)$$

Chebyshev II

Because of the nature of the derivation of the inverse Chebyshev approximation function from the standard Chebyshev approximation the calculation of the order (??) is the same.

14.2.2 Quarter wavelength filters

Quarter wavelength filters are based on the fact that short and open $\lambda/4$ stubs are equivalent to series and parallel resonant tanks, respectively. Once given the normalized g_i coefficients of the filter prototype, it is quite straightforward to obtain the impedance of the short/open stubs, depending on the desired filter mask (bandpass/notch). Figure ?? illustrates the topology of a quarter wavelength filter.

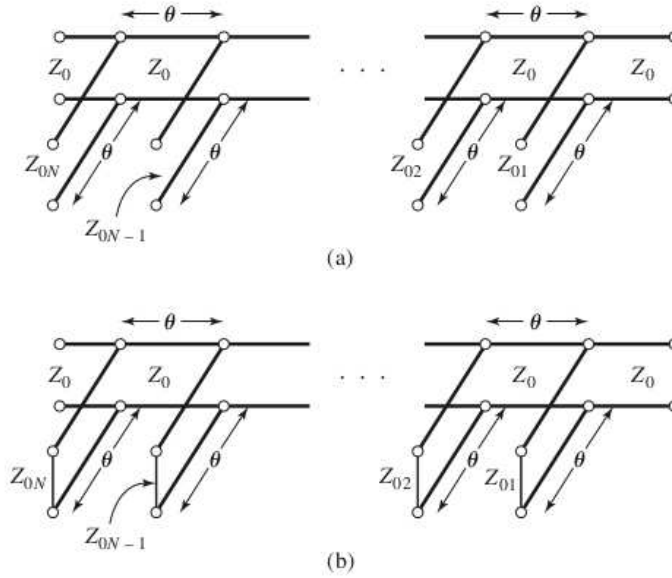


Figure 14.3: Quarter wavelength filters topology. a) Notch filter b) Bandpass filter [?]

The impedance of the i -th stub, Z_i , is given by:

- Bandpass (short circuit) $Z_i = \frac{\pi Z_0 \Delta}{4g_i}$
- Notch (open circuit) $Z_i = \frac{4Z_0}{g_i \pi \Delta}$

where Z_0 is the characteristic impedance of the filter (typically, 50Ω) and Δ is the relative bandwidth defined as $\Delta = \frac{f_2 - f_1}{f_c}$. The demonstration [?] is not included in this document in favour of greater simplicity.

It has to be said that quarter wavelength filters are narrowband by nature. Besides this, these kind of filters may not be feasible because of the stubs may require impractical widths (i.e. impedances) for proper operation.

14.3 Matching newtworks

The purpose of a matching network is to transform an impedance into another (e.g. convert the impedance of a device to the characteristic impedance of a transmission line, Z_0) and play an important role in RF circuit design. The matching networks can be implemented with lumped components (typically in applications below $< 1\text{GHz}$), but also with transmission lines (more common in the microwave band). The following sections show the basics of the most common matching network synthesis techniques as well as the design equations implemented in the Qucs impedance matching tool.

14.3.1 L-section

The *L-section* method combines two reactances (shunt and series) to match any complex load impedance, $Z_L = R_L + jX_L$ to a source impedance, Z_0 , (typically, $Z_0 = 50\Omega$) at a certain frequency. As shown in ??, eight different combinations of inductor and capacitor can be used depending on where Z_L is located on the Smith chart.

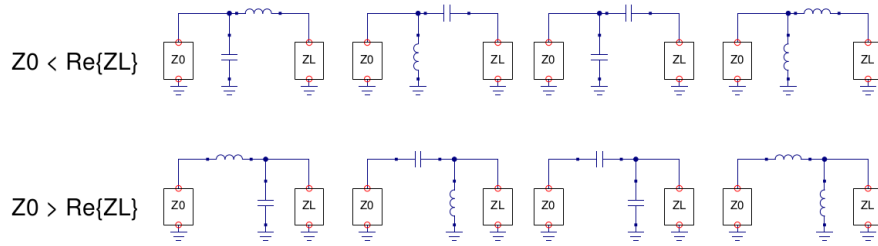


Figure 14.4: L-section topology

The series and the shunt reactances move Z_L towards Z_0 through the impedance and admittance circles respectively. As shown at ??, these eight combinations of series and shunt reactances provide the required degrees of freedom to move Z_L to any other point in the Smith chart.

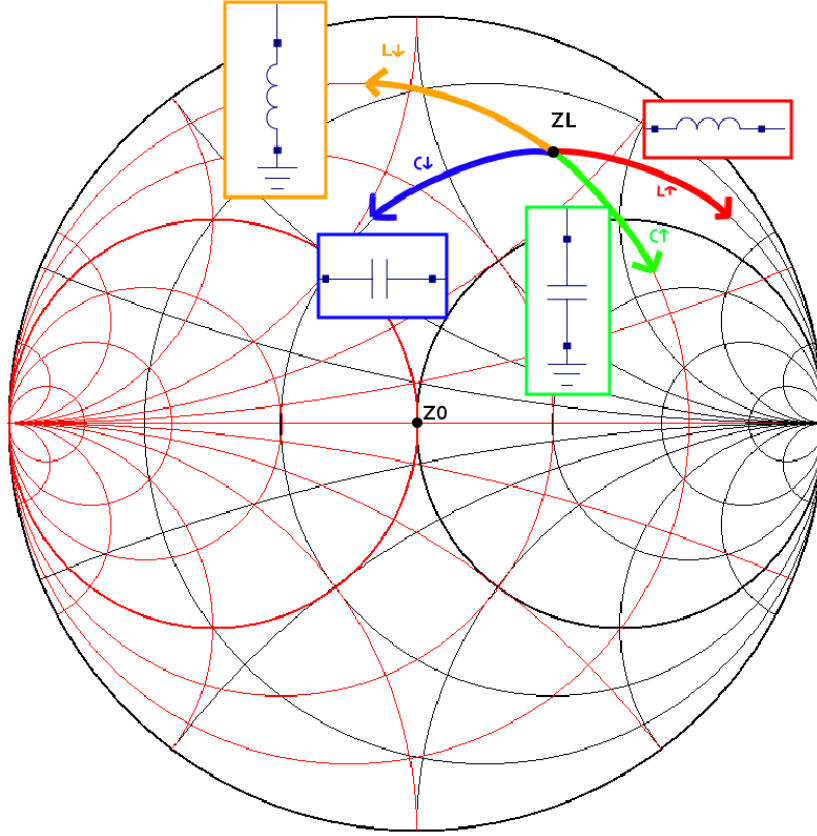


Figure 14.5: Impedance traslation produced by a series/shunt capacitor and inductor

The design equations for the *L-section* method are the following:

if $\frac{R_L}{Z_0} > 1$

$$B = \frac{X_L \pm \sqrt{\frac{R_L}{Z_0} \cdot \sqrt{R_L^2 + X_L^2} - Z_0 \cdot R_L}{R_L^2 + X_L^2} \quad (14.42)$$

$$X = \frac{1}{B} + \frac{X_L \cdot Z_0}{R_L} - \frac{Z_0}{B \cdot R_L} \quad (14.43)$$

on the counterpart, if $\frac{R_L}{Z_0} < 1$

$$X = \pm \sqrt{R_L \cdot (Z_0 - R_L)} - X_L \quad (14.44)$$

$$B = \pm \frac{(Z_0 - R_L)/R_L}{Z_0} \quad (14.45)$$

Reference: [?], pages 229-233.

14.3.2 Single stub

At microwave frequencies, transmission lines are usually preferred rather than lumped components. The simplest method to achieve narrowband matching with transmission lines is to use a stub and a transmission line.

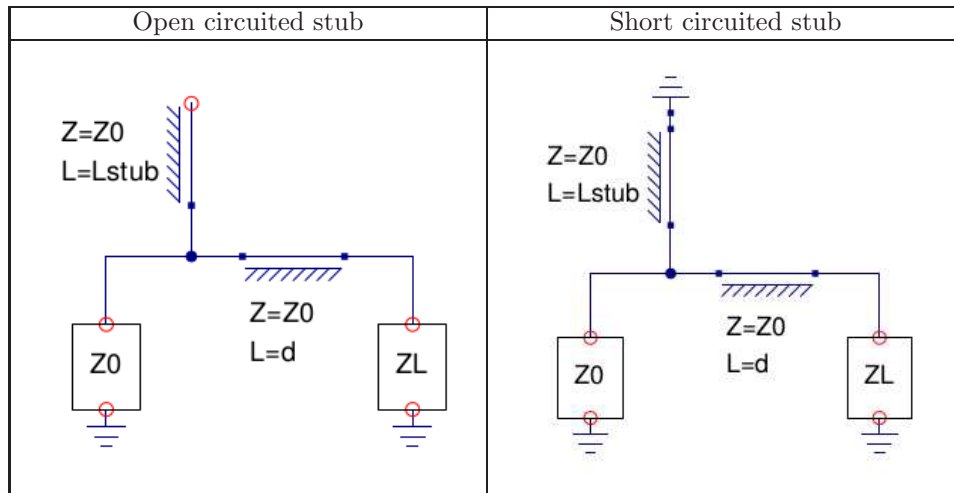


Table 14.1: Single stub matching

As ?? shows, a transmission line describes a circle around Z_0 whereas an open and a short circuited stub behave like a shunt capacitor and a shunt inductor respectively, which gives the necessary degrees of freedom to match any complex load, Z_L to Z_0 :

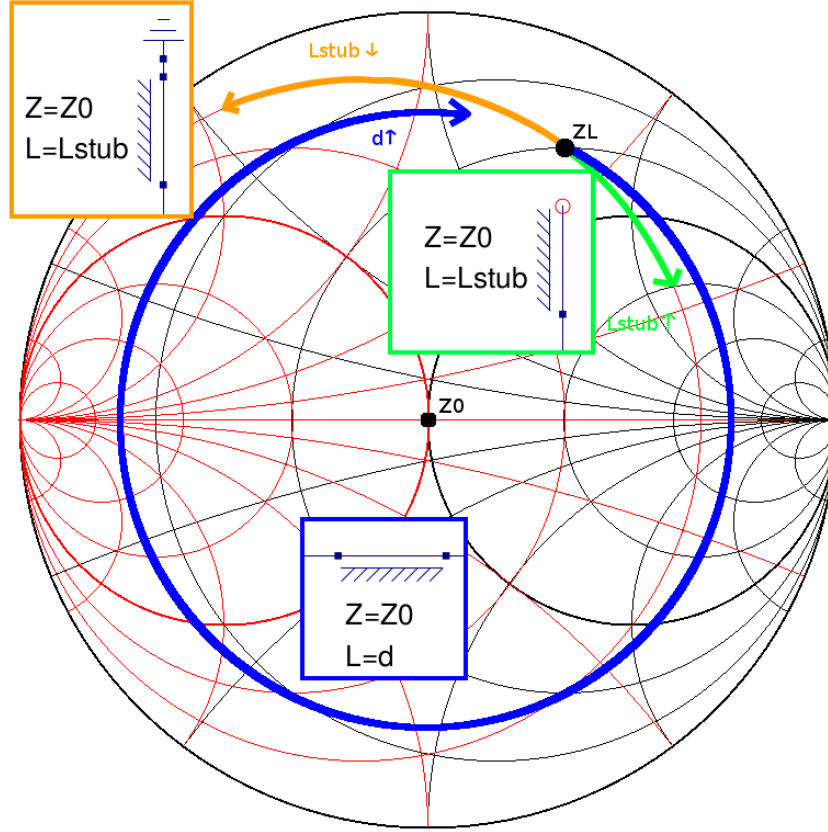


Figure 14.6: Impedance traslation caused by a transmission line and an open/short circuited stub

Below are detailed the design equations for the *single stub* technique, where the variables are the distance between the stub and the load, and the length of the stub. The impedance seen at the stub, $Z(d)$ is given by:

$$Z(d) = Z_0 \frac{Z_L + j \cdot Z_0 \cdot \tan(\beta \cdot d)}{Z_0 + j \cdot Z_L \cdot \tan(\beta \cdot d)} \quad (14.46)$$

where $\beta = \frac{2\pi}{\lambda}$ is the wave number and d the distance between the load and the stub. The admittance at the stub is: $Y(d) = G(d) + j \cdot B(d)$, where $G(d)$ is the conductance and $B(d)$ the susceptance at d . It can be shown that:

$$G(d) = \frac{R_L \cdot (1 + \tan^2(\beta \cdot d))}{R_L^2 + (X_L + Z_0 \cdot \tan(\beta \cdot d))^2} \quad (14.47)$$

and,

$$B(d) = \frac{R_L^2 \cdot \tan(\beta \cdot d) - (Z_0 - X_L \cdot \tan(\beta \cdot d)) \cdot (X_L + Z_0 \cdot \tan(\beta \cdot d))}{Z_0 \cdot (R_L^2 + (X_L + Z_0 \cdot \tan(\beta \cdot d))^2)} \quad (14.48)$$

The conductance at the stub should be $G(d) = 1/Z_0$ so as to match the real part of the load to Z_0 . Then, applying this condition and solving the previous equation for $t = \tan(\beta \cdot d)$, it gives:

$$R_L \neq Z_0 \quad t = \frac{X_L \pm \sqrt{(R_L/Z_0) \cdot ((Z_0 - R_L)^2 + X_L^2)}}{R_L - Z_0} \quad (14.49)$$

$$R_L = Z_0 \quad t = -X_L/(2 \cdot Z_0) \quad (14.50)$$

Now, the susceptance of the stub should be calculated to cancel the reactive part of the load impedance, so $B_{stub} = -B(d)$.

Since $t = \tan(\beta \cdot d)$, d can be calculated as:

$$\text{if } t \geq 0 \quad d = \frac{\lambda}{2\pi} \operatorname{atan}(t) \quad (14.51)$$

$$\text{if } t < 0 \quad d = \frac{\lambda}{2\pi} (\pi + \operatorname{atan}(t)) \quad (14.52)$$

The length of the stub is given by:

$$\text{Open circuited stub} \quad l_{open} = \frac{-\lambda}{2\pi} \operatorname{atan}(Z_0 \cdot B(d)) \quad (14.53)$$

$$\text{Short circuited stub} \quad l_{short} = \frac{\lambda}{2\pi} \operatorname{atan}\left(\frac{1}{Z_0 \cdot B(d)}\right) \quad (14.54)$$

Finally, if needed, the stubs may be splitted into balanced stubs as follows:

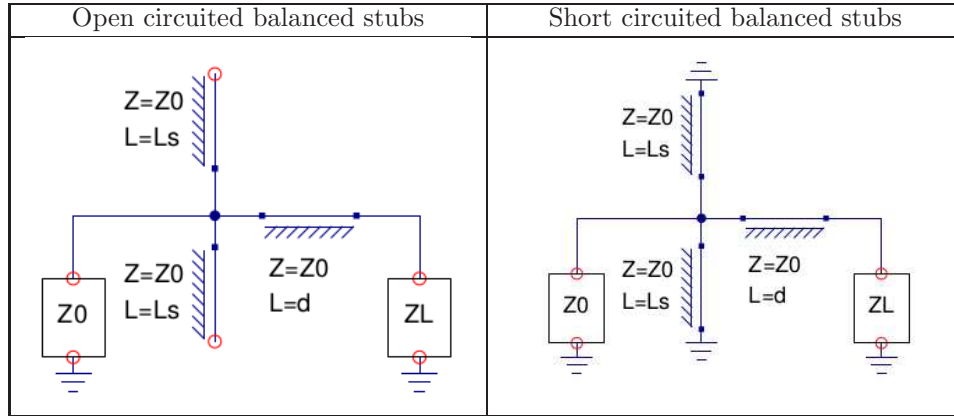


Table 14.2: Single stub matching (balanced stubs)

The length of the balanced stubs is given by the following equations [?]:

$$\text{Open circuited stub:} \quad l_{sb} = \frac{\lambda}{2\pi} \tan^{-1} \left(2 \tan \left(\frac{2\pi l_s}{\lambda} \right) \right) \quad (14.55)$$

$$\text{Short circuited stub:} \quad l_{sb} = \frac{\lambda}{2\pi} \tan^{-1} \left(\frac{1}{2} \tan \left(\frac{2\pi l_s}{\lambda} \right) \right) \quad (14.56)$$

where l_{sb} and l_s are the lengths of the unbalanced and balanced stubs respectively.

Reference: [?], pages 234-241.

14.3.3 Double stub

The distance between different blocks (e.g. amplifier stages, filters,...) in a board is determined by the transmission line, so the its length is not easily tunable in practice. Obviously, this fact locks one degree of freedom in the *single stub* technique, making it less suitable. On the counterpart, the *double stub* method keeps the length of the transmission line fixed, so tuning can be done simply by adjusting the length of the stubs. This is usually far more convenient since the stubs can be trimmed or elongated as needed. However, this technique comes with the disadvantage of having restrictions about the load to be matched, Z_L .

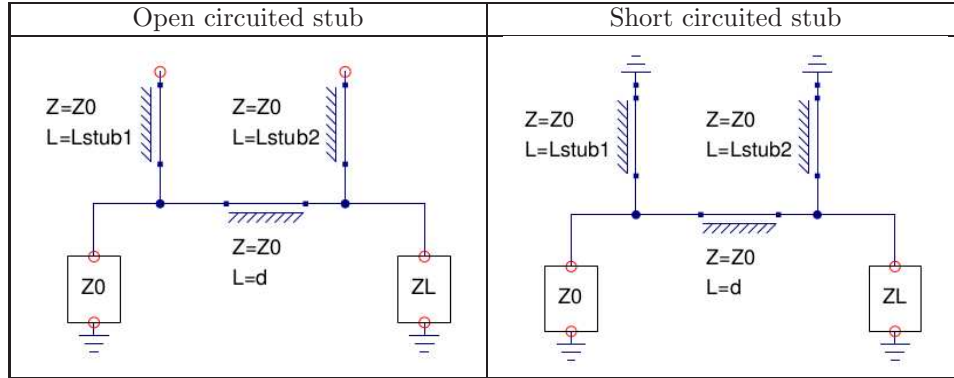


Table 14.3: Double stub matching

The admittance at the first stub is $Y_1 = G_L + j(B_L + B_{stub1})$, where Y_L , B_L and G_L have the same meaning as in the previous section. B_{stub1} is the susceptance of the first stub. The distance between the stubs, d , is typically fixed to $\lambda/8$. Under this condition, the admittance at the second stub is given by:

$$Y_2 = Y_0 \frac{G_L + j(B_L + B_{stub1} + Y_0 \cdot \tan(\beta \cdot d))}{Y_0 + j(G_L + j(B_L + B_{stub1})) \cdot \tan(\beta \cdot d)} \quad (14.57)$$

The impedance at the second stub must be Z_0 , so equating $Real\{Y_2\} = 1/Z_0$ it gives:

$$G_L = Y_0 \frac{1 + \tan^2(\beta \cdot d)}{2 \tan^2(\beta \cdot d)} \left\{ 1 \pm \sqrt{1 - \frac{4 \tan^2(\beta \cdot d) (Y_0 - (B_L + B_{stub1}) \cdot \tan(\beta \cdot d))^2}{Y_0^2 (1 + \tan^2(\beta \cdot d))^2}} \right\} \quad (14.58)$$

Notice here that those conductances which make the factor inside the square root negative cannot be matched. So, the *double stub* method is limited to load conductances which meet the following condition:

$$G_L \leq \frac{Y_0}{\sin^2(\beta \cdot d)} \quad (14.59)$$

The susceptance of the first stub can be calculated as:

$$B_{stub1} = -B_L + \frac{Y_0 \pm \sqrt{(1 + \tan^2(\beta \cdot d)) G_L \cdot Y_0 - (G_L \cdot \tan(\beta \cdot d))^2}}{\tan(\beta \cdot d)} \quad (14.60)$$

Similarly, the susceptance of the second stub is:

$$B_{stub_2} = \frac{\pm \sqrt{Y_0 \cdot G_L (1 + \tan^2(\beta \cdot d)) - (G_L \cdot \tan(\beta \cdot d))^2} + G_L \cdot Y_0}{G_L \cdot \tan(\beta \cdot d)} \quad (14.61)$$

Finally, the length of the stubs is given by:

$$l_{open} = \frac{\lambda}{2\pi} \text{atan}(Z_0 \cdot B_{stub_i}) \quad \text{Open stub} \quad (14.62)$$

$$l_{short} = \frac{-\lambda}{2\pi} \text{atan}\left(\frac{1}{Z_0 \cdot B_{stub_i}}\right) \quad \text{Short circuited stub} \quad (14.63)$$

where $i = \{1, 2\}$ denotes the number of the stub.

As in *single stub* method, the use of balanced stubs is possible:

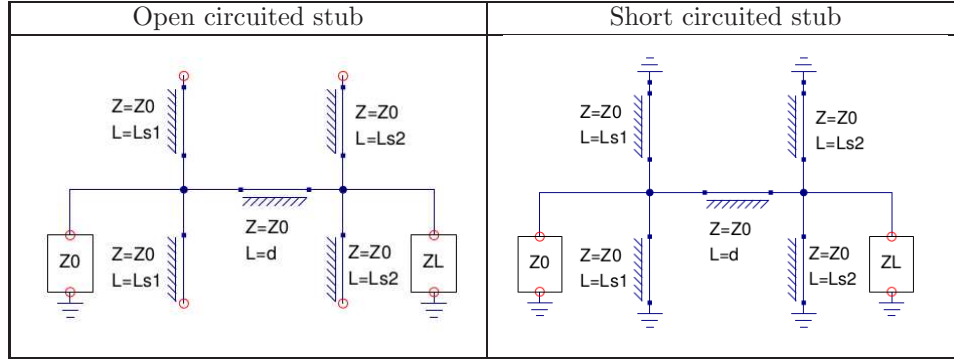


Table 14.4: Double stub matching (balanced stubs)

Reference [?], pages 241-246.

14.3.4 $\lambda/4 + \lambda/8$ matching

Another narrowband matching network technique is to combine a $\lambda/4$ and $\lambda/8$ transmission to match a complex impedance, Z_L to real source impedance. The $\lambda/8$ line converts the complex impedance to a real one and the $\lambda/4$ line converts that intermediate real impedance into the source impedance, Z_0 .

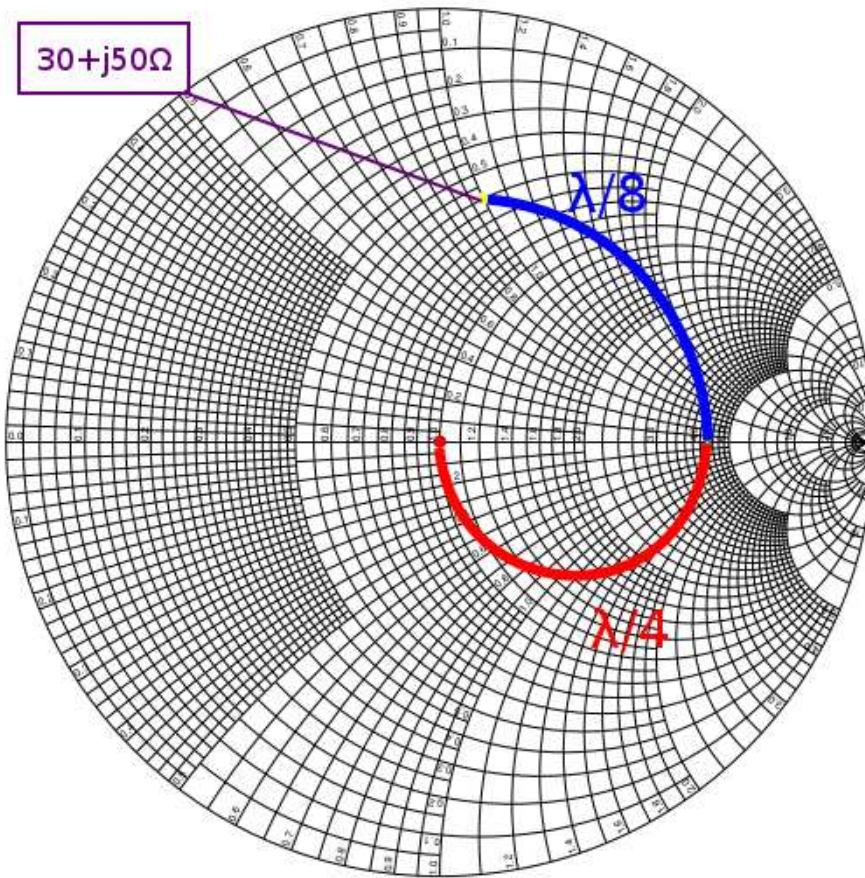


Figure 14.7: Example of $\lambda/4 + \lambda/8$ matching: $Z_S = 50\Omega$, $Z_L = 30 + j50\Omega$, $Z_{\lambda/4} = 102.6\Omega$, $Z_{\lambda/4} = 58.3\Omega$

Design equations:

$$Z_{\lambda/8} = \sqrt{R_L^2 + X_L^2} \quad (14.64)$$

$$Z_{\lambda/4} = \sqrt{\frac{R_S \cdot R_L \cdot Z_{\lambda/8}}{Z_{\lambda/8} - X_L}} \quad (14.65)$$

Reference: [?], pages 159-161.

14.3.5 Multistage $\lambda/4$ sections

Typically, broadband matching between real impedances at microwave frequencies is done with multisection transformers. This matching technique relies on the *Theory of the small reflections* which states that the overall reflection coefficient of N cascaded transmission lines of the same length can be approximated as:

$$\Gamma(\theta) = \sum_{n=0}^N \Gamma_n \cdot e^{-j \cdot 2n\theta} \quad (14.66)$$

where θ is the electrical length of a single transmission line (typically $\lambda/4$). The Γ_n weights define the shape of $\Gamma(\theta)$ and they are tailored so as to get a binomial or a Chebyshev overall response.

Once given $\Gamma_n \forall n \in [0, N]$, it is possible to calculate the impedance of each transmission line:

$$Z_{n+1} \sim Z_n \cdot e^{2 \cdot \Gamma_n} \quad (14.67)$$

where n indicates the section number.

Binomial weighting

Γ_n is defined as:

$$\Gamma_n = A \cdot C_n^N \quad \forall n \in [0, N] \quad (14.68)$$

where $A = 2^{-N} \frac{Z_L - Z_0}{Z_L + Z_0}$ and $C_n^N = \frac{N!}{(N-n)!n!}$

Chebyshev weighting

The overall reflection coefficient is set to:

$$\Gamma(\theta) = A \cdot e^{-jN\theta} \cdot T_N(\sec(\theta_m)\cos(\theta)) \quad (14.69)$$

where:

$$A = \frac{Z_L - Z_0}{Z_L + Z_0} \cdot \frac{1}{T_N(\sec(\theta_m))} \quad (14.70)$$

$$\sec(\theta_m) = \cosh\left(\frac{1}{N} \operatorname{acosh}\left(\frac{1}{\Gamma_m} \left|\frac{Z_L - Z_0}{Z_L + Z_0}\right|\right)\right) \quad (14.71)$$

and $T_N(x)$ is the Chebyshev polynomial of order N . The Chebyshev polynomials are defined such that:

$$T_n(x) = 2xT_{n-1}(x) - T_{n-2}(x) \quad (14.72)$$

$$T_1(x) = x \quad (14.73)$$

$$T_0(x) = 1 \quad (14.74)$$

$T_N(\sec(\theta_m)\cos(\theta))$ involves terms such as $\cos^N(x)$, $\cos^{N-1}(x)$,... In order to equate this terms to the Γ_n terms in ??, it is necessary to expand $\cos^N(x)$ terms into sums of terms such as $\cos((n-2m)\theta)$

Reference: [?], pages 250-260.

14.3.6 Cascaded L-sections

An alternative method to achieve broadband matching between real impedances is to cascade N (lowpass) *L-section* stages. In this sense, the *Cascaded L-sections* method does $N-1$ matches between intermediate real impedances such that $R_i > R_{i+1} \forall i \in N$.

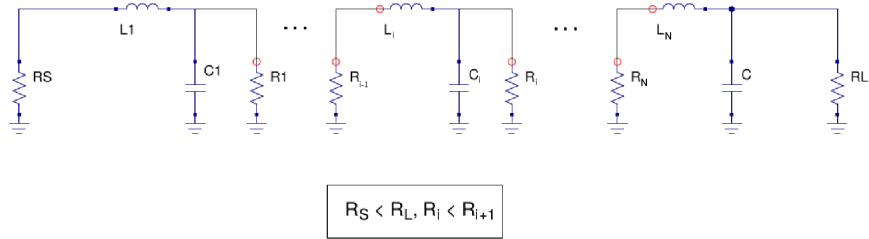


Figure 14.8: Cascaded L-sections method

Depending on the values of the intermediate impedances, the bandwidth of the matching network will vary. In order to maximize it, the ratios of consecutive virtual resistances should be equal:

$$\frac{R_S}{R_1} = \dots = \frac{R_i}{R_{i+1}} = \dots = \frac{R_{N-1}}{R_L} \quad (14.75)$$

Thus, the intermediate virtual resistance is given by:

$$R_i = R_S^{\frac{N-i}{N}} \cdot R_L^{\frac{i}{N}}, \forall i \in [1, N-1] \quad (14.76)$$

Once having R_i , the values of the inductors and the can be determined:

$$Q_i = \sqrt{\frac{R_{i-1}}{R_i}}, \quad (14.77)$$

$$C_i = \frac{Q_i}{\omega R_{i-1}}, \quad (14.78)$$

$$L_i = \frac{Q_i R_i}{\omega}, \forall i \in [1, N-1] \quad (14.79)$$

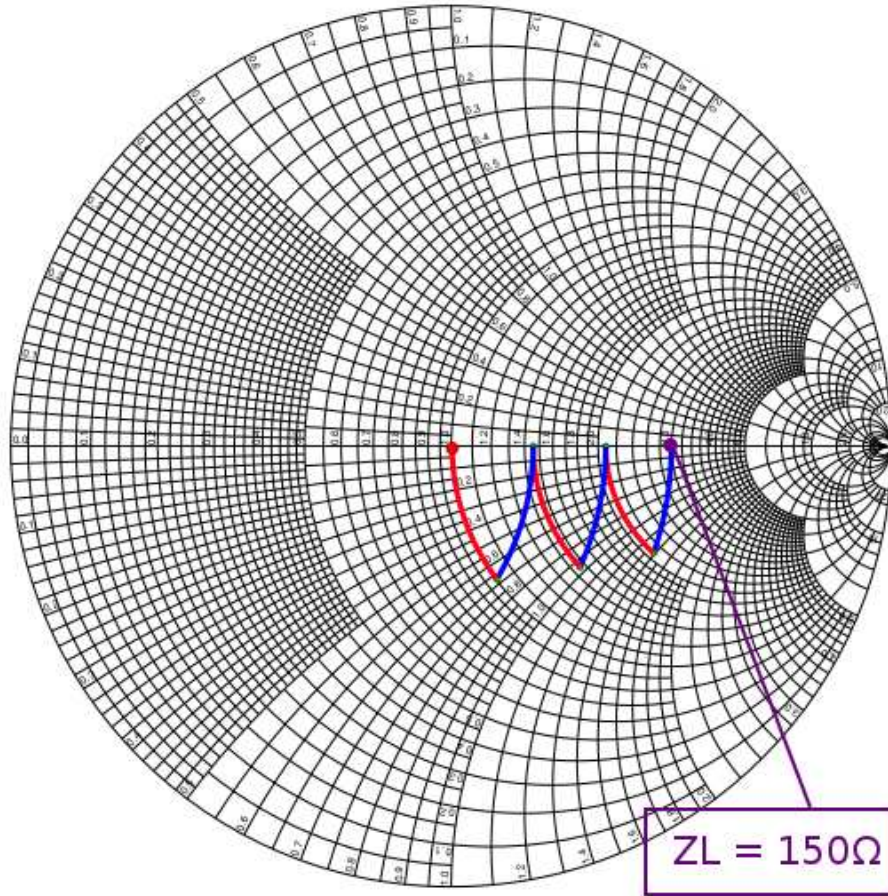


Figure 14.9: Example of *cascaded L-sections* matching: $Z_S = 50\Omega$, $Z_L = 150\Omega$, $L_i = \{5.3, 7.6, 11\}nH$, $C_i = \{1.47, 1, 0.7\}pF$

Reference: [?], pages: 169-172.

Chapter 15

Mathematical background

15.1 N-port matrix conversions

When dealing with n-port parameters it may be necessary or convenient to convert them into other matrix representations used in electrical engineering. The following matrices and notations are used in the transformation equations.

$$[\underline{X}]^{-1} = \text{inverted matrix of } [\underline{X}]$$

$$[\underline{X}]^* = \text{complex conjugated matrix of } [\underline{X}]$$

$$[\underline{E}] = \begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{bmatrix} \text{identity matrix}$$

$$[\underline{S}] = \text{S-parameter matrix}$$

$$[\underline{Z}] = \text{impedance matrix}$$

$$[\underline{Y}] = \text{admittance matrix}$$

$$[\underline{Z}_{ref}] = \begin{pmatrix} \underline{Z}_{0,1} & 0 & \dots & 0 \\ 0 & \underline{Z}_{0,2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \underline{Z}_{0,N} \end{pmatrix}$$

$$\underline{Z}_{0,n} = \text{reference impedance of port } n$$

$$[\underline{G}_{ref}] = \begin{pmatrix} G_1 & 0 & \dots & 0 \\ 0 & G_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & G_N \end{pmatrix}$$

$$G_n = \frac{1}{\sqrt{\text{Re} |\underline{Z}_{0,n}|}}$$

15.1.1 Renormalization of S-parameters to different port impedances

During S-parameter usage it sometimes appears to have not all components in a circuit normalized to the same impedance. But calculations can only be performed with all ports being normalized to the same impedance. In the field of high frequency techniques this is usually 50Ω . In order to transform to different port impedances, the following computation must be applied to the resulting S-parameter matrix (valid only for real reference impedances).

$$[\underline{S}'] = [\underline{A}]^{-1} \cdot ([\underline{S}] - [\underline{R}]) \cdot ([\underline{E}] - [\underline{R}] \cdot [\underline{S}])^{-1} \cdot [\underline{A}] \quad (15.1)$$

With

Z_n = reference impedance of port n after the normalizing process
 $Z_{n,before}$ = reference impedance of port n before the normalizing process
 $[S]$ = original S-parameter matrix
 $[S']$ = recalculated scattering matrix

$$[R] = \begin{pmatrix} \underline{r}(Z_1) & 0 & \dots & 0 \\ 0 & \underline{r}(Z_2) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \underline{r}(Z_N) \end{pmatrix} \text{ reflection coefficient matrix}$$

$$\underline{r}(Z_n) = \frac{Z_n - Z_{n,before}}{Z_n + Z_{n,before}}$$

$$[A] = \begin{pmatrix} A_1 & 0 & \dots & 0 \\ 0 & A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & A_N \end{pmatrix}$$

$$A_n = \sqrt{\frac{Z_n}{Z_{n,before}}} \cdot \frac{2Z_{n,before}}{Z_n + Z_{n,before}}$$

15.1.2 Transformations of n-Port matrices

S-parameter, admittance and impedance matrices are not limited to One- or Two-Port definitions. They are defined for an arbitrary number of ports. The following section contains transformation formulas forth and back each matrix representation.

Converting a scattering parameter matrix to an impedance matrix is done by the following formula.

$$[Z] = [G_{ref}]^{-1} \cdot ([E] - [S])^{-1} \cdot ([S] \cdot [Z_{ref}] + [Z_{ref}]) \cdot [G_{ref}] \quad (15.2)$$

$$= [G_{ref}]^{-1} \cdot ([E] - [S])^{-1} \cdot ([S] + [E]) \cdot [Z_{ref}] \cdot [G_{ref}] \quad (15.3)$$

Converting a scattering parameter matrix to an admittance matrix can be achieved by computing the following formula.

$$[Y] = [G_{ref}]^{-1} \cdot ([S] \cdot [Z_{ref}] + [Z_{ref}])^{-1} \cdot ([E] - [S]) \cdot [G_{ref}] \quad (15.4)$$

$$= [G_{ref}]^{-1} \cdot [Z_{ref}]^{-1} \cdot ([S] + [E])^{-1} \cdot ([E] - [S]) \cdot [G_{ref}] \quad (15.5)$$

Converting an impedance matrix to a scattering parameter matrix is done by the following formula.

$$[S] = [G_{ref}] \cdot ([Z] - [Z_{ref}]) \cdot ([Z] + [Z_{ref}])^{-1} \cdot [G_{ref}]^{-1} \quad (15.6)$$

Converting an admittance matrix to a scattering parameter matrix is done by the following formula.

$$[\underline{S}] = [\underline{G}_{ref}] \cdot ([\underline{E}] - [\underline{Z}_{ref}] \cdot [\underline{Y}]) \cdot ([\underline{E}] + [\underline{Z}_{ref}] \cdot [\underline{Y}])^{-1} \cdot [\underline{G}_{ref}]^{-1} \quad (15.7)$$

Converting an impedance matrix to an admittance matrix is done by the following simple formula.

$$[\underline{Y}] = [\underline{Z}]^{-1} \quad (15.8)$$

Converting an admittance matrix to an impedance matrix is done by the following simple formula.

$$[\underline{Z}] = [\underline{Y}]^{-1} \quad (15.9)$$

15.1.3 Two-Port transformations

Two-Port matrix conversion based on current and voltage

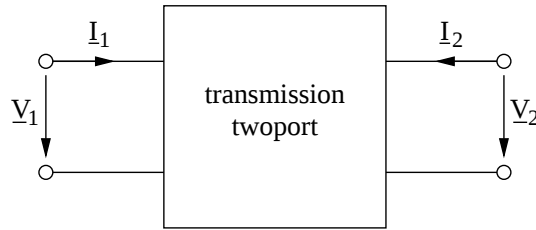


Figure 15.1: twoport definition using current and voltage

There are five different matrix forms for the correlations between the quantities at the transmission twoport shown in fig. ??, each having its special meaning when connecting twoports with each other.

- Y-parameters (also called admittance parameters)

$$\begin{pmatrix} I_1 \\ I_2 \end{pmatrix} = \begin{pmatrix} Y_{11} & Y_{12} \\ Y_{21} & Y_{22} \end{pmatrix} \cdot \begin{pmatrix} V_1 \\ V_2 \end{pmatrix} \quad (15.10)$$

- Z-parameters (also called impedance parameters)

$$\begin{pmatrix} V_1 \\ V_2 \end{pmatrix} = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix} \cdot \begin{pmatrix} I_1 \\ I_2 \end{pmatrix} \quad (15.11)$$

- H-parameters (also called hybrid parameters)

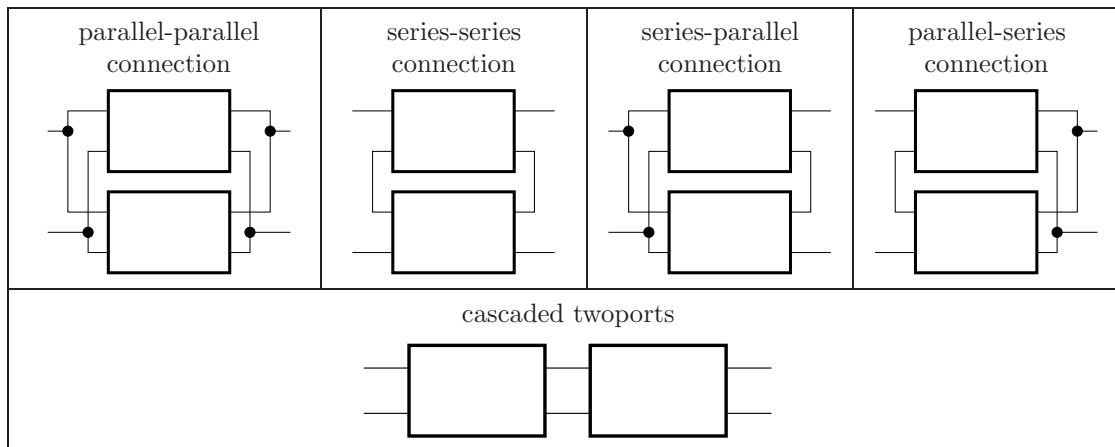
$$\begin{pmatrix} V_1 \\ I_2 \end{pmatrix} = \begin{pmatrix} H_{11} & H_{12} \\ H_{21} & H_{22} \end{pmatrix} \cdot \begin{pmatrix} I_1 \\ V_2 \end{pmatrix} \quad (15.12)$$

- G-parameters (also called C-parameters)

$$\begin{pmatrix} I_1 \\ V_2 \end{pmatrix} = \begin{pmatrix} G_{11} & G_{12} \\ G_{21} & G_{22} \end{pmatrix} \cdot \begin{pmatrix} V_1 \\ I_2 \end{pmatrix} \quad (15.13)$$

- A-parameters (also called chain or ABCD-parameters)

$$\begin{pmatrix} V_1 \\ I_1 \end{pmatrix} = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \cdot \begin{pmatrix} V_2 \\ -I_2 \end{pmatrix} \quad (15.14)$$



Basically there are five different kinds of twoport connections. Using the corresponding twoport matrix representations, complicated networks can be analysed by connecting elementary twoports. The linear correlations between the complex currents and voltages rms values of a twoport are described by four complex twoport parameters (i.e. the twoport matrix). These parameters are used to describe the AC behaviour of the twoport.

- parallel-parallel connection: use Y-parameters: $Y = Y_1 + Y_2$
- series-series connection: use Z-parameters: $Z = Z_1 + Z_2$
- series-parallel connection: use H-parameters: $H = H_1 + H_2$
- parallel-series connection: use G-parameters: $G = G_1 + G_2$
- chain connection: use A-parameters: $A = A_1 \cdot A_2$

	A	Y	Z	H	G
A	$\begin{matrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{matrix}$	$\begin{matrix} \frac{-Y_{22}}{Y_{21}} & \frac{-1}{Y_{21}} \\ \frac{-\Delta Y}{Y_{21}} & \frac{-Y_{11}}{Y_{21}} \end{matrix}$	$\begin{matrix} \frac{Z_{11}}{Z_{21}} & \frac{\Delta Z}{Z_{21}} \\ 1 & \frac{Z_{22}}{Z_{21}} \end{matrix}$	$\begin{matrix} \frac{-\Delta H}{H_{21}} & \frac{-H_{11}}{H_{21}} \\ \frac{-H_{22}}{H_{21}} & \frac{-1}{H_{21}} \end{matrix}$	$\begin{matrix} \frac{1}{G_{21}} & \frac{G_{22}}{G_{21}} \\ \frac{G_{11}}{G_{21}} & \frac{\Delta G}{G_{21}} \end{matrix}$
Y	$\begin{matrix} \frac{A_{22}}{A_{12}} & \frac{-\Delta A}{A_{12}} \\ -1 & \frac{A_{11}}{A_{12}} \\ \frac{A_{22}}{A_{12}} & \frac{A_{11}}{A_{12}} \end{matrix}$	$\begin{matrix} Y_{11} & Y_{12} \\ Y_{21} & Y_{22} \end{matrix}$	$\begin{matrix} \frac{Z_{22}}{\Delta Z} & \frac{-Z_{12}}{\Delta Z} \\ -\frac{Z_{21}}{\Delta Z} & \frac{Z_{11}}{\Delta Z} \end{matrix}$	$\begin{matrix} \frac{1}{H_{11}} & \frac{-H_{12}}{H_{11}} \\ \frac{H_{21}}{H_{11}} & \frac{\Delta H}{H_{11}} \end{matrix}$	$\begin{matrix} \frac{\Delta G}{G_{22}} & \frac{G_{12}}{G_{22}} \\ -\frac{G_{21}}{G_{22}} & \frac{1}{G_{22}} \end{matrix}$
Z	$\begin{matrix} \frac{A_{11}}{A_{21}} & \frac{\Delta A}{A_{21}} \\ 1 & \frac{A_{22}}{A_{21}} \\ \frac{A_{11}}{A_{21}} & \frac{\Delta A}{A_{21}} \end{matrix}$	$\begin{matrix} \frac{Y_{22}}{\Delta Y} & \frac{-Y_{12}}{\Delta Y} \\ -\frac{Y_{21}}{\Delta Y} & \frac{Y_{11}}{\Delta Y} \end{matrix}$	$\begin{matrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{matrix}$	$\begin{matrix} \frac{\Delta H}{H_{22}} & \frac{H_{12}}{H_{22}} \\ \frac{-H_{21}}{H_{22}} & \frac{1}{H_{22}} \end{matrix}$	$\begin{matrix} \frac{1}{G_{11}} & \frac{-G_{12}}{G_{11}} \\ \frac{G_{21}}{G_{11}} & \frac{\Delta G}{G_{11}} \end{matrix}$
H	$\begin{matrix} \frac{A_{12}}{A_{22}} & \frac{\Delta A}{A_{22}} \\ -1 & \frac{A_{21}}{A_{22}} \\ \frac{A_{12}}{A_{22}} & \frac{\Delta A}{A_{22}} \end{matrix}$	$\begin{matrix} \frac{1}{Y_{11}} & \frac{-Y_{12}}{Y_{11}} \\ \frac{Y_{21}}{Y_{11}} & \frac{\Delta Y}{Y_{11}} \end{matrix}$	$\begin{matrix} \frac{\Delta Z}{Z_{22}} & \frac{Z_{12}}{Z_{22}} \\ -\frac{Z_{21}}{Z_{22}} & \frac{1}{Z_{22}} \end{matrix}$	$\begin{matrix} H_{11} & H_{12} \\ H_{21} & H_{22} \end{matrix}$	$\begin{matrix} \frac{G_{22}}{\Delta G} & \frac{-G_{12}}{\Delta G} \\ -\frac{G_{21}}{\Delta G} & \frac{G_{11}}{\Delta G} \end{matrix}$
G	$\begin{matrix} \frac{A_{21}}{A_{11}} & \frac{-\Delta A}{A_{11}} \\ 1 & \frac{A_{12}}{A_{11}} \\ \frac{A_{21}}{A_{11}} & \frac{A_{12}}{A_{11}} \end{matrix}$	$\begin{matrix} \frac{\Delta Y}{Y_{22}} & \frac{Y_{12}}{Y_{22}} \\ -\frac{Y_{21}}{Y_{22}} & \frac{1}{Y_{22}} \end{matrix}$	$\begin{matrix} \frac{1}{Z_{11}} & \frac{-Z_{12}}{Z_{11}} \\ \frac{Z_{21}}{Z_{11}} & \frac{\Delta Z}{Z_{11}} \end{matrix}$	$\begin{matrix} \frac{H_{22}}{\Delta H} & \frac{-H_{12}}{\Delta H} \\ \frac{-H_{21}}{\Delta H} & \frac{H_{11}}{\Delta H} \end{matrix}$	$\begin{matrix} G_{11} & G_{12} \\ G_{21} & G_{22} \end{matrix}$

Two-Port matrix conversion based on signal waves

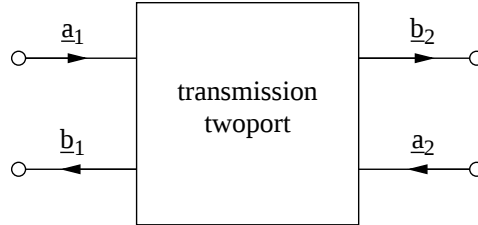


Figure 15.2: twoport definition using signal waves

There are two main different matrix forms for the correlations between the quantities at the transmission twoport shown in fig. ??.

- S-parameters (also called scattering parameters)

$$\begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} \underline{S}_{11} & \underline{S}_{12} \\ \underline{S}_{21} & \underline{S}_{22} \end{pmatrix} \cdot \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \quad (15.15)$$

the S-parameters matrix relates the output waves b_n (dependent variables) and the input waves a_n (independent variables) and is customarily defined with the waves ports in ascending order.

- T-parameters (also called transfer scattering parameters)
one definition of the T-parameters is

$$\begin{pmatrix} \underline{b}_1 \\ \underline{a}_1 \end{pmatrix} = \begin{pmatrix} \underline{T}_{11} & \underline{T}_{12} \\ \underline{T}_{21} & \underline{T}_{22} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_2 \\ \underline{b}_2 \end{pmatrix} \quad (15.16)$$

here the T-parameters relate the input *port* waves and the output *port* waves. Note that there are other possible definitions, depending on the choice of the dependent and independent variables and on the ports ordering used.

According to Janusz A. Dobrowolski [?] the following table contains the matrix transformation formulas for the T parameters definition above.

	S	T
S	$\begin{matrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{matrix}$	$\begin{matrix} \frac{T_{12}}{T_{22}} & \frac{\Delta T}{T_{22}} \\ \frac{1}{T_{22}} & \frac{-T_{21}}{T_{22}} \end{matrix}$
T	$\begin{matrix} \frac{-\Delta S}{S_{21}} & \frac{S_{11}}{S_{21}} \\ \frac{-S_{22}}{S_{21}} & \frac{1}{S_{21}} \end{matrix}$	$\begin{matrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{matrix}$

The T-parameters matrix allows to easily obtain the overall network response of a network cascade, a process called *embedding*. For the two cascaded twoports of fig. ??

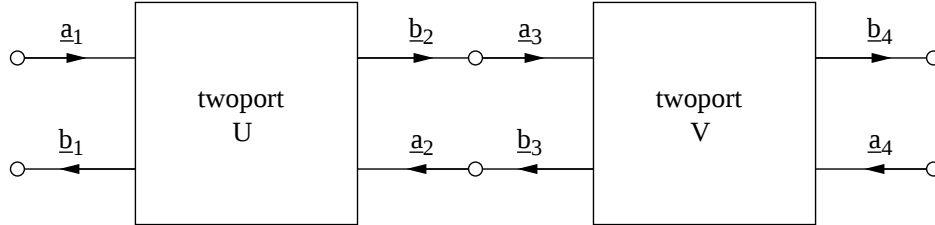


Figure 15.3: twoport cascade

the relationships between the various waves at the networks ports can be written as

$$\begin{pmatrix} \underline{b}_1 \\ \underline{a}_1 \end{pmatrix} = \begin{pmatrix} \underline{T}_{U11} & \underline{T}_{U12} \\ \underline{T}_{U21} & \underline{T}_{U22} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_2 \\ \underline{b}_2 \end{pmatrix} \quad \text{and} \quad \begin{pmatrix} \underline{b}_3 \\ \underline{a}_3 \end{pmatrix} = \begin{pmatrix} \underline{T}_{V11} & \underline{T}_{V12} \\ \underline{T}_{V21} & \underline{T}_{V22} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_4 \\ \underline{b}_4 \end{pmatrix} \quad (15.17)$$

since the output waves of the twoport U are the input waves of the twoport V, that is

$$\begin{pmatrix} \underline{a}_2 \\ \underline{b}_2 \end{pmatrix} = \begin{pmatrix} \underline{b}_3 \\ \underline{a}_3 \end{pmatrix} \quad (15.18)$$

the relationship between the input port waves and the output port waves of the cascade connection can be written

$$\begin{pmatrix} \underline{b}_1 \\ \underline{a}_1 \end{pmatrix} = \begin{pmatrix} \underline{T}_{U11} & \underline{T}_{U12} \\ \underline{T}_{U21} & \underline{T}_{U22} \end{pmatrix} \cdot \begin{pmatrix} \underline{T}_{V11} & \underline{T}_{V12} \\ \underline{T}_{V21} & \underline{T}_{V22} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_4 \\ \underline{b}_4 \end{pmatrix} \quad (15.19)$$

from which follows that the overall transfer scattering parameters matrix T is given by

$$\underline{T} = \underline{T}_U \cdot \underline{T}_V \quad (15.20)$$

Similarly, for the cascade of any number of twoports the overall T-parameters matrix is simply given by the product of the individual networks T parameters.

Incidentally, note that the cascade response can also be computed directly from the individual networks S parameters using the Redheffer *star product* [?].

De-embedding

In practice is also often useful to be able to remove the effect of a network in a network cascade, a process which is called *de-embedding*. This can also be done using the T-parameters matrix: from the overall transfer scattering parameters $\underline{T}$ given by eq. (??) the T-parameters matrix of the network V can be obtained from

$$\underline{T}_V = \underline{T}_U^{-1} \cdot \underline{T} \quad (15.21)$$

which can be seen as the result of cascading the inverse network (or *anti-network*) of U with the original composite network.

For a twoport, the S-parameters matrix of the inverse network of U is

$$\underline{S}_U^{(inv)} = \frac{1}{\Delta S_U} \begin{pmatrix} \underline{S}_{U11} & -\underline{S}_{U21} \\ -\underline{S}_{U12} & \underline{S}_{U22} \end{pmatrix} \quad (15.22)$$

Four- and N-port matrix conversion based on signal waves

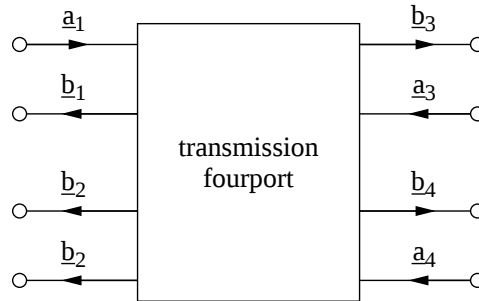


Figure 15.4: four-port network definition using signal waves

Similarly to the two-port case, two main matrix forms are used to express the relationships between the signal waves for network with four or more ports:

- S-parameters (scattering parameters), relating the output waves b_n to the input waves a_n

$$\begin{pmatrix} \underline{b}_1 \\ \underline{b}_2 \\ \underline{b}_3 \\ \underline{b}_4 \end{pmatrix} = \begin{pmatrix} \underline{S}_{11} & \underline{S}_{12} & \underline{S}_{13} & \underline{S}_{14} \\ \underline{S}_{21} & \underline{S}_{22} & \underline{S}_{23} & \underline{S}_{24} \\ \underline{S}_{31} & \underline{S}_{32} & \underline{S}_{33} & \underline{S}_{34} \\ \underline{S}_{41} & \underline{S}_{42} & \underline{S}_{43} & \underline{S}_{44} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_1 \\ \underline{a}_2 \\ \underline{a}_3 \\ \underline{a}_4 \end{pmatrix} \quad (15.23)$$

customarily defined with the waves ports in ascending order.

- T-parameters (transfer scattering parameters), relating the input *ports* waves and the output *ports* waves

one definition of the T-parameters is

$$\begin{pmatrix} \underline{b}_1 \\ \underline{b}_2 \\ \underline{a}_1 \\ \underline{a}_2 \end{pmatrix} = \begin{pmatrix} \underline{T}_{11} & \underline{T}_{12} & \underline{T}_{13} & \underline{T}_{14} \\ \underline{T}_{21} & \underline{T}_{22} & \underline{T}_{23} & \underline{T}_{24} \\ \underline{T}_{31} & \underline{T}_{32} & \underline{T}_{33} & \underline{T}_{34} \\ \underline{T}_{41} & \underline{T}_{42} & \underline{T}_{43} & \underline{T}_{44} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_3 \\ \underline{a}_4 \\ \underline{b}_3 \\ \underline{b}_4 \end{pmatrix} \quad (15.24)$$

this is just one of the possible definitions for a transfer scattering parameter matrix; a different choice of the dependent and independent variables and of the ports ordering will give a different definition.

To convert between the S-parameters and T-parameters matrices the left- and right-side waves are first grouped together

$$\underline{b}_l = \begin{pmatrix} \underline{b}_1 \\ \underline{b}_2 \end{pmatrix}, \quad \underline{b}_r = \begin{pmatrix} \underline{b}_3 \\ \underline{b}_4 \end{pmatrix}, \quad \underline{a}_l = \begin{pmatrix} \underline{a}_1 \\ \underline{a}_2 \end{pmatrix}, \quad \underline{a}_r = \begin{pmatrix} \underline{a}_3 \\ \underline{a}_4 \end{pmatrix} \quad (15.25)$$

so that the S-parameters matrix can be written also as

$$\begin{pmatrix} \underline{b}_l \\ \underline{b}_r \end{pmatrix} = \begin{pmatrix} \underline{S}_{ll} & \underline{S}_{lr} \\ \underline{S}_{rl} & \underline{S}_{rr} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_l \\ \underline{a}_r \end{pmatrix} \quad (15.26)$$

By grouping the left and right *ports*, the transfer scattering parameters matrix is obtained

$$\begin{pmatrix} \underline{b}_l \\ \underline{a}_l \end{pmatrix} = \begin{pmatrix} \underline{T}_{11} & \underline{T}_{12} \\ \underline{T}_{21} & \underline{T}_{22} \end{pmatrix} \cdot \begin{pmatrix} \underline{a}_r \\ \underline{b}_r \end{pmatrix} \quad (15.27)$$

and from eq. (??) and eq. (??) the relationship between the two matrix representation can be derived.

$$\begin{pmatrix} \underline{T}_{11} & \underline{T}_{12} \\ \underline{T}_{21} & \underline{T}_{22} \end{pmatrix} = \begin{pmatrix} \underline{S}_{lr} - \underline{S}_{ll}\underline{S}_{rl}^{-1}\underline{S}_{rr} & \underline{S}_{ll}\underline{S}_{rl}^{-1} \\ -\underline{S}_{rl}^{-1}\underline{S}_{rr} & \underline{S}_{rl}^{-1} \end{pmatrix} \quad (15.28)$$

and

$$\begin{pmatrix} \underline{S}_{ll} & \underline{S}_{lr} \\ \underline{S}_{rl} & \underline{S}_{rr} \end{pmatrix} = \begin{pmatrix} \underline{T}_{12}\underline{T}_{22}^{-1} & \underline{T}_{11} - \underline{T}_{12}\underline{T}_{22}^{-1}\underline{T}_{21} \\ \underline{T}_{22}^{-1} & -\underline{T}_{22}^{-1}\underline{T}_{21} \end{pmatrix} \quad (15.29)$$

The S-parameters matrix of the anti-network can be computed by first inverting blockwise the T matrix (??) and then calculating the corresponding S-parameters matrix using eq. (??). The resulting expression is

$$\underline{S}^{(inv)} = \begin{pmatrix} \underline{S}_{ll}^{(inv)} & \underline{S}_{lr}^{(inv)} \\ \underline{S}_{rl}^{(inv)} & \underline{S}_{rr}^{(inv)} \end{pmatrix} = \begin{pmatrix} (\underline{S}_{rr} - \underline{S}_{rl}\underline{S}_{ll}^{-1}\underline{S}_{lr})^{-1} & (\underline{S}_{lr} - \underline{S}_{ll}\underline{S}_{rl}^{-1}\underline{S}_{rr})^{-1} \\ (\underline{S}_{rl} - \underline{S}_{rr}\underline{S}_{ll}^{-1}\underline{S}_{ll})^{-1} & (\underline{S}_{ll} - \underline{S}_{lr}\underline{S}_{rr}^{-1}\underline{S}_{rl})^{-1} \end{pmatrix} \quad (15.30)$$

The above expressions represent the general formulas for converting between multiport S and T parameters and for computing the S parameters of the anti-network. Note that these reduce to the formulas shown previously for the two-port case and can be extended to any network having an even number of ports.

Mixed Two-Port matrix conversions

Sometimes it may be useful to have a twoport matrix representation based on signal waves in a representation based on voltage and current and the other way around. There are two more parameters involved in this case: The reference impedance at port 1 (denoted as Z_1) and the reference impedance at port 2 (denoted as Z_2).

Converting from scattering parameters to chain parameters results in

$$A_{11} = \frac{Z_1^* + Z_1 \cdot S_{11} - Z_1^* \cdot S_{22} - Z_1 \cdot \Delta S}{2 \cdot S_{21} \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}} \quad (15.31)$$

$$A_{12} = \frac{Z_1^* \cdot Z_2^* + Z_1 \cdot Z_2^* \cdot S_{11} + Z_1^* \cdot Z_2 \cdot S_{22} + Z_1 \cdot Z_2 \cdot \Delta S}{2 \cdot S_{21} \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}} \quad (15.32)$$

$$A_{21} = \frac{1 - S_{11} - S_{22} + \Delta S}{2 \cdot S_{21} \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}} \quad (15.33)$$

$$A_{22} = \frac{Z_2^* - Z_2^* \cdot S_{11} + Z_2 \cdot S_{22} - Z_2 \cdot \Delta S}{2 \cdot S_{21} \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}} \quad (15.34)$$

Converting from chain parameters to scattering parameters results in

$$S_{11} = \frac{A_{11} \cdot Z_2 + A_{12} - A_{21} \cdot Z_1^* \cdot Z_2 - A_{22} \cdot Z_1^*}{A_{11} \cdot Z_2 + A_{12} + A_{21} \cdot Z_1 \cdot Z_2 + A_{22} \cdot Z_1} \quad (15.35)$$

$$S_{12} = \frac{\Delta A \cdot 2 \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}}{A_{11} \cdot Z_2 + A_{12} + A_{21} \cdot Z_1 \cdot Z_2 + A_{22} \cdot Z_1} \quad (15.36)$$

$$S_{21} = \frac{2 \cdot \sqrt{\text{Re}(Z_1) \cdot \text{Re}(Z_2)}}{A_{11} \cdot Z_2 + A_{12} + A_{21} \cdot Z_1 \cdot Z_2 + A_{22} \cdot Z_1} \quad (15.37)$$

$$S_{22} = \frac{-A_{11} \cdot Z_2^* + A_{12} - A_{21} \cdot Z_1 \cdot Z_2^* + A_{22} \cdot Z_1}{A_{11} \cdot Z_2 + A_{12} + A_{21} \cdot Z_1 \cdot Z_2 + A_{22} \cdot Z_1} \quad (15.38)$$

Converting from scattering parameters to hybrid parameters results in

$$H_{11} = Z_1 \cdot \frac{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}}{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}} \quad (15.39)$$

$$H_{12} = \sqrt{\frac{Z_1}{Z_2}} \cdot \frac{2 \cdot S_{12}}{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}} \quad (15.40)$$

$$H_{21} = \sqrt{\frac{Z_1}{Z_2}} \cdot \frac{-2 \cdot S_{21}}{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}} \quad (15.41)$$

$$H_{22} = \frac{1}{Z_2} \cdot \frac{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}}{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}} \quad (15.42)$$

Converting from hybrid parameters to scattering parameters results in

$$S_{11} = \frac{(H_{11} - Z_1) \cdot (1 + Z_2 \cdot H_{22}) - Z_2 \cdot H_{12} \cdot H_{21}}{(H_{11} + Z_1) \cdot (1 + Z_2 \cdot H_{22}) - Z_2 \cdot H_{12} \cdot H_{21}} \quad (15.43)$$

$$S_{12} = \frac{2 \cdot H_{12} \cdot \sqrt{Z_1 \cdot Z_2}}{(H_{11} + Z_1) \cdot (1 + Z_2 \cdot H_{22}) - Z_2 \cdot H_{12} \cdot H_{21}} \quad (15.44)$$

$$S_{21} = \frac{-2 \cdot H_{21} \cdot \sqrt{Z_1 \cdot Z_2}}{(H_{11} + Z_1) \cdot (1 + Z_2 \cdot H_{22}) - Z_2 \cdot H_{12} \cdot H_{21}} \quad (15.45)$$

$$S_{22} = \frac{(H_{11} + Z_1) \cdot (1 - Z_2 \cdot H_{22}) + Z_2 \cdot H_{12} \cdot H_{21}}{(H_{11} + Z_1) \cdot (1 + Z_2 \cdot H_{22}) - Z_2 \cdot H_{12} \cdot H_{21}} \quad (15.46)$$

Converting from scattering parameters to the second type of hybrid parameters results in

$$G_{11} = \frac{1}{Z_1} \cdot \frac{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}}{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}} \quad (15.47)$$

$$G_{12} = \sqrt{\frac{Z_2}{Z_1}} \cdot \frac{-2 \cdot S_{12}}{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}} \quad (15.48)$$

$$G_{21} = \sqrt{\frac{Z_2}{Z_1}} \cdot \frac{2 \cdot S_{21}}{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}} \quad (15.49)$$

$$G_{22} = Z_2 \cdot \frac{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}}{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}} \quad (15.50)$$

Converting from the second type of hybrid parameters to scattering parameters results in

$$S_{11} = \frac{(1 - G_{11} \cdot Z_1) \cdot (G_{22} + Z_2) + Z_1 \cdot G_{12} \cdot G_{21}}{(1 + G_{11} \cdot Z_1) \cdot (G_{22} + Z_2) - Z_1 \cdot G_{12} \cdot G_{21}} \quad (15.51)$$

$$S_{12} = \frac{-2 \cdot G_{12} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 + G_{11} \cdot Z_1) \cdot (G_{22} + Z_2) - Z_1 \cdot G_{12} \cdot G_{21}} \quad (15.52)$$

$$S_{21} = \frac{2 \cdot G_{21} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 + G_{11} \cdot Z_1) \cdot (G_{22} + Z_2) - Z_1 \cdot G_{12} \cdot G_{21}} \quad (15.53)$$

$$S_{22} = \frac{(1 + G_{11} \cdot Z_1) \cdot (G_{22} - Z_2) - Z_1 \cdot G_{12} \cdot G_{21}}{(1 + G_{11} \cdot Z_1) \cdot (G_{22} + Z_2) - Z_1 \cdot G_{12} \cdot G_{21}} \quad (15.54)$$

Converting from scattering parameters to Y-parameters results in

$$Y_{11} = \frac{1}{Z_1} \cdot \frac{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}}{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}} \quad (15.55)$$

$$Y_{12} = \sqrt{\frac{1}{Z_1 \cdot Z_2}} \cdot \frac{-2 \cdot S_{12}}{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}} \quad (15.56)$$

$$Y_{21} = \sqrt{\frac{1}{Z_1 \cdot Z_2}} \cdot \frac{-2 \cdot S_{21}}{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}} \quad (15.57)$$

$$Y_{22} = \frac{1}{Z_2} \cdot \frac{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}}{(1 + S_{11}) \cdot (1 + S_{22}) - S_{12} \cdot S_{21}} \quad (15.58)$$

Converting from Y-parameters to scattering parameters results in

$$S_{11} = \frac{(1 - Y_{11} \cdot Z_1) \cdot (1 + Y_{22} \cdot Z_2) + Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2}{(1 + Y_{11} \cdot Z_1) \cdot (1 + Y_{22} \cdot Z_2) - Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2} \quad (15.59)$$

$$S_{12} = \frac{-2 \cdot Y_{12} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 + Y_{11} \cdot Z_1) \cdot (1 + Y_{22} \cdot Z_2) - Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2} \quad (15.60)$$

$$S_{21} = \frac{-2 \cdot Y_{21} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 + Y_{11} \cdot Z_1) \cdot (1 + Y_{22} \cdot Z_2) - Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2} \quad (15.61)$$

$$S_{22} = \frac{(1 + Y_{11} \cdot Z_1) \cdot (1 - Y_{22} \cdot Z_2) + Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2}{(1 + Y_{11} \cdot Z_1) \cdot (1 + Y_{22} \cdot Z_2) - Y_{12} \cdot Z_1 \cdot Y_{21} \cdot Z_2} \quad (15.62)$$

Converting from scattering parameters to Z-parameters results in

$$Z_{11} = Z_1 \cdot \frac{(1 + S_{11}) \cdot (1 - S_{22}) + S_{12} \cdot S_{21}}{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}} \quad (15.63)$$

$$Z_{12} = \frac{2 \cdot S_{12} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}} \quad (15.64)$$

$$Z_{21} = \frac{2 \cdot S_{21} \cdot \sqrt{Z_1 \cdot Z_2}}{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}} \quad (15.65)$$

$$Z_{22} = Z_2 \cdot \frac{(1 - S_{11}) \cdot (1 + S_{22}) + S_{12} \cdot S_{21}}{(1 - S_{11}) \cdot (1 - S_{22}) - S_{12} \cdot S_{21}} \quad (15.66)$$

Converting from Z-parameters to scattering parameters results in

$$S_{11} = \frac{(Z_{11} - Z_1) \cdot (Z_{22} + Z_2) - Z_{12} \cdot Z_{21}}{(Z_{11} + Z_1) \cdot (Z_{22} + Z_2) - Z_{12} \cdot Z_{21}} \quad (15.67)$$

$$S_{12} = \sqrt{\frac{Z_2}{Z_1}} \cdot \frac{2 \cdot Z_{12} \cdot Z_1}{(Z_{11} + Z_1) \cdot (Z_{22} + Z_2) - Z_{12} \cdot Z_{21}} \quad (15.68)$$

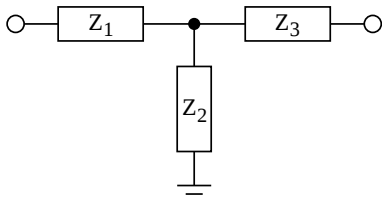
$$S_{21} = \sqrt{\frac{Z_1}{Z_2}} \cdot \frac{2 \cdot Z_{21} \cdot Z_2}{(Z_{11} + Z_1) \cdot (Z_{22} + Z_2) - Z_{12} \cdot Z_{21}} \quad (15.69)$$

$$S_{22} = \frac{(Z_{11} + Z_1) \cdot (Z_{22} - Z_2) - Z_{12} \cdot Z_{21}}{(Z_{11} + Z_1) \cdot (Z_{22} + Z_2) - Z_{12} \cdot Z_{21}} \quad (15.70)$$

Two-Port parameters of passive devices

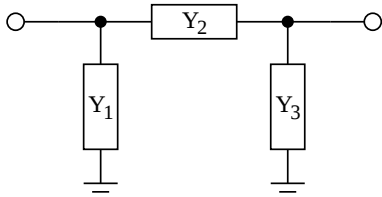
Basically the twoport parameters of passive twoports can be determined using Kirchhoff's voltage law and Kirchhoff's current law or by applying the definition equations of the twoport parameters. This has been done [?] for some example circuits.

- T-topology



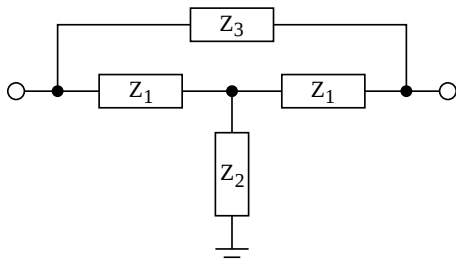
$$Z = \begin{bmatrix} Z_1 + Z_2 & Z_2 \\ Z_2 & Z_2 + Z_3 \end{bmatrix}$$

- π -topology



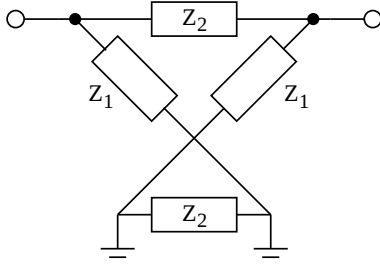
$$Y = \begin{bmatrix} Y_1 + Y_2 & -Y_2 \\ -Y_2 & Y_2 + Y_3 \end{bmatrix}$$

- symmetric T-bridge



$$Z = \begin{bmatrix} \frac{Z_1^2 + Z_1 \cdot Z_3}{2 \cdot Z_1 + Z_3} + Z_2 & \frac{Z_1^2}{2 \cdot Z_1 + Z_3} + Z_2 \\ \frac{Z_1^2}{2 \cdot Z_1 + Z_3} + Z_2 & \frac{Z_1^2 + Z_1 \cdot Z_3}{2 \cdot Z_1 + Z_3} + Z_2 \end{bmatrix}$$

- symmetric X-topology



$$Z = \frac{1}{2} \begin{bmatrix} Z_1 + Z_2 & Z_1 - Z_2 \\ Z_1 - Z_2 & Z_1 + Z_2 \end{bmatrix}$$

15.2 Solving linear equation systems

When dealing with non-linear networks the number of equation systems to be solved depends on the required precision of the solution and the average necessary iterations until the solution is stable. This emphasizes the meaning of the solving procedures choice for different problems.

The equation systems

$$[A] \cdot [x] = [z] \quad (15.71)$$

solution can be written as

$$[x] = [A]^{-1} \cdot [z] \quad (15.72)$$

15.2.1 Matrix inversion

The elements $\beta_{\mu\nu}$ of the inverse of the matrix A are

$$\beta_{\mu\nu} = \frac{A_{\mu\nu}}{\det A} \quad (15.73)$$

whereas $A_{\mu\nu}$ is the matrix elements $a_{\mu\nu}$ cofactor. The cofactor is the sub determinant (i.e. the minor) of the element $a_{\mu\nu}$ multiplied with $(-1)^{\mu+\nu}$. The determinant of a square matrix can be recursively computed by either of the following equations.

$$\det A = \sum_{\mu=1}^n a_{\mu\nu} \cdot A_{\mu\nu} \quad \text{using the } \nu\text{-th column} \quad (15.74)$$

$$\det A = \sum_{\nu=1}^n a_{\mu\nu} \cdot A_{\mu\nu} \quad \text{using the } \mu\text{-th row} \quad (15.75)$$

This method is called the Laplace expansion. In order to save computing time the row or column with the most zeros in it is used for the expansion expressed in the above equations. A sub determinant $(n-1)$ -th order of a matrix's element $a_{\mu\nu}$ of n -th order is the determinant which is computed by cancelling the μ -th row and ν -th column. The following example demonstrates calculating the determinant of a 4th order matrix with the elements of the 3rd row.

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{vmatrix} = a_{31} \begin{vmatrix} a_{12} & a_{13} & a_{14} \\ a_{22} & a_{23} & a_{24} \\ a_{42} & a_{43} & a_{44} \end{vmatrix} - a_{32} \begin{vmatrix} a_{11} & a_{13} & a_{14} \\ a_{21} & a_{23} & a_{24} \\ a_{41} & a_{43} & a_{44} \end{vmatrix} \\ + a_{33} \begin{vmatrix} a_{11} & a_{12} & a_{14} \\ a_{21} & a_{22} & a_{24} \\ a_{41} & a_{42} & a_{44} \end{vmatrix} - a_{34} \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{41} & a_{42} & a_{43} \end{vmatrix} \quad (15.76)$$

This recursive process for computing the inverse of a matrix is most easiest to be implemented but as well the slowest algorithm. It requires approximately $n!$ operations.

15.2.2 Gaussian elimination

The Gaussian algorithm for solving a linear equation system is done in two parts: forward elimination and backward substitution. During forward elimination the matrix A is transformed into an upper triangular equivalent matrix. Elementary transformations due to an equation system having the same solutions for the unknowns as the original system.

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix} \rightarrow \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & a_{nn} \end{bmatrix} \quad (15.77)$$

The modifications applied to the matrix A in order to achieve this transformations are limited to the following set of operations.

- multiplication of a row with a scalar factor
- addition or subtraction of multiples of rows
- exchanging two rows of a matrix

Step 1: Forward elimination

The transformation of the matrix A is done in $n - 1$ elimination steps. The new matrix elements of the k -th step with $k = 1, \dots, n - 1$ are computed with the following recursive formulas.

$$a_{ij} = 0 \quad i = k + 1, \dots, n \text{ and } j = k \quad (15.78)$$

$$a_{ij} = a_{ij} - a_{kj} \cdot a_{ik} / a_{kk} \quad i = k + 1, \dots, n \text{ and } j = k + 1, \dots, n \quad (15.79)$$

$$z_i = z_i - z_k \cdot a_{ik} / a_{kk} \quad i = k + 1, \dots, n \quad (15.80)$$

The triangulated matrix can be used to calculate the determinant very easily. The determinant of a triangulated matrix is the product of the diagonal elements. If the determinant $\det A$ is non-zero the equation system has a solution. Otherwise the matrix A is singular.

$$\det A = a_{11} \cdot a_{22} \cdot \dots \cdot a_{nn} = \prod_{i=1}^n a_{ii} \quad (15.81)$$

When using row and/or column pivoting the resulting determinant may differ in its sign and must be multiplied with $(-1)^m$ whereas m is the number of row and column substitutions.

Finding an appropriate pivot element

The Gaussian elimination fails if the pivot element a_{kk} turns to be zero (division by zero). That is why row and/or column pivoting must be used before each elimination step. If a diagonal element $a_{kk} = 0$, then exchange the pivot row k with the row $m > k$ having the coefficient with the largest absolute value. The new pivot row is m and the new pivot element is going to be a_{mk} . If no such pivot row can be found the matrix is singular.

Total pivoting looks for the element with the largest absolute value within the matrix and exchanges rows and columns. When exchanging columns in equation systems the unknowns get reordered as well. For the numerical solution of equation systems with Gaussian elimination column pivoting is clever, and total pivoting recommended.

In order to improve numerical stability pivoting should also be applied if $a_{kk} \neq 0$ because division by small diagonal elements propagates numerical (rounding) errors. This appears especially with poorly conditioned (the two dimensional case: two lines with nearly the same slope) equation systems.

Step 2: Backward substitution

The solutions in the vector x are obtained by backward substituting into the triangulated matrix. The elements of the solution vector x are computed by the following recursive equations.

$$x_n = \frac{z_n}{a_{nn}} \quad (15.82)$$

$$x_i = \frac{z_i}{a_{ii}} - \sum_{k=i+1}^n x_k \cdot \frac{a_{ik}}{a_{ii}} \quad i = n-1, \dots, 1 \quad (15.83)$$

The forward elimination in the Gaussian algorithm requires approximately $n^3/3$, the backward substitution $n^2/2$ operations.

15.2.3 Gauss-Jordan method

The Gauss-Jordan method is a modification of the Gaussian elimination. In each k -th elimination step the elements of the k -th column get zero except the diagonal element which gets 1. When the right hand side vector z is included in each step it contains the solution vector x afterwards.

The following recursive formulas must be applied to get the new matrix elements for the k -th elimination step. The k -th row must be computed first.

$$a_{kj} = a_{kj}/a_{kk} \quad j = 1 \dots n \quad (15.84)$$

$$z_k = z_k/a_{kk} \quad (15.85)$$

Then the other rows can be calculated with the following formulas.

$$a_{ij} = a_{ij} - a_{ik} \cdot a_{kj} \quad j = 1, \dots, n \text{ and } i = 1, \dots, n \text{ with } i \neq k \quad (15.86)$$

$$z_i = z_i - a_{ik} \cdot z_k \quad i = 1, \dots, n \text{ with } i \neq k \quad (15.87)$$

Column pivoting may be necessary in order to avoid division by zero. The solution vector x is not harmed by row substitutions. When the Gauss-Jordan algorithm has been finished the original matrix has been transformed into the identity matrix. If each operation during this process is applied to an identity matrix the resulting matrix is the inverse matrix of the original matrix. This means that the Gauss-Jordan method can be used to compute the inverse of a matrix.

Though this elimination method is easy to implement the number of required operations is larger than within the Gaussian elimination. The Gauss-Jordan method requires approximately $N^3/2 + N^2/2$ operations.

15.2.4 LU decomposition

LU decomposition (decomposition into a lower and upper triangular matrix) is recommended when dealing with equation systems where the matrix A does not alter but the right hand side (the vector z) does. Both the Gaussian elimination and the Gauss-Jordan method involve both the right hand side and the matrix in their algorithm. Consecutive solutions of an equation system with an altering right hand side can be computed faster with LU decomposition.

The LU decomposition splits a matrix A into a product of a lower triangular matrix L with an upper triangular matrix U .

$$A = LU \quad \text{with} \quad L = \begin{bmatrix} l_{11} & 0 & \dots & 0 \\ l_{21} & l_{22} & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ l_{n1} & \dots & \dots & l_{nn} \end{bmatrix} \quad \text{and} \quad U = \begin{bmatrix} u_{11} & u_{12} & \dots & u_{1n} \\ 0 & u_{22} & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & u_{nn} \end{bmatrix} \quad (15.88)$$

The algorithm for solving the linear equation system $Ax = z$ involves three steps:

- LU decomposition of the coefficient matrix A
 $\rightarrow Ax = LUx = z$
- introduction of an (unknown) arbitrary vector y and solving the equation system $Ly = z$ by forward substitution
 $\rightarrow y = Ux = L^{-1}z$
- solving the equation system $Ux = y$ by backward substitution
 $\rightarrow x = U^{-1}y$

The decomposition of the matrix A into a lower and upper triangular matrix is not unique. The most important decompositions, based on Gaussian elimination, are the Doolittle, the Crout and the Cholesky decomposition.

If pivoting is necessary during these algorithms they do not decompose the matrix A but the product with an arbitrary matrix PA (a permutation of the matrix A). When exchanging rows and columns the order of the unknowns as represented by the vector z changes as well and must be saved during this process for the forward substitution in the algorithms second step.

Step 1: LU decomposition

Using the decomposition according to Crout the coefficients of the L and U matrices can be stored in place the original matrix A . The upper triangular matrix U has the form

$$U = \begin{bmatrix} 1 & u_{12} & \dots & u_{1n} \\ 0 & 1 & & \vdots \\ \vdots & \ddots & \ddots & u_{n-1,n} \\ 0 & \dots & 0 & 1 \end{bmatrix} \quad (15.89)$$

The diagonal elements u_{jj} are ones and thus the determinant $\det U$ is one as well. The elements of the new coefficient matrix LU for the k -th elimination step with $k = 1, \dots, n$ compute as

follows:

$$u_{jk} = \frac{1}{l_{jj}} \left(a_{jk} - \sum_{r=1}^{j-1} l_{jr} u_{rk} \right) \quad j = 1, \dots, k-1 \quad (15.90)$$

$$l_{jk} = a_{jk} - \sum_{r=1}^{k-1} l_{jr} u_{rk} \quad j = k, \dots, n \quad (15.91)$$

Pivoting may be necessary as you are going to divide by the diagonal element l_{jj} .

Step 2: Forward substitution

The solutions in the arbitrary vector y are obtained by forward substituting into the triangulated L matrix. At this stage you need to remember the order of unknowns in the vector z as changed by pivoting. The elements of the solution vector y are computed by the following recursive equation.

$$y_i = \frac{z_i}{l_{ii}} - \sum_{k=1}^{i-1} y_k \cdot \frac{l_{ik}}{l_{ii}} \quad i = 1, \dots, n \quad (15.92)$$

Step 3: Backward substitution

The solutions in the vector x are obtained by backward substituting into the triangulated U matrix. The elements of the solution vector x are computed by the following recursive equation.

$$x_i = y_i - \sum_{k=i+1}^n x_k \cdot u_{ik} \quad i = n, \dots, 1 \quad (15.93)$$

The division by the diagonal elements of the matrix U is not necessary because of Crout's definition in eq. (??) with $u_{ii} = 1$.

The LU decomposition requires approximately $n^3/3 + n^2 - n/3$ operations for solving a linear equation system. For M consecutive solutions the method requires $n^3/3 + Mn^2 - n/3$ operations.

15.2.5 QR decomposition

Singular matrices actually having a solution are over- or under-determined. These types of matrices can be handled by three different types of decompositions: Householder, Jacobi (Givens rotation) and singular value decomposition. Householder decomposition factors a matrix A into the product of an orthonormal matrix Q and an upper triangular matrix R , such that:

$$A = Q \cdot R \quad (15.94)$$

The Householder decomposition is based on the fact that for any two different vectors, v and w , with $\|v\| = \|w\|$, i.e. different vectors of equal length, a reflection matrix H exists such that:

$$H \cdot v = w \quad (15.95)$$

To obtain the matrix H , the vector u is defined by:

$$u = \frac{v - w}{\|v - w\|} \quad (15.96)$$

The matrix H defined by

$$H = I - 2 \cdot u \cdot u^T \quad (15.97)$$

is then the required reflection matrix.

The equation system

$$A \cdot x = z \quad \text{is transformed into} \quad QR \cdot x = z \quad (15.98)$$

With $Q^T \cdot Q = I$ this yields

$$Q^T QR \cdot x = Q^T z \quad \rightarrow \quad R \cdot x = Q^T z \quad (15.99)$$

Since R is triangular the equation system is solved by a simple matrix-vector multiplication on the right hand side and backward substitution.

Step 1: QR decomposition

Starting with $A_1 = A$, let v_1 = the first column of A_1 , and $w_1^T = (\pm\|v_1\|, 0, \dots, 0)$, i.e. a column vector whose first component is the norm of v_1 with the remaining components equal to 0. The Householder transformation $H_1 = I - 2 \cdot u_1 \cdot u_1^T$ with $u_1 = v_1 - w_1 / \|v_1 - w_1\|$ will turn the first column of A_1 into w_1 as with $H_1 \cdot A_1 = A_2$. At each stage k , v_k = the k th column of A_k on and below the diagonal with all other components equal to 0, and w_k 's k th component equals the norm of v_k with all other components equal to 0. Letting $H_k \cdot A_k = A_{k+1}$, the components of the k th column of A_{k+1} below the diagonal are each 0. These calculations are listed below for each stage for the matrix A .

$$\begin{aligned} v_1 = \begin{bmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{n1} \end{bmatrix} \quad w_1 = \begin{bmatrix} \pm\|v_1\| \\ 0 \\ \vdots \\ 0 \end{bmatrix} \quad u_1 = \frac{v_1 - w_1}{\|v_1 - w_1\|} = \begin{bmatrix} u_{11} \\ u_{21} \\ \vdots \\ u_{n1} \end{bmatrix} \\ H_1 = I - 2 \cdot u_1 \cdot u_1^T \quad \rightarrow \quad H_1 \cdot A_1 = A_2 = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & a_{n2} & \dots & a_{nn} \end{bmatrix} \end{aligned} \quad (15.100)$$

With this first step the upper left diagonal element of the R matrix, $a_{11} = \pm\|v_1\|$, has been generated. The elements below are zeroed out. Since H_1 can be generated from u_1 stored in place of the first column of A_1 the multiplication $H_1 \cdot A_1$ can be performed without actually generating H_1 .

$$\begin{aligned}
v_2 &= \begin{bmatrix} 0 \\ a_{22} \\ \vdots \\ a_{n2} \end{bmatrix} & w_1 &= \begin{bmatrix} 0 \\ \pm \|v_2\| \\ \vdots \\ 0 \end{bmatrix} & u_2 &= \frac{v_2 - w_2}{\|v_2 - w_2\|} = \begin{bmatrix} 0 \\ u_{22} \\ \vdots \\ u_{n2} \end{bmatrix} \\
H_2 &= I - 2 \cdot u_2 \cdot u_2^T \rightarrow H_2 \cdot A_2 = A_3 = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & 0 & \ddots & \vdots \\ 0 & 0 & & a_{nn} \end{bmatrix}
\end{aligned} \tag{15.101}$$

These elimination steps generate the R matrix because Q is orthonormal, i.e.

$$\begin{aligned}
A &= Q \cdot R \rightarrow Q^T A = Q^T Q \cdot R \rightarrow Q^T A = R \\
&\rightarrow H_n \cdot \dots \cdot H_2 \cdot H_1 \cdot A = R
\end{aligned} \tag{15.102}$$

After n elimination steps the original matrix A contains the upper triangular matrix R , except for the diagonal elements which can be stored in some vector. The lower triangular matrix contains the Householder vectors $u_1 \dots u_n$.

$$A = \begin{bmatrix} u_{11} & r_{12} & \dots & r_{1n} \\ u_{21} & u_{22} & & r_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ u_{n1} & u_{n2} & \dots & u_{nn} \end{bmatrix} \quad R_{diag} = \begin{bmatrix} r_{11} \\ r_{22} \\ \vdots \\ r_{nn} \end{bmatrix} \tag{15.103}$$

With $Q^T = H_1 \cdot H_2 \cdot \dots \cdot H_n$ this representation contains both the Q and R matrix, in a packed form, of course: Q as a composition of Householder vectors and R in the upper triangular part and its diagonal vector R_{diag} .

Step 2: Forming the new right hand side

In order to form the right hand side $Q^T z$ let remember eq. (??) denoting the reflection matrices used to compute Q^T .

$$H_n \cdot \dots \cdot H_2 \cdot H_1 = Q^T \tag{15.104}$$

Thus it is possible to replace the original right hand side vector z by

$$H_n \cdot \dots \cdot H_2 \cdot H_1 \cdot z = Q^T \cdot z \tag{15.105}$$

which yields for each $k = 1 \dots n$ the following expression:

$$H_k \cdot z = (I - 2 \cdot u_k \cdot u_k^T) \cdot z = z - 2 \cdot u_k \cdot u_k^T \cdot z \tag{15.106}$$

The latter $u_k^T \cdot z$ is a simple scalar product of two vectors. Performing eq. (??) for each Householder vector finally results in the new right hand side vector $Q^T z$.

Step 3: Backward substitution

The solutions in the vector x are obtained by backward substituting into the triangulated R matrix. The elements of the solution vector x are computed by the following recursive equation.

$$x_i = \frac{z_i}{r_{ii}} - \sum_{k=i+1}^n x_k \cdot \frac{r_{ik}}{r_{ii}} \quad i = n, \dots, 1 \tag{15.107}$$

Motivation

Though the QR decomposition has an operation count of $2n^3 + 3n^2$ (which is about six times more than the LU decomposition) it has its advantages. The QR factorization method is said to be unconditional stable and more accurate. Also it can be used to obtain the minimum-norm (or least square) solution of under-determined equation systems.

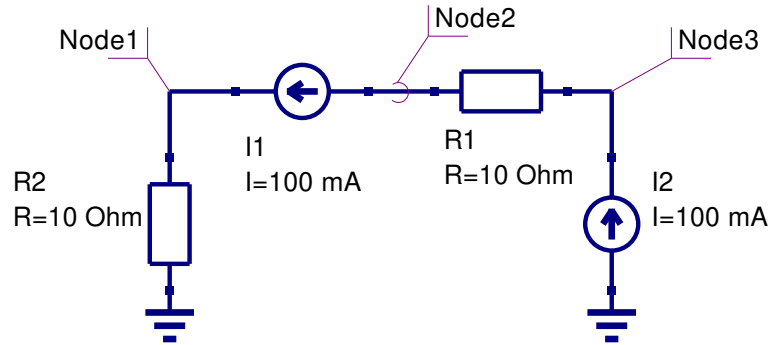


Figure 15.5: circuit with singular modified nodal analysis matrix

The circuit in fig. ?? has the following MNA representation:

$$Ax = \begin{bmatrix} \frac{1}{R_2} & 0 & 0 \\ 0 & \frac{1}{R_1} & -\frac{1}{R_1} \\ 0 & -\frac{1}{R_1} & \frac{1}{R_1} \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \end{bmatrix} = \begin{bmatrix} 0.1 & 0 & 0 \\ 0 & 0.1 & -0.1 \\ 0 & -0.1 & 0.1 \end{bmatrix} \cdot \begin{bmatrix} V_1 \\ V_2 \\ V_3 \end{bmatrix} = \begin{bmatrix} I_1 \\ -I_1 \\ I_2 \end{bmatrix} = \begin{bmatrix} 0.1 \\ -0.1 \\ 0.1 \end{bmatrix} = z \quad (15.108)$$

The second and third row of the matrix A are linear dependent and the matrix is singular because its determinant is zero. Depending on the right hand side z , the equation system has none or unlimited solutions. This is called an under-determined system. The discussed QR decomposition easily computes a valid solution without reducing accuracy. The LU decomposition would probably fail because of the singularity.

QR decomposition with column pivoting

Least norm problem

With some more effort it is possible to obtain the minimum-norm solution of this problem. The algorithm as described here would probably yield the following solution:

$$x = \begin{bmatrix} V_1 \\ V_2 \\ V_3 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} \quad (15.109)$$

This is one out of unlimited solutions. The following short description shows how it is possible to obtain the minimum-norm solution. When decomposing the transposed problem

$$A^T = Q \cdot R \quad (15.110)$$

the minimum-norm solution $\hat{x}$ is obtained by forward substitution of

$$R^T \cdot x = z \quad (15.111)$$

and multiplying the result with Q .

$$\hat{x} = Q \cdot x \quad (15.112)$$

In the example above this algorithm results in a solution vector with the least vector norm possible:

$$\hat{x} = \begin{bmatrix} V_1 \\ V_2 \\ V_3 \end{bmatrix} = \begin{bmatrix} 1 \\ -0.5 \\ 0.5 \end{bmatrix} \quad (15.113)$$

This algorithm outline is also sometimes called LQ decomposition because of R^T being a lower triangular matrix used by the forward substitution.

15.2.6 Singular value decomposition

Very bad conditioned (ratio between largest and smallest eigenvalue) matrices, i.e. nearly singular, or even singular matrices (over- or under-determined equation systems) can be handled by the singular value decomposition (SVD). This type of decomposition is defined by

$$A = U \cdot \Sigma \cdot V^H \quad (15.114)$$

where the U matrix consists of the orthonormalized eigenvectors associated with the eigenvalues of $A \cdot A^H$, V consists of the orthonormalized eigenvectors of $A^H \cdot A$ and Σ is a matrix with the singular values of A (non-negative square roots of the eigenvalues of $A^H \cdot A$) on its diagonal and zeros otherwise.

$$\Sigma = \begin{bmatrix} \sigma_1 & 0 & \cdots & 0 \\ 0 & \sigma_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_n \end{bmatrix} \quad (15.115)$$

The singular value decomposition can be used to solve linear equation systems by simple substitutions

$$A \cdot x = z \quad (15.116)$$

$$U \cdot \Sigma \cdot V^H \cdot x = z \quad (15.117)$$

$$\Sigma \cdot V^H \cdot x = U^H \cdot z \quad (15.118)$$

since

$$U^H \cdot U = V^H \cdot V = V \cdot V^H = I \quad (15.119)$$

To obtain the decomposition stated in eq. (??) Householder vectors are computed and their transformations are applied from the left-hand side and right-hand side to obtain an upper bidiagonal matrix B which has the same singular values as the original A matrix because all of the transformations introduced are orthogonal.

$$U_B^{H(n)} \cdot \dots \cdot U_B^{H(1)} \cdot A \cdot V_B^{(1)} \cdot \dots \cdot V_B^{(n-2)} = U_B^H \cdot A \cdot V_B = B^{(0)} \quad (15.120)$$

Specifically, $U_B^{H(i)}$ annihilates the subdiagonal elements in column i and $V_B^{(j)}$ zeros out the appropriate elements in row j .

$$B^{(0)} = \begin{bmatrix} \beta_1 & \delta_2 & 0 & \cdots & 0 \\ 0 & \beta_2 & \delta_3 & 0 & 0 \\ \vdots & 0 & \ddots & \ddots & 0 \\ 0 & 0 & 0 & \beta_{n-1} & \delta_n \\ 0 & 0 & 0 & 0 & \beta_n \end{bmatrix} \quad (15.121)$$

Afterwards an iterative process (which turns out to be a QR iteration) is used to transform the bidiagonal matrix B into a diagonal form by applying successive Givens transformations (therefore orthogonal as well) to the bidiagonal matrix. This iteration is said to have cubic convergence and yields the final singular values of the matrix A .

$$B^{(0)} \rightarrow B^{(1)} \rightarrow \dots \rightarrow \Sigma \quad (15.122)$$

$$B^{(k+1)} = \left(\tilde{U}^{(k)} \right)^H \cdot B^{(k)} \cdot \tilde{V}^{(k)} \quad (15.123)$$

Each of the transformations applied to the bidiagonal matrix is also applied to the matrices U_B and V_B^H which finally yield the U and V^H matrices after convergence.

So far for the algorithm outline. Without the very details the following sections briefly describe each part of the singular value decomposition.

Notation

Beforehand some notation marks are going to be defined.

- Conjugate transposed (or adjoint):

$$A \rightarrow (A^T)^* = (A^*)^T = A^H$$

- Euclidean norm:

$$\|x\| = \sqrt{\sum_{i=1}^n x_i \cdot x_i^*} = \sqrt{\sum_{i=1}^n |x_i|^2} = \sqrt{|x_1|^2 + \dots + |x_n|^2} = \sqrt{x \cdot x^H}$$

- Hermitian (or self adjoint):

$$A = A^H$$

whereas A^H denotes the conjugate transposed matrix of A . In the real case the matrix A is then said to be “symmetric”.

- Unitary:

$$A \cdot A^H = A^H \cdot A = I$$

Real matrices A with this property are called “orthogonal”.

Householder reflector

A Householder matrix is an elementary unitary matrix that is Hermitian. Its fundamental use is their ability to transform a vector x to a multiple of $\vec{e}_1$, the first column of the identity matrix. The elementary Hermitian (i.e. the Householder matrix) is defined as

$$H = I - 2 \cdot u \cdot u^H \quad \text{where} \quad u^H \cdot u = 1 \quad (15.124)$$

Beside excellent numerical properties, their application demonstrates their efficiency. If A is a matrix, then

$$\begin{aligned} H \cdot A &= A - 2 \cdot u \cdot u^H \cdot A \\ &= A - 2 \cdot u \cdot (A^H \cdot u)^H \end{aligned} \quad (15.125)$$

and hence explicit formation and storage of H is not required. Also columns (or rows) can be transformed individually exploiting the fact that $u^H \cdot A$ yields a scalar product for single columns or rows.

Specific case In order to reduce a 4×4 matrix A to upper triangular form successive Householder reflectors must be applied.

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} \quad (15.126)$$

In the first step the diagonal element a_{11} gets replaced and its below elements get annihilated by the multiplication with an appropriate Householder vector, also the remaining right-hand columns get modified.

$$u_1 = \begin{bmatrix} u_{11} \\ u_{21} \\ u_{31} \\ u_{41} \end{bmatrix} \quad H_1 = I - 2 \cdot u_1 \cdot u_1^H \quad \rightarrow A_1 = H_1 \cdot A = \begin{bmatrix} \beta_1 & a_{12}^{(1)} & a_{13}^{(1)} & a_{14}^{(1)} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & a_{24}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} & a_{34}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} & a_{44}^{(1)} \end{bmatrix} \quad (15.127)$$

This process must be repeated

$$u_2 = \begin{bmatrix} 0 \\ u_{22} \\ u_{32} \\ u_{42} \end{bmatrix} \quad H_2 = I - 2 \cdot u_2 \cdot u_2^H \quad \rightarrow A_2 = H_2 \cdot A_1 = \begin{bmatrix} \beta_1 & a_{12}^{(2)} & a_{13}^{(2)} & a_{14}^{(2)} \\ 0 & \beta_2 & a_{23}^{(2)} & a_{24}^{(2)} \\ 0 & 0 & a_{33}^{(2)} & a_{34}^{(2)} \\ 0 & 0 & a_{43}^{(2)} & a_{44}^{(2)} \end{bmatrix} \quad (15.128)$$

$$u_3 = \begin{bmatrix} 0 \\ 0 \\ u_{33} \\ u_{43} \end{bmatrix} \quad H_3 = I - 2 \cdot u_3 \cdot u_3^H \quad \rightarrow A_3 = H_3 \cdot A_2 = \begin{bmatrix} \beta_1 & a_{12}^{(3)} & a_{13}^{(3)} & a_{14}^{(3)} \\ 0 & \beta_2 & a_{23}^{(3)} & a_{24}^{(3)} \\ 0 & 0 & \beta_3 & a_{34}^{(3)} \\ 0 & 0 & 0 & a_{44}^{(3)} \end{bmatrix} \quad (15.129)$$

$$u_4 = \begin{bmatrix} 0 \\ 0 \\ 0 \\ u_{44} \end{bmatrix} \quad H_4 = I - 2 \cdot u_4 \cdot u_4^H \quad \rightarrow A_4 = H_4 \cdot A_3 = \begin{bmatrix} \beta_1 & a_{12}^{(4)} & a_{13}^{(4)} & a_{14}^{(4)} \\ 0 & \beta_2 & a_{23}^{(4)} & a_{24}^{(4)} \\ 0 & 0 & \beta_3 & a_{34}^{(4)} \\ 0 & 0 & 0 & \beta_4 \end{bmatrix} \quad (15.130)$$

until the matrix A contains an upper triangular matrix R . The matrix Q can be expressed as the the product of the Householder vectors. The performed operations deliver

$$H_4^H \cdot H_3^H \cdot H_2^H \cdot H_1^H \cdot A = Q^H \cdot A = R \quad \rightarrow \quad A = Q \cdot R \quad (15.131)$$

since Q is unitary. The matrix Q itself can be expressed in terms of H_i using the following transformation.

$$Q^H = H_4^H \cdot H_3^H \cdot H_2^H \cdot H_1^H \quad (15.132)$$

$$(Q^H)^H = (H_4^H \cdot H_3^H \cdot H_2^H \cdot H_1^H)^H \quad (15.133)$$

$$Q = H_1 \cdot H_2 \cdot H_3 \cdot H_4 \quad (15.134)$$

The eqn. (??)-(??) are necessary to be mentioned only in case Q is not Hermitian, but still unitary. Otherwise there is no difference computing Q or Q^H using the Householder vectors. No care must be taken in choosing forward or backward accumulation.

General case In the general case it is necessary to find an elementary unitary matrix

$$H = I - \tau \cdot u \cdot u^H \quad (15.135)$$

that satisfies the following three conditions.

$$|\tau|^2 \cdot u^H \cdot u = \tau + \tau^* = 2 \cdot \text{Re} \{ \tau \} \quad , \quad H^H \cdot x = \gamma \cdot \|x\| \cdot \vec{e}_1 \quad , \quad |\gamma| = 1 \quad (15.136)$$

When choosing the elements $u_{ii} = 1$ it is possible to store the Householder vectors as well as the upper triangular matrix R in the same storage of the matrix A . The Householder matrices H_i can be completely restored from the Householder vectors.

$$A = \begin{bmatrix} \beta_1 & a_{12} & a_{13} & a_{14} \\ u_{21} & \beta_2 & a_{23} & a_{24} \\ u_{31} & u_{32} & \beta_3 & a_{34} \\ u_{41} & u_{42} & u_{43} & \beta_4 \end{bmatrix} \quad (15.137)$$

There exist several approaches to meet the conditions expressed in eq. (??). For fewer computational effort it may be convenient to choose γ to be real valued. With the notation

$$H^H \cdot x = H^H \cdot \begin{bmatrix} \alpha \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} = \begin{bmatrix} \beta \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (15.138)$$

one possibility is to define the following calculation rules.

$$\nu = \text{sign}(\text{Re} \{ \alpha \}) \cdot \|x\| \quad (15.139)$$

$$\tau = \frac{\alpha + \nu}{\nu} \quad (15.140)$$

$$\gamma = -1 \quad (15.141)$$

$$\beta = \gamma \cdot \|x\| = -\|x\| \quad \rightarrow \quad \text{real valued} \quad (15.142)$$

$$u = \frac{x + \nu \cdot \vec{e}_1}{\alpha + \nu} \quad \rightarrow \quad u_{ii} = 1 \quad (15.143)$$

These definitions yield a complex τ , thus H is no more Hermitian but still unitary.

$$H = I - \tau \cdot u \cdot u^H \quad \rightarrow \quad H^H = I - \tau^* \cdot u \cdot u^H \quad (15.144)$$

Givens rotation

A Givens rotation is a plane rotation matrix. Such a plane rotation matrix is an orthogonal matrix that is different from the identity matrix only in four elements.

$$M = \begin{bmatrix} 1 & 0 & \cdots & & & \cdots & 0 \\ 0 & \ddots & & & & & \vdots \\ \vdots & & 1 & & & & \\ & & & +c & 0 & \cdots & 0 & +s \\ & & & 0 & 1 & & & 0 \\ & & & \vdots & & \ddots & & \vdots \\ & & & 0 & & & 1 & 0 \\ & & & -s & 0 & \cdots & 0 & +c \\ & & & & & & & 1 & \vdots \\ \vdots & & & & & & & & \ddots & 0 \\ 0 & \cdots & & & & & \cdots & 0 & & 1 \end{bmatrix} \quad (15.145)$$

The elements are usually chosen so that

$$R = \begin{bmatrix} c & s \\ -s & c \end{bmatrix} \quad c = \cos \theta, \quad s = \sin \theta \quad \rightarrow \quad |c|^2 + |s|^2 = 1 \quad (15.146)$$

The most common use of such a plane rotation is to choose c and s such that for a given a and b

$$R = \begin{bmatrix} c & s \\ -s & c \end{bmatrix} \cdot \begin{bmatrix} a \\ b \end{bmatrix} = \begin{bmatrix} d \\ 0 \end{bmatrix} \quad (15.147)$$

multiplication annihilates the lower vector entry. In such an application the matrix R is often termed “Givens rotation” matrix. The following equations satisfy eq. (??) for a given a and b exploiting the conditions given in eq. (??).

$$c = \frac{\pm a}{\sqrt{|a|^2 + |b|^2}} \quad \text{and} \quad s = \frac{\pm b}{\sqrt{|a|^2 + |b|^2}} \quad (15.148)$$

$$d = \sqrt{|a|^2 + |b|^2} \quad (15.149)$$

Eigenvalues of a 2-by-2 matrix

The eigenvalues of a 2-by-2 matrix

$$A = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \quad (15.150)$$

can be obtained directly from the quadratic formula. The characteristic polynomial is

$$\begin{aligned} \det(A - \mu I) &= \det \begin{bmatrix} a - \mu & b \\ c & d - \mu \end{bmatrix} = (a - \mu) \cdot (d - \mu) - bc \\ 0 &= \mu^2 - (a + d) \cdot \mu + (ad - bc) \end{aligned} \quad (15.151)$$

The polynomial yields the two eigenvalues.

$$\mu_{1,2} = \frac{a+d}{2} \pm \sqrt{\left(\frac{a+d}{2}\right)^2 + bc - ad} \quad (15.152)$$

For a symmetric matrix A (i.e. $b = c$) eq.(??) can be rewritten to:

$$\mu_{1,2} = e + d \pm \sqrt{e^2 + b^2} \quad \text{with} \quad e = \frac{a-d}{2} \quad (15.153)$$

Step 1: Bidiagonalization

In the first step the original matrix A is bidiagonalized by the application of Householder reflections from the left and right hand side. The matrices U_B^H and V_B can each be determined as a product of Householder matrices.

$$\underbrace{U_B^{H(n)} \cdot \dots \cdot U_B^{H(1)}}_{U_B^H} \cdot A \cdot \underbrace{V_B^{(1)} \cdot \dots \cdot V_B^{(n-2)}}_{V_B} = U_B^H \cdot A \cdot V_B = B^{(0)} \quad (15.154)$$

Each of the required Householder vectors are created and applied as previously defined. Suppose a $n \times n$ matrix, then applying the first Householder vector from the left hand side eliminates the first column and yields

$$U_B^{H(1)} \cdot A = \begin{bmatrix} \beta_1 & a_{12}^{(1)} & a_{13}^{(1)} & \cdots & a_{1n}^{(1)} \\ u_{21} & a_{22}^{(1)} & a_{23}^{(1)} & & a_{2n}^{(1)} \\ u_{31} & a_{32}^{(1)} & a_{33}^{(1)} & & a_{3n}^{(1)} \\ \vdots & & & \ddots & \vdots \\ u_{n1} & a_{n2}^{(1)} & a_{n3}^{(1)} & \cdots & a_{nn}^{(1)} \end{bmatrix} \quad (15.155)$$

Next, a Householder vector is applied from the right hand side to annihilate the first row.

$$U_B^{H(1)} \cdot A \cdot V_B^{(1)} = \begin{bmatrix} \beta_1 & \delta_2 & v_{13} & \cdots & v_{1n} \\ u_{21} & a_{22}^{(2)} & a_{23}^{(2)} & & a_{2n}^{(2)} \\ u_{31} & a_{32}^{(2)} & a_{33}^{(2)} & & a_{3n}^{(2)} \\ \vdots & & & \ddots & \vdots \\ u_{n1} & a_{n2}^{(2)} & a_{n3}^{(2)} & \cdots & a_{nn}^{(2)} \end{bmatrix} \quad (15.156)$$

Again, a Householder vector is applied from the left hand side to annihilate the second column.

$$U_B^{H(2)} \cdot U_B^{H(1)} \cdot A \cdot V_B^{(1)} = \begin{bmatrix} \beta_1 & \delta_2 & v_{13} & \cdots & v_{1n} \\ u_{21} & \beta_2 & a_{23}^{(3)} & & a_{2n}^{(3)} \\ u_{31} & u_{32} & a_{33}^{(3)} & & a_{3n}^{(3)} \\ \vdots & \vdots & & \ddots & \vdots \\ u_{n1} & u_{n2} & a_{n3}^{(3)} & \cdots & a_{nn}^{(3)} \end{bmatrix} \quad (15.157)$$

This process is continued until

$$U_B^H \cdot A \cdot V_B = \begin{bmatrix} \beta_1 & \delta_2 & v_{13} & \cdots & v_{1n} \\ u_{21} & \beta_2 & \delta_3 & & v_{2n} \\ u_{31} & u_{32} & \ddots & \ddots & \\ \vdots & \vdots & & \beta_{n-1} & \delta_n \\ u_{n1} & u_{n2} & u_{n3} & & \beta_n \end{bmatrix} \quad (15.158)$$

For each of the Householder transformations from the left and right hand side the appropriate τ values must be stored in separate vectors.

Step 2: Matrix reconstructions

Using the Householder vectors stored in place of the original A matrix and the appropriate τ value vectors it is now necessary to unpack the U_B and V_B^H matrices. The diagonal vector β and the super-diagonal vector δ can be saved in separate vectors previously. Thus the U_B matrix can be unpacked in place of the A matrix and the V_B^H matrix is unpacked in a separate matrix.

There are two possible algorithms for computing the Householder product matrices, i.e. forward accumulation and backward accumulation. Both start with the identity matrix which is successively multiplied by the Householder matrices either from the left or right.

$$U_B^H = H_{U_n}^H \cdot \dots \cdot H_{U_2}^H \cdot H_{U_1}^H \cdot I \quad (15.159)$$

$$\rightarrow U_B = I \cdot H_{U_n} \cdot \dots \cdot H_{U_2} \cdot H_{U_1} \quad (15.160)$$

Recall that the leading portion of each Householder matrix is the identity except the first. Thus, at the beginning of backward accumulation, U_B is “mostly the identity” and it gradually becomes full as the iteration progresses. This pattern can be exploited to reduce the number of required flops. In contrast, U_B^H is full in forward accumulation after the first step. For this reason, backward accumulation is cheaper and the strategy of choice. When unpacking the U_B matrix in place of the original A matrix it is necessary to choose backward accumulation anyway.

$$V_B = I \cdot H_{V_1}^H \cdot H_{V_2}^H \cdot \dots \cdot H_{V_n}^H \quad (15.161)$$

$$\rightarrow V_B^H = I \cdot H_{V_n} \cdot \dots \cdot H_{V_2} \cdot H_{V_1} \quad (15.162)$$

Unpacking the V_B^H matrix is done in a similar way also performing successive Householder matrix multiplications using backward accumulation.

Step 3: Diagonalization – shifted QR iteration

At this stage the matrices U_B and V_B^H exist in unfactored form. Also there are the diagonal vector β and the super-diagonal vector δ . Both vectors are real valued. Thus the following algorithm can be applied even though solving a complex equation system.

$$B^{(0)} = \begin{bmatrix} \beta_1 & \delta_2 & 0 & \cdots & 0 \\ 0 & \beta_2 & \delta_3 & & 0 \\ \vdots & 0 & \ddots & \ddots & 0 \\ 0 & 0 & 0 & \beta_{n-1} & \delta_n \\ 0 & 0 & 0 & 0 & \beta_n \end{bmatrix} \quad (15.163)$$

The remaining problem is thus to compute the SVD of the matrix B . This is done applying an implicit-shift QR step to the tridiagonal matrix $T = B^T B$ which is a symmetric. The matrix T is not explicitly formed that is why a QR iteration with implicit shifts is applied.

After bidiagonalization we have a bidiagonal matrix $B^{(0)}$:

$$B^{(0)} = U_B^H \cdot A \cdot V_B \quad (15.164)$$

The presented method turns $B^{(k)}$ into a matrix $B^{(k+1)}$ by applying a set of orthogonal transforms

$$B^{(k+1)} = \tilde{U}^H \cdot B^{(k)} \cdot \tilde{V} \quad (15.165)$$

The orthogonal matrices $\tilde{U}$ and $\tilde{V}$ are chosen so that $B^{(k+1)}$ is also a bidiagonal matrix, but with the super-diagonal elements smaller than those of $B^{(k)}$. The eq.(??) is repeated until the non-diagonal elements of $B^{(k+1)}$ become smaller than ε and can be disregarded.

The matrices $\tilde{U}$ and $\tilde{V}$ are constructed as

$$\tilde{U} = \tilde{U}_1 \cdot \tilde{U}_2 \cdot \tilde{U}_3 \cdot \dots \cdot \tilde{U}_{n-1} \quad (15.166)$$

and similarly $\tilde{V}$ where $\tilde{V}_i$ and $\tilde{U}_i$ are matrices of simple rotations as given in eq.(??). Both $\tilde{V}$ and $\tilde{U}$ are products of Givens rotations and thus perform orthogonal transforms.

Single shifted QR step. The left multiplication of $B^{(k)}$ by $\tilde{U}_i^H$ replaces two rows of $B^{(k)}$ by their linear combinations. The rest of $B^{(k)}$ is unaffected. Right multiplication of $B^{(k)}$ by $\tilde{V}_i$ similarly changes only two columns of $B^{(k)}$.

A matrix $\tilde{V}_1$ is chosen the way that

$$B_1^{(k)} = B_0^{(k)} \cdot \tilde{V}_1 \quad (15.167)$$

is a QR transform with a shift. Note that multiplying $B^{(k)}$ by $\tilde{V}_1$ gives rise to a non-zero element which is below the main diagonal.

$$B_0^{(k)} \cdot \tilde{V}_1 = \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ \boxed{\otimes} & \times & \times & 0 & 0 & 0 \\ 0 & 0 & \times & \times & 0 & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} \quad (15.168)$$

A new rotation angle is then chosen so that multiplication by $\tilde{U}_1^H$ gets rid of that element. But this will create a non-zero element which is right beside the super-diagonal.

$$\tilde{U}_1^H \cdot B_1^{(k)} = \begin{bmatrix} \times & \times & \boxed{\otimes} & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & 0 & \times & \times & 0 & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} \quad (15.169)$$

Then $\tilde{V}_2$ is made to make it disappear, but this leads to another non-zero element below the diagonal, etc.

$$B_2^{(k)} \cdot \tilde{V}_2 = \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & \boxed{\otimes} & \times & \times & 0 & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} \quad (15.170)$$

In the end, the matrix $\tilde{U}^H B \tilde{V}$ becomes bidiagonal again. However, because of a special choice of $\tilde{V}_1$ (QR algorithm), its non-diagonal elements are smaller than those of B .

Please note that each of the transforms must also be applied to the unfactored U_B^H and V_B matrices which turns them successively into U^H and V

Computation of the Wilkinson shift. For a single QR step the computation of the eigenvalue μ of the trailing 2-by-2 submatrix of $T_n = B_n^T \cdot B_n$ that is closer to the t_{22} matrix element is required.

$$T_n = \begin{bmatrix} t_{11} & t_{12} \\ t_{21} & t_{22} \end{bmatrix} = B_n^T \cdot B_n = \begin{bmatrix} \delta_{n-1} & \beta_{n-1} & 0 \\ 0 & \delta_n & \beta_n \end{bmatrix} \cdot \begin{bmatrix} \delta_{n-1} & 0 \\ \beta_{n-1} & \delta_n \\ 0 & \beta_n \end{bmatrix} \quad (15.171)$$

$$= \begin{bmatrix} \beta_{n-1}^2 + \delta_{n-1}^2 & \delta_n \cdot \beta_{n-1} \\ \delta_n \cdot \beta_{n-1} & \beta_n^2 + \delta_n^2 \end{bmatrix} \quad (15.172)$$

The required eigenvalue is called Wilkinson shift, see eq.(??) for details. The sign for the eigenvalue is chosen such that it is closer to t_{22} .

$$\mu = t_{22} + d - \text{sign}(d) \cdot \sqrt{d^2 + t_{12}^2} \quad (15.173)$$

$$= t_{22} + d - t_{12} \cdot \text{sign}\left(\frac{d}{t_{12}}\right) \cdot \sqrt{\left(\frac{d}{t_{12}}\right)^2 + 1} \quad (15.174)$$

$$= t_{22} - \frac{t_{12}^2}{d + t_{12} \cdot \text{sign}\left(\frac{d}{t_{12}}\right) \cdot \sqrt{\left(\frac{d}{t_{12}}\right)^2 + 1}} \quad (15.175)$$

whereas

$$d = \frac{t_{11} - t_{22}}{2} \quad (15.176)$$

The Givens rotation $\tilde{V}_1$ is chosen such that

$$\begin{bmatrix} c_1 & s_1 \\ -s_1 & c_1 \end{bmatrix}^T \cdot \begin{bmatrix} (\beta_1^2 - \mu)/\beta_1 \\ \delta_2 \end{bmatrix} = \begin{bmatrix} \times \\ 0 \end{bmatrix} \quad (15.177)$$

The special choice of this first rotation in the single QR step ensures that the super-diagonal matrix entries get smaller. Typically, after a few of these QR steps, the super-diagonal entry δ_n becomes negligible.

Zeros on the diagonal or super-diagonal. The QR iteration described above claims to hold if the underlying bidiagonal matrix is unreduced, i.e. has no zeros neither on the diagonal nor on the super-diagonal.

When there is a zero along the diagonal, then premultiplication by a sequence of Givens transformations can zero the right-hand super-diagonal entry as well. The inverse rotations must be applied to the U_B^H matrix.

$$\begin{aligned}
 B = \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & 0 & 0 & \times & 0 & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} & \rightarrow \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & \otimes & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} \\
 \rightarrow \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \otimes \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix} & \rightarrow \begin{bmatrix} \times & \times & 0 & 0 & 0 & 0 \\ 0 & \times & \times & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \times & \times & 0 \\ 0 & 0 & 0 & 0 & \times & \times \\ 0 & 0 & 0 & 0 & 0 & \times \end{bmatrix}
 \end{aligned}$$

Thus the problem can be decoupled into two smaller matrices B_1 and B_2 . The diagonal matrix B_3 is successively getting larger for each super-diagonal entry being neglected after the QR iterations.

$$\begin{bmatrix} B_1 & 0 & 0 \\ 0 & B_2 & 0 \\ 0 & 0 & B_3 \end{bmatrix} \quad (15.178)$$

Matrix B_2 has non-zero super-diagonal entries. If there is any zero diagonal entry in B_2 , then the super-diagonal entry can be annihilated as just described. Otherwise the QR iteration algorithm can be applied to B_2 .

When there are only B_3 matrix entries left (diagonal entries only) the algorithm is finished, then the B matrix has been transformed into the singular value matrix Σ .

Step 4: Solving the equation system

It is straight-forward to solve a given equation system once having the singular value decomposition computed.

$$A \cdot x = z \quad (15.179)$$

$$U \Sigma V^H \cdot x = z \quad (15.180)$$

$$\Sigma V^H \cdot x = U^H \cdot z \quad (15.181)$$

$$V^H \cdot x = \Sigma^{-1} U^H \cdot z \quad (15.182)$$

$$x = V \Sigma^{-1} U^H \cdot z \quad (15.183)$$

The inverse of the diagonal matrix Σ yields

$$\Sigma^{-1} = \begin{bmatrix} 1/\sigma_1 & 0 & \cdots & 0 \\ 0 & 1/\sigma_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1/\sigma_n \end{bmatrix} \quad (15.184)$$

With v_i being the i -th row of the matrix V , u_i the i -th column of the matrix U and σ_i the i -th singular value eq. (??) can be rewritten to

$$x = \sum_{i=1}^n \frac{u_i^H \cdot z}{\sigma_i} \cdot v_i \quad (15.185)$$

It must be mentioned that very small singular values σ_i corrupt the complete result. Such values indicate (nearly) singular (ill-conditioned) matrices A . In such cases, the solution vector x obtained by zeroing the small σ_i 's and then using equation (??) is better than direct-method solutions (such as LU decomposition or Gaussian elimination) and the SVD solution where the small σ_i 's are left non-zero. It may seem paradoxical that this can be so, since zeroing a singular value corresponds to throwing away one linear combination of the set of equations that is going to be solved. The resolution of the paradox is that a combination of equations that is so corrupted by roundoff error is thrown away precisely as to be at best useless; usually it is worse than useless since it "pulls" the solution vector way off towards infinity along some direction that is almost a nullspace vector.

15.2.7 Jacobi method

This method quite simply involves rearranging each equation to make each variable a function of the other variables. Then make an initial guess for each solution and iterate. For this method it is necessary to ensure that all the diagonal matrix elements a_{ii} are non-zero. This is given for the nodal analysis and almostly given for the modified nodal analysis. If the linear equation system is solvable this can always be achieved by rows substitutions.

The algorithm for performing the iteration step $k + 1$ writes as follows.

$$x_i^{(k+1)} = \frac{1}{a_{ii}} \left(z_i - \sum_{j=1}^{i-1} a_{ij} x_j^{(k)} - \sum_{j=i+1}^n a_{ij} x_j^{(k)} \right) \quad \text{for } i = 1, \dots, n \quad (15.186)$$

This has to repeated until the new solution vectors $x^{(k+1)}$ deviation from the previous one $x^{(k)}$ is sufficiently small.

The initial guess has no effect on whether the iterative method converges or not, but with a good initial guess (as possibly given in consecutive Newton-Raphson iterations) it converges faster (if it converges). To ensure convergence the condition

$$\sum_{j=1, j \neq i}^n |a_{ij}| \leq |a_{ii}| \quad \text{for } i = 1, \dots, n \quad (15.187)$$

and at least one case

$$\sum_{i=1, i \neq j}^n |a_{ij}| \leq |a_{ii}| \quad (15.188)$$

must apply. If these conditions are not met, the iterative equations may still converge. If these conditions are met the iterative equations will definitely converge.

Another simple approach to a convergence criteria for iterative algorithms is the Schmidt and v. Mises criteria.

$$\sqrt{\sum_{i=1}^n \sum_{j=1, j \neq i}^n \left| \frac{a_{ij}}{a_{ii}} \right|^2} < 1 \quad (15.189)$$

15.2.8 Gauss-Seidel method

The Gauss-Seidel algorithm is a modification of the Jacobi method. It uses the previously computed values in the solution vector of the same iteration step. That is why this iterative method is expected to converge faster than the Jacobi method.

The slightly modified algorithm for performing the $k + 1$ iteration step writes as follows.

$$x_i^{(k+1)} = \frac{1}{a_{ii}} \left(z_i - \sum_{j=1}^{i-1} a_{ij} x_j^{(k+1)} - \sum_{j=i+1}^n a_{ij} x_j^{(k)} \right) \quad \text{for } i = 1, \dots, n \quad (15.190)$$

The remarks about the initial guess $x^{(0)}$ as well as the convergence criteria noted in the section about the Jacobi method apply to the Gauss-Seidel algorithm as well.

15.2.9 A comparison

There are direct and iterative methods (algorithms) for solving linear equation systems. Equation systems with large and sparse matrices should rather be solved with iterative methods.

method	precision	application	programming effort	computing complexity	notes
Laplace expansion	numerical errors	general	straight forward	$n!$	very time consuming
Gaussian elimination	numerical errors	general	intermediate	$\frac{n^3}{3} + \frac{n^2}{2}$	
Gauss-Jordan	numerical errors	general	intermediate	$\frac{n^3}{3} + \frac{n^2}{3}$	computes the inverse besides
LU decomposition	numerical errors	general	intermediate	$\frac{n^3}{3} + \frac{n^2}{3}$	useful for consecutive solutions
QR decomposition	good	general	high	$2n^3 + 3n^3$	
Singular value decomposition	good	general	very high	$2n^3 + 4n^3$	ill-conditioned matrices can be handled
Jacobi	very good	diagonally dominant systems	easy	n^2 in each iteration step	possibly no convergence
Gauss-Seidel	very good	diagonally dominant systems	easy	n^2 in each iteration step	possibly no convergence

15.3 Polynomial approximations

15.3.1 Cubic splines

15.4 Frequency-Time Domain Transformation

Any signal can completely be described in time or in frequency domain. As both representations are equivalent, it is possible to transform them into each other. This is done by the so-called Fourier Transformation and the inverse Fourier Transformation, respectively:

$$\text{Fourier Transformation:} \quad \underline{U}(j\omega) = \int_{-\infty}^{\infty} u(t) \cdot e^{-j\omega \cdot t} dt \quad (15.191)$$

$$\text{inverse Fourier Transformation:} \quad u(t) = \frac{1}{2\pi} \cdot \int_{-\infty}^{\infty} \underline{U}(j\omega) \cdot e^{j\omega \cdot t} d\omega \quad (15.192)$$

In digital systems the data $u(t)$ or $\underline{U}(j\omega)$, respectively, consists of a finite number N of samples u_k and $\underline{U}_n$. This leads to the discrete Fourier Transformation (DFT) and its inverse operation (IDFT):

$$\text{DFT:} \quad \underline{U}_n = \sum_{k=0}^{N-1} u_k \cdot \exp\left(-j \cdot n \frac{2\pi \cdot k}{N}\right) \quad (15.193)$$

$$\text{IDFT:} \quad u_k = \frac{1}{N} \cdot \sum_{n=0}^{N-1} \underline{U}_n \cdot \exp\left(j \cdot k \frac{2\pi \cdot n}{N}\right) \quad (15.194)$$

The absolute time and frequency values do not appear anymore in the DFT. They depend on the sample frequency f_T and the number of samples N .

$$\Delta f = \frac{1}{N \cdot \Delta t} = \frac{f_T}{N} \quad (15.195)$$

Where Δt is distance between time samples and Δf distance between frequency samples.

With DFT the N time samples are transformed into N frequency samples. This also holds if the time data are real numbers, as is always the case in "real life": The complex frequency samples are conjugate complex symmetrical and so equalizing the score:

$$\underline{U}_{N-n} = \underline{U}_n^* \quad (15.196)$$

That is, knowing the input data has no imaginary part, only half of the Fourier data must be computed.

15.4.1 Fast Fourier Transformation

As can be seen in equation ?? the computing time of the DFT rises with N^2 . This is really huge, so it is very important to reduce the time consumption. Using a strongly optimized algorithm, the so-called Fast Fourier Transformation (FFT), the DFT is reduced to an $N \cdot \log_2 N$ time rise. The following information stems from [?], where the theoretical background is explained comprehensively.

The fundamental trick of the FFT is to cut the DFT into two parts, one with data having even indices and the other with odd indices:

$$\underline{U}_n = \sum_{k=0}^{N-1} u_k \cdot \exp\left(-j \cdot n \frac{2\pi \cdot k}{N}\right) \quad (15.197)$$

$$= \sum_{k=0}^{N/2-1} u_{2k} \cdot \exp\left(-j \cdot n \frac{2\pi \cdot 2k}{N}\right) + \sum_{k=0}^{N/2-1} u_{2k+1} \cdot \exp\left(-j \cdot n \frac{2\pi \cdot (2k+1)}{N}\right) \quad (15.198)$$

$$= \underbrace{\sum_{k=0}^{N/2-1} u_{2k} \cdot \exp\left(-j \cdot n \frac{2\pi \cdot k}{N/2}\right)}_{F_{even}} + W_{n,N} \cdot \underbrace{\sum_{k=0}^{N/2-1} u_{2k+1} \cdot \exp\left(-j \cdot n \frac{2\pi \cdot k}{N/2}\right)}_{F_{odd}} \quad (15.199)$$

$$\text{with } W_{n,N} = \exp\left(2\pi \cdot j \cdot \frac{n}{N}\right) \quad (\text{so-called 'twiddle factor'}) \quad (15.200)$$

The new formula shows no speed advantages. The important thing is that the even as well as the odd part each is again a Fourier series. Thus the same procedure can be repeated again and again until the equation consists of N terms. Then, each term contains only one data u_k with factor $e^0 = 1$. This works if the number of data is a power of two (2, 4, 8, 16, 32, ...). So finally, the FFT method performs $\log_2 N$ times the operation

$$u_{k1,even} + W_{n,x} \cdot u_{k2,odd} \quad (15.201)$$

to get one data of $\underline{U}_n$. This is called the Danielson-Lanzcos algorithm. The question now arises which data values of u_k needs to be combined according to equation (??). The answer is quite

easy. The data array must be reordered by the bit-reversal method. This means the value at index k_1 is swapped with the value at index k_2 where k_2 is obtained by mirroring the binary number k_1 , i.e. the most significant bit becomes the least significant one and so on. Example for $N = 8$:

000	$\leftrightarrow$	000	011	$\leftrightarrow$	110	110	$\leftrightarrow$	011
001	$\leftrightarrow$	100	100	$\leftrightarrow$	001	111	$\leftrightarrow$	111
010	$\leftrightarrow$	010	101	$\leftrightarrow$	101			

Having this new indexing, the values to combine according to equation ?? are the adjacent values. So, performing the Danielson-Lanczos algorithm has now become very easy.

Figure ?? illustrates the whole FFT algorithm starting with the input data u_k and ending with one value of the output data $\underline{U}_n$.

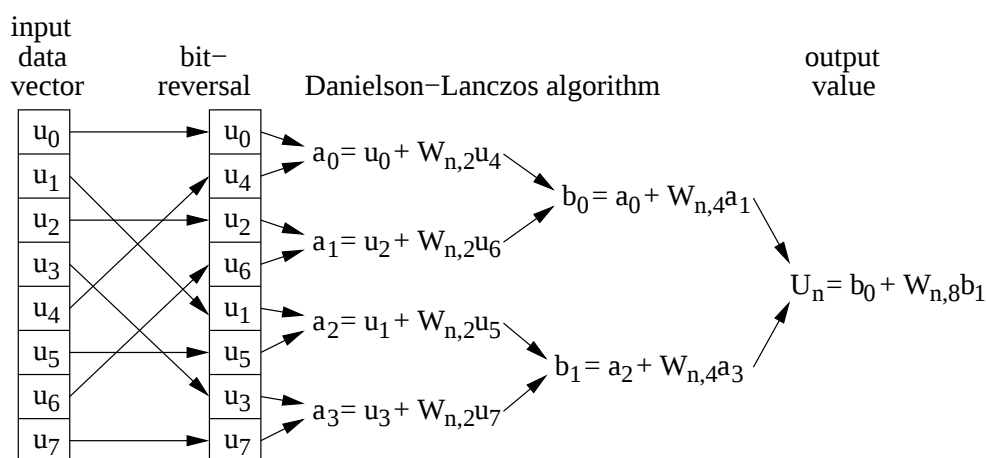


Figure 15.6: principle of a FFT with data length 8

This scheme alone gives no advantage. But it can compute all output values within, i.e. no temporary memory is needed and the periodicity of $W_{n,N}$ is best exploited. To understand this, let's have a look on the first Danielson-Lanczos step in figure ?. The four multiplications and additions have to be performed for each output value (here 8 times!). But indeed this is not true, because $W_{n,2}$ is 2-periodical in n and furthermore $W_{n,2} = -W_{n+1,2}$. So now, $u_0 + W_{0,2} \cdot u_4$ replaces the old u_0 value and $u_0 - W_{0,2} \cdot u_4$ replaces the old u_4 value. Doing this for all values, four multiplications and eight additions were performed in order to calculate the first Danielson-Lanczos step for all (!!!) output values. This goes on, as $W_{n,4}$ is 4-periodical in n and $W_{n,4} = -W_{n+2,4}$. So this time, two loop iterations (for $W_{n,4}$ and for $W_{n+1,4}$) are necessary to compute the current Danielson-Lanczos step for all output values. This concept continues until the last step.

Finally, a complete FFT source code in C should be presented. The original version was taken from [?]. It is a radix-2 algorithm, known as the Cooley-Tukey Algorithm. Here, several changes were made that gain about 10% speed improvement.

Listing 15.1: 1D-FFT algorithm in C

```

1 // *****
2 // Parameters:
3 // num      - number of complex samples
4 // data[]   - array containing the data samples, real and imaginary
5 //            part in alternating order (length: 2*num)
6 // isign    - is 1 to calculate FFT and -1 to calculate inverse FFT
7 // *****
8
9 #define SWAP(a,b) { wr = a; a = b; b = wr; }
10
11 void fft_radix2(int num, double *data, int isign)
12 {
13     double wt, theta, wr, wi, wpr, wpi, tempi, tempr;
14     int i, j, m, n;
15     n = 2*num;
16     j = 0;
17
18     // bit reversal method
19     // 1) index 0 need not to be swapped
20     //    -> start at i=2
21     // 2) swap scheme is symmetrical
22     //    -> swap first and second half in one iteration
23     for (i=2; i<num; i+=2) {
24         m = num;
25         while (j >= m) { // calculate swap index
26             j -= m;
27             m >>= 1;
28         }
29         j += m;
30
31         if (j > i) { // was index already swapped ?
32             SWAP(data[j], data[i]); // swap real part
33             SWAP(data[j+1], data[i+1]); // swap imaginary part
34
35             if (j < num) { // swap second half ?
36                 SWAP (data[n-j-2], data[n-i-2]); // swap real part
37                 SWAP (data[n-j-1], data[n-i-1]); // swap imaginary part
38             }
39         }
40     }
41
42     // Danielson-Lanzcos algorithm
43     int mmax, istep;
44     mmax = 2;
45     while (n > mmax) { // each Danielson-Lanzcos step
46         istep = mmax << 1;
47         theta = isign * (2.0 * PI / mmax);
48         wpr = cos(theta);
49         wpi = sin(theta);
50         wr = 1.0;

```

```

51     wi = 0.0;
52     for (m=1; m<mmax; m+=2) {
53         for (i=m; i<=n; i+=istep) {
54             j = i+mmax;
55             tempr = wr*data[j-1] + wi*data[j];
56             tempi = wr*data[j] - wi*data[j-1];
57             data[j-1] = data[i-1] - tempr;
58             data[j] = data[i] - tempi;
59             data[i-1] += tempr;
60             data[i] += tempi;
61         }
62         wt = wr;
63         wr = wt*wpr - wi*wpi;
64         wi = wi*wpr + wt*wpi;
65     }
66     mmax = istep;
67 }
68
69 if (istep == -1)           // perform inverse FFT ?
70     for (i=0; i<num; i++)
71         data[i] /= num;    // normalize result
72 }

```

There are many other FFT algorithms mainly aiming at higher speed (radix-8 FFT, split-radix FFT, Winograd FFT). These algorithms are much more complex, but on modern processors with numerical co-processors they gain no or hardly no speed advantage, because the reduced FLOPS are equalled by the far more complex indexing.

15.4.2 Real-Valued FFT

All physical systems are real-valued in time domain. As already mentioned above, this fact leads to a symmetry in frequency domain, which can be exploited to save 50% memory usage and about 30% computation time. Rewriting the C listing from above to a real-valued FFT routine creates the following function. As this scheme is not symmetric anymore, an extra procedure for the inverse transformation is needed. It is also depicted below.

Listing 15.2: real-valued FFT algorithm in C

```

1  // *****
2  // Parameters:
3  // num      - number of real-valued samples
4  // data[]   - array containing the data samples (length: num)
5  //
6  // Output:
7  // data[]   - r(0), r(1), i(1), ..., r(N/2-1), i(N/2-1), r(N/2)
8  // *****
9
10 #define SWAP(a,b) { wr = a; a = b; b = wr; }
11
12 void real_fft_radix2(int num, double *data)
13 {

```

```

14  int i, j, k, l, n1 = num >> 1, n2 = 1;
15  double t1, t2, t3, wr, wi, wpr, wpi;
16
17  // bit reversal method
18  // 1) index 0 need not to be swapped
19  //    -> start at i=1
20  // 2) swap scheme is symmetrical
21  //    -> swap first and second half in one iteration
22  j = 0;
23  for (i=1; i<n1; i++) {
24      k = n1;
25      while (j >= k) { // calculate swap index
26          j -= k;
27          k >>= 1;
28      }
29      j += k;
30
31      if (j > i) { // was index already swapped ?
32          SWAP(data[j], data[i]);
33
34          if (j < n1) // swap second half ?
35              SWAP (data[num-j-1], data[num-i-1]);
36      }
37  }
38
39  // length two butterflies
40  for (i=0; i<num; i+=2) {
41      t1 = data[i+1];
42      data[i+1] = data[i] - t1;
43      data[i] += t1;
44  }
45
46  while (n1 < num) {
47      n2 <=<= 1; // half a butterfly
48      n1 = n2 << 1; // length of a butterfly
49
50      for (i=0; i<num; i+=n1) {
51          t1 = data[i+n2];
52          data[i+n2] = -data[i+n1-1];
53          data[i+n1-1] = data[i] - t1;
54          data[i] += t1;
55      }
56
57      t1 = 2.0*M_PI / ((double)n1);
58      wpr = cos(t1); // real part of twiddle factor
59      wpi = sin(t1); // imaginary part of twiddle factor
60      wr = 1.0; // start of twiddle factor
61      wi = 0.0;
62
63      for (j=3; j<n2; j+=2) { // all complex lines of a butterfly
64          t1 = wr;
65          wr = t1*wpr - wi*wpi; // calculate next twiddle factor
66          wi = wi*wpr + t1*wpi;

```

```

67
68     for (i=0; i<num; i+=n1) { // through all butterflies
69         k = i + j - 2;
70         l = i + n1 - j;
71         t1 = data[l]*wr + data[l+1]*wi;
72         t3 = data[k+1];
73         t2 = data[l+1]*wr - data[l]*wi;
74         data[l] = data[k];
75
76         if((i & n1) != 0) { // index swap ?
77             t1 = -t1;
78             t3 = -t3;
79         }
80
81         data[k] += t1;
82         data[k+1] = t2 + t3;
83         data[l] -= t1;
84         data[l+1] = t2 - t3;
85     }
86 }
87 }
88 }

```

Listing 15.3: real-valued inverse FFT algorithm in C

```

1 // *****
2 // Parameters:
3 // num - count of numbers in data
4 // data[] - r(0), r(1), i(1), .... , r(N/2-1), i(N/2-1), r(N/2)
5 //
6 // Output:
7 // data[] - array containing the data samples (length: num)
8 // *****
9
10 #define SWAP(a,b) { wr = a; a = b; b = wr; }
11
12 void real_ifft_order_speed(int num, double *data)
13 {
14     int i, j, k, l, n1, n2 = num;
15     double t1, t2, t3, wr, wi, wpr, wpi;
16
17     while (n2 > 2) {
18         n1 = n2; // length of a butterfly
19         n2 >>= 1; // half a butterfly
20
21         for (i=0; i<num; i+=n1) { // through all butterflies
22             t1 = data[i+n1-1];
23             data[i+n1-1] = -2.0 * data[i+n2];
24             data[i+n2-1] *= 2.0;
25             data[i+n2] = data[i] - t1;
26             data[i] += t1;
27         }
28     }

```

```

29     t1 = 2.0*M_PI / ((double)n1);
30     wpr = cos(t1); // real part of twiddle factor
31     wpi = sin(t1); // imaginary part of twiddle factor
32     wr = 1.0; // start of twiddle factor
33     wi = 0.0;
34
35     for (j=3; j<n2; j+=2) { // all complex lines of a butterfly
36         t1 = wr;
37         wr = t1*wpr + wi*wpi; // calculate next twiddle factor
38         wi = wi*wpr - t1*wpi;
39
40         for (i=0; i<num; i+=n1) { // through all butterflies
41             k = i + j - 2;
42             l = i + n1 - j;
43             t1 = data[l] - data[k];
44             t2 = data[l+1] + data[k+1];
45             t3 = data[k+1] - data[l+1];
46             data[k] += data[l];
47
48             if((i & n1) != 0) {
49                 t1 = -t1;
50                 t3 = -t3;
51             }
52
53             data[k+1] = t3;
54             data[l] = t2*wi - t1*wr;
55             data[l+1] = t2*wr + t1*wi;
56         }
57     }
58 }
59
60 // length two butterflies
61 for (i=0; i<num; i+=2) {
62     t1 = data[i+1];
63     data[i+1] = (data[i] - t1) / num;
64     data[i] = (data[i] + t1) / num;
65 }
66
67 // bit reversal method
68 // 1) index 0 need not to be swapped
69 // -> start at i=1
70 // 2) swap scheme is symmetrical
71 // -> swap first and second half in one iteration
72 j = 0;
73 n1 = num >> 1;
74 for (i=1; i<n1; i++) {
75     k = n1;
76     while (j >= k) { // calculate swap index
77         j -= k;
78         k >>= 1;
79     }
80     j += k;
81

```

```

82     if(j > i) { // was index already swapped ?
83         SWAP(data[j], data[i]);
84
85         if(j < n1) // swap second half ?
86             SWAP (data[num-j-1], data[num-i-1]);
87     }
88 }
89 }

```

15.4.3 More-Dimensional FFT

A standard Fourier Transformation is not useful in harmonic balance methods, because with multi-tone excitation many mixing products appear. The best way to cope with this problem is to use multi-dimensional FFT.

Fourier Transformations in more than one dimension soon become very time consuming. Using FFT mechanisms is therefore mandatory. A more-dimensional Fourier Transformation consists of many one-dimensional Fourier Transformations (1D-FFT). First, 1D-FFTs are performed for the data of the first dimension at every index of the second dimension. The results are used as input data for the second dimension that is performed the same way with respect to the third dimension. This procedure is continued until all dimensions are calculated. The following equations shows this for two dimensions.

$$\underline{U}_{n1,n2} = \sum_{k_2=0}^{N_2-1} \sum_{k_1=0}^{N_1-1} u_{k_1,k_2} \cdot \exp\left(-j \cdot n_1 \frac{2\pi \cdot k_1}{N_1}\right) \cdot \exp\left(-j \cdot n_2 \frac{2\pi \cdot k_2}{N_2}\right) \quad (15.202)$$

$$= \sum_{k_2=0}^{N_2-1} \exp\left(-j \cdot n_2 \frac{2\pi \cdot k_2}{N_2}\right) \cdot \underbrace{\sum_{k_1=0}^{N_1-1} u_{k_1,k_2} \cdot \exp\left(-j \cdot n_1 \frac{2\pi \cdot k_1}{N_1}\right)}_{\text{1D-FFT}} \quad (15.203)$$

Finally, a complete n -dimensional FFT source code should be presented. It was taken from [?] and somewhat speed improved.

Parameters:

- ndim** - number of dimensions
- num[]** - array containing the number of complex samples for every dimension
- data[]** - array containing the data samples,
real and imaginary part in alternating order (length: 2*sum of num[]),
going through the array, the first dimension changes least rapidly !
all subscripts range from 1 to maximum value !
- isign** - is 1 to calculate FFT and -1 to calculate inverse FFT

Listing 15.4: multidimensional FFT algorithm in C

```

1  int idim, i1, i2, i3, i2rev, i3rev, ip1, ip2, ip3, ifp1, ifp2;
2  int ibit, k1, k2, n, nprev, nrem, ntot;
3  double tempi, tempr, wt, theta, wr, wi, wpi, wpr;
4

```

```

5  ntot = 1;
6  for (idim=0; idim<ndim; idim++) // compute total number of complex values
7      ntot *= num[idim];
8
9  nprev = 1;
10 for (idim=ndim-1; idim>=0; idim--) { // main loop over the dimensions
11     n = num[idim];
12     nrem = ntot/(n*nprev);
13     ip1 = nprev << 1;
14     ip2 = ip1*n;
15     ip3 = ip2*nrem;
16     i2rev = 1;
17
18     for (i2=1; i2<=ip2; i2+=ip1) { // bit-reversal method
19         if (i2 < i2rev) {
20             for (i1=i2; i1<=i2+ip1-2; i1+=2) {
21                 for (i3=i1; i3<=ip3; i3+=ip2) {
22                     i3rev = i2rev+i3-i2;
23                     SWAP(data[i3-1],data[i3rev-1]);
24                     SWAP(data[i3],data[i3rev]);
25                 }
26             }
27         }
28         ibit=ip2 >> 1;
29         while (ibit >= ip1 && i2rev > ibit) {
30             i2rev -= ibit;
31             ibit >>= 1;
32         }
33         i2rev += ibit;
34     }
35
36     ifp1 = ip1;
37     while (ifp1 < ip2) { // Danielson-Lanzcos algorithm
38         ifp2 = ifp1 << 1;
39         theta = isign*2*pi/(ifp2/ip1);
40         wpr = cos(theta);
41         wpi = sin(theta);
42         wr = 1.0; wi = 0.0;
43         for (i3=1; i3<=ifp1; i3+=ip1) {
44             for (i1=i3; i1<=i3+ip1-2; i1+=2) {
45                 for (i2=i1; i2<=ip3; i2+=ifp2) {
46                     k1 = i2;
47                     k2 = k1+ifp1;
48                     tempr = wr*data[k2-1] - wi*data[k2];
49                     tempi = wr*data[k2] + wi*data[k2-1];
50                     data[k2-1] = data[k1-1] - tempr;
51                     data[k2] = data[k1] - tempi;
52                     data[k1-1] += tempr; data[k1] += tempi;
53                 }
54             }
55             wt = wr;
56             wr = wt*wpr - wi*wpi;
57             wi = wi*wpr + wt*wpi;

```

```
58     }  
59     ifp1 = ifp2;  
60     }  
61     nprev *= n;  
62 }
```

Appendix A

Qucs file formats

Qucs uses plain-text (ASCII) files as its input and transfer format for netlists and data. This appendix explains the file formats by describing the grammars of the languages used. The files are generally line-oriented but arbitrary whitespace between the token is allowed. You can also use the backslash (\) to continue a line on the next line. This works almost everywhere but in comment lines.

The grammars are presented using a version of the Extended Backus-Naur Form (EBNF) which works as follows:

- $A \rightarrow B$ Nonterminal A produces sentential form B .
- $B|C$ Produces B or C .
- $\{A\}$ Arbitrary repetition of form A . No repetition is allowed as well (“Kleene operator”).
- $[A]$ Form A is optional.
- (A) Grouping, stands for A itself.

Nonterminal symbols are set in normal font, terminal symbols are in bold font. Terminal symbols in single quotes are literally found in the input while the other terminal symbols are compositions. See below for definition of composed terminal symbols.

A.1 Qucs netlist grammar

Syntactic Structure

Input → { InputLine }
InputLine → DefinitionLine
 | SubcircuitBody
 | EquationLine
 | ActionLine
 | ['#' (*Entire line is comment*) / **Eol**

A netlist is read in a line-based fashion. There are several types of lines.

Definition

DefinitionLine → **Identifier** ':' **Identifier** { **Identifier** } PairList **Eol**

Components of the circuit are defined by providing its nodes and parameters for a property.

Subcircuits

SubcircuitBody → DefBegin { DefBodyLine } DefEnd
DefBegin → '.' **Def** ':' **Identifier** { **Identifier** } **Eol**
DefBodyLine → DefinitionLine
 | SubcircuitBody
 | **Eol**
DefEnd → '.' **Def** ':' **End** **Eol**

Subcircuits are recursively defined by blocks of component definitions.

Action

ActionLine → '.' **Identifier** ':' **Identifier** PairList **Eol**

Defines what to simulate with the circuit.

Equation

EquationLine → '**Eqn**' ':' **Identifier** Equation { Equation } **Eol**
Equation → **Identifier** '=' '“' Expression '”'

Named equation definitions consist of a list of assignments with expressions on their right hand side.

Declarations

PairList	→ { Identifier '=' Value }
Value	→ PropertyValue "" PropertyValue ""
PropertyValue	→ Identifier PropertyReal '[' PropertyReal [{ ';' PropertyReal }] '['
PropertyReal	→ Real [Scale [Unit]]

The above constructs are used to define properties (parameters) of components and actions.

Mathematical Expressions

Expression	→ Constant Reference Application '(' Expression ')'
Constant	→ Real Imag Character String Range
Range	→ [Real] ':' [Real]
Reference	→ Identifier
Application	→ Identifier '(' Expression [{ ';' Expression }] ')' Reference '[' Expression [{ ';' Expression }] '[' ('+' '-') Expression Expression ('+' '-' '*' '/' '%' '^') Expression

Operator precedence works as expected in common mathematical expressions.

Lexical structure

Identifier	→ Alpha { AlphaNum } { '.' Alpha { AlphaNum } }
Alpha	→ 'a' ... 'z' 'A' ... 'Z' '_'
AlphaNum	→ 'a' ... 'z' 'A' ... 'Z' '0' ... '9' '_'
Real	→ ['+' '-'] [Num] '.' Num [('e' 'E') ['+' '-'] Num]
Num	→ Digit { Digit }
Digit	→ '0' ... '9'
Imag	→ ['+' '-'] ('i' 'j') [Num] '.' Num [('e' 'E') ['+' '-'] Num]
Character	→ "" (Any character but newline and ") ""
String	→ "" { (Any character but newline and ") } ""
Scale	→ ('dBm' 'dB' 'T' 'G' 'M' 'k' 'm' 'u' 'n' 'p' 'f' 'a')
Unit	→ ('Ohm' 'S' 's' 'K' 'H' 'F' 'Hz' 'V' 'A' 'W' 'm')
Eol	→ ['\r'] '\n'

This defines the composed terminal symbols.

A.2 Qucs dataset grammar

Syntactic structure

Input → VersionLine { Variable }
VersionLine → '<Qucs Dataset ' Num '.' Num '.' Num '>' Eol

The file consists of a header line indicating the software version and a number of variables.

Data

Variable → '<dep' Identifier { Identifier } '>' { Float } '</dep>'
| '<indep' Identifier Integer '>' { Float } '</indep>'
| '#' (*Entire line is comment*) Eol

The **Float** values itself may be scattered over several lines. An independent variable denotes a list of real or complex values. The dependent variables denote lists of real or complex values depending on independent variables, i.e. a function of some other variable, a $f(x, \dots)$ in mathematical terms.

Lexical structure

Float → Real
| Imag
| Complex
Real → ['+' | '-'] [Num] '.' Num [('e' | 'E') ['+' | '-'] Num]
Imag → ['+' | '-'] ('i' | 'j') [Num] '.' Num [('e' | 'E') ['+' | '-'] Num]
Complex → Real Imag
Num → Digit { Digit }
Digit → '0' ... '9'
Integer → ['+' | '-'] Num
Eol → ['\r'] '\n'

Appendix B

Bibliography