

Skript zur Vorlesung EMV

Elektromagnetische Verträglichkeit an der HTWG Konstanz, Prof. Heinz Rebholz.

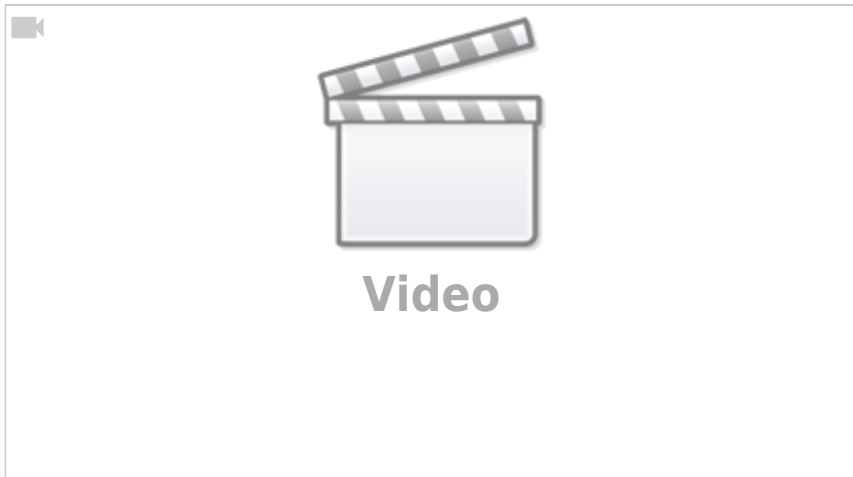
1.0 Einleitung

Vorwort

Die Elektromagnetische Verträglichkeit, kurz EMV, ist für viele junge Entwicklungsabteilungen zunächst ein Schreckgespenst. Grundsätzlich ist allen Beteiligten klar, dass Produkte eigensicher sein müssen und ihr Umfeld nicht unzulässig beeinflussen dürfen – allerdings oft nicht gerade jetzt, kurz vor der ersten Abgabe oder in einer Phase, in der die eigentliche Funktionsentwicklung deutlich mehr Aufmerksamkeit und Motivation bindet. Doch der Markt und insbesondere der Kunde fordern zu Recht ein konformes Produkt mit CE-Kennzeichnung. Damit bleibt letztlich keine Alternative, als sich strukturiert mit den Anforderungen der EMV auseinanderzusetzen und auf bewährte Literatur oder dieses Skript zurückzugreifen.

Dabei zeigt die Praxis schnell, dass EMV kein nachgelagerter Prüfpunkt ist, sondern ein integraler Bestandteil der gesamten Produktentwicklung. Frühzeitig berücksichtigt, lassen sich viele Probleme mit vergleichsweise geringem Aufwand vermeiden. Wird die EMV hingegen erst spät betrachtet, steigen sowohl Entwicklungsaufwand als auch Kosten erheblich, da grundlegende Designentscheidungen häufig erneut hinterfragt und angepasst werden müssen.

Das nachfolgende Skript versucht, einen kompakten Überblick über wesentliche Aspekte der EMV-Entwicklung zu geben, wie so oft ohne den Anspruch auf Vollständigkeit. Im Fokus stehen dabei praxisrelevante Zusammenhänge, typische Problemstellungen sowie bewährte Lösungsansätze, die den Einstieg in die Thematik erleichtern und eine strukturierte Herangehensweise unterstützen.



Inhalt

Ich vermute, Sie haben dieses Skript nicht in der Hand, weil Sie Spaß an komplizierter Elektrotechnik haben. Vielmehr geht es darum, möglichst reibungslos durch die EMV-Klausur zu kommen oder durch die Zertifizierung Ihres Produkts. Beides sind ehrenwerte Ziele, die ich im Folgenden zu unterstützen versuche. Das Skript enthält neben allgemeinen Hinweisen zur EMV viele praktische Tipps zur Hardwaregestaltung und zu EMV-Messverfahren sowie deren Einsatz. An manchen Stellen kann ich allerdings nicht widerstehen und werde auf einige Formeln zurückgreifen. Für Studierende ist es wichtig, diese Abschnitte nicht zu überspringen. Entwickler in Eile, die ihr Produkt in Serie bringen müssen, können in Ruhe dazu zurückkehren, sobald die groben Arbeiten erledigt sind oder, wie man auch sagt, die Kuh vom Eis ist. Eventuell sind diese Inhalte für Sie erst beim zweiten oder dritten Projekt von Interesse.

Spätestens nach der ersten Serienentwicklung ist man eigentlich schon ein EMV-Profi, denn man hat die wichtigste Regel der EMV im Tal der Tränen aufgesammelt:

Die EMV-Arbeit beginnt mit dem ersten Tag der Hardwareentwicklung. Besser noch, mit dem Entwurf des Blockschaltbilds.

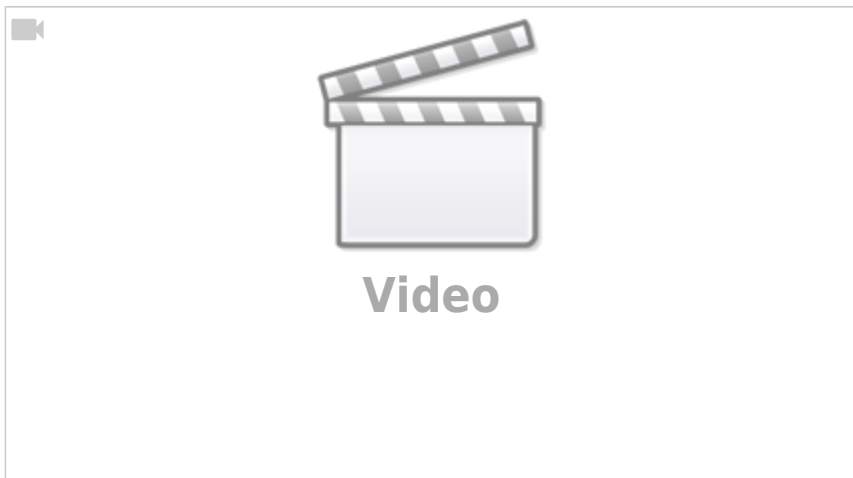
Je später mit Überlegungen zur EMV begonnen wird, desto schwieriger wird es, EMV-Maßnahmen umzusetzen. Denn egal, welche EMV-Maßnahme Sie einbringen,

sie muss getestet werden. Nicht nur hinsichtlich der EMV-Grenzwerte, sondern noch schlimmer hinsichtlich:

- Lebensdauer
- Temperaturabhängigkeit
- Gesamtfunktionalität

Je näher die angedachte Serienproduktion rückt, desto mehr sinkt somit das Verständnis für notwendige EMV-Maßnahmen. Die gängigste Frage: „Kann man hier nicht auch etwas mit Software machen?“ muss dann leider in den allermeisten Fällen verneint werden. In späteren Kapiteln werden wir sehen, dass Software durchaus einen Einfluss auf die EMV-Performance Ihres Produktes haben kann, allerdings üblicherweise in überschaubaren Größenordnungen.

1.1 Was versteht man unter EMV



Elektromagnetische Verträglichkeit (EMV) ist ein Sammelbegriff, dem ganz unterschiedliche Themenfelder aus der Elektrotechnik zugeordnet werden. Vereinfacht könnte man sagen, die EMV beschreibt einerseits alle elektromagnetischen Emissionen, welche von meiner Schaltung ausgehen, und bewertet andererseits die Festigkeit gegenüber von außen auf mein Produkt einwirkenden elektrischen Störgrößen.

Manchmal werden zusätzlich alle unerklärlichen Eigenschaften einer elektronischen Schaltung der EMV zugeschrieben, auch wenn es sich eventuell

um funktionale Probleme handelt (in diesem Fall sprechen wir von der Eigenstörfestigkeit). Der Gesetzgeber reguliert über die EMV-Richtlinie die maximal zulässigen Emissionen eines Produktes und schreibt eine definierte Festigkeit gegenüber Störgrößen vor. \ Die Einhaltung der EMV-Richtlinie ist meist eine von mehreren anzuwendenden Richtlinien zur CE-Zertifizierung.\

EMV:

- Bewertung der von meinem Produkt ausgehenden Strörgrößen (Emissionen)
- Festigkeitsprüfung gegenüber einwirkender elektrischer Störgrößen (Immission)
- Sicherstellen, dass sich mein Gerät nicht selber stört (Eigenstörfestigkeit)

Das hat nichts mit EMV zu tun

Aus Gesprächen mit jungen Gründern sind mir einige Missverständnisse aufgefallen, die oft im Zusammenhang mit der EMV oder der CE-Kennzeichnung auftreten.

Vielleicht erkennen Sie das eine oder andere Vorurteil aus Ihrem Umfeld:

- Wir verwenden nur bleifreie und zertifizierte Bauteile. Damit ist unser Gesamtprodukt automatisch konform mit den EU-Richtlinien
- Wir kaufen unsere Baugruppen nur bei nach ISO9001 zertifizierten Betrieben und lassen danach fertigen. Damit ist auch unser Produkt konform mit der Richtlinie.
- CE-Kennzeichnung gilt nur für Geräte mit einphasigem Netzanschluss. 12V Anwendungen sind davon nicht betroffen.
- Für geringe Stückzahlen ist keine CE-Kennzeichnung notwendig.
- Bei der der Kennzeichnung handelt es sich um eine Selbstzertifizierung und zeigt nur an, dass das Produkt in der EU hergestellt wurde.

Alle diese Behauptungen zur Zertifizierung sind falsch.

Damit Produkte in den Handel gebracht werden dürfen, sind eine Reihe an Vorschriften und Regulierungen zu beachten.

Welche Normen für Ihr Produkt relevant sind, wird in den nachfolgenden Kapiteln

besprochen. Um es vorwegzunehmen, eine Normenrecherche ist nichts für schwache Nerven und bedarf einiges an Geduld und Durchhaltevermögen.

1.2 Mess- und Bewertungsverfahren

Bevor wir weitergehen, sollten wir die Begriffe **Emission** und **elektrische Störgrößen** (**Immission**) genauer betrachten. Elektromagnetische Emission (lat. emittire „aussenden“) bedeutet, dass es möglich ist, außerhalb unserer Komponente physikalisch messbare Größen zu erfassen, auch wenn wir das funktional gar nicht vorgesehen haben. Hier wird deutlich, dass es sich bei dem Begriff „elektromagnetische Emission“ um einen Oberbegriff handelt. Die Emissionen einer mit elektrischer Energie betriebenen Komponente können weit vielfältiger sein als die beschriebenen elektrischen und magnetischen Felder. Folgende physikalischen Größen werden bei Emissionsmessungen bewertet:

Beschreibung	Einheit
Konstante und niederfrequente magnetische Felder	A/m
Konstante und niederfrequente elektrische Felder	V/m
Hochfrequente elektromagnetisch Felder	V/m
Hochfrequente Störströme bzw. Störspannungen	A bzw. V
Oberschwingungen / Netzurückwirkung	A oder V
Transiente Störgrößen z.B. Überspannungen aufgrund Blitzeinwirkung	V

Elektrische Störgrößen können in zwei unterschiedliche Bereiche gegliedert werden. Es ist davon auszugehen, dass sich unser Produkt später in einer Welt umringt von zahlreichen weiteren elektronischen Geräten aufhalten wird. Daher ist es logisch, dass uns die Emissionen der umliegenden Geräte nichts anhaben dürfen. Als ein Entwickler, der besonders gute Produkte herstellen möchte, würden Sie sicher zustimmen, wenn ich sage, unser Produkt sollte sogar noch etwas robuster sein als die Emissionen, welche die umliegenden Geräte emittieren dürfen. Man könnte auch sagen, wir müssen einen fiktiven Schutzschirm um unser Gerät aufspannen, welcher die externen Störgrößen blockiert.

Dank der einheitlichen CE-Kennzeichnung sind die erlaubten Emissionen in einer spezifizierten Umgebung für alle Geräte identisch oder sehr ähnlich, womit wir sehr genau sagen können, welche Störgrößen zu erwarten sind. Allerdings dürfen wir uns auf den spezifizierten Grenzwerten nicht ausruhen. Stellen Sie sich vor, neben Ihrem Gerät steht eine Anlage, welche aufgrund eines Defekts plötzlich mehr emittiert. Sollte Ihr Gerät dadurch ebenfalls Schaden nehmen, wäre das für den Ruf Ihrer Produkte sicher nicht hilfreich.

Neben den Störgrößen, welche durch benachbarte Geräte generiert werden, gibt es eine Reihe von natürlichen elektrischen Phänomenen, gegenüber denen man sich schützen muss. Die bekannteste natürliche Störgröße ist eine Blitzentladung. Sie wirkt sich bei einer lokalen Entladung durch eine massive Spannungsüberhöhung im Versorgungsnetz aus. Weit entfernte Entladungen machen sich in Form elektromagnetischer Wellen bemerkbar, welche in die Geräte einkoppeln. Sicherlich sind Ihnen bereits beide Phänomene bekannt, sei es durch zerstörte Elektrogeräte oder durch wahrnehmbare Knackgeräusche im Radio während entfernter Gewitter. Gegen einen, wenn auch sehr selten auftretenden, direkten Blitzeinschlag hilft meist auch keine EMV-Maßnahme mehr. Allerdings ist es möglich, sich gegenüber den auftretenden Überspannungen und dem Einkoppeln elektromagnetischer Wellen zu schützen.

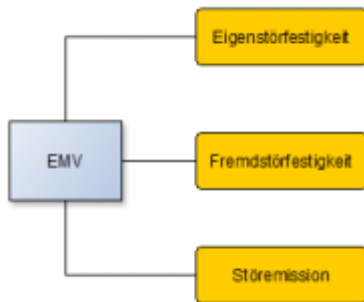
Im Gegensatz zu einem direkten Blitzeinschlag ist es wichtig, sich gegen direkte ESD (Electrostatic Discharge) Entladungen zu schützen. Jeder von uns hat diese kleinen, lästigen elektrischen Entladungen bereits am eigenen Körper gespürt. Allerdings liegt die Wahrnehmungsgrenze bei ca. 2–3 kV, womit die Dunkelziffer der vorhandenen Entladungen sehr hoch sein dürfte. Das bedeutet, elektronische Schaltungen sind ständig Überspannungsimpulsen ausgesetzt, gegen die es sich zu schützen gilt. Die maximal messbaren Entladungen liegen bei sehr ungünstigen Verhältnissen (materialabhängig, Luftfeuchtigkeit, ...) bei bis zu 30 kV. Mit der ESD- und Blitzentladung haben wir die zwei wichtigsten natürlichen Störgrößen betrachtet, vor denen es sich zu schützen gilt. Eine weitere natürliche Störquelle ist die vorhandene Höhenstrahlung, welche allerdings in der Gesetzgebung keine Anwendung findet und somit nur für Spezialanwendungen unter dem Thema Eigenstörfestigkeit (siehe nächster Abschnitt) von Bedeutung sein sollte (Beachten Sie diesen Punkt, falls Sie Ihr Produkt in Höhen über 2000 m NN einsetzen möchten).

1.3 Gruppierung in Arbeitsfelder

Die EMV kann in drei Arbeitsgebiete unterteilt werden, der

- Eigenstörfestigkeit
- Fremdstörfestigkeit
- Störaussendung

Gesetzliche Anforderungen existieren nur für die beiden letzten Punkte.



Die **Eigenstörfestigkeit** bewertet die funktionale Seite der Produkte. Werden interne Komponenten durch benachbarte Schaltungsteile gestört so spricht man von einer internen Beeinflussung oder eben einer fehlenden Eigenstörfestigkeit. Bemerkbar macht sich eine fehlende Eigenstörfestigkeit falls in einem bestimmten Betriebspunkt einzelne Funktionen nicht mehr gegeben sind.

Eventuell ist Ihnen so ein Verhalten bei der Inbetriebnahme Ihrer Geräte schon aufgefallen. Oft treten solche Phänomene beim Umschalten zwischen verschiedener Betriebspunkte (z.B. Motor wird gestartet / stark abgebremst) oder beim Zuschalten von Leistungsendstufen und Hochstromverbrauchern auf. Die häufigste Ursache liegt in einer fehlenden Signalintegrität bzw. einem fehlenden Signal-zu-Rauschabstand einzelner Signale Ihrer Schaltung. Die Eigenstörfestigkeit ist natürlich nicht reguliert und obliegt keinen gesetzlichen Rahmenbedingungen. Sie selbst haben daran Interesse, zuverlässige und sichere Produkte zu entwickeln und herzustellen.

Die **Fremdstörfestigkeit** oder elektromagnetische Immission hingegen ist sehr wohl reglementiert und schreibt vor, dass die Produkte gegenüber von außen auftretenden Störgrößen resistent sein müssen. Oder um im Bild zu bleiben, wie resistent muss der von uns aufgespannte Schutzschirm sein. Die Fremdstörfestigkeit unterscheidet dabei eine Vielzahl an physikalischen Phänomenen welche zur Bewertung herangezogen werden. Bei der bekanntesten Prüfung wird untersucht, ob sich unser Gerät durch elektromagnetische Felder beeinflusst wird. Dabei simuliert eine Antenne die Anwesenheit weitere Geräte. Weniger bekannt ist die Prüfung gegenüber Burst- und Surge-Impulsen, welche das Gerät fit machen gegenüber den am Netzanschluss auftretenden Störgrößen.

Zuletzt die **Störemission**. Der Begriff ist identisch mit der elektromagnetischen Emission und beschreibt die Summe aller außerhalb meiner Komponente messbarer physikalischen Größen. Die bekannteste Messung ist die Emissionsmessung bei der mit einer Antenne die emittierte Feldstärke gemessen wird. Weniger bekannt ist, dass auch die Emissionen auf den Versorgungsleitungen reglementiert sind. Werden während des Produktdesigns alle Themenfelder berücksichtigt, steht der Einhaltung der EMV-Richtlinie nichts im Weg. Vermutlich haben Sie als Hardwareentwickler bereits die eine oder andere EMV-Maßnahme umgesetzt, ohne es zu wissen.

1.4 Landläufige Definition der EMV

Leider hat es sich im Sprachgebrauch eingebürgert, die EMV auf die Einhaltung der gesetzlichen Vorschriften bezüglich der Emission und Immissionen zu begrenzen. Die moderne EMV-Arbeit vermag jedoch viel mehr. Als ambitionierter Entwickler liegt es uns am Herzen, zuverlässige und robuste Geräte zu entwickeln. EMV-Maßnahmen sind immer zuerst auch qualitätssichernde Maßnahmen, welche Ihr Produkt aufwerten. Mit dem Schreckgespenst der EMV ist es in etwa wie mit einer Klausur im Studium. Wird der Vorlesungsstoff missachtet und mit zu viel Mut zur Lücke gelernt, kann das böse Erwachen kommen. Je nach Dozent kommt man mit einem blauen Auge davon. Falls man sich entscheidet, konsequent dem Vorlesungsinhalt zu folgen und alle Punkte zu bearbeiten, ist es kein Problem, die Klausur respektive die EMV-Prüfung zu bestehen.

1.5 Ziele der EMV-Arbeit

Zusammengefasst ist es Ihr Ziel als EMV-Ingenieur oder Hardwareentwickler, zum einen sichere und zuverlässige Produkte herzustellen, welche sich nicht selber stören, und auf der anderen Seite alle gesetzlichen Rahmenbedingungen einzuhalten. Der Gesetzgeber regelt mit der EMV-Richtlinie sehr detailliert, welche Schritte und Prüfverfahren dazu notwendig sind. Falls Sie als Zulieferer tätig sind, der eine Unterbaugruppe herstellt, erhalten Sie die notwendigen Anforderungen an die Elektronik im Lastenheft von Ihrem Auftraggeber. Sie vertreiben dann in der Regel das Produkt nicht selber, sondern sind Teil eines Gesamtaufbaus. Der Inverkehrbringer des Gesamtprodukts ist in diesem Fall für die Einhaltung der nationalen Gesetze verantwortlich. Problematisch ist hier, die Schnittstellen und Prüfaufbauten sauber zu definieren, um Zuständigkeiten und Entwicklungsverantwortungen auch firmenübergreifend sauber zu trennen. Ein Beispiel für die Aufteilung der Aufgaben und die Schnittstellendefinition ist im Kapitel Blockschaltbild zu finden.

1.6 Woher kommen die Grenzwerte?

Für die Eigenstörfestigkeit gibt es natürlich keine gesetzlichen Grenzwerte. Man könnte auch sagen, sind die Geräte funktional in Ordnung, braucht man sich um diesen Punkt keine Sorgen zu machen. Ganz so einfach ist es allerdings nicht. Es wird sich im weiteren Verlauf zeigen, dass je robuster die Geräte gegen interne Störungen sind, das Einhalten der Grenzwerte sowohl für Emission und Immission deutlich einfacher ist. Man könnte auch sagen: „Eine durchdachte (natürlich ist

EMV-durchdacht gemeint) Hardware wird in allen drei Themenfeldern der EMV gleich gut abschneiden!“ Aber wie kommt man nun an die gültigen Grenzwerte? Um diese Frage richtig zu beantworten, muss ich leider etwas weiter ausholen. Aus dem täglichen Leben ist Ihnen vielleicht bekannt, dass weltweit nicht alle Länder harmonisierte bzw. abgestimmte technische Standards haben. Jedes Land oder besser Wirtschaftsraum unterwirft Maschinen und elektrische Geräte einer Vielzahl an Vorschriften und Reglementierungen. Vielleicht haben Sie sich schon einmal gewundert, dass bei US-amerikanischen Fahrzeugen die Blinker in Rot ausgeführt sind oder sogar die Bremsleuchte dazu verwendet wird. Oder noch banaler: Es muss ein Gesetz geben, das Warnhinweise und die Betriebsanleitung in der jeweiligen Landessprache fordert!

Für den europäischen Wirtschaftsraum gibt es einheitliche Standards in Form von EU-Richtlinien. Für häufig anzuwendende Richtlinien, wie die Maschinen- oder EMV-Richtlinien, gibt es jeweils Leitfäden, welche bei der Interpretation behilflich sind. Sämtliche Richtlinien verfolgen im Wesentlichen zwei Ziele. Zum einen soll ein uneingeschränkter Warenaustausch innerhalb der Wirtschaftsunion erfolgen können, zum anderen sollen die Anwender maximal vor möglichen Gefahren geschützt werden. Aktuell laufen verschiedene Verhandlungen, den europäischen Wirtschaftsraum über Handelsabkommen zu erweitern, siehe TTIP oder CETA. Ein viel diskutierter Punkt ist genau der erwähnte Verbraucherschutz, da die Handelsabkommen gegenseitig die lokalen Richtlinien anerkennen (Stichwort Chlorhühnchen). Europäische Richtlinien müssen von den jeweiligen Mitgliedsstaaten in nationales Recht übersetzt werden. Das bedeutet, dass es zu jeder Richtlinie ein deutsches Pendant geben muss, welches in den meisten Fällen die EU-Richtlinie wiedergibt. Erst durch die nationale Gesetzgebung wird aus der Richtlinie für Sie ein anzuwendendes Gesetz.

- 2001/95/EG allgemeine Produktsicherheit
→ ProdSG Produktsicherheitsgesetz
- 2014/30/EU EMV → Gesetz über die elektromagnetische Verträglichkeit von Betriebsmitteln EMVG
- 2014/35/EU Niederspannungsrichtlinie → ProdSV Produktsicherheit elektrischer Betriebsmittel
- und viele Weitere (Richtlinie über Aufzüge, Bauprodukte, Druckbehälter, ...)
- ...

Spätestens wenn man die Liste zum ersten Mal durcharbeitet, wünscht man sich,

man hätte nie damit angefangen, da die Gesetze teilweise aufeinander verweisen oder geschachtelt sind. Auf jeden Fall eine schweißtreibende Arbeit Bisher sind wir unseren Grenzwerten noch kein Stück nähergekommen, und es wird noch schlimmer. Schauen wir uns dazu das für uns wichtige EMV-Gesetz (EMVG) genauer an. Der Gesetzestext ist einzusehen unter

http://www.gesetze-im-internet.de/emvg_2016/ Die für uns Anwender/Entwickler wichtigsten Aussagen des EMVG besagen:\ §1 (1)

Dieses Gesetz gilt für alle Betriebsmittel, die elektromagnetische Störungen verursachen können oder deren Betrieb durch elektromagnetische Störungen beeinträchtigt werden kann.\ §4

Betriebsmittel müssen nach dem Stand der Technik so entworfen und hergestellt sein, dass

1. die von ihnen verursachten elektromagnetischen Störungen keinen Pegel erreichen, bei dem ein bestimmungsgemäßer Betrieb von Funk- und Telekommunikationsgeräten oder anderen Betriebsmitteln nicht möglich ist;
2. sie gegen die bei bestimmungsgemäßem Betrieb zu erwartenden elektromagnetischen Störungen hinreichend unempfindlich sind, um ohne unzumutbare Beeinträchtigung bestimmungsgemäß arbeiten zu können.

Die Forderungen in §4 entsprechen exakt unserer bisherigen Definition der EMV für die elektromagnetische Emission und Immission. Ein Verweis auf einzuhaltende physikalische Größen ist allerdings nur versteckt zu finden in der Erklärung zum Stand der Technik. Dort heißt es: „.....„Stand der Technik“ der allgemein anerkannte Stand der Technik in Bezug auf die elektromagnetische Verträglichkeit entsprechend den harmonisierten Normen;“. Genau an dieser Stelle werden wir verwiesen auf harmonisierte, gleichbedeutend mit international abgestimmte, Normen.



Neuentwicklungen müssen sich stets am Stand der Technik orientieren und somit die relevanten Normen berücksichtigen

Normen entstehen in einem öffentlichen Normungsprozess unter Mitwirkung von Vertretern aus der Wissenschaft und der Industrie. Da sich der „Stand der Technik“ kontinuierlich weiterentwickelt, sind Normen selten statische Dokumente, sondern ebenfalls laufend dem Wandel unterworfen in Form von Revisionen oder neuen Ausgaben. Organisiert wird der Normungsprozess durch

eine der drei europäischen Komitees

- CEN Europäisches Komitee für Normung
- CENELEC Europäisches Institut für elektrische Normung
- ETSI Europäisches Institut für Telekommunikationsnormen

Zu erkennen sind die europäischen Normen am Zusatz EN XX, z.B. EN 61000-6-1. Ein rein formaler Akt ist die Ratifizierung (Übernahme) der europäischen Norm in das nationale Normenwerk. Die Übernahme erfolgt mit dem Zusatz DIN-EN. Weitere Infos zu der Arbeit von Normenausschüssen sind unter www.din.de zu finden. Sie werden es bereits errahnen, dass es sich bei den harmonisierten Normen zur EMV nicht nur um ein Dokument handelt, sondern eine ganze Reihe an unterschiedlichen Schriften. Um eine Systematik in die Normenwelt zu bekommen, ist es möglich, drei verschiedene Arten von Normen zu unterscheiden:

Eine **Fachgrundnorm** enthält die übergeordneten Anforderungen an unsere Geräte hinsichtlich Störfestigkeit und Aussendung. **Grundnormen** hingegen erläutern die Anforderungen an physikalische Messverfahren. Sie sind die Basis dafür, dass einheitlich und reproduzierbar gemessen und geprüft werden kann. An letzter Stelle stehen die **Produktnormen**. Hierin wird gezielt eine Unterscheidung in der Anwendung getroffen. So existieren Produktnormen für Hausgeräte, Radio- und TV-Geräte, Leuchten und weitere.

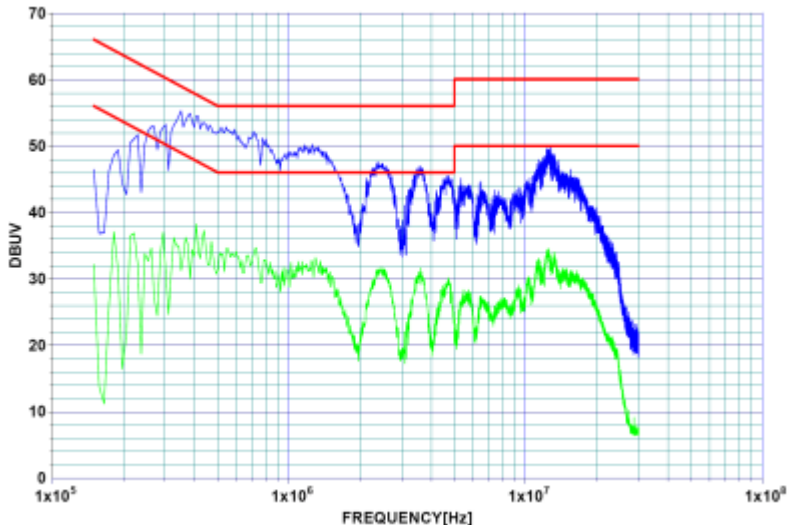
In der Konformitätserklärung zu Ihrem Produkt reicht prinzipiell für die EMV ein Verweis auf die Fachgrundnorm, da diese wiederum auf die entsprechende Produktnorm verweist. Die Fachgrundnorm ist somit für uns die erste Anlaufstelle, um herauszufinden, welche EMV-Anforderungen nach dem Stand der Technik aktuell gültig sind.

Fachgrundnorm	Bezeichnung
Störfestigkeit für Wohnbereich, Geschäfts- und Gewerbebereiche sowie Kleinbetriebe	EN 61000-6-1
Störfestigkeit für Industriebereiche	EN 61000-6-2
Störaussendung für Wohnbereich, Geschäfts- und Gewerbebereiche sowie Kleinbetriebe	EN 61000-6-3
Störaussendung für Industriebereiche	EN 61000-6-4

Leider sind die Normen der DIN-EN XX Reihe nicht frei verfügbar, sondern müssen über den VDE bzw. Beuth-Verlag erworben werden www.beuth.de. Bibliotheken, typischerweise von Hochschulen und Universitäten, sind meistens sogenannte Normen-Infopoints. Alle Normen können hier kostenlos eingesehen werden.

Allerdings besteht meistens keine Möglichkeit zum Druck oder Speichern der Dokumente, womit man sich nur einen groben Überblick verschaffen kann.

Und ja, endlich, nach langen Umwegen sind wir am Ziel angekommen. Wir haben jetzt endlich ein erstes Dokument gefunden, das uns verbindliche Ziele vorschreibt nach denen wir entwickeln können:



Das Bild zeigt eine Messung der leitungsgebundenen Emission auf den Versorgungsleitungen mit dazugehörigen Grenzwerten.

Die Fachgrundnormen, oder falls für unser Produkt verfügbar die Produktnorm, verweisen auf ca. 5- 10 Grundnormen, in welchen die Messverfahren und Versuchsaufbauten der Einzelmessungen definiert sind. Alleine der Gedanke daran, wie viele Seiten Normentext auf uns zum Lesen zukommt, kann einem den Angstschweiß in den Nacken treiben. Aber so schlimm kommt es nicht. Sie werden feststellen, dass sich vieles wiederholt und schon nach ein wenig Übung ein Gefühl dafür aufkommt, welche Passagen wirklich wichtig sind. Ich hoffe schwer, dass Sie als Techniker / Ingenieur stets darum bemüht sind, sichere und funktionsfähige Produkte zu entwickeln. Sehen Sie Normen daher als Hilfsmittel, gute Produkte herzustellen, und nicht als Last . Die angegebenen Grenzwerte sind jeweils das Mindestmaß dessen, was nach dem Stand der Technik ein Gerät erfüllen muss. Sie können jederzeit natürlich noch bessere Produkte herstellen. In der Automobilindustrie werden typischerweise um den Faktor drei strengere Anforderungen an die Fahrzeuge und elektrischen Komponenten angelegt. Das hat banale Gründe: Denken Sie an den Imageschaden einer Marke, falls die

Fahrzeuge reihenweise bei Gewitter oder beim Überqueren von Bahnübergängen stehen blieben.

Normen können allerdings nicht alle technischen Details erfassen. Dazu sind unsere Systeme viel zu komplex. Das bedeutet, Normen dürfen auch nicht blind angewendet werden, sondern jeweils für die entsprechende Situation bewertet und ggf. in modifizierter Form angewendet werden. Falls Sie nicht sicher sind, wenden Sie sich am einfachsten an eine benannte Stelle oder suchen sich Rat bei Dienstleistern, die sich auf die Normenauslegung und rechtliche Absicherung spezialisiert haben.

1.7 Wie werden Grenzwerte festgelegt?

Zu Beginn habe ich die Frage gestellt: „Woher kommen die Grenzwerte?“. Klar, wir wissen jetzt, in welchen Dokumenten wir suchen müssen. Wir wissen sogar, wer sie geschrieben hat, allerdings noch nicht, warum sie so sind, wie sie sind! Wichtig ist es, sich an dieser Stelle klarzumachen, dass Grenzwerte immer Kompromisse sind, welche von unterschiedlichen Interessensvertretern ausgehandelt wurden. Einerseits darf die EMV-Regulierung nicht zu strenge Anforderungen stellen, da ansonsten die Kosten massiv steigen oder sogar zusätzlicher Bauraum für zum Beispiel Filterschaltungen notwendig würde. Auf der anderen Seite dürfen die Anforderungen auch nicht zu schwach formuliert sein, da die Geräte je nach Einsatzort im ihrem Funktionsumfang eingeschränkt wären. Die gesetzten Grenzwerte für Emission und Immissionsprüfungen bewegen sich somit zwischen diesen beiden Anforderungen.



Ein CE-Kennzeichen bietet keinen 100% Schutz gegenüber elektrischen Störgrößen!

Das bedeutet aber auch, dass die in der Realität auftretenden elektrischen Störgrößen deutlich größer sein können als die geprüften Grenzwerte. Die in der Grundnorm definierten Prüfaufbauten gewährleisten zwar ein einheitliches Prüfverfahren, können allerdings natürlich nicht für jedes Gerät das Kundenverhalten mitberücksichtigen. Wollen Sie das Ausfallrisiko aufgrund externer Störgrößen weiter minimieren, bleibt Ihnen keine andere Möglichkeit, als das Kundenverhalten bzw. deren Umgebung näher zu untersuchen und entsprechend zu handeln. Fahrzeughersteller rüsten zum Beispiel Fahrzeuge mit Empfangsantennen aus und messen die im Fahrbetrieb höchste von außen kommende Feldstärke.

1.8 Fazit - Was nun? Wie geht es weiter?

Mit dem Wissen, welche Normen für uns von Bedeutung sind, haben wir bereits eine große Hürde in Richtung CE-Kennzeichnung bzw. Einhalten der EMV-Richtlinie genommen. Spannender wird jetzt die Frage, wie geht es weiter? Leider können Sie mit den vorgegebenen Grenzwerten und Immunitätsprüfungen (noch) nichts anfangen. Falls Sie ein neues Produkt entwickeln oder ein junges Start-up-Unternehmen sind, fehlt Ihnen vermutlich jegliche Referenz für Ihr Produkt. Da Sie zum jetzigen Zeitpunkt noch nicht wissen, ob Ihr Produkt die Anforderungen locker einhält oder vielleicht sehr weit davon entfernt ist, können wir auf die absoluten Werte noch verzichten. Viel wichtiger ist es, während der Produktentwicklung die EMV-Anforderungen mit zu berücksichtigen, kritische Pfade zu erkennen und mit Maßnahmen zu hinterlegen. EMV-Maßnahmen in einem frühen Entwicklungsstadium sind meistens kostenneutral und benötigen nur wenig Volumen. Vielleicht haben Sie Glück und auf die eine oder andere Maßnahme kann wieder verzichtet werden.

Im nächsten Schritt schauen wir uns die Grundlagen der EMV-Arbeit für Ihr Produkt an und überlegen, an welcher Stelle EMV-Maßnahmen eingesetzt werden müssen.

1.10 EMV Entwicklung

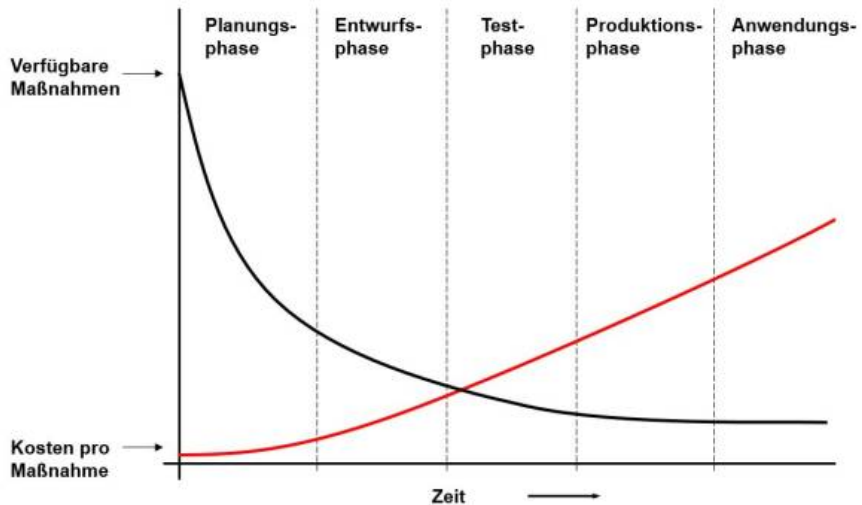
Die Berücksichtigung der EMV-Anforderungen an das Produkt sollte bereits bei der Systementwicklung beginnen. Das bedeutet, bereits im Blockschaltbild ist es notwendig, erste EMV-Ansätze zu integrieren. Mit den Informationen aus den vorangehenden Kapiteln ist es zum Beispiel naheliegend, keine langen Leitungen mit hochfrequenten Signalanteilen über lange Leitungsverbindungen durch das Gerät zu ziehen. Zumindest nicht ohne ausreichende Maßnahmen. Solche einfachen Kernaussagen lassen sich ohne Mühe in den Produktentstehungsprozess integrieren. Die Integration zusätzlicher Schaltungsteile, zur Erhöhung der Störfestigkeit oder für eine reduzierte Emission, ist in einer frühen Projektphase mit sehr geringen Kosten verbunden. Vermutlich in Summe nur wenige Euro. Je später die Maßnahmen integriert werden, desto teurer und schwieriger ist eine Integration.



Ladegerät für die Hochvoltbatterie eines Elektro- / Hybridfahrzeugs mit zahlreichen Ferritkernen und EMV-Maßnahmen

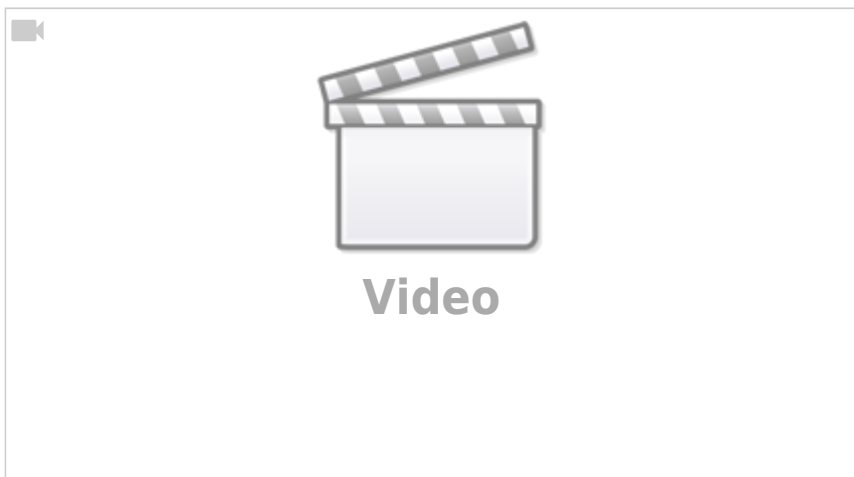
Ein Beispiel, das jeder kennt, sind Geräte mit Ferritkernen über der Leitung. Ferrite wirken wie Induktivitäten und hemmen hochfrequente Ströme in ihrer Ausbreitung. Sie sind generell ein beliebtes Element, da sie während der EMV-Tests einfach angebracht werden können, und mit etwas Glück bleibt man damit unter den Grenzwerten. Allerdings weisen die Ferrite meist auf einen harten Kampf zur Einhaltung der Grenzwerte hin, bei dem zu allen Hilfsmitteln gegriffen werden musste. Ferrite kosten ein Vielfaches mehr als auf der Leiterplatte integrierte Lösungen. Zusätzlich müssen sie in der Produktion händisch angebracht werden, wohingegen zusätzliche Elemente auf der Leiterplatte nahezu kostenneutral im Automaten bestückt werden. Genau um dies zu verhindern, versuchen wir im Folgenden eine systematische EMV-Entwicklungsarbeit aufzubauen.

Wie jedes gute EMV-Skript komme auch ich nicht um die Grafik von Henry W. Ott herum, der bereits 1988 in seinem Buch „Noise Reduction Techniques in Electronic Systems“ den bekannten Zusammenhang zwischen der Entwicklungsdauer und den EMV-Kosten bzw. Lösungsoptionen hergestellt hat.



Die Grafik spricht für sich und beschreibt das Dilemma, welches einem widerfährt, ohne rechtzeitige Berücksichtigung der EMV ziemlich genau. Bevor wir uns näher mit der Entwicklung beschäftigen, schauen wir uns allerdings die wichtigsten EMV-Messverfahren genauer an. Nur dann bekommen wir ein Gefühl dafür, vor was wir unsere Schaltungen schützen möchten bzw. wie die Emissionen reduziert werden können.

2.0 EMV Mess- und Bewertungsverfahren



Bereits im ersten Kapitel haben wir uns mit der elektromagnetischen Emission beschäftigt und uns über die Messgrößen / Einheiten Gedanken gemacht. Beginnen wir mit den beiden wichtigsten Verfahren zur Messung der elektromagnetischen Emission:

2.1 Messung der gestrahlten Emission

Das bekannteste Messverfahren ist die Bewertung der von einer Komponente ausgehenden elektromagnetischen Strahlung. Formal richtig wäre zu sagen: „Wir messen die Amplitude der von den Geräten ausgehenden elektromagnetischen Wellen im Fernfeld“. Mit dem Zusatz „Fernfeld“ wird deutlich gemacht, dass es sich eben nicht um rein elektrische oder rein magnetische Felder handelt. In Wirklichkeit handelt es sich bei der elektromagnetischen Welle um eine Kombination aus elektrischen und magnetischen Feldern, die 90° zueinander angeordnet sind und sich im Raum ausbreiten. Der Fachbegriff lautet transversalelektromagnetische Welle (TEM-Welle), spielt aber für unsere EMV-Messung nur eine untergeordnete Rolle. Obwohl es sich um eine Kombination beider Feldarten handelt, wird das Messergebnis der emittierten Welle in V/m (sprich: Volt pro Meter) ausgedrückt, also in der Einheit für die elektrische Feldstärke. Dieser Umstand braucht uns nicht weiter zu beunruhigen, da es im Fernfeld eine feste Beziehung zwischen den beiden Feldern gibt. Man hätte sich ebenso gut für die magnetische Feldstärke entscheiden können.



Bikonische Antenne der Firma ETS-Lindgren

[Link](#) für Messungen im Frequenzbereich von 20 MHz – 200 MHz

Das Ziel der Messung ist es, über einen sehr weiten Frequenzbereich eine Aussage über die emittierte Feldstärke zu erhalten. Die ermittelten Messgrößen unterliegen den vorgegebenen Grenzwerten, die es einzuhalten gilt. Zur Messung benötigen wir eine Empfangsantenne und ein geeignetes Messgerät.

Empfangsantennen arbeiten meist nur in einem eingeschränkten Frequenzbereich, weshalb meist mehrere Antennen zum Einsatz kommen müssen. Außerhalb der Arbeitsbereiche sind die Antennen zu unempfindlich. Typische Empfangsantennen zur EMV-Messung sind bikonische und logarithmisch-periodische Antennen für den Einsatz im MHz- bis GHz-Bereich.

Frequenzen kleiner 1 MHz können mit klassischen Monopolen, sehr hohe Frequenzen (> 1 GHz) mit Hornantennen bewertet werden. Als Messinstrument wird ein Messempfänger verwendet. Verglichen mit einem Oszilloskop sind Messungen mit einem Messempfänger leider wenig intuitiv und es bedarf einiger Übung, die Geräte richtig zu bedienen. Das liegt vermutlich daran, dass wir es eher gewohnt sind, Abläufe im Zeitbereich anstatt im Frequenzbereich zu beschreiben. Das Erste, was wir machen müssen, ist, dem Messempfänger zu sagen, in welchem Frequenzbereich er messen soll. Der Frequenzbereich muss natürlich zur entsprechenden Antenne passen. Im nächsten Schritt wird festgelegt, wie fein die Frequenzauflösung sein soll. Damit wären schon die wichtigsten Einstellungen vorgenommen. Zusätzlich können wir festlegen, wie lange der Messempfänger auf einer Frequenz verweilen soll.

Wird die Messung gestartet, setzt sich der Messempfänger auf den ersten Frequenzpunkt und misst das Antennensignal. Danach springt er zum nächsten Frequenzpunkt und wiederholt dieses Spiel, bis er den gesamten Frequenzbereich abgefahren hat.

Eine wichtige Einstellung haben wir allerdings noch nicht beachtet: Wollen wir das korrekte Frequenzspektrum ermitteln, analog dem Ergebnis einer Fourier-Reihe, müssten wir zum einen unendlich viele Einzelmessungen (Frequenzschritte) durchführen und zum anderen wirklich exakt an einem Frequenzpunkt messen. Dazu ist ein Bandpass mit unendlich hoher Güte notwendig oder eben ein idealer Bandpass. Je breitbandiger der Bandpass ist, desto mehr werden benachbarte Frequenzen das Ergebnis an dem eigentlichen Frequenzpunkt verfälschen. Da wir weder unendlich lange warten können, bis das Ergebnis vorliegt, noch ideale Bandpässe herstellen können, behilft man sich mit einer einfachen Lösung. Die Grundnormen geben innerhalb der Messvorschrift auch die Bandbreite an, mit der gemessen werden muss. Typische Messbandbreiten sind 9 kHz oder 120 kHz. Um den Messfehler möglichst klein zu halten und dennoch nicht zu langsam zu messen, wird für die Frequenzauflösung meist die halbe Messbandbreite gewählt.

Für ein korrektes Ergebnis sind noch zwei weitere Dinge zu beachten. Die physikalische Ausgangsgröße einer Antenne ist eine Spannung bzw. die Fußpunktspannung. Somit ist der Messempfänger vom Prinzip her ein Voltmeter, das im Frequenzbereich messen kann. Vom Antennenhersteller erhalten Sie eine Korrekturkurve, die sogenannte Antennenkorrektur, mit der das gemessene Ergebnis multipliziert werden muss.

Gestrahlte Emission [V/m] = Messwert [V] $\times$ Antennenkorrektur [1/m]

Neben der Antennenkorrektur muss die Dämpfung der hochfrequenten Signale auf dem Signalkabel zwischen der Antenne und dem Messempfänger

berücksichtigt werden. Generell gilt: Je hochfrequenter das Signal, desto größer ist die auftretende Dämpfung. Dieses Verhalten ist über das einfache LC-Ersatzschaltbild mit dessen Tiefpassstruktur leicht herzuleiten. Die Kabeldämpfung ist eine einheitslose Größe, die ebenfalls im Ergebnis des Messempfängers berücksichtigt werden muss. Mit dem vorhandenen Setup und den Korrekturfaktoren steht einer erfolgreichen Messung nichts mehr im Weg. Moderne Messempfänger übernehmen die Korrektur der Antenne und der Kabeldämpfung für uns. Oft ist es sogar möglich, die Kabeldämpfung mit den Geräten zu ermitteln und direkt in das Setup einzupflegen. Zu beachten ist, dass für die Antennenkorrektur meist zwei Faktoren angegeben werden, je nach Ausrichtung der Antenne. Unterschieden wird zwischen der horizontalen und vertikalen Polarisation der Antenne (Ausrichtung).

2.2 Was versteckt sich hinter der Einheit dB?

Vielleicht ist Ihnen schon aufgefallen, dass Messgrößen und Grenzwerte in der EMV in dB ausgedrückt werden. Die Einheit dB ist ein Kunstwort, welches sich aus zwei Abkürzungen zusammensetzt. In Wirklichkeit handelt es sich sogar um gar keine „richtige“ Einheit, Wikipedia spricht sogar zurecht von einem - Pseudomaß! „d“ steht hier für die Multiplikation mit dem Faktor 10 (Dezi) „B“ wurde zu Ehren des [Alexander Graham Bell](#), welcher als Erster die Industrialisierung des Telefons vorangetrieben hat, gesetzt. Ähnlich wie bei den trigonometrischen Funktionen muss das Argument des Logarithmus stets einheitenlos sein! So ist es ja zum Beispiel nicht möglich, den Sinus von 5V $\sin(5V)$ bzw. den Logarithmus von 20A $\log_{10}(20A)$. Um die Einheiten loszuwerden, müssen wir uns auf eine Referenz beziehen, welche beliebig gewählt werden darf. Damit auch unser Bürokollege weiß, was bei einer Umrechnung gemeint ist, wird am Ende der dB-Angabe die Bezugseinheit gesetzt.

Bezugsgröße 1V	Ergebnis in dBV (sprich: dB volt)
Bezugsgröße 1mV	Ergebnis in dBmV (sprich: dB Millivolt)
Bezugsgröße 1 μ V/m	Ergebnis in dB μ V/m (sprich: dB Mikrovolt pro Meter)
Bezugsgröße 1mW	Ergebnis in dBm (sprich dBm)

Die Tabelle ist für alle in der Elektrotechnik verwendeten Einheiten erweiterbar, z.B. dBA, dB μ A, Achtung: Die einzige Ausnahme stellt die Angabe der Leistung in mW dar. Hier hat es sich durchgesetzt, anstatt dBmW nur dBm zu verwenden, siehe Tabelle.

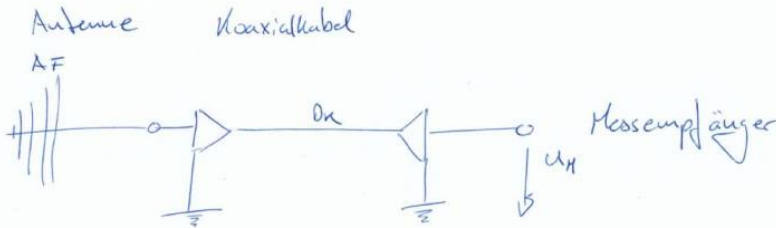
Mit der Normierung auf den Referenzwert ist es jetzt möglich, den Logarithmus zu bilden:

Leistungspegel	$p_{dB} = 10 \cdot \lg \frac{P_1}{P_0} dB_m$	Bezugsgröße: $P_0 = 1 \text{ mW}$
Spannungspegel	$u_{dB} = 20 \cdot \lg \frac{U_x}{U_0} dB_{\mu V}$	Bezugsgröße: $U_0 = 1 \mu V$
Strompegel	$i_{dB} = 20 \cdot \lg \frac{I_x}{I_0} dB_{\mu A}$	Bezugsgröße: $I_0 = 1 \mu A$
E-Feldstärkepegel	$E_{dB} = 20 \cdot \lg \frac{E_x}{E_0} dB_{\mu V/m}$	Bezugsgröße: $E_0 = 1 \mu V/m$
H-Feldstärkepegel	$H_{dB} = 20 \cdot \lg \frac{H_x}{H_0} dB_{\mu A/m}$	Bezugsgröße: $H_0 = 1 \mu A/m$

Aufpassen muss man bei der Berechnung lediglich darauf, welcher Faktor anzusetzen ist. Am einfachsten ist es, man merkt sich die Formel für die Leistungspegel mit dem Alleinstellungsmerkmal des Faktors zehn. Die restlichen Ergebnisse des Logarithmus werden mit zwanzig multipliziert für Spannung, Strom, Feldstärken usw.

Logarithmische Größen haben den Vorteil, dass sich auch Ergebnisse über einen weiten Messbereich darstellen lassen. Allerdings führt das auch dazu, einzelne Amplituden massiv zu unterschätzen.

Die Aussage: "Das Ergebnis ist nur 6dBµV über dem erlaubten Grenzwert" bedeutet durch die logarithmische Darstellung, dass der vorhandene Spannungspegel doppelt so hoch ist als erlaubt! Die Angabe aller Werte in dB hat einen weiteren entscheidenden Vorteil. Wir haben uns bei der Betrachtung zum Messaufbau der gestrahlten Emission die Korrekturfaktoren für die Antenne und die Koaxialleitungen angeschaut. Beide Korrekturfaktoren mussten mit dem Messergebnis des Messempfängers multipliziert werden. Durch die Rechenregeln des Logarithmus bzw. der 10^x Funktion wird aus der Multiplikation eine einfache Addition. Das bedeutet, dass zur Korrektur der Werte lediglich alle Einzelergebnisse addiert werden.



Feldstärke

$$E \left[\frac{\text{dB}\mu\text{V/m}}{\text{m}} \right] = U_M \left[\text{dB}\mu\text{V} \right] + \underset{\substack{\uparrow \\ \text{Antennefaktor} \\ \text{(Korrektur)}}}{AF \left[\text{dB} \frac{1}{\text{m}} \right]} + \underset{\substack{\uparrow \\ \text{Kabeldämpfung}}}{D_K \left[\text{dB} \right]}$$

dB und Kopfrechnen

Die Umrechnung zwischen linearen und logarithmischen Größenordnungen kann einfach als Kopfrechnung überschlagen werden. Voraussetzung ist, dass man den Grundverlauf der Logarithmusfunktion kennt.

Der wichtigste Punkt für uns ist $\log(1) = 0$.

Damit ergibt sich sofort:

- $1\text{V} = 0\text{dBV}$
- $1\mu\text{V} = 0\text{dB}\mu\text{V}$
- $1\text{V/m} = 0\text{dBV/m}$
- $1\text{mW} = 0\text{dBm}$ usw.

Entspricht der zu berechnende Wert einer Zehnerpotenz der Bezugsgröße, kann das Ergebnis direkt angegeben werden, da $\log(10) = 1$, $\log(100) = 2$, $\log(1000) = 3$ usw..

So erhält man zum Beispiel:

- $10\text{V} = 20\text{dBV}$
- $100\text{V} = 40\text{dBV}$
- $1000\text{V} = 60\text{dBV}$

Vorsicht ist geboten bei der Berechnung von Leistungspegeln, da hier mit dem Faktor 10 vor dem Logarithmus gearbeitet wird:

- $10\text{mW} = 10\text{dBm}$
- $100\text{mW} = 20\text{dBm}$
- $1000\text{W} = 30\text{dBm}$

Die Zwischengrößen lassen sich ohne Taschenrechner nicht ohne Weiteres berechnen. \ Allerdings kann in vielen Fällen eine Abschätzung für eine Verdoppelung des linearen Wertes herangezogen werden.

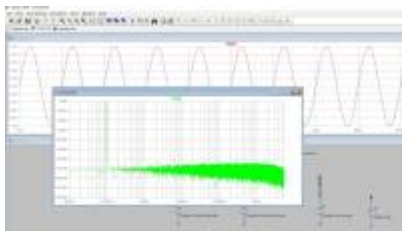
+3dB entspricht näherungsweise einer Verdoppelung der Leistung,
+6dB entspricht näherungsweise einer Verdoppelung der Spannung / Strom / ..

Beispiele:

- $1\text{V} = 120\text{dB}\mu\text{V}$
- $2\text{V} \approx 126\text{dB}\mu\text{V}$ (Exakter Wert $126,02\text{dB}\mu\text{V}$)
- $4\text{V} \approx 132\text{dB}\mu\text{V}$ (Exakter Wert $132,04\text{dB}\mu\text{V}$)
- $8\text{V} \approx 138\text{dB}\mu\text{V}$ (Exakter Wert $138,06\text{dB}\mu\text{V}$)
- $10\text{V} = 140\text{dB}\mu\text{V}$

Wir sehen, der Fehler durch die 6dB Regel ist vernachlässigbar und kann bei einer Überschlagsrechnung in Kauf genommen werden. \

Neben Messempfängern sind auch Oszilloskope in der Lage, das Frequenzspektrum eines Signals darzustellen. Auch hier kommt meist eine Anzeige bzw. Berechnung in dB zum Einsatz. Da Simulationsprogramme die Realität möglichst genau abbilden möchten, wird auch hier auf die dB-Darstellung zurückgegriffen. \ Alle Messgeräte und Simulationstools haben gemein, dass man zuerst herausfinden muss, welche Werte man angezeigt bekommt. Entweder man bemüht das Handbuch oder erzeugt sich entsprechende Testsignale.



Das Beispiel zeigt das Ergebnis im Frequenzbereich einer periodischen Sinusfunktion mit 1V Amplitude und einer Frequenz von 100kHz, erstellt mit einer LT-Spice-Simulation. Deutlich zu sehen ist ein Peak bei den erwarteten 100kHz. Es gilt jetzt nur noch herauszufinden, ob es sich um einen dargestellten Spitzen- oder Effektivwert handelt sowie, was sich hinter der Darstellung dB versteckt. \ \ LTSpice hält sich leider nicht an die gebräuchliche Darstellung und man erkennt, dass mit der Achsenbeschriftung dB eigentlich gemeint ist dBV, da der Peak bis knapp an die 0dB heranreicht. Der Zoom zeigt, dass die Amplitude eben gerade nicht genau bis an die 0dB heranreicht, sondern nur bis -3dB. Dahinter versteckt sich, dass uns das Programm nicht den Spitzenwert (eigentliches Ergebnis der Fourierreihe) ausgibt, sondern den Effektivwert. In der Simulationsdatei zum Download sind verschiedene Impulse angelegt, welche sich besonders zur Untersuchung im Frequenzbereich eignen.

impulse.zip



Die FFT-Funktion in LT-Spice gibt als Ergebnis den Effektivwert in dBV aus!
Ströme entsprechend als Effektivwert in dBA, usw.

Eine gute Zusammenfassung und viele weitere Informationen zum Thema logarithmische Pegel finden Sie in den Application Notes von Rhode&Schwarz [Link](#)
(Natürlich gibt es dazu auch verschiedene Rechner für Smartphones, siehe [Link](#))

2.3 Messung der leitungsgebundenen Emission

Das zweite wichtige Messverfahren, das wir uns anschauen, ist die leitungsgebundene Emission. Nahezu jedes elektrische Gerät besitzt einen Anschluss für die Versorgungsleitung, sei es für 230VAC oder im Fahrzeug an die Batterieklemmen. Das bedeutet aber auch, dass Störströme auf den Versorgungsleitungen ungehindert jeden Teilnehmer im Netz erreichen und beeinflussen können. Aus diesem Grund hat die Messung und Reglementierung der leitungsgebundenen Emission einen hohen Stellenwert. Wir werden später sehen, dass zur Reduktion der gestrahlten Emission viele geometrische und mechanische Aspekte mit eingehen. Die Reduktion der leitungsgebundenen Emissionen hingegen ist meist nur eine Herausforderung an die Hardwareentwickler.



Gemessen wird je nach Anwendung in einem Frequenzbereich von 150kHz bis 110MHz, wobei die exakten Werte je nach Anwendung variieren können (Automotive, Industrie, Rüstung, ...). Als Ergebnis der Messung erhält man typischerweise eine Störspannung in dB μ V.

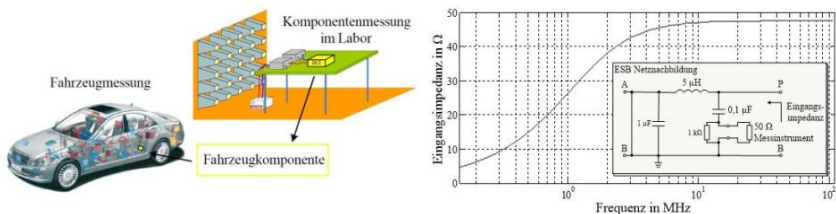
Der Messaufbau zur Messung leitungsgebundener Störungen besteht aus einer Netznachbildung, dem Messkabel und einem Messempfänger oder Spektrumanalysator. Das zentrale Element der Messung ist die Netznachbildung, welche eine Auskopplung der Störungen ermöglicht. Betrachtet man den Schaltungsaufbau einer Netznachbildung, wäre es besser zu sagen, die Messung koppelt den Störstrom auf den Versorgungsleitungen aus, welcher dann am 50 Ω Widerstand des Messempfängers eine messbare Störspannung hervorruft. Netznachbildungen gibt es in verschiedenen Ausführungsformen. Die dargestellte, gebräuchlichste Form (Bild: HTWG) ist die V-Netznachbildung, welche die Spannung von jeder Versorgungsleitung (L1, L2, L3, N) zur Referenzmasse bewertet. Die Netznachbildungen gibt es natürlich auch in kompakter Bauform für einphasige Anwendungen.\



Neben der Auskopplung der Störströme hat die Netznachbildung, entsprechend ihrer Bezeichnung, die Aufgabe, das spätere Versorgungsnetz nachzubilden. Das ist natürlich keine einfache Aufgabe, da je nachdem, an welcher Stelle im Netz

die Komponente betrieben wird, sich unterschiedliche Netzimpedanzen ergeben. Das bedeutet, eine normgerechte Impedanznachbildung kann nur eine typische Versorgungsstruktur abbilden, jedoch nie die tatsächlich in der Realität auftretenden Werte. Dadurch besteht immer die Möglichkeit, dass ein die Grenzwerte einhaltendes Gerät in der Anwendung Probleme verursacht. Am häufigsten ist dieser Effekt aufgrund der hohen Elektronikdichte in Kraftfahrzeugen anzufinden. Als i.O geprüfte Komponenten im Labor können plötzlich im Fahrzeug auffällig werden und benachbarte Teilnehmer im Versorgungsnetz beeinflussen.

Im Aufbau unterscheiden sich die Netznachbildungen hauptsächlich über die eingesetzte Induktivität. Geht man im Fahrzeug von einer maximalen Anschlusslänge zur Batterie von 5m aus, ergeben sich mit der Daumenregel der Leitungsinduktivität von $1\mu\text{H}/\text{m}$ die eingesetzten $5\mu\text{H}$. Da für die Netznachbildung kein Korrekturfaktor berücksichtigt werden muss, ist es lediglich notwendig, die Dämpfung der Messleitung zum Messempfänger zu berücksichtigen.



Für Haushalts- und Industrielängen geht man von einer deutlich höheren Anschlusslänge aus und damit von einer Induktivität von $50\mu\text{H}$ pro Leitung, also separat für L und N. Als Ergebnis erhält man zwei Messungen, jeweils gegen die Referenzmasse, in diesem Fall PE. Im Fahrzeug ist die Lage nicht so eindeutig wie an der Steckdose. Unsere Energieversorgung entspricht an der Versorgungsstelle (Steckdose) einem echten Dreileitersystem. Im Fahrzeug gibt es normalerweise nur zwei Bezugspunkte, Kl. 30 und Kl. 31. Wir wissen aber, dass das Fahrzeugchassis meist auf Kl. 31 (Batterie Minuspol) angeschlossen wird und damit als Rückleiter verwendet wird. Man spart sich damit zusätzliche Verkabelungswege. Dies kann aber im Fahrzeug jetzt zu der identischen Situation wie zu Hause führen. Auch dort wird im Keller der Rückleiter aufgetrennt in einen Schutzleiter (Referenzerde oder Referenzmasse) und den Neutraleiter. An der Trennstelle besitzen beide Leiter das identische Potential. Mit zunehmender Entfernung jedoch hängt das Potenzial von den Betriebsparametern wie Stromfluss, Frequenz des Stromes, Leitungsquerschnitt usw. ab. Bei genauer Betrachtung stellen wir fest, dass die Situation in Fahrzeugen identisch ist. Allerdings existiert nicht nur ein Aufspaltungspunkt, sondern in jeder

Fahrzeugpartie meist mehrere Punkte, den sogenannten Massebolzen. Die Bolzen verteilen die Masse an die einzelnen Steuergeräte.

2.4 Weitere Messverfahren zur Emissionsmessung

Mit der gestrahlten und leitungsgebundenen Emission kennen wir bereits die wichtigsten Messverfahren. Je nach Produkt können weitere Messverfahren herangezogen werden, zum Beispiel leitungsgebundene Emissionen auf Telekommunikationsleitungen oder Sensorleitungen mit Hilfe einer Stromzange (siehe Stromzangenmessverfahren CISPR25). Alle diese Messverfahren eint, dass sie versuchen, die hochfrequenten Emissionen der Prüflinge zu ermitteln. Die EMV betrachtet neben den hochfrequenten Störungen auch Störungen auf dem AC-Versorgungsnetz. Bisher wurde die Rückwirkung auf die Versorgungsleitungen von Geräten weitestgehend ignoriert oder fand nur bei Anwendungen mit einer hohen Anschlussleistung Beachtung. Durch den vermehrten Einsatz von nicht linearen Verbrauchern (z.B. einphasige LED-Beleuchtung) ist das Thema Netzurückwirkung mehr in den Fokus gerückt. Aber von vorne, worum geht es?

2.4.1 Messung von Oberschwingungen

Nur noch wenige elektrische Geräte am einphasigen Netz verwenden die Netzspannung (50Hz, 230V AC) als Arbeitsspannung, wie zum Beispiel Gebläse, Staubsauger, Heizlüfter, Backofen. Die Mehrzahl arbeitet mit sogenannten Gleichspannungszwischenkreisen, nachgeschaltet an Brückengleichrichter. Typische Anwendungen sind LED-Vorschaltgeräte, Umrichter für Motoren kleiner und mittlerer Leistung, Leistungsverstärker, Ladegeräte usw.



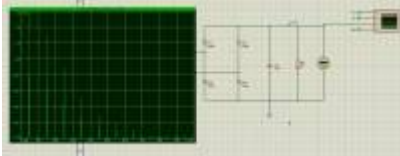
Die Stromaufnahme der Gleichspannungszwischenkreise ist allerdings bei weitem nicht mehr sinusförmig. Die resultierende Stromform ist abhängig von der Zwischenkreiskapazität und der Leistungsaufnahme der anschließenden Applikation. Eine einfache Simulation des Gleichrichters im Zeitbereich zeigt die Stromaufnahme der Schaltung. Deutlich zu erkennen ist, dass der Kurvenverlauf nichts mehr mit der Sinus- Kosinusfunktion zu tun hat. Prinzipiell kein Problem

und man könnte argumentieren: Dem Netz doch egal.....

Stimmt, wären da nicht die generierten Oberschwingungen (Vielfache der Grundschwingung).

Wir wissen: Nach Fourier lassen sich periodische Signale zerlegen in eine unendliche Folge sinus- und cosinusförmiger Signale, welche sich in ihrer Amplitude, Frequenz und Phasenlage unterscheiden.

In der EMV interessieren uns normalerweise nur die ersten beiden Parameter. Die Phaseninformation wird nicht weiter beachtet, da es ja hinsichtlich der Störbeeinflussung egal ist, in welcher Relation die Störungen zueinander stehen. Wichtig ist, ob sie uns stören (Amplitude) und in welchem Frequenzbereich.

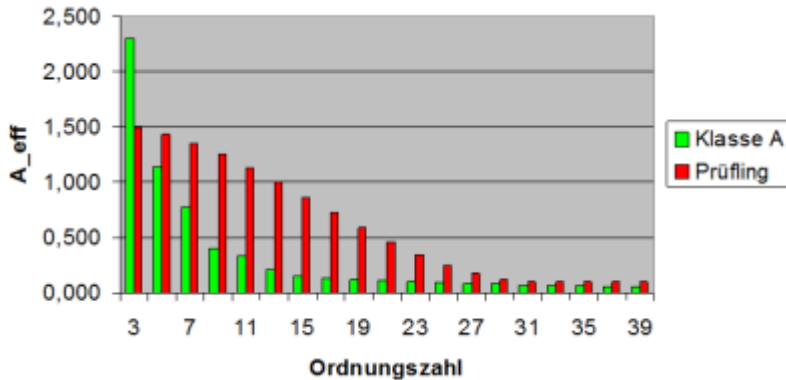


Die Simulationen zu den Oberschwingungen wurden mit dem Programm Proteus der Fa. Labcenter erstellt. Auch hier gilt es vorab zu klären, welche Informationen angezeigt werden. Gewählt wurde hier eine lineare Achsenaufteilung, die Amplitude entspricht

dem Effektivwert.

Die Fourierzerlegung des Stromverlaufs zeigt einen hohen Anteil an Oberschwingungen, selbst für Frequenzanteile, welche einem Vielfachen der Grundfrequenz von 50Hz entsprechen. Die generierten Oberschwingungen sind für unser System kein Problem, allerdings für die im Netzverbund angeschlossenen weiteren Verbraucher. Die Spannung im Netz erfährt durch die Leitungsimpedanz eine hohe Verzerrung und kann dazu führen, dass die gemessene Netzspannung eher einem Rechteckverlauf ähnelt als einem sinusförmigen Verlauf.

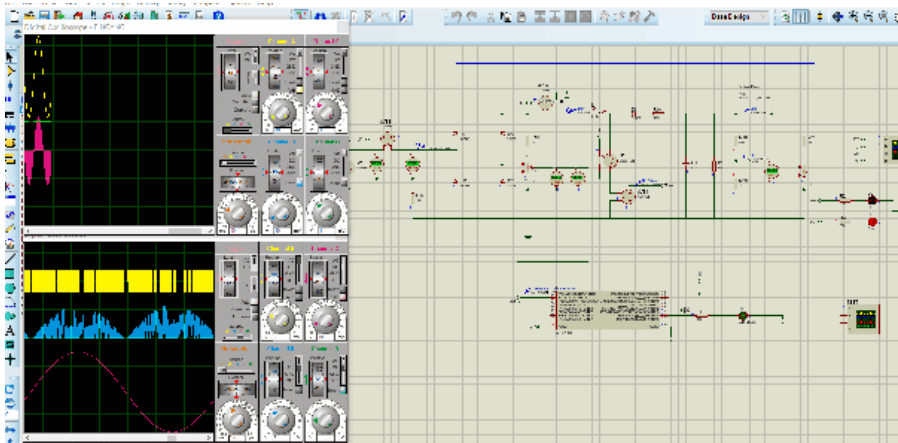
Ein weiterer ungewollter Effekt der Oberschwingungen besteht darin, dass nur die 50Hz-Grundschwingung in der Lage ist, Wirkleistung zu übertragen. Sämtliche Oberschwingungen erzeugen reine Blindleistung, welche durch den Blindstrom die Übertragungsnetze unnötig belasten.



Für das dargestellte Beispiel der Oberschwingungen beträgt die Wirkleistungsaufnahme der Schaltung ca. 350W. Die Scheinleistung, definiert als $U_{eff} \cdot I_{eff}$, beträgt 880VA. Es besteht somit ein krasses Missverhältnis und dementsprechend ein schlechter $\cos(\phi)$ von 0,4 (Verhältnis von Wirk- zu Blindleistung). Die zulässigen Grenzwerte für die Oberschwingungen auf den Versorgungsleitungen sind in EN61000-3-2 geregelt (Limits for harmonic current emissions). Hierin werden Grenzwerte für die Oberschwingungen definiert für Geräte mit einer Leistungsaufnahme $>75W$. Es stehen verschiedene Lösungsmöglichkeiten zur Verfügung, den hohen Anteil an Blindströmen und damit eine hohe Blindleistung zu reduzieren. Alle Verfahren haben gemeinsam, dass sie versuchen, einer sinusförmigen Stromaufnahme möglichst nahe zu kommen. Die einfachste Möglichkeit, hochfrequente Signale zu blockieren, besteht in deren Filterung. Mithilfe einer Induktivität wird der Stromverlauf geglättet und damit die Oberschwingungen gedämpft. Um einen ausreichend hohen Effekt festzustellen, sind allerdings meist Induktivitätswerte im mH-Bereich notwendig, wodurch die Produkte groß und schwer werden. Weiterhin sind große Spulen aufgrund des hohen Kupferanteils im Vergleich zu den restlichen Elektronikkomponenten sehr teuer. Diese Art der Leistungsfaktorkorrektur bietet sich an für Produkte mit kleiner Stückzahl, bei denen Gewicht und Volumen keine entscheidenden Merkmale darstellen. Diese Art der Oberschwingungskorrektur wird meist als passive PFC (Power Factor Correction) bezeichnet.

Deutlich bessere Ergebnisse, sowohl beim erzielten $\cos(\phi)$ als auch bei den verbleibenden Oberschwingungen, erhält man durch den Einsatz einer aktiven PFC-Schaltung. Die aktive PFC bildet mit Hilfe eines Hochsetzstellers den sinusförmigen Eingangsstrom nach. Proportional zur anliegenden Eingangsspannung (AC-seitig) wird der Strom aufgenommen. Das Versorgungsnetz wird damit analog einer realen Last belastet. Die Nachteile der aktiven PFC sind direkt aus der Prinzipschaltung ablesbar.

Man könnte die Funktion auch so beschreiben, dass hier ein Hochsetzsteller verwendet wird, der nicht auf den Ausgang geregelt wird, sondern auf den Eingangsstrom.

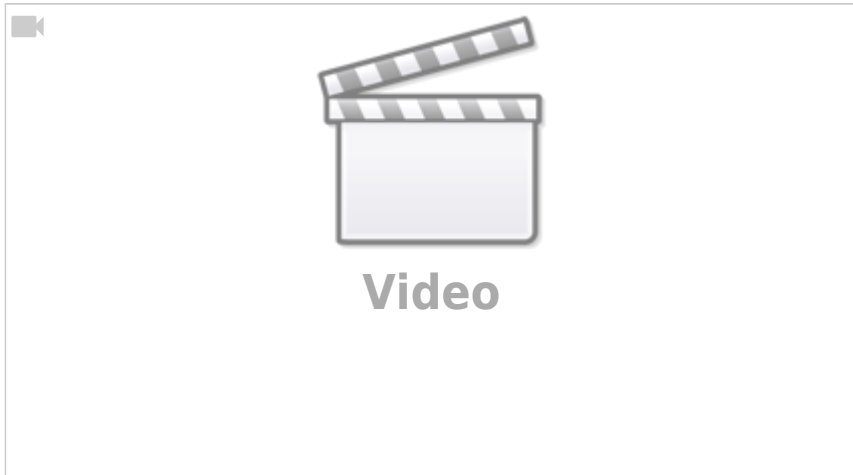


Das animierte Gif öffnet sich durch Anklicken. Die Hardware ist Bestandteil der Vorlesung PFC und wird als praktische Arbeit zuerst simuliert und anschließend aufgebaut. Durch den getakteten Hochsetzsteller können die Oberschwingungen reduziert werden, allerdings holt man sich durch die Schaltung selber eine EMV-Quelle in das Gerät. Weiterhin bedarf es einer deutlich höheren Entwicklungsarbeit, die Schaltung zu dimensionieren und für den Serieneinsatz zu qualifizieren. Auf dem Halbleitermarkt sind verschiedene integrierte Reglerbausteine verfügbar, z.B. TDA4863-2, welche mit nur wenigen externen Bauteilen auskommen. Die von Ihrem Produkt ausgehenden Oberschwingungen lassen sich auf einfache Weise im Labor mit Hilfe einer Stromzange und einem Oszilloskop ermitteln. Über die FFT-Funktion können die Oberschwingungen abgeschätzt werden. Für Abnahmemessungen sind kombinierte Prüfgeräte verfügbar, welche das gesamte Spektrum der DIN EN61000-3-2 abdecken.

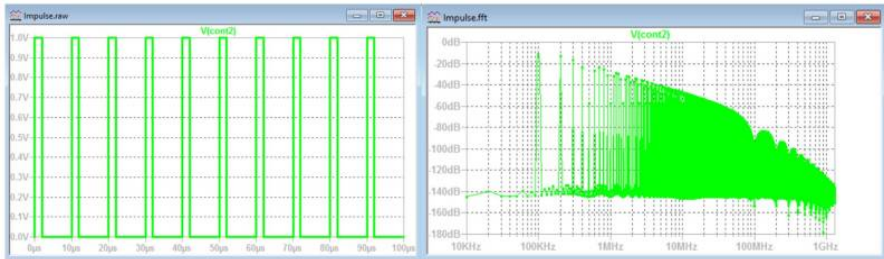
3.0 Woher kommen die Störgrößen?

Vielleicht ist Ihnen aufgefallen, dass die unterschiedlichen Messverfahren verschiedene Frequenzbereiche abdecken. Oberschwingungen betrachten ein geringes Vielfaches der Netzfrequenz, wohingegen leitungsgebundene Störgrößen bis zu 100MHz gemessen werden. Gestrahlte Emissionen müssen

sogar bis weit in den GHz-Bereich betrachtet werden. Aber woher kommen denn überhaupt die Störgrößen? Wer stört? Schaut man sich die Historie der EMV an, stellt man fest, dass das Thema erst seit den 80er Jahren so richtig Fahrt aufgenommen hat. Hätten sich Ingenieure vor etwas mehr als hundert Jahren noch gewünscht, sie wären in der Lage, elektromagnetische Emissionen (Wellen) zu generieren, haben wir heute mehr als genug davon, meist mehr ungewollte als gewollte



Ursächlich dafür sind die schnell taktenden Prozessorsysteme und die mit immer höheren Strömen schaltenden leistungselektronischen Systeme. Ob Prozessor oder die Leistungselektronik, sie alle erzeugen periodische Stromflüsse. [Joseph Fourier](#) hat uns gezeigt, dass jedes periodische Signal in eine unendliche Reihe an periodischen Sinus- und Kosinusschwingungen zerlegt werden kann. Die Schwingungen unterscheiden sich in deren Amplitude, Frequenz und Phasenlage. Wie so oft in der EMV interessiert uns die relative Lage der Signale zueinander, also die Phasenverschiebung, meistens nicht. Einzig die Amplitude entscheidet über die Intensität, mit der gestört wird, wohingegen die Frequenz festlegt, welche Anwendung sich stören lässt.

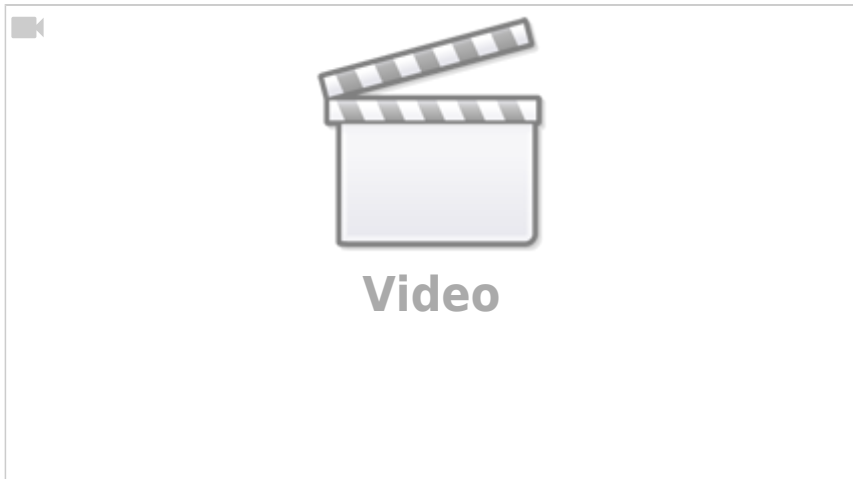


Das Beispiel zeigt ein einfaches Rechtecksignal im Zeitbereich mit der Grundfrequenz von 100kHz (links) mit den dazugehörigen Oberschwingungen im Frequenzbereich. Wir sehen, die Signale sind bis weit in den GHz-Bereich erfassbar bzw. verschwinden erst ab ca. 2GHz im mathematischen Rauschen der Berechnung in LT-Spice. Genau diese Signale sind es, die uns meist Kopferbrechen machen.



Unsere Aufgabe als EMV-Ingenieur ist es somit, dafür zu sorgen, dass die unerwünschten Oberschwingungen das Gerät nicht verlassen.

Der Download `Impulse.zip` (siehe oben) enthält das gezeigte Beispiel im Zeit- und Frequenzbereich. Zusätzlich sind weitere Impulse enthalten, mit denen sich Änderungen im Zeitbereich (Anstiegszeit, Periodendauer, usw.) im Frequenzbereich beobachten lassen. Das nachfolgende Video beinhaltet eine kurze Einführung in LTSpice zum Aufbau und Simulationssetup der Rechteckimpulse. Zum Abschluss werden die Fourierkoeffizienten händisch berechnet und mit den Werten aus LTSpice verglichen.



3.1 EMV- Tafel

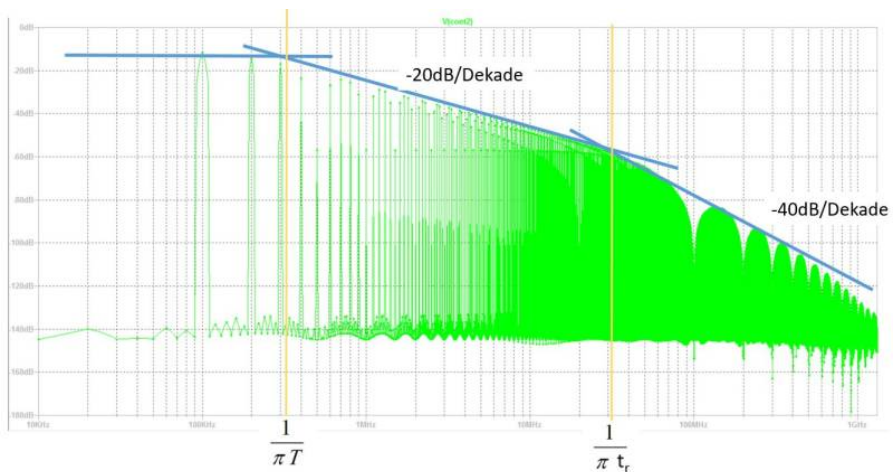
Die EMV-Arbeit kennt typischerweise zwei Arbeitsabschnitte. Am wichtigsten ist während der Entwicklungsarbeit zum neuen Produkt die Vorhersage, welche Auswirkungen Einzelmaßnahmen (z.B. Leitungslängen, Filterschaltungen, Gehäusekonstruktionen, usw.) auf das spätere EMV-Verhalten haben. Wie so oft bei nicht greifbaren Dingen tun wir uns mit Vorhersagen ziemlich schwer. Es ist zwar heute möglich, über Feldberechnungsprogramme einzelne Teilaspekte zu bewerten, allerdings ist dies eine mühsame und zeitraubende Angelegenheit. Sie müssen dazu ein 3D-Modell der Gesamtanordnung besitzen mit der genauen (elektrischen) Beschreibung aller Materialien. Da während der Entwicklung meist weder die Materialien im Detail bekannt sind noch die finale Geometrie, bleibt also nichts anderes, als sich auf die Erfahrung zu verlassen.

Der zweite EMV-Phase beschäftigt sich dann mit den Musterständen im EMV-Labor und versucht, die während der Entwicklung getroffenen Annahmen zu bestätigen oder entsprechend die Störquellen zu lokalisieren. Hier ist detektivische Arbeit gefordert, um sozusagen das Übel an der Quelle zu bekämpfen.

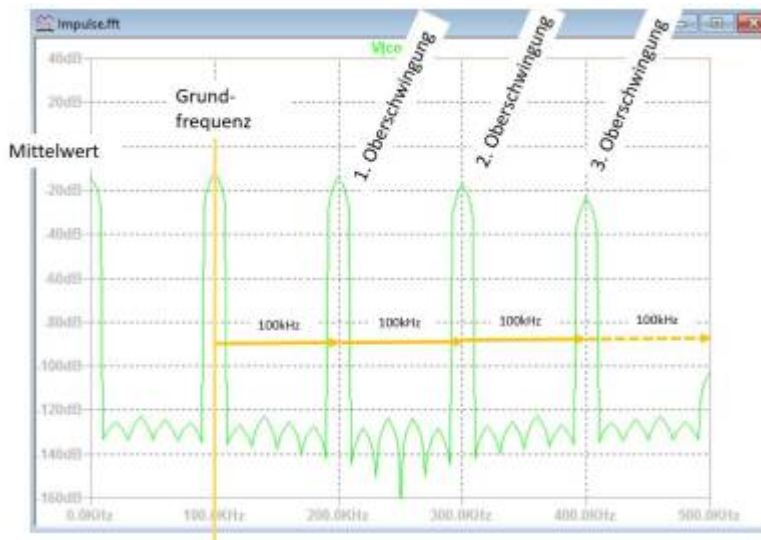
In beiden Arbeitsschritten ist die EMV- Tafel ein wichtiges Hilfsmittel. Die EMV-Tafel versucht, aus einem bekannten trapezförmigen Impuls im Zeitbereich die auftretenden Störsignale im Frequenzbereich abzuschätzen. Oder vielmehr deren Dämpfung über der Frequenz. Wie gesehen erzeugt ein einfaches Rechtecksignal unendlich viele Oberschwingungen. Die gute Nachricht ist, dass die erzeugten

Oberschwingungen mit steigender Frequenz beständig abnehmen und niemals anwachsen. Das bedeutet, dass wir beruhigt sagen können, dass sobald wir die Oberschwingungen eines Signales nicht mehr messen können (verschwinden im Rauschen), dann auch oberhalb davon keine neuen Oberschwingungen mehr hervortreten. Das bedeutet, wir müssen uns lediglich überlegen, ab wann die Signale so stark gedämpft werden, dass sie nicht mehr messbar sind bzw. für die Grenzwerte keinerlei Bedeutung mehr haben.

Genau hier setzt die EMV- Tafel an. Das zuvor gezeigte Beispiel des Rechteckimpulses zeigt einen typischen Verlauf im Frequenzbereich. Da Rechtecksignale sehr oft auftreten (μC , Prozessortakte, Signaltakte, Kommunikation), ist es interessant, sich den „Fußabdruck“ dieser Signale genauer anzuschauen. Die EMV-Tafel lässt sich herleiten über die Spektraldichtefunktion eines einmalig auftretenden Trapezimpulses. Die Spektraldichte lässt sich mit Hilfe des Fourierintegrals berechnen, welches an dieser Stelle allerdings zu weit führt. Für periodische Signale kann die EMV-Tafel als Worst-Case-Abschätzung herangezogen werden.



Charakteristisch sind die beiden Eckfrequenzen, welche durch die Einschaltdauer T_e und die Flankenanstiegs- und Abfallzeit t_r definiert werden. Je länger die Periodendauer gewählt wird bei gleichbleibender Impulslänge, desto weiter erfolgt der Übergang von der Spektralfunktion zum Amplitudendichtespektrum. Das bedeutet, die einzelnen Spektrallinien gehen über in unsere Worst-Case-Abschätzung, bei der keine diskreten Frequenzanteile mehr vorhanden sind, sondern ein sogenanntes breitbandiges Spektrum bilden.



Typischerweise werden Taktsignale oder PWM ([Pulsweitenmodulation](https://de.wikipedia.org/wiki/Pulsweitenmodulation))-Signale] (<https://de.wikipedia.org/wiki/Pulsdauernmodulation>] %29-Signale) zur Leistungssteuerung eingesetzt, bei denen die Einschaltdauer in der Größenordnung der Periodendauer gewählt wird. Zum Beispiel 50% an, 50% aus für symmetrische Rechtecksignale oder 10% an, 90% aus, um die Ausgangsspannung eines Tiefsetzstellers auf 10% zu setzen. Hier können wir die **Schmalbandigkeit** des Störspektrums für unsere detektivische Arbeit ausnutzen. Für periodische Signale wissen wir aus der Fourierreihe, dass sich das Ursprungssignal aus der Grundschwingung plus (Ober-)Schwingungen zusammensetzt, deren Frequenzen stets ein ganzzahliges Vielfaches der Grundfrequenz entsprechen. Das bedeutet, durch den Abstand zweier schmalbandiger Störungen im Frequenzbereich können wir oft direkt auf die Grundfrequenz schließen. Das Blockschaltbild der Anwendung gibt uns dann Aufschluss auf den Ursprung der Taktgenerierung.

Ich hoffe, Sie erheben Einspruch zu dieser These, da Sie in anderen Vorlesungen gelernt haben, dass bei idealen symmetrischen Rechtecksignalen nur die ungeraden Vielfachen der Grundfrequenz vorhanden sind. Die geradzahligten werden zu Null. Stimmt, Sie haben recht! Allerdings treten die geradzahligten Oberschwingungen aus dem Rauschen heraus, falls die Anstiegs- und Abfallzeiten ca. $1/1000$ der Pulsbreite entsprechen.

Schwerer macht uns das Leben somit eher, dass die Störungen zeitlich nicht konstant sind. Kommunikationssignale haben als Abfolge von digitalen Signalen niemals eine strenge Periodizität. Auch PWM-Signale zur Leistungsregelung werden gegebenenfalls ständig nachgeregelt, um Sollwerte zu halten. Für solche

Fälle empfiehlt es sich, den Softwarekollegen zu bitten, eine separate EMV-Software zu schreiben, mit deaktivierten Regelschleifen und falls möglich mit deaktivierter oder eingeschränkter Kommunikation.

3.2 Funktion und EMV

Die elektromagnetischen Eigenschaften sind natürlich immer mit der Funktion einer Komponente verknüpft. Keine Funktion, keine Störungen. \ Das wichtigste Dokument bei der Produktentwicklung ist ein übersichtliches Blockschaltbild. Mit Hilfe des Blockschaltbilds erhält man einen schnellen Überblick über die Gesamtfunktionalität, weiterhin ist es möglich, jeder Anforderung aus dem Lastenheft einen funktionalen Block zuzuordnen. Beispiele von funktionalen Blöcken sind

- Spannungs- / Strommessung
- Sensorsignale
- Endstufen bzw. leistungselektronische Schaltungen
- Spannungsversorgung, typ. DC/DC-Wandler
- Gleichrichter
- Zukaufteile wie z.B. Netzteile
- usw.

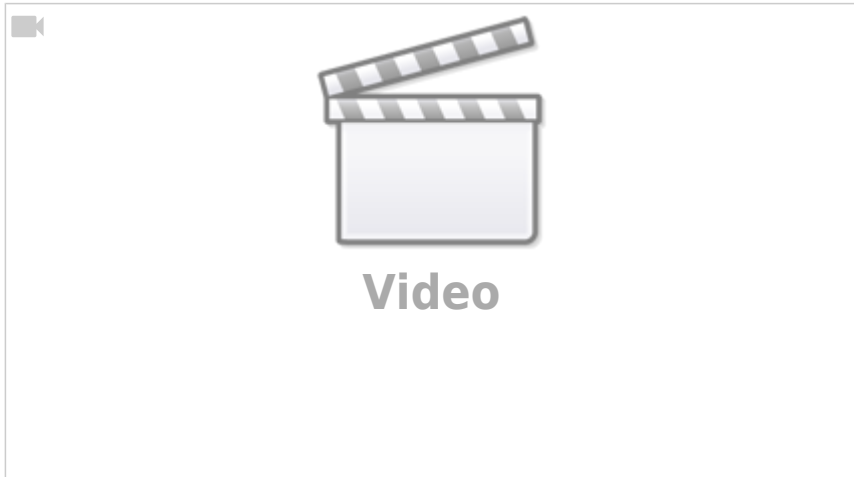
Das Blockschaltbild kann auch als Basis für die EMV-Arbeit dienen. Hier ist es wichtig, generell **alle** Schnittstellen zu bewerten. Das bedeutet, jedem Übergang zwischen zwei Blöcken muss eine EMV-Maßnahme zugeordnet werden. Schnittstellen nach außen sind in gleicher Weise zu betrachten. In der frühen Entwicklungsphase ist es noch nicht entscheidend, wie die Maßnahme umgesetzt oder ausgeführt wird. Wichtig ist, sie vorzusehen und im späteren Entwicklungsablauf umzusetzen.



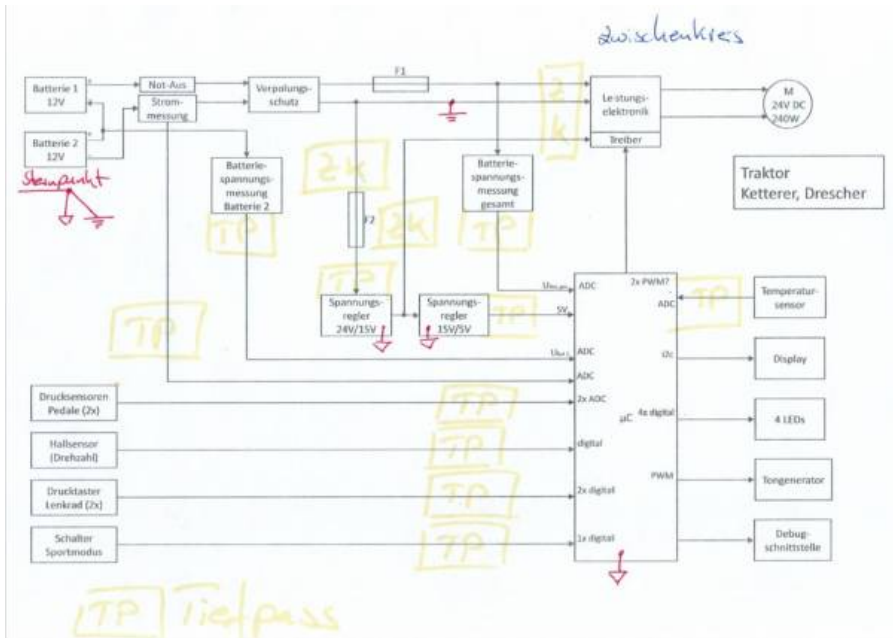
Für alle Ein- und Ausgänge der funktionalen Blöcke müssen EMV-Maßnahmen vorgesehen werden.

Natürlich ist es ebenso wichtig, die EMV-Eigenschaften innerhalb der einzelnen Blöcke zu berücksichtigen. Hier lässt sich in gleicher Weise verfahren. Die Funktion der Blöcke wird über ein Blockschaltbild beschrieben, welches seinerseits wieder EMV-Maßnahmen enthält.

3.3 Blockschaltbild / EMV-Konzept



Nicht ohne Grund wird ein Blockschaltbild mit integrierten EMV-Maßnahmen als EMV-Konzept bezeichnet. Hier laufen alle EMV-Betrachtungen von der Einzelmaßnahme bis hin zur Masseanbindung bzw. dem Massekonzept zusammen. Zu Beginn des EMV-Konzepts muss davon ausgegangen werden, dass zwischen allen Blöcken sämtliche möglichen Koppelmechanismen bestehen können.



Je nach zu erwartender Störung innerhalb eines Blockes werden die entsprechenden Maßnahmen ausgewählt. Ohne tief in die Theorie der Filterschaltungen einzusteigen, ist es naheliegend, dass auf allen Leitungen, auf denen ein reiner DC-Wert zu erwarten ist (z.B. DC-Spannungsversorgungen, DC-Spannungsmessung, generell reine DC-Größen), auch ein wechsellspannungsfreies Signal anliegen sollte. In der Praxis bedeutet das, dass sämtliche DC-Leitungen mit Kondensatoren in Parallelschaltung gefiltert werden können. Diese Art der Filterschaltung entspricht bei als ideal angenommenen Kondensatoren der einfachsten Tiefpassfilterschaltung. Die Auswahl der Kondensatoren ist dabei abhängig davon, welche hochfrequenten Signale unterdrückt werden sollen (siehe Kapitel: Kondensatorauswahl). Meist ist hier nicht der Kapazitätswert ausschlaggebend, sondern der parasitäre Serienwiderstand bzw. die parasitäre Serieninduktivität. Eine weitere naheliegende Lösung liegt darin, bereits im Blockdiagramm mechanische Maßnahmen wie Schirmbleche und Masseanbindungen vorzusehen. Je nach Funktion gilt es entweder, sensible Bereiche abzuschirmen (z.B. Audioverstärker, Magnetfeldsensoren, Sende- und Empfangseinrichtungen) oder besonders emittierende Bereiche von der restlichen Schaltung zu trennen. Wie immer besteht die Schwierigkeit darin, zu identifizieren, welche Bereiche besonderen Schutz benötigen bzw. welche Schaltungsteile separiert werden müssen. Natürlich obliegt diese Einschätzung einer gewissen Berufserfahrung, welche

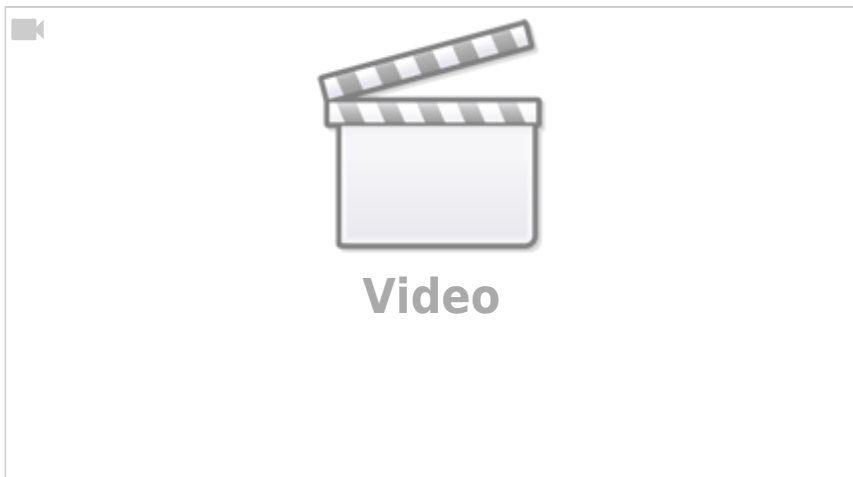
aber auch durch einfache Abschätzungen ermittelt werden kann. Vorsicht ist generell geboten bei sehr kleinen Nutzsignalen wie bei Audio- oder HF-Kommunikationssignalen. Die eingesetzten Verstärker arbeiten meist mit sehr hohen Verstärkungsfaktoren und sind damit sensibel gegenüber eindringenden Störgrößen. Auf der anderen Seite müssen Bereiche mit getakteten Strömen bei leistungselektronischen Schaltungen gesondert betrachtet werden. Hier ist es ggf. notwendig, die „Störquelle“ gegenüber den restlichen Schaltungsteilen zu separieren.

Das Blockschaltbild im dargestellten Beispiel wurde ursprünglich für eine Elektronik entwickelt, welche einen DC-Motor ansteuert. Mit der Überlegung, die Signale aller Blöcke zu filtern sowie die Massepfade zwischen Leistungs- und Digitalteil getrennt zu halten, wird daraus ein EMV-Konzept. Im nächsten Schritt geht es dann darum, die mit TP für Tiefpassfilter und ZK für Zwischenkreis bezeichneten Blöcke mit Leben zu füllen.

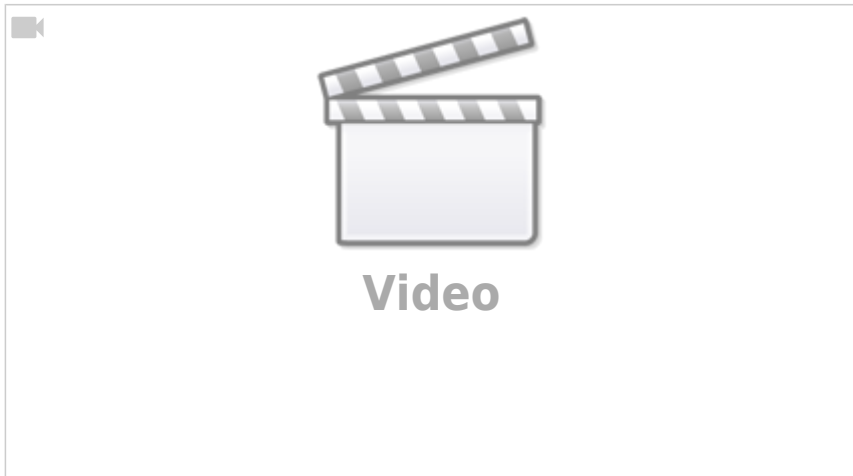
4.0 Koppelmechanismen

Was hat ein Phasenprüfer mit der kapazitiven Kopplung zu tun? Wie werden Störungen übertragen?

Die Antworten dazu im Video:



5.0 Filterschaltungen



Unsere schärfste Waffe, um vorhandene Störungen zu unterdrücken, ist der Einsatz von Filterschaltungen. Im Blockschaltbild ist es daher sinnvoll, alle Leitungen (Verbindungen zwischen einzelnen Blöcken und nach außen) mit einem Filter zu versehen. Hierbei geht es sowohl darum, die von den Blöcken generierte Emission zu unterdrücken bzw. nicht an den angrenzenden Block weiterzugeben, als auch die Störfestigkeit gegenüber von außen eindringenden Größen zu erhöhen. In der ersten Version des Blockschaltbilds ist es auch unerheblich zu wissen, wie der Filter aufgebaut ist und welchen Frequenzgang er aufweist. Wichtig ist hier lediglich, ihn zu berücksichtigen.

5.1 Filterauslegung



Bild: Fa. Schaffner Die Auslegung von Filterschaltungen ist leider nicht trivial und basiert in vielen Fällen auf dem Trial-and-Error-Prinzip. Wird die notwendige Störunterdrückung nicht erreicht, kommt einfach eine weitere oder modifizierte Filterschaltung zum Einsatz. Am bekanntesten und von außen sichtbar sind Netzfilter für Geräte mit einem ein- oder mehrphasigen Schukostecker. Die eingesetzten Filter werden zur Unterdrückung der emittierten

leitungsgebundenen Emission eingesetzt und sind in vielen unterschiedlichen Varianten erhältlich. In unserem Blockschaltbild ist der Netzfilter vor dem Steckeranschluss und dem ersten Funktionsblock zu finden.

Die klassische, passive Filterschaltung besteht im Wesentlichen aus drei Bauteilen, die in unterschiedlichen Anordnungen eingesetzt werden. Am häufigsten kommen **Kondensatoren** zum Einsatz. Dies hat den einfachen Grund, dass Kondensatoren in zahlreichen Bauformen und Technologien zur Verfügung stehen. **Induktivitäten** hingegen sind deutlich schwieriger einzusetzen. Hier ist vor allem darauf zu achten, dass die Sättigungsgrenze nicht überschritten wird sowie der Einsatz einer geschlossenen Bauform, damit auftretende Streufelder nicht zu Störungen auf der eigenen Platine führen. Von Stabkerndrosseln ist generell abzuraten, da sich die magnetischen Feldlinien stets außerhalb des Kerns bewegen. Bitte beachten Sie die im Sprachgebrauch übliche Bezeichnung für eine Induktivität als Spule oder Drossel. In der Literatur werden beide Begriffe zwischenzeitlich vermischt und nicht einheitlich verwendet. Ab und zu wird auch der Begriff Drosselspule verwendet. Von Drosseln spricht man überwiegend, sobald es darum geht, Ströme zu bremsen bzw. zu reduzieren.

Kombinierte Spulen mit mehr als einer Wicklung werden in der EMV als stromkompensierte Drosseln eingesetzt. Hier wird ausgenutzt, dass sich bei gegensinnigem Wicklungssinn das Magnetfeld auslöscht. Eine kompensierte Drossel oder engl. Common-Mode-Choke ist somit nur für gleichgerichtete Ströme oder besser, Gleichtaktströme wirksam. Die Unterscheidung in Gleich- und Gegentaktströme erfolgt in den nächsten Kapiteln.

Das Funktionsprinzip aller Filterschaltungen beruht darauf, entweder hochfrequente Störgrößen auf die Referenzmasse abzuleiten oder die Ausbreitung über eine hochohmige Impedanz zu verhindern. Vielleicht haben Sie bereits die Erfahrung gemacht, dass Filterschaltungen umso besser funktionieren, je besser die Masseanbindung ist. Dieser Effekt ist auf das Prinzip der Ableitung der Ströme zur Masse zurückzuführen und erklärt die großflächigen Metallaschen

am gezeigten Netzfilter der Firma Schaffner.

Bitte beachten Sie, dass die Befilterung von Versorgungsleitungen nichts mit der bekannten Filterauslegung aus der Signaltheorie zu tun hat, bei der Sie Tiefpässe, Bandpässe, Bandsperren und Hochpässe auslegen. Für DC-Versorgungsleitungen versuchen wir natürlich, sämtliche hochfrequenten Anteile herauszufiltern, für AC-Leitungen alles oberhalb der 50Hz- Netzfrequenz.



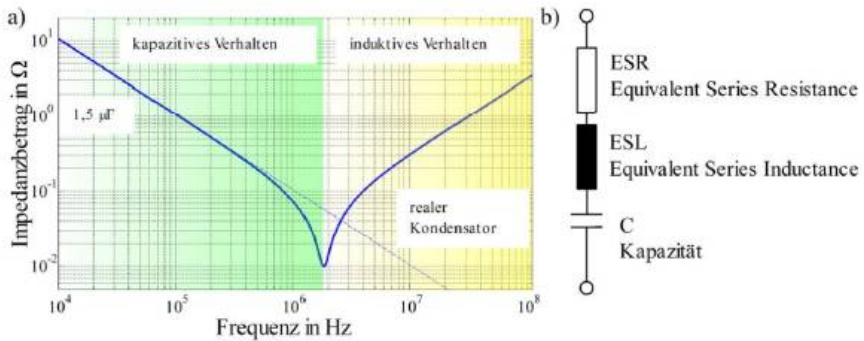
Bezogen auf unsere Störgrößen würden wir in der Signaltheorie von Tiefpassfiltern mit einer Grenzfrequenz von 0Hz sprechen.

Noch ein Unterschied zur Filterauslegung aus Vorlesungen wie Elektrotechnik II oder Signale und Systeme ist die Lastimpedanz. Während in Grundlagenvorlesungen stets von unbelasteten Filterschaltungen ausgegangen wird, setzen Ingenieure in der Hochfrequenztechnik sowohl an den Ein- als auch an den Ausgang 50-Ohm-Widerstände. Bei der Filterauswahl für Versorgungsleitungen ist im allgemeinen Fall sowohl die Ein- als auch die Ausgangsimpedanz (Belastung rechts und links des Filters) nicht bekannt. Auf der Geräteseite ist die Impedanz stets abhängig vom Betriebspunkt/Arbeitspunkt der Anwendung. Auf der Netzseite ist die Impedanz durch die Netznachbildungen während der EMV-Messung definiert. Im Betrieb hängt die Impedanz der Netzseite von der Leitungslänge zum (Batterie-)Speicher bzw. der Netzversorgung ab.

Damit sind im allgemeinen Fall die Impedanzverhältnisse auf beiden Seiten des Filters nicht bekannt, womit entweder eine Worst-Case-Analyse durchgeführt werden muss.

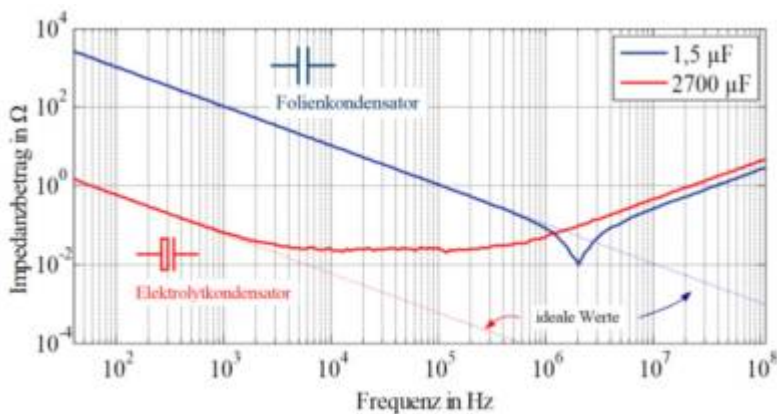
5.2 Reales Verhalten von Filterbauteilen

Reale Bauelemente verhalten sich nur in einem sehr engen Betriebsbereich wie im Schaltbild angenommen. Im überwiegenden Einsatz dominieren sogar parasitäre Elemente und bestimmen das Verhalten der Bauelemente. Das geht so weit, dass Kondensatoren oberhalb der Resonanzfrequenz induktives Verhalten zeigen, anstatt wie eine Kapazität zu wirken.



a) Impedanzverlauf eines 1,5 μF Folienkondensators, b) HF-Ersatzschaltbild.

Die Ursache dieser parasitären Eigenschaften ist im geometrischen Aufbau der Elemente zu suchen. Über Anschlussdrähte (Pins) werden bedrahtete Bauelemente auf der Platine angeschlossen. Da jedes noch so kleine Drahtstück sowohl einen realen Widerstand als auch eine Induktivität darstellt, muss das Schaltbild des Kondensators um diese zwei Eigenschaften erweitert werden. Im Datenblatt sind die Werte unter ESR (Equivalent Series Resistance) und ESL (Equivalent Series Inductance) zu finden, welche jeweils die parasitären Eigenschaften des Gesamtaufbaus berücksichtigen. Mit dem realen Ersatzschaltbild lässt sich das Verhalten der Kondensatoren über der Frequenz erklären. Generell gilt: Je größer (Bauform) der Kondensator, desto ausgeprägter sind die parasitären Eigenschaften. Ein 4700 μF (60V) Elektrolyt-Kondensator kann für Frequenzen im MHz-Bereich nichts ausrichten, da er längst induktives Verhalten zeigt.



Auch wenn es der Kollege in der Produktion oder Logistik nicht so gerne hört, ist es besser, Kondensatoren unterschiedlicher Bauform einzusetzen, um möglichst für ein breites Frequenzband eine niederimpedante Impedanz darzustellen (wir wollen die hochfrequenten Störgrößen ja gegen Masse kurzschließen). Induktivitäten sollen eine möglichst hohe Impedanz darstellen, um die hochfrequenten Störgrößen am Fließen zu hindern. Das bedeutet, wir möchten typischerweise bereits bei tiefen Frequenzen einen hohen Impedanzwert erzielen. Leider zeigen reale Induktivitäten nicht den idealisierten Verlauf der Impedanz mit $Z = \omega \cdot L$.

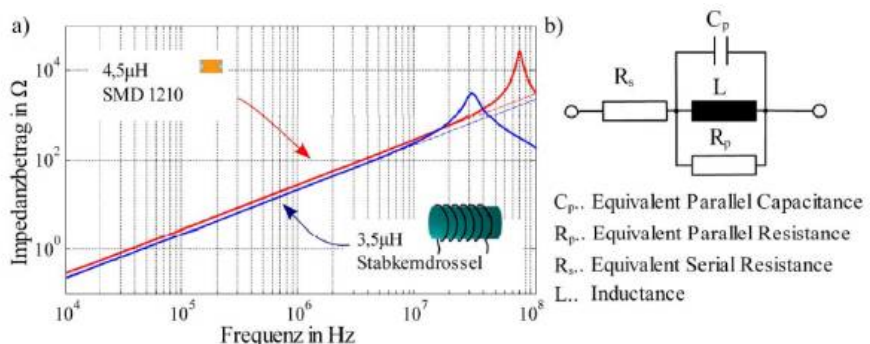
Induktivitäten in Form von Spulen kommen immer dann zum Einsatz, wenn die Ausbreitung der Störströme auf den Leitungen behindert werden soll. Die Auslegung richtet sich nach der zulässigen Stromtragfähigkeit der Spule sowie der maximalen Magnetisierung des Kerns, ohne dass er dabei in Sättigung gerät. Das HF-Ersatzschaltbild einer Spule wird aus deren Gleichstromwiderstand R_s , parasitären Kapazitäten C_p über den einzelnen Windungen sowie der eigentlichen Induktivität L gebildet, wie in Bild 3.2 b) dargestellt. Die verteilten Kapazitäten zwischen den Windungen werden zu einer konzentrierten Kapazität parallel zur Induktivität zusammengefasst. Damit bildet sich mit der Induktivität der Spule und der Ersatzkapazität ein Parallelschwingkreis aus. Die Güte der Parallelresonanz wird begrenzt über dem parallel zur Spule eingesetzten Widerstand R_p . Je nach Bauform der Spule treten im Frequenzbereich von 100–1000 MHz weitere Resonanzpunkte auf, die sich aus der Teilkapazität und Teilinduktivität der Wicklung ergeben.



Messen lassen sich die parasitären Eigenschaften einzelner Bauelemente elegant mit Hilfe eines Impedanzanalysators. Der Analysator macht vom Prinzip das, was Sie vermutlich auch machen würden, um die komplexe Impedanz einer Schaltung zu ermitteln. Er legt eine bekannte Spannung (bekannt nach Amplitude und Frequenz) an und misst den sich einstellenden Strom. Teilt man die Beträge durcheinander, ergibt sich die Impedanz für diese Frequenz. Wird zusätzlich noch die Phasenverschiebung des Stromes zur angelegten Spannung gemessen, kann die komplexe Impedanz sofort in eulerscher Form nach Betrag und Phase notiert werden. Dieses Vorgehen wiederholen Sie jetzt für die von Ihnen gewünschte Anzahl an Frequenzpunkten, fertig sind die gezeigten Impedanzverläufe. Falls Sie

sich für die automatisierte Variante mit dem Impedanzanalysator entscheiden, sollten Sie für die Anschaffung 20–30 T€ einplanen. Die Abbildung zeigt unseren Impedanzanalysator im Labor für Leistungselektronik.

Impedanzanalysatoren arbeiten in einem Frequenzbereich bis ca. 100 MHz. Natürlich ist es möglich, auch weit über diesen Frequenzbereich Bauteile und Schaltungen zu analysieren. Dazu werden meist Netzwerkanalysatoren herangezogen, welche in der Lage sind, bis in den hohen GHz- oder sogar THz-Bereich zu messen. Zu beachten ist jedoch, dass bereits ab dem zweistelligen MHz-Bereich der Messaufbau das Ergebnis maßgeblich beeinflusst.

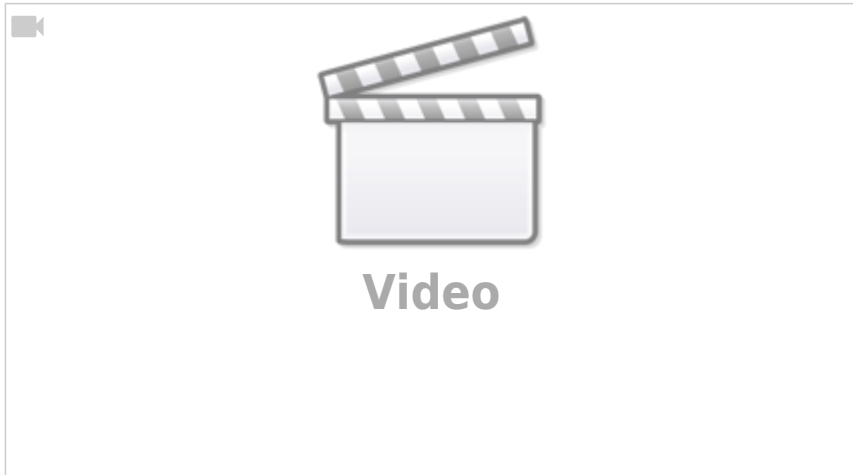


a) Impedanz verschiedener Spulenformen, b) HF-Ersatzschaltbild.

Die Abbildung zeigt den typischen Impedanzverlauf einer Stabkerninduktivität und einer Spule in SMD-Ausführung. Als Vergleich sind die idealen Impedanzwerte der Spulen mit gestrichelten Kurven dargestellt. Die Stabkerninduktivität besitzt eine Parallelresonanz bei 30 MHz, wohingegen die Resonanzfrequenz der SMD-Spule erst bei 80 MHz auftritt. Je kleiner die Bauform, desto eher nähern sich die Eigenschaften der Spule idealen Werten im Frequenzbereich bis 110 MHz an. Allerdings sinkt damit die verfügbare Stromtragfähigkeit der Elemente. Im Beispiel aus Bild 3.2 beträgt die maximale Strombelastung der SMD-Spule 220 mA, wohingegen die Stabkerninduktivität Ströme von über 50 A tragen kann, bei einem ähnlichen Induktivitätswert beider Spulen.)

Ich habe vorhin bei der Einleitung zu den Filterschaltungen von **drei unterschiedlichen** Bauteilen gesprochen, bisher aber nur zwei vorgestellt. Es fehlt noch die sogenannte Gleichtaktinduktivität oder engl. Common-Mode Choke. Der Name verrät es bereits, es geht um ein Bauelement, das dazu verwendet wird, **Gleichtaktstörungen** zu unterdrücken. Dazu müssen wir allerdings zuerst klären, was Gleichtaktstörgrößen überhaupt sind.

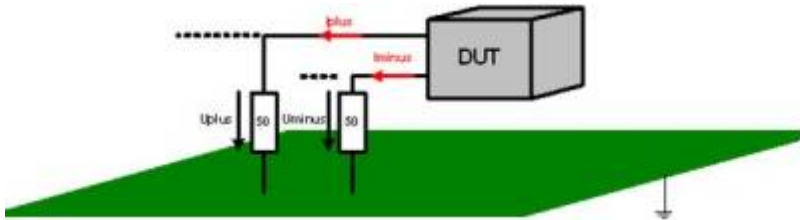
5.3 Gleich- Gegentaktzerlegung



Wir haben uns vorhin die Messung leitungsgebundener Emissionen mit Hilfe von Netznachbildungen angeschaut. Genauer gesagt haben wir überlegt, wie wir die auf den Leitungen fließenden Störströme auskoppeln und messbar machen. Die Messung wird bei einphasigen Systemen typischerweise für beide Anschlussleitungen durchgeführt, also L1 und N jeweils gegen die Referenzmasse (Erde). Bei genauerer Betrachtung stellen wir fest, dass es sich ja eigentlich um ein Dreileitersystem handelt. Bisher haben wir uns in keiner Weise über die Differenzspannung zwischen den Punkten L - N Gedanken gemacht. Wozu auch, sie spielt normativ keine Rolle! Alle Grenzwerte beziehen sich auf Spannungen gegen die Referenzmasse. Das macht auch Sinn, da sich die Störungen ausgehend von jedem Potenzialpunkt ausbreiten können und falls das Radio rauscht, ist es uns ja egal, von welchem Potenzialpunkt die Störungen ausgehen. Sie sind nun mal da.

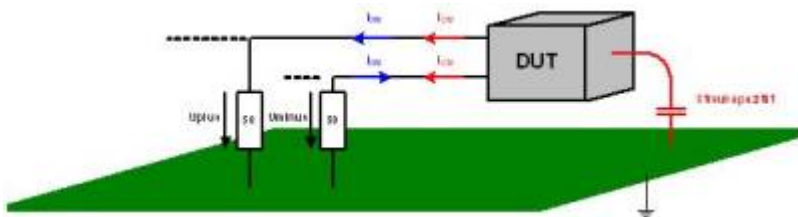
Bei der Bekämpfung der Störung spielt es allerdings eine entscheidende Rolle, um welche Störungen es sich handelt.

In einem idealen System, in dem sich lediglich Störungen entsprechend dem Nutzstrom ausbreiten, fließen keine Ströme über das Gehäuse zur Referenzmasse ab.



Die Störströme schließen sich differenziell über die Netzschaltungen. Für diesen Fall gilt weiterhin, dass I_{plus} um 180° phasenverschoben zu I_{minus} ist. Es existiert somit lediglich eine **Gegentaktstörung**.

In der realen Messumgebung werden neben den differenziellen Störgrößen Gleichtaktstörungen auftreten. Gleichtaktstörungen entstehen aufgrund der unsymmetrischen Kopplung der Versorgungsleitungen zur Referenzmasse bzw. der kapazitiven Kopplung des DUT zur unterliegenden Massefläche.



Das Ersatzschaltbild muss somit um den Anteil der Gleichtaktstörungen erweitert werden.

Die gemessenen Störspannungen an den Netzschaltungen ergeben sich zu

$$U_{plus} = 50\Omega \cdot I_{plus}$$

$$U_{minus} = 50\Omega \cdot I_{minus}$$

Die Spannungen an den Netzschaltungen setzen sich nun zusammen aus:

$$U_{plus} = 50\Omega \cdot (I_{DM} + I_{CM})$$

$$U_{minus} = 50\Omega \cdot (I_{CM} - I_{DM})$$

Gemessen wird somit jeweils eine Kombination aus Gleich- und Gegenaktstörungen, die sich aus einer Addition bzw. Subtraktion der einzelnen Teilströme ergeben.

Die nodalen Spannungen (gegen Masse bezogen) lassen sich darstellen als:

$$U_{plus} = U_{CM} + U_{DM}$$

$$|U_{minus} = U_{CM} - U_{DM}$$

Daraus ergeben sich die gesuchten modalen Störspannungen zu:

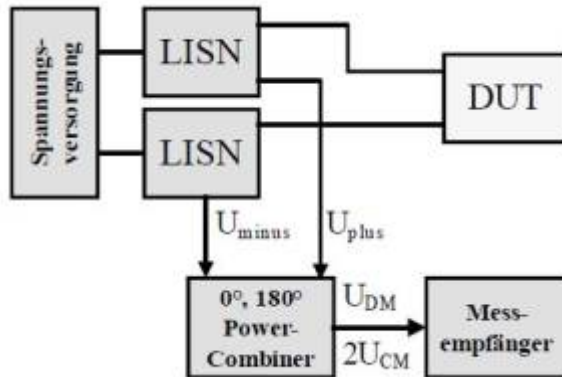
$$U_{CM} = \frac{1}{2} \cdot (U_{plus} + U_{minus})$$

$$U_{DM} = \frac{1}{2} \cdot (U_{plus} - U_{minus})$$

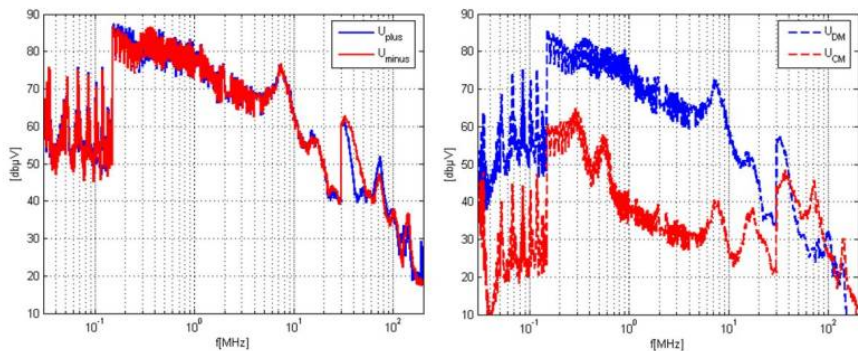
Messempfänger und Spektrumanalysatoren sind lediglich in der Lage, den Betrag der gemessenen Störspannung zu erfassen. Die Phaseninformation geht dabei verloren. Zur Messung der differenziellen Störspannungen muss allerdings eine phasenrichtige Addition bzw. Subtraktion der nodalen Störspannungen durchgeführt werden, die vor dem Messempfänger mit einer zusätzlichen analogen Schaltung erfolgt. Dafür werden Power-Splitter verwendet, die zwischen der Netznachbildung und dem Messempfänger eingesetzt werden. Die Power-Splitter bestehen aus HF-Transformatoren, die entsprechend ihrem Wicklungssinn die anliegenden Spannungen addieren bzw. subtrahieren.



Das Bild zeigt zwei Splitter von [Mini-Circuits](#). Der **ZFSC-2-4** ist ein 0° Power-Combiner, der die Spannungen an den Eingängen 1 und 2 addiert. Der **ZFSCJ-2-1** addiert zwar auch, allerdings mit 180° Phasenverschiebung, welches dann einer Subtraktion entspricht.



Als Messbeispiel werden die Power-Splitter für die Messung des Störspektrums eines klassischen Tiefsetzstellers mit Halbbrücke eingesetzt. Im ersten Schritt werden die Störungen in herkömmlicher Weise mittels nodaler Störspannungen auf den Versorgungsleitungen U_{plus} und U_{minus} bewertet. Danach erfolgt eine Aufteilung der Störungen in modale Störspannungen hinsichtlich Gleich- und Gegentaktanteile.



Wie erwartet, ist das Ergebnis der Störspannung an beiden Netznachbildungen für U_{plus} und U_{minus} nahezu identisch. Eine Aufteilung in Gleich- und Gegentaktanteile zeigt allerdings einen dominierenden Gegentaktanteil der Störungen bis ca. 50 MHz. Danach überwiegt der Gleichtaktanteil entsprechend der roten Kurve (rechts).

Die Messung zeigt das typische Bild der Gleich- Gegentaktzerlegung. Je größer die geometrischen Abmessungen, Kühlflächen etc. bzw. die Kopplung zur Referenzmasse, desto früher im Frequenzbereich dominieren die

Gleichtaktstörungen. Für geometrisch kleine Prüflinge ohne leitfähige Gehäuseanordnung (einfacher PCB-Aufbau) zeigt der Laborversuch einen überwiegenden Gleichtaktanteil lediglich ab ca. 80–100 MHz. Mit Hilfe der dominierenden Störgröße ist es nun möglich, das Filternetzwerk entsprechend der auftretenden Störgröße anzupassen.

5.4 Zusammenhang zwischen nodalen und modalen Größen

Da bei der EMV-Messung die Grenzwerte für die nodalen Größen (U_p , U_n) definiert sind, ist vor allem der Zusammenhang zwischen den modalen Größen U_{cm} und U_{dm} sowie den nodalen Größen U_p und U_n von Interesse. Ausgehend von der herkömmlichen, nodalen Messung, können drei unterschiedliche Fälle definiert werden.

Entscheidend dabei ist der Faktor

$$I_{cm}/2 \pm I_{dm}$$

wobei zu beachten ist, dass es sich jeweils um komplexe Größen handelt!

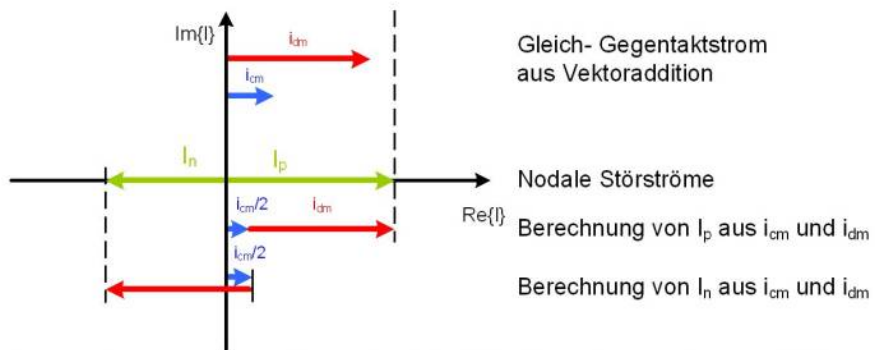
Da bei der Messung jeweils nur Beträge der Störspannungen vorliegen, müssen die einzelnen Spezialfälle betrachtet werden:

I):	$ I_{cm} \gg I_{dm} $	$U_n \approx 50 \times I_{cm}$
	Nodale Größen verhalten sich dabei entsprechend:	$U_p \approx 50 \times I_{cm}$
II):	$ I_{dm} \gg I_{cm} $	$U_n \approx 50 \times I_{dm}$
	Nodale Größen verhalten sich dabei entsprechend:	$U_p \approx 50 \times I_{dm}$

Fall I und II beschreiben die Fälle, in denen jeweils ein Störmodus überwiegt und als dominante Störung auftritt. Die nodalen Störgrößen sind dabei in ihrer Amplitude fast identisch. Ähnliche Störspannungen sind somit ein direktes Zeichen für einen dominanten Störmodus.

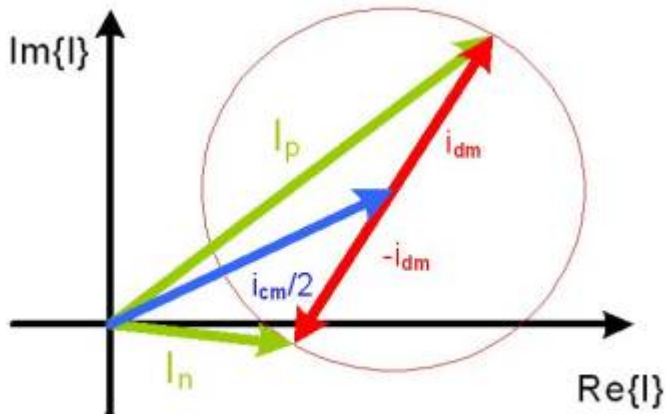
III): $|I_{dm}| \approx |I_{cm}|$

Besitzt die Gleichtaktstörung die Größenordnung der Gegentaktstörung, ist aufgrund der fehlenden Phaseninformation eine eindeutige Aussage über die nodalen Größen nicht mehr möglich.



Zusammenhang zwischen nodalen und modalen Störströmen für rein **reale** Lastverhältnisse.

Die Abbildung zeigt den Zusammenhang zwischen nodalen und modalen Störströmen unter der Annahme rein **realer** Lastimpedanzen. Dabei ist es möglich, zwischen beiden Systemen zu wechseln, ohne dass Information über die Amplitude einzelner Teilströme verloren geht. Somit kann aus den modalen Strömen i_{dm} und i_{cm} auf die Leiterströme I_p und I_n geschlossen werden. Leider tritt dieser Fall in der Realität nur selten bzw. vermutlich nie auf.

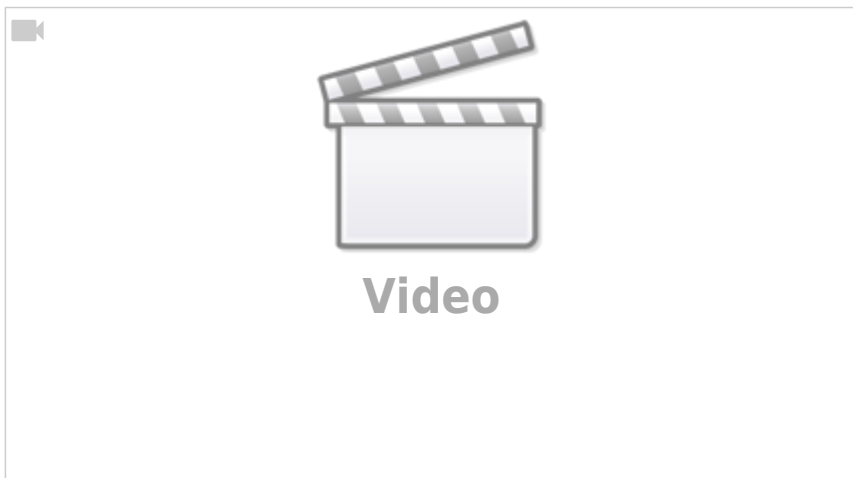


Die Gegentaktstörströme verhalten sich meist wie induktive Ströme. Die Gleichtaktstörströme sind durch die Verknüpfung der Leitungsinduktivität und der Streukapazität geprägt. Die Abbildung zeigt eine beliebig angenommene Phasenlage der modalen Störströme i_{cm} und i_{dm} und die daraus möglichen resultierenden nodalen Ströme I_p und I_n . Aufgrund der unbekannten Phasenlage

der modalen Störströme zueinander bewegen sich die nodalen Ströme entlang des durch ihm aufgespannten Kreisumfangs. Damit ist eine Aussage über die Amplitude der Störgrößen nicht mehr eindeutig möglich.

6.0 Noch einmal Filterschaltungen

Da wir jetzt wissen welche Störgrößen in welchem Frequenzbereich dominieren, sind wir in der Lage Filterelemente und Schaltungen zur Reduktion der dominanten Störgröße einzusetzen.



6.1 Gleichtaktdrossel

Neben den klassischen Kondensatoren und Spulen kennt die EMV noch ein weiteres Filterelement, das häufig eingesetzt wird. Gleichtaktdrosseln oder engl. (Common-Mode Choke CMC) werden hauptsächlich dazu verwendet, der Name



lässt es vermuten, Gleichtaktstörgrößen zu unterdrücken. Es ist allerdings auch möglich, kombinierte Drosseln zu erhalten, welche sowohl Gleich- als auch Gegentaktstörgrößen unterdrücken können.

Das Funktionsprinzip ist simpel. Zwei magnetisch gekoppelte Spulen teilen sich einen gemeinsamen Kern, meist einen Ringkern. Wir erinnern uns, **Gegentaktströme** (oder funktionale Ströme) fließen in entgegengesetzter Richtung. Bei entsprechendem Wicklungssinn der beiden Spulen löscht sich der magnetische Fluss im Kern aus. Die induktive Wirkung beider Wicklungen wird aufgehoben. Kein magnetischer Fluss → keine Induktivität!
Bild: WE-Online.de

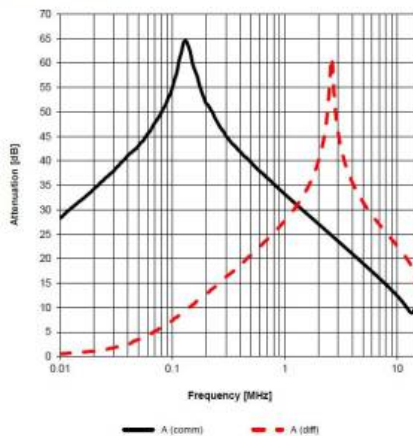
Von Vorteil ist, dass sich der magnetische Kern durch die Auslöschung der Magnetfelder nicht in Sättigung bringen lässt. Ein Problem weniger bei der



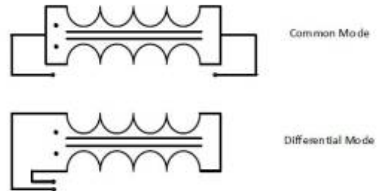
Spulenauslegung

Der magnetische Fluss, welcher durch die **Gleichtaktströme** generiert wird, addiert sich entsprechend dem Wicklungssinn. Die ideale Gleichtaktdrossel wirkt sich somit nur auf Gleichtaktgrößen und nicht auf Gegentaktgrößen aus. Würth Elektronik bietet eine große Auswahl an Gleichtaktdrosseln. Würth
Ein Blick in das Datenblatt der oben gezeigten Drossel zeigt den Impedanzverlauf der Gleichtaktdrossel über der Frequenz.

Typical Insertion Loss:



Test Setup:



Das Datenblatt zeigt zwei Messkurven, getrennt nach Gleichtaktimpedanz (comm) und Gegentaktimpedanz (diff). Gemessen wird der Impedanzverlauf typischerweise mit einem Impedanzanalysator mithilfe zweier Messschaltungen, dargestellt im Bild rechts.

Bei der Messung der **Gegentaktimpedanz** müssen die Ströme den identischen Weg nehmen, analog zu den späteren Störströmen. Sie bilden somit den Nutzstrompfad nach, bei dem die auf der oberen Seite einfließenden Ströme auf der unteren Seite zurückkehren. Gleichtaktströme fließen wie gesehen auf beiden Leitungen in die gleiche Richtung. Damit ist es naheliegend, zur Messung der **Gleichtaktimpedanz** beide Enden des Bauteils zu verbinden und einen parallel fließenden Stromfluss zu ermöglichen.

Das Ergebnis beider Impedanzmessungen ist der Betrag der Impedanz der Gleichtaktdrossel. Auf die Phaseninformation wird, wie in der EMV üblich, kein Wert gelegt – wozu auch ... Bei genauer Betrachtung ist für beide Impedanzverläufe das induktive Verhalten, aufgrund der über einen weiten Frequenzbereich linear steigenden Impedanz ($Z = R + j\omega L$), ersichtlich. Oberhalb der Resonanzfrequenz dominiert jeweils das kapazitive Verhalten.

6.2 Filterauswahl im Dreileitersystem

Bei der Filterauswahl im vorangegangenen Kapitel sind wir stillschweigend davon ausgegangen, dass unser Gerät nur zwei Anschlussleitungen besitzt bzw. wir uns bei der Filterauswahl stets auf eine Referenzmasse beziehen können.

Bei der Auswahl von geeigneten Netzfiltern können wir leider nie von einem

einfachen Zweileitersystem ausgehen. Ihr Gerät mag vielleicht ohne einen PE-Anschluss auskommen, da keine leitfähigen Gehäuseteile vorhanden sind, dennoch lassen sich alle auftretenden Störgrößen auf diese Referenz beziehen. Auch im Fahrzeug handelt es sich meistens um ein echtes Dreileitersystem aus

- Kl. 30 (Versorgungsleitung)
- Kl. 31 (Rückstrompfad)
- Chassis (Referenzmasse, eigentlicher Rückstrompfad)

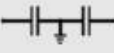
Jetzt könnte man natürlich davon ausgehen, dass Kl. 31 sowieso auf das Chassis gezogen wird, womit sich das Problem wieder auf ein Zweileitersystem reduzieren lässt. Das stimmt, hilft uns aber hier nicht weiter.

Stellen Sie sich vor, das Steuergerät holt sich seine Kl. 31 über einen 1 m entfernten Verteiler. Der durch die Leitung eingebrachte DC-Widerstand, vermutlich kleiner 1 mOhm, spielt für die Funktion der Energieversorgung des Geräts eine untergeordnete Rolle. Die vorhandene Leitungsinduktivität von ca. 1 $\mu\text{H/m}$ durchkreuzt allerdings in vielen Fällen unsere EMV-Pläne. Zur Veranschaulichung: Wir erhalten über die Beziehung $Z = \omega \cdot L$ bereits bei 1 MHz eine Impedanz von 2π , also mehr als 6 Ohm, nur durch diese ein Meter lange Leitung. Dadurch heben wir das für die Filterauslegung so wichtige Referenzpotenzial an, womit alle unsere Pläne, die Störgrößen auf die Referenzmasse abzuleiten, zum Scheitern verurteilt sind. Spätestens jetzt sind wir auch hier wieder in einem echten Dreileitersystem angekommen. Nachfolgende Tabelle erweitert die zuvor erstellte Tabelle zur Filterauswahl für Dreileitersysteme:

Maßnahme / Ergebnis Grundstellung	Filtertyp	Gleichst.- Dämpfung	Gegentakt- Dämpfung	Dämpfung U_{ges}	Dämpfung U_{aus}
$U_{\text{dies}} > U_{\text{om}}$		—	Hohe Dämpfung für Zlast > Zfilter	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{om}} > U_{\text{dies}}$		—	Hohe Dämpfung für Zlast > Zfilter	—	—
$U_{\text{dies}} > U_{\text{om}}$		Hohe Dämpfung für Zlast > Zfilter	Hohe Dämpfung für Zlast > Zfilter	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{om}} > U_{\text{dies}}$		Hohe Dämpfung für Zlast > Zfilter	Hohe Dämpfung für Zlast > Zfilter	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{dies}} > U_{\text{om}}$		Hohe Dämpfung für Zlast < Zfilter	Hohe Dämpfung für Zlast < Zfilter	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{om}} > U_{\text{dies}}$		Hohe Dämpfung für Zlast < Zfilter	Keine Dämpfung, Verstärkung der Störspannung	Dämpfung entsprechend Impedanz- verhältnisse	—
$U_{\text{dies}} > U_{\text{om}}$		Hohe Dämpfung für Zlast < Zfilter	Hohe Dämpfung für Zlast < Zfilter	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{om}} > U_{\text{dies}}$		Hohe Dämpfung für Zlast < Zfilter	Keine Dämpfung, Verstärkung der Störspannung	—	Dämpfung entsprechend Impedanz- verhältnisse
$U_{\text{dies}} > U_{\text{om}}$	 Ideale CMC	Hohe Dämpfung für Zlast < Zfilter	—	—	—
$U_{\text{om}} > U_{\text{dies}}$	 Ideale CMC	Hohe Dämpfung für Zlast < Zfilter	—	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{dies}} > U_{\text{om}}$	Allgemeiner unsymmetrischer Filter	Entsprechend Topologie und Elemente	Entsprechend Topologie und Elemente	Entsprechend Gegentakt- dämpfung	Entsprechend Gegentakt- dämpfung
$U_{\text{om}} > U_{\text{dies}}$	Allgemeiner unsymmetrischer Filter	Entsprechend Topologie und Elemente	Keine Dämpfung, evtl. Verstärkung entsprechend Modenverteilung	Abhängig von Symmetrie des DUT	Abhängig von Symmetrie des DUT

Wie bereits gesehen, kann die Störspannung an der Netznachbildung nur reduziert werden, wenn es gelingt, die dominante Störgröße zu unterdrücken. Als

Filterelemente kommen für herkömmliche passive Filter lediglich vier Elemente bzw. Grundschaltungen zum Einsatz:

Filterschaltung	Einfluss auf die Filterdämpfung			Anwendung	
	DM	CM	CM→DM	Vorteil	Nachteil
	✓			- Einfach einsetzbar - Hohe Verfügbarkeit	Hohe Abhängigkeit von der Lastimpedanz
	✓	✓	✓	Hohe Verfügbarkeit	- Modenwandlung - Stromtragfähigkeit
	✓	✓		Gleich- und Gegentakt-dämpfung	Nur in Systemen mit vorhandener Masse einsetzbar
		✓		Auch ohne Massesystem anwendbar	- Dimension - Stromtragfähigkeit

\ Für die Grundelemente wird folgende Bezeichnung gewählt, in Reihenfolge obiger Tabelle:

X-Kondensator: Kondensator zwischen Kl. 30 – Kl. 31 bzw. L - N

Längsdrossel L: Drossel als Serielement in Kl. 30 oder Kl. 31

Y-Kondensator: Kondensatorpaar mit Massebezug im Mittelabgriff

CMC-Drossel: Gleichtaktdrossel, Common-Mode-Choke oder stromkompensierte Drossel

CM: Common-Mode, Gleichtaktstörung

DM: Differential-Mode, Gegentaktstörung

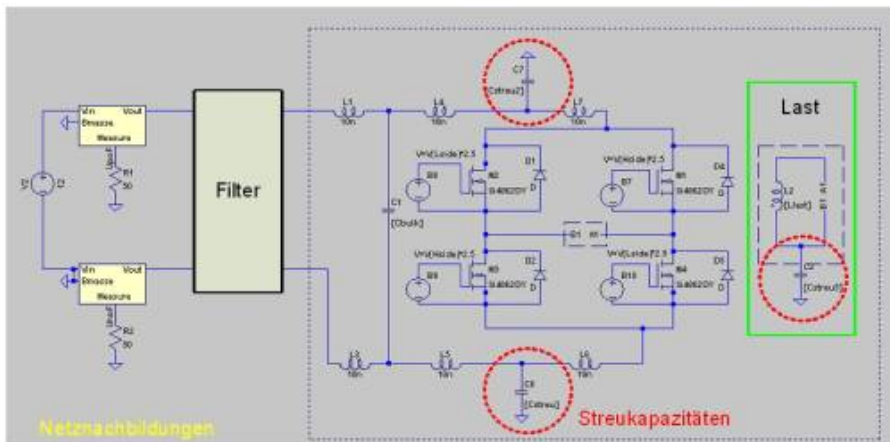
Die Tabelle zeigt den Zusammenhang zwischen den Grundelementen und dem Einfluss auf die jeweilige Störgröße. Y-Kondensatorpaare sind ideal geeignet, Gleichtaktgrößen gegen die Masse zu schließen sowie Gegentaktströme im Sinne eines X-Kondensators mit $C/2$ zu bedämpfen. Die Gleichtaktunterdrückung kann natürlich nur funktionieren, falls eine ausreichend niederohmige Masseverbindung vorhanden ist. Für die Seriendrossel gilt es, auf eine Besonderheit zu achten: Neben der Dämpfung von Gegentaktstörungen kann eine Modenwandlung auftreten. Eine durch das Gerät erzeugte Gleichtaktstörung kann somit nach der Drossel als Gegentaktssignal auftreten.

CM → DM: Transformation von Common- nach Differential-Mode

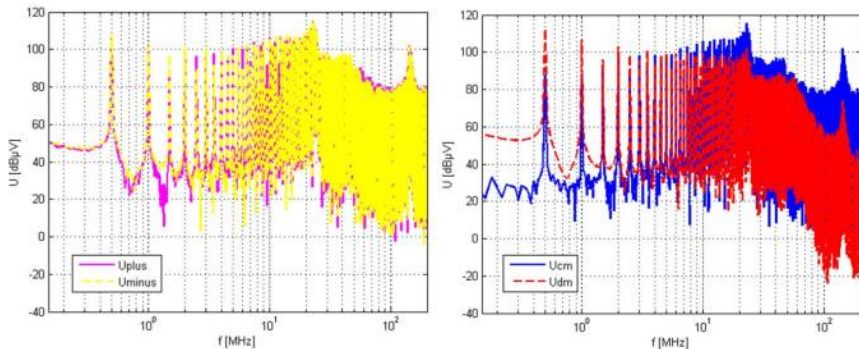
6.3 Beispiel zur Filterwirkung

Als Beispiel wird die Simulation einer H-Brücke (Endstufe) in LTSpice herangezogen und hinsichtlich der Störspannungen an den Netznachbildungen ausgewertet. Zur Nachbildung auftretender Gleichtaktströme werden Streukapazitäten zur Referenzmasse eingefügt. Im Folgenden wird der Einfluss der einzelnen Filterelemente auf die Störgrößen untersucht.

Streukapazitäten bilden sich generell von allen leitfähigen Strukturen zur Referenzmasse aus. Ganz allgemein gilt natürlich: Je größer (geometrisch größer), desto größer ist die sich ausbildende Kapazität. Deren Wert liegt typischerweise im fF- (Femto-) bis pF-Bereich. Kann für große Kühlkörperstrukturen auch bis zweistellige pF-Werte annehmen.



Entsprechend dem Vorgehen in der Messung wird im ersten Schritt der Grundstörpegel (Störgrößen ohne jegliche Filtermaßnahme, Baseline) der Schaltung ohne Filtermaßnahme betrachtet. Wir vergleichen zuerst die nodalen (gegen Referenzmasse bezogenen Größen) mit den differentiellen Gleich- und Gegentaktgrößen. Gleichtakt = Common-Mode, Gegentakt = Differential-Mode. Der dargestellte Frequenzbereich endet bei 110 MHz. Welcher Frequenzbereich von Interesse ist, hängt von der Anwendung und damit von den geltenden Normen ab. Die gewählten 110 MHz beziehen sich auf CISPR25 leitungsgebundene Emission (conducted emission) für Automotive-Anwendungen.



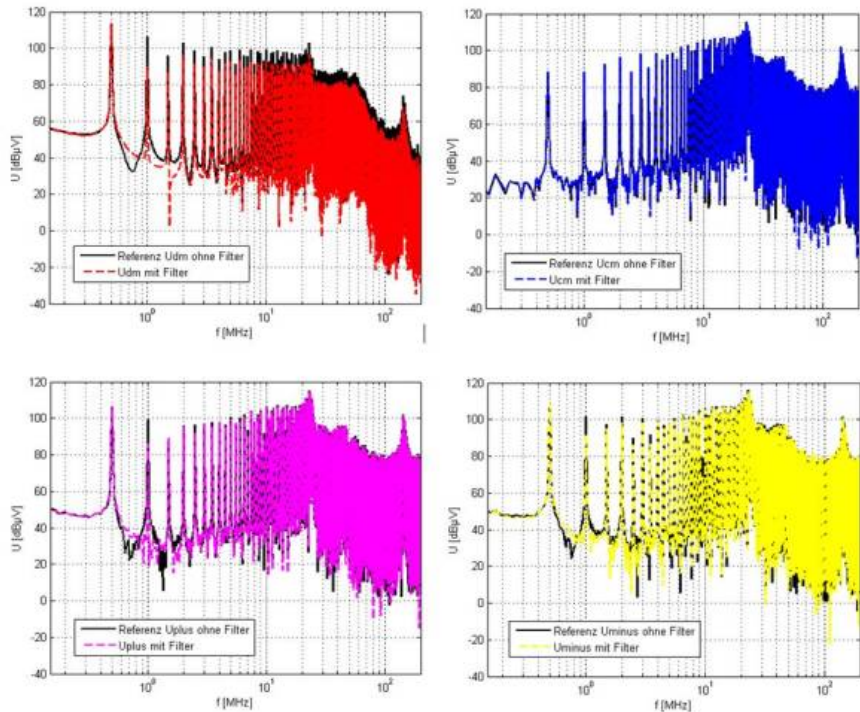
Links: Nodale Störspannungen Uplus und Uminus

Rechts: Modale Störspannungen Ucm und Udm

Das Störspektrum wird in der Simulation analog zur Messung an den Netznachbildungen ermittelt. Der Hauptunterschied zur Messung besteht in der Bewertung der Störgröße im Zeitbereich, welche dann mittels FFT (Fast Fourier Transformation) in den Frequenzbereich überführt wird. Während die nodale Auswertung keine Rückschlüsse auf die Art der Störungen zulässt, zeigt sich in der modalen Auswertung ein überwiegender Gleichtaktanteil bereits ab ca. 1 MHz. Das bedeutet, dass die Gleichtaktstörungen für die gewählte Konfiguration überwiegen und in erster Linie gedämpft werden müssen.

6.3.1 Einfluss eines X-Kondensators auf die Störgrößen

Als Filterelement wird ein einfacher X-Kondensator verwendet mit 1,5 μ F Kapazität, 20 nH Serieninduktivität und 20 m Ω Serienwiderstand. Die Ersatzelemente des Kondensators wurden mit Hilfe eines Impedanzanalysators im Frequenzbereich von 40 Hz – 110 MHz ermittelt.

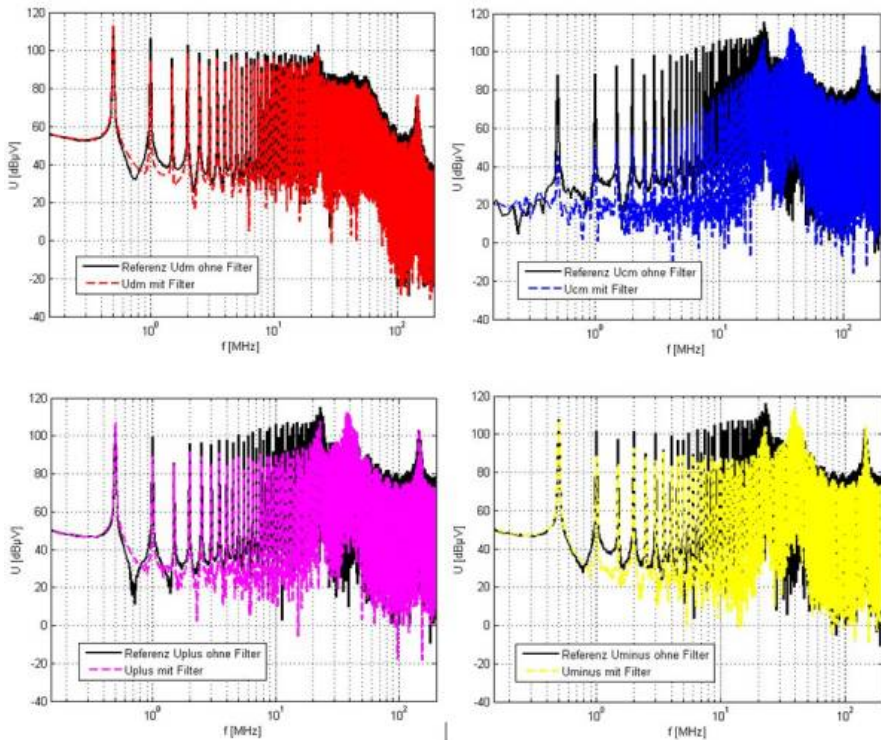


Oben: Links: Differential-Mode Rechts: Common-Mode
 Unten: Links: Uplus Rechts: Uminus

Entsprechend der Abbildung zeigt die X-Kondensatorschaltung keinen Einfluss auf die Gleichtaktstörspannung. Lediglich der Gegentakt-Mode wird gedämpft. Wie erwartet zeigen sich die nodalen Störspannungen von der Filtermaßnahme nicht beeindruckt, da die dominierende Gleichtaktstörung nicht gedämpft wird.

6.3.2 Einfluss eines Y-Kondensators auf die Störgrößen

Als Filterelement wird nun die Erweiterung der einfachen X-Struktur zur Y-Struktur untersucht, mit zwei gegen die Referenzmasse geschalteten Kondensatoren. Als Filterelement kommt der bereits im vorhergehenden Abschnitt vorgestellte Kondensator in zweifacher Ausführung zum Einsatz.



Oben: Links: Differential-Mode Rechts: Common-Mode
 Unten: Links: Uplus Rechts: Uminus

Die Y-Struktur zeigt wie erwartet eine Dämpfung sowohl für die Gleichtakt- als auch für die Gegentaktstörung. Die Gegentaktdämpfung fällt durch die Serienschaltung der beiden Kondensatoren in der Y-Struktur entsprechend geringer aus als bei der X-Anordnung. Aufgrund der dominierenden Gleichtaktstörung zeigt die erzielte Dämpfung des Gleichtaktstromes durch die Y-Struktur nun auch Wirkung auf die nodalen Störspannungen U_{plus} und U_{minus} . Die Wirkung der Y-Kondensatoren nimmt ab ca. 10 MHz stark ab. Dieses Verhalten ist auf die parasitären Eigenschaften (Serienwiderstand und Serieninduktivität) zurückzuführen. In der realen Schaltung nimmt die Wirkung ggf. noch stärker ab aufgrund der nicht idealen Masseanbindung unserer Y-Kondensatoren. Man beachte: Auf einer Leiterkarte entsprechen 1 cm ca. 10 nH Anschlussinduktivität (grobe Abschätzung), welche allerdings bei 10 MHz bereits mit 0,6 Ohm zu Buche schlägt. Bei 100 MHz würde das bereits 6 Ohm entsprechen und somit zur Ableitung von Störgrößen gänzlich ungeeignet sein. Daher gilt die allgemeine Regel:

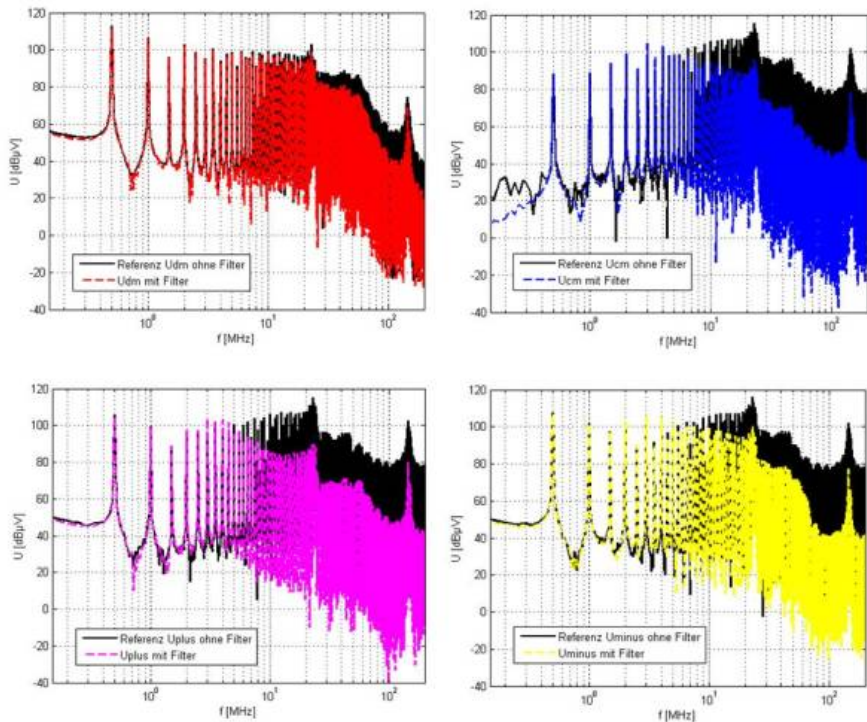
Masseanbindungen stets so niederohmig wie möglich ausführen:



- Breite Leiterbahnen oder besser, Masseflächen einsetzen
- Leiterbahnen zur Referenzmasse so kurz wie möglich
- Nutzstrompfade (Rückleiter) sind keine Masseverbindungen!

6.3.3 Einfluss einer Gleichtaktdrossel auf die Störgrößen

Die Gleichtaktdrossel oder stromkompensierte Drossel basiert auf dem Prinzip der Feldauslöschung bei gegenphasiger Stromanregung der beiden Wicklungen. Sie ist generell unsere stärkste Waffe gegen Gleichtaktstörgrößen. Falls wir keine Möglichkeit haben, die hochfrequenten Störströme zur Referenzmasse abzuleiten, ist sie sogar die einzige Möglichkeit, Gleichtaktgrößen zu dämpfen. Die ideale Gleichtaktdrossel stellt für differenzielle Ströme keine Induktivität dar, wobei bei Gleichtaktströmen durch die Feldaddition eine induktive Wirkung eintritt. Zur Untersuchung der Dämpfungswirkung der Gleichtaktdrossel werden zwei Drosseln mit je $2\ \mu\text{H}$ Induktivität als Filterelement eingefügt. Die Drosseln sind in der Simulation nicht ideal gekoppelt, sondern erhalten einen Kopplungsfaktor von 0,8, um die Wirkung einer realen Gleichtaktdrossel zu simulieren.



Oben: Links: Differential-Mode Rechts: Common-Mode
 Unten: Links: Uplus Rechts: Uminus

Die Gleichtaktdrossel zeigt wie erwartet eine hohe Gleichtaktdämpfung sowie eine aufgrund der nicht idealen Kopplung der beiden Spulenteile geringe Gegentaktdämpfung. Da die dominierende Störgröße U_{cm} deutlich reduziert wird, erhält man ebenfalls eine hohe Dämpfung der leitungsgebundenen Störspannungen U_{plus} und U_{minus} .

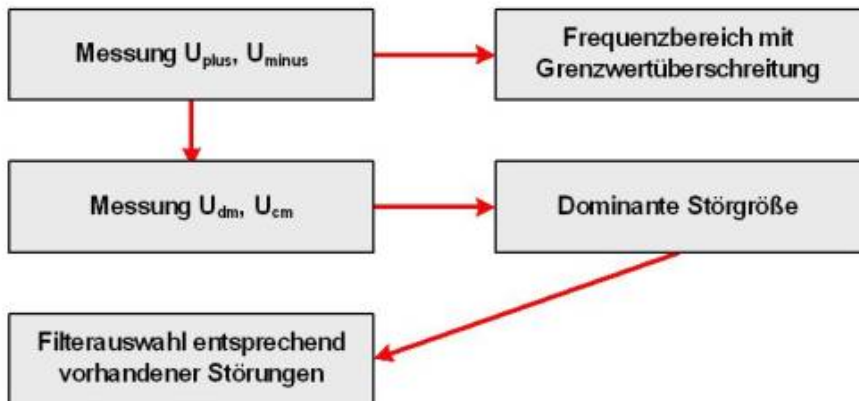
6.3.4 Fazit

In der Simulation konnte der Zusammenhang zwischen den einzelnen Filter-Grundtypen und der jeweiligen Modendämpfung entsprechend der vorgestellten Tabelle verifiziert werden. Betrachtet man neben den modalen Störspannungen die nodalen Größen, so zeigt sich, dass eine Verringerung der nodalen Störspannungen nur erreicht werden kann, falls die dominierende modale Störgröße gedämpft wird. Dies ist deutlich ersichtlich in der Simulation zum X-

Kondensator. Die X-Struktur wirkt sich lediglich auf die differentiellen Störströme aus, welche entsprechend gedämpft werden. Die nodalen Größen zeigen allerdings keine Reaktion auf die Maßnahme. Selbst eine Erhöhung der Kapazität oder eine bessere Anbindung (niederinduktiver) würde keine Verbesserung in den nodalen Störspannungen erzielen. Lediglich über eine Gleichtaktdämpfung, mittels Gleichtaktdrossel oder Y-Struktur, zeigt sich die gewünschte Dämpfung im nodalen System.

6.4 Ablauf zur Filterauslegung

Nachfolgende Abbildung zeigt die prinzipielle Vorgehensweise zur Identifikation der dominanten Störgröße. Die Messung erfolgt dabei im Grundzustand ohne angeschlossenes Filterelement. Vorhandene Speicherkondensatoren bzw. Bulk-Kondensatoren (oder auch Zwischenkreiskondensatoren genannt), die für die Funktion des DUT notwendig sind, werden nicht dem Filter angerechnet, sondern der Funktion der Anwendung zugeordnet.



Die Auswahl der Filterelemente ist in den meisten Fällen nicht frei wählbar, sondern obliegt den Randbedingungen wie:

- Beschränkte Anzahl an möglichen Elementen
- Bauraumbeschränkung (Abmessungen)
- Maximale Belastung (Temperatur, Klima, etc.)
- Anbindung an die Umgebung (Stanzgitter)

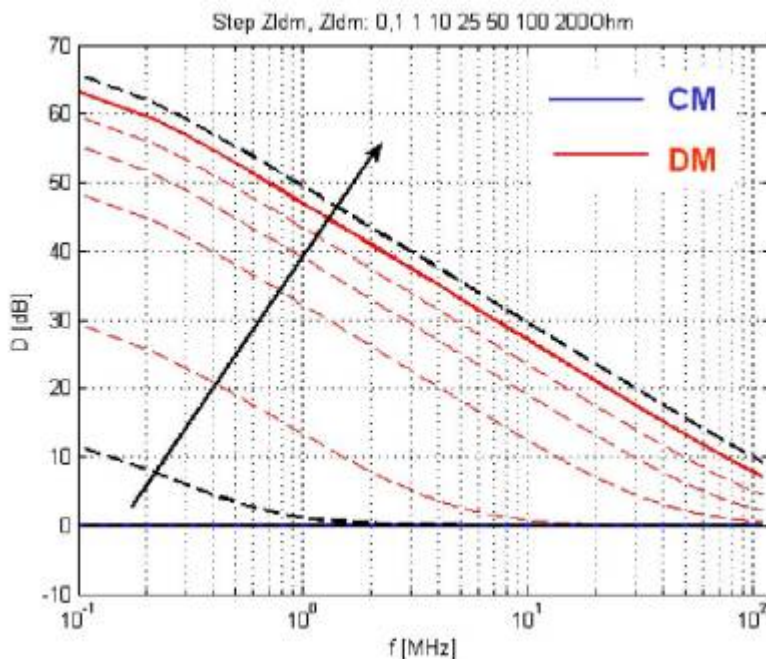
Neben den physikalischen / projektspezifischen Randbedingungen bei der

Auswahl der Elemente spielen die nur sehr schwer abschätzbaren Parameter wie:

- Kopplung der Bauelemente
- Kopplung des Stanzgitters / Leiterbahnen auf einzelne Elemente
- Modenwandlung aufgrund eines unsymmetrischen Aufbaus
- Eingangsimpedanz der Anwendung

eine erhebliche Rolle für die erreichbare Dämpfung des Filters.

6.5 Simulation möglicher Filterkomponenten / Kombinationen



In vielen Fällen stehen aufgrund der physikalischen Randbedingungen nur zwei bis drei Schaltungsvarianten zur Verfügung, zum Beispiel eine C-L-

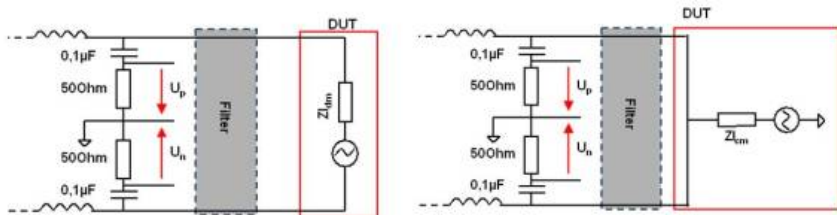
oder eine C-L-L-Kombination. Die Filterelemente ergeben sich aus der maximalen Stromtragfähigkeit der Spulen bzw. aus der benötigten Kapazität / dem Impedanzverlauf der Kondensatoren.

Die Auswahl der Filterelemente wird weiterhin durch die unbekannte Eingangsimpedanz der Anwendung erschwert, da die Dämpfung der Filterschaltung stets abhängig von der an den Filter angeschlossenen Impedanz

ist! Die Abbildung zeigt die Gleich- und Gegentaktdämpfung eines 1000 μF Kondensators als X-Kondensator in Abhängigkeit von der Eingangsimpedanz. Dabei wird von einer möglichen Lastimpedanz für den Gleich- als auch den Gegentaktfall von 0,1–200 Ω ausgegangen. Die schwarzen Kurven entsprechen der maximal und minimal angenommenen Lastimpedanz.

Wie erwartet kann mit dem X-Kondensator keine Gleichtaktdämpfung erzielt werden. Die Gegentaktdämpfung hingegen zeigt eine starke Abhängigkeit der Dämpfungswerte von der angenommenen Eingangsimpedanz. Sie variiert um bis zu 50 dB, je nach angeschlossener Lastimpedanz.

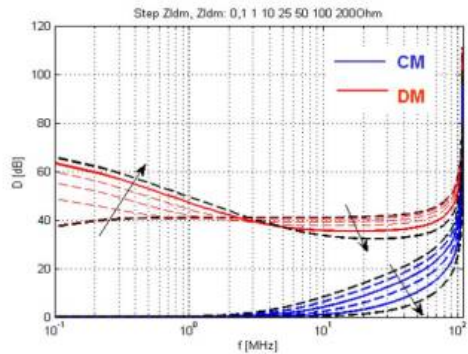
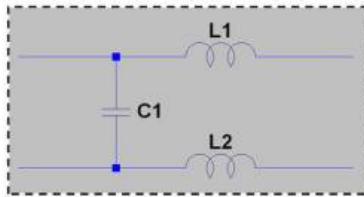
Sind die möglichen Kombinationen der Filterschaltung und Filterelemente bekannt, kann die Filterwirkung sowie die Abhängigkeit der Dämpfung von verschiedenen Lastfällen untersucht werden. Im ersten Schritt werden die Filterelemente mit Hilfe eines Impedanzanalysators über der Frequenz charakterisiert. Zur Berechnung der Gleich- und Gegentaktdämpfung können die Ersatzschaltbilder entsprechend nachfolgender Abbildung in einer Simulation herangezogen werden.



Für die modalen Dämpfungsparameter errechnet sich die Dämpfung analog der bekannten Definition der Filterdämpfung zu:

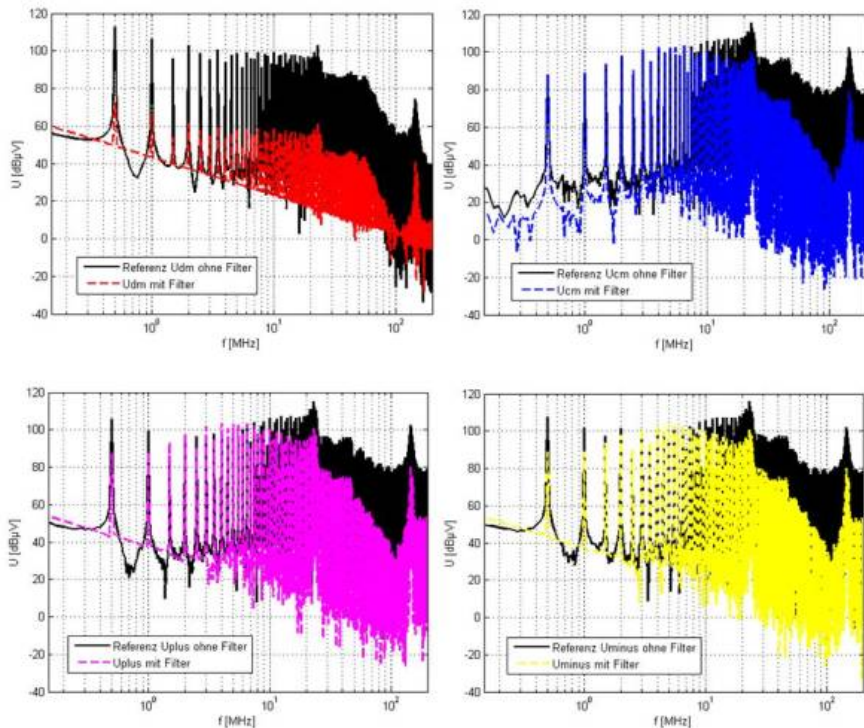
$$\text{Dämpfung} = \frac{\text{Störspannung ohne Filter}}{\text{Störspannung mit Filter}}$$

Als Beispiel wird eine Filterschaltung bestehend aus zwei Spulen und einem X-Kondensator betrachtet:



Der Filter ist in seinem Grundaufbau symmetrisch. Leichte Unsymmetrien können entstehen, falls bei der Anordnung unterschiedliche Zuleitungslängen für die Einzelelemente verwendet werden. Die Abbildung zeigt die berechnete Dämpfung der Filterstruktur, dargestellt für Gleich- und Gegentaktsignale, wobei Leitungsinduktivitäten und eventuell vorhandene Unsymmetrien vernachlässigt werden. Erneut zeigt ein Parametersweep die Dämpfung in Abhängigkeit der gewählten Lastimpedanz. Die Abhängigkeit der Einfügedämpfung von der Lastimpedanz ist deutlich geringer als für den reinen X-Kondensator.

Zur Verifikation der Filterdämpfung wird die Filterstruktur in die Leistungsbrücke aus der vorangegangenen Simulation integriert und der Einfluss auf die Störspannungen ausgewertet. Zu beachten ist, dass die Filterelemente lediglich bis 110 MHz verifiziert sind, durch die limitierte Frequenzauflösung des verwendeten Impedanzanalysators.



Oben: Links: Differential-Mode Rechts: Common-Mode
 Unten: Links: Uplus Rechts: Uminus

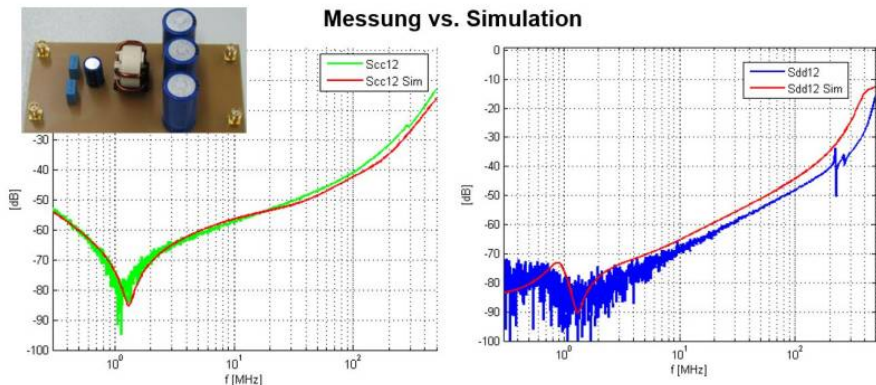
Die Simulation zeigt die Filterwirkung auf die auftretenden nodalen und modalen Störspannungen der Leistungsbrücke. Die breitbandige 40 dB Gegentaktdämpfung sowie die erst ab ca. 10 MHz einsetzende Gleichtaktdämpfung stimmen sehr gut mit dem Ergebnis der simulierten Filterdämpfung überein. Erneut zeigt sich, dass trotz der sehr hohen Gegentaktdämpfung die nodalen Störspannungsgrößen dem Verlauf der Gleichtaktdämpfung folgen, aufgrund des überwiegenden Gleichtaktanteils der Grundstörungen.

6.6 Filtersimulation vs. Messung

Filterschaltungen lassen sich für Anwendungen bis 100 MHz besonders gut in einer Simulation abbilden. Oberhalb dieser Frequenz fangen die Eigenschaften des PCBs sowie die direkte Kopplung der Bauelemente an zu dominieren.

Bezahlbare Impedanzanalysatoren haben typischerweise einen Messbereich bis ca. 100 MHz, weiterhin definiert die CISPR25 die maximal zu messende Frequenz bei 108 MHz. Es spricht also viel dafür, die Filtersimulation bis zu dieser Frequenz durchzuführen.

Im nachfolgenden Beispiel wird eine typische Filterschaltung auf einem Test-PCB aufgebaut und mit Hilfe von BNC-Buchsen mit einem 4-Port-Netzwerkanalysator vermessen. Mit Hilfe einer einfachen Matrixtransformation können die bereits beschriebenen Dämpfungswerte für Gleich- und Gegentakt berechnet werden. Die Simulation hingegen kennt nur die Impedanzwerte nach Betrag und Phase, welche wir mit einem Impedanzanalysator gemessen haben. Die daraus generierten Ersatzschaltbilder lassen sich direkt z. B. in LT-Spice einbinden. Die Filterschaltung besteht aus einer $C_y - C_x - CMC - C_{Bulke}$ (von den Versorgungsklemmen zum Gerät) Kombination, wie sie häufig zu finden ist.



Die Abbildung vergleicht, getrennt nach Gleichtaktdämpfung und Gegentaktdämpfung, jeweils die Messergebnisse (Netzwerkanalysator) mit der LT-Spice-Simulation. Die Ergebnisse sind mehr als ausreichend genau und können für weiterführende Simulationen herangezogen werden. Im nächsten Schritt würde es sich anbieten, die zuvor beschriebene Sensitivität hinsichtlich der Lastimpedanz zu bestimmen.

6.7 Das oder der perfekte Filter

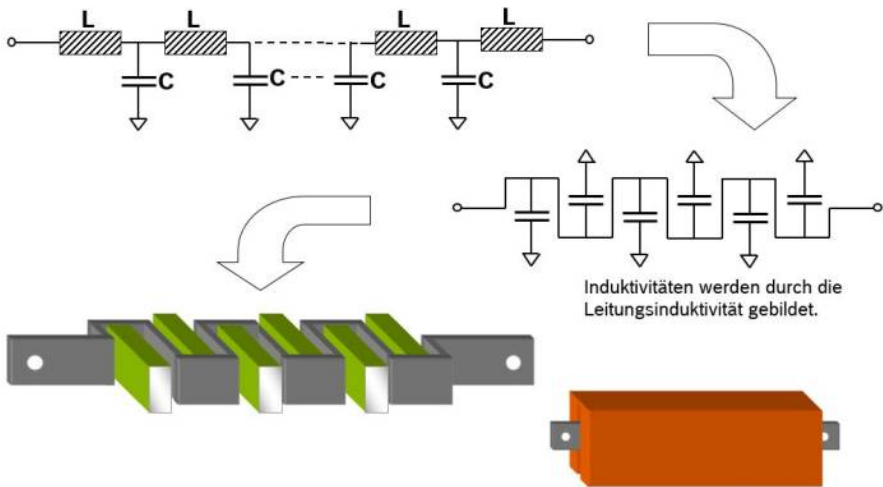
Vielleicht reden wir an dieser Stelle doch kurz über die Rechtschreibung und den Artikel. Mein schwäbisches Bauchgefühl sagt mir es muss „der Filter“ heißen. Eine Abfrage bei Duden ergibt es ist beides möglich. Fühlen sie sich also frei

Aber zurück zum Thema! Welche Eigenschaften würden wir einem perfekten Filter zuschreiben? Mir fallen folgende Punkte ein:

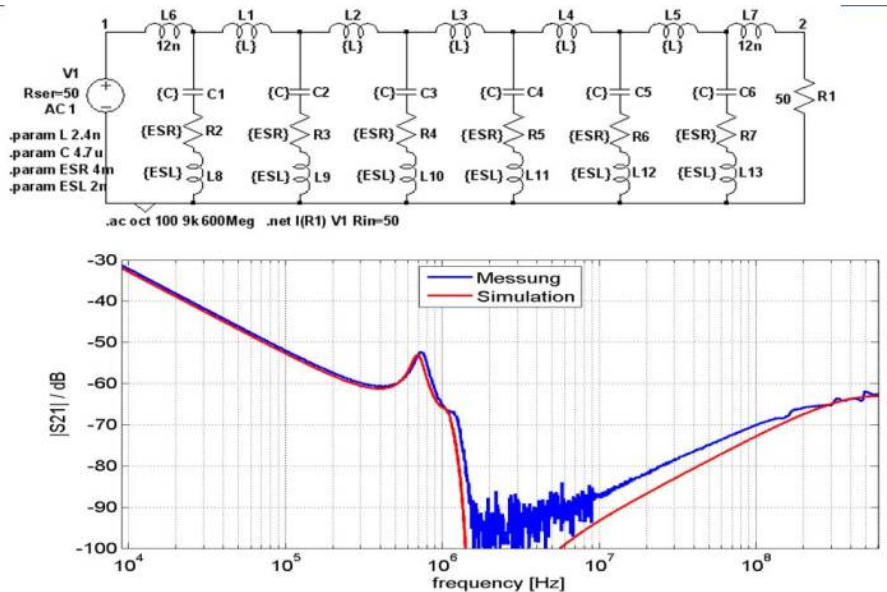
- Keinen Einfluss auf das Nutzsignal
- Sehr hohe Dämpfung aller hochfrequenten Anteile
- Möglichst hohe Stromtragfähigkeit

Die Eigenschaft „sehr hoch“ ist natürlich relativ. In der Praxis reichen uns Dämpfungswerte aus welche unser Störgrößen unterhalb die Grenzwerte drücken. Als perfekt würde ich einen Filter bezeichnen der Dämpfungswerte erreicht die mit einem Netzwerkanalysator nicht mehr messbar sind. Je nach Kalibrierung liegt die Rauschgrenze zwischen -80 und -100dB. Für -100dB bedeutet das, dass ein Eingangssignal um den Faktor 10^5 verringert wird oder anschaulich: Aus einem 1V Eingangspegel verbleiben noch 10 μ V. Überlegen sie sich an der stelle kurz welche Auflösung ihr Oszilloskop hat wenn die Amplitude auf höchster Auflösung eingestellt ist. Eine möglichst hohe Stromtragfähigkeit ist noch relativer. Die Herausforderung besteht darin Filter hoher Stromtragfähigkeit aufzubauen, da die Auswahl an verfügbaren Induktivitäten sich schlagartig reduziert bei Strömen >10A. Besonders im Bereiche > 100A wird die Luft sehr dünn und es können nur noch Ferrite eingesetzt werden. Damit keine Sättigung auftritt müssen die Ferritkerne zusätzlich mit einem Luftspalt von einander getrennt werden.

Nachfolgend ein Beispiel für ein 12V Filter welches bei einem Strom von 200A maximal mögliche Dämpfungswerte erzielen soll. Theoretisch benötigen wir um hohe Dämpfungswerte zu erzielen keine mehrstufigen Filter. Die Mehrstufigkeit verwendet man typischerweise um steile Flanken im Dämpfungsverlauf zu erhalten. Sie erinnern sich, der Dämpfungsverlauf einstufiger Tiefpassfilter sinkt mit 20dB/Dekade, zweistufige Tiefpassfilter mit 40dB/Dekade und so weiter. In der EMV verwenden wir die Mehrstufigkeit um das nicht ideale Verhalten der Bauteile zu kompensieren. Der aufgebaute Filter besteht aus sechs L-C Stufen. Die Besonderheit liegt darin, dass der induktive Anteil nur durch die Kupferleitung entsteht. Die Werte sind mit 12nH natürlich entsprechend klein. Als Daumenregel kann als Induktivitätsbelag ein Werte von 10nH/cm angenommen werden.

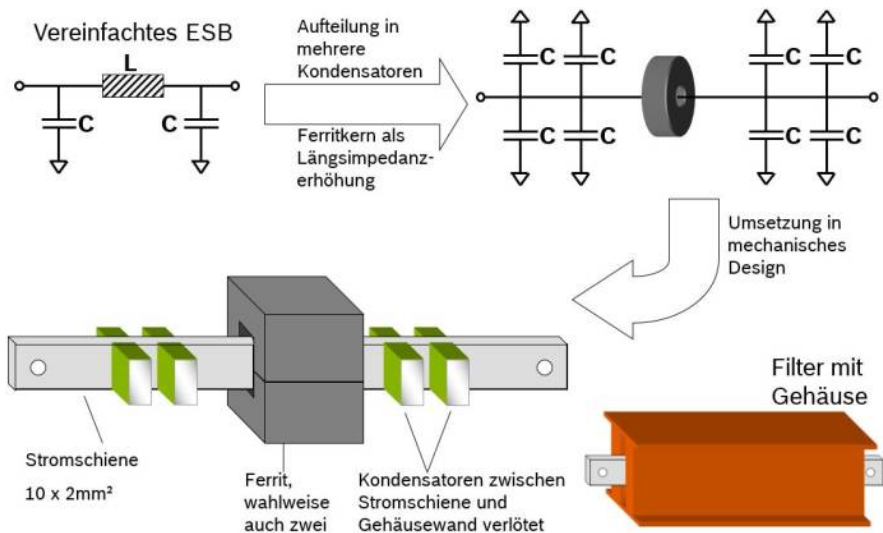


Der Filter ist mäanderförmig um die Kondensatoren aufgebaut welche durch das Kupfergehäuse sehr gut an die Masse angebunden sind. Im nächsten Bild sehen wir das Ersatzschaltbild des Filters mit den Einzelbauteilen. Zugegeben, wenn wir im nächsten Bild die Messung mit der Simulation vergleichen dann sind die Werte „frisiert“. Das bedeutet zuerst war die Messung vorhanden und erst danach wurde die Simulation auf das Ergebnis angepasst. Zu meiner Verteidigung muss ich allerdings sagen, dass nur der Induktivitätsbelag angepasst wurde. Die Ersatzwerte der Kondensatoren wurden zuvor mit einem Imepdedanzanalysator ermittelt.

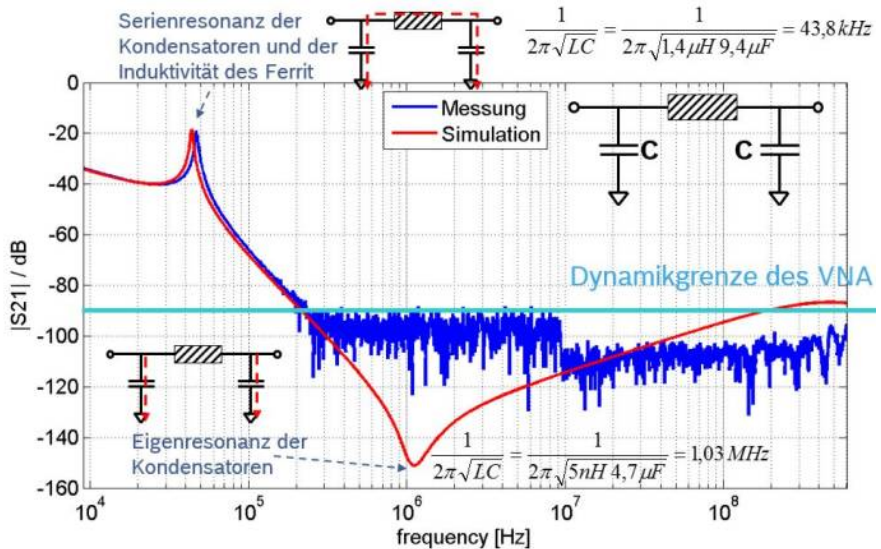


Die Simulation passt erstaunlich gut zur Messung. Man könnte fast annehmen es wurde noch mehr frisiert: Ja, das stimmt. Diesmal allerdings nicht die Bauteilwerte, sondern die Messmethode. Auch durch die Messung können sich Fehler einschleichen. Ironischerweise handelt es dabei um ein EMV-Problem Die Messmethode möchte ich anhand der nächsten Filterschaltung vorstellen.

Dass es keiner sechsstufigen Filter bedarf um hohe Dämpfungswerte zu erzeugen möchte ich im nächsten Beispiel zeigen. Zur Erhöhung der Induktivität beiße ich in den sauren Apfel und setzen ein permeables Material, in Form eines Ferritkerns ein. Wichtig ist, dass der Kern nicht in Sättigung gerät da ansonsten die Induktivität schlagartig verschwindet. Bei der Schaltung handelt es sich um die weit verbreitete π -Stuktur. Da bedeutet nichts anderes als Kondensator - Induktivität - Kondensator. Da es für derart hohe Ströme keine klassischen gewickelten Spulen mehr gibt müssen wir die Induktivität zusammen mit dem Ferritkern und dem Induktivitätsbelag der Stromschiene bilden. Zur Reduktion der parasitären Eigenschaften der Kondensatoren werden jeweils vier in Parallelschaltung verwendet. Nachfolgende Abbildung zeigt den prinzipiellen Aufbau und das Ersatzschaltbild des Filters.

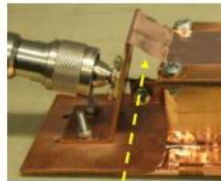


Gehen wir gleich weiter zum Ergebnis der Filterdämpfung. Es mag überraschen: Obwohl wir deutlich weniger Filterstufen einsetzen erzielen wir über einen weiten Frequenzbereich eine deutlich höhere Dämpfung. Die Reduktion der Filterstufen macht sich allerdings nicht in der Baugröße bemerkbar. Sie erinnern sich, der magnetische Fluß ergibt sich aus der magnetischen Induktion multipliziert mit der vom Fluß durchsetzten Fläche. Oder anders ausgedrückt: Die Größe des Ferritkerns geht direkt in den erzielten Induktivitätswert ein.

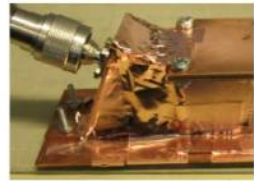


Man könnte sagen, wir haben unser perfektes Filter gefunden! Ok, mit der Einschränkung für Frequenzen oberhalb von 200kHz. Ganz allgemein gilt, dass Frequenzen kleiner 100kHz nur mit „Größe“ zu bekämpfen sind. Größe bedeutet hier im Sinne von geometrischer Größe. Schuld daran sind die frequenzabhängigen komplexen Impedanzen unserer komplexen Bauteilen. Damit eine Induktivität bereits bei geringen Frequenzen eine hohe Impedanz aufweist müssen wir hohe Induktivitätswerte verwenden da der ω Wert ($2\pi f$) entsprechend noch klein ist in der Gleichung $Z_L = \omega L$. Bei einem Kondensator haben wir die identische Problematik entsprechend reziprok. Störungen in diesem tiefen Frequenzbereich werden meist durch das Takten leistungselektronischer Systeme hervorgerufen. Wie sich die Eingangskondensatoren (Bulk- oder Zwischenkreiskondensatoren) berechnen lassen sind in einem eigenen Kapitel beschrieben.

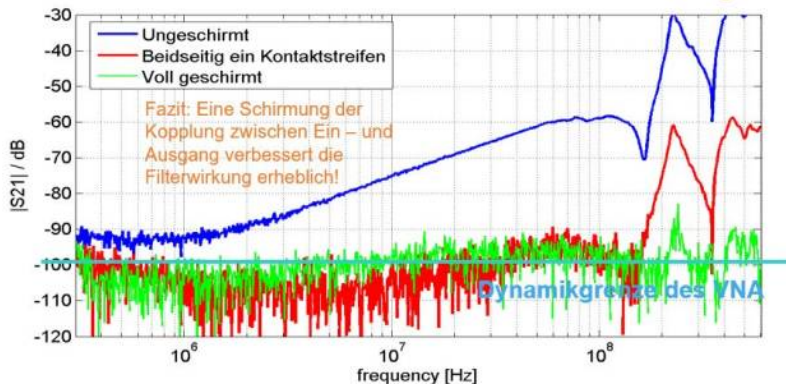
Verbindung des Anschlusswinkels mit einem Kontaktstreifen oben auf beiden Seiten



Kontaktstreifen



Voll geschirmt



Jetzt aber wie zuvor besprochen zum Trick mit der Messung. Auch wenn wir in diesem Kapitel nur leitungsgebundene Störgrößen betrachten können wir nicht ausschließen, dass es nicht auch bei unserem Filter zu einer gestrahlten Emission kommt. Netzwerkanalysatoren haben die für uns positive Eigenschaft, dass wir lediglich die Transmission zwischen den zwei verwendeten Ports, den S21 Parameter, verwenden müssen da dieser direkt der Einfügedämpfung des Filters entspricht. Aber wie den Filter kontaktieren? Das Messgerät besitzt typischerweise BNC- oder N-Buchsen für koaxiale Leiteranschlüsse. Diese Situation hat natürlich nichts mit dem späteren Einsatzfeld des Filters zu tun. Wer kontaktiert schon einen 200A Leistungsausgang über N-Buchse, abgesehen davon dass dies nicht möglich ist

Und genau hier liegt unser Dilemma. Obige Abbildung zeigt sehr gut den Einfluss zusätzlicher Massestreifen bzw. die vollflächige Schirmung der Anschlussbleche. Zur Bestimmung der Einfügedämpfung wurde eigens für diesen Filter eine Testvorrichtung erstellt. Sie besteht aus einer massiven durchgängigen Massefläche welche für eine vollflächige niederohmige Anbindung des Gehäuses sorgt. Die Signal Ein- und Ausgänge sind mit N-Buchsen realisiert welche über zusätzliche Kupferwinkel an die Referenzmasse angebunden sind. Da zuerst die Simulationsergebnisse so gar nicht zu den Messergebnissen passen wollten mussten wir auf Fehlersuche gehen. Die Ursache war schnell gefunden. Die Grafik zeigt die hohe Abhängigkeit der Dämpfungsergebnisse von der Schirmung der

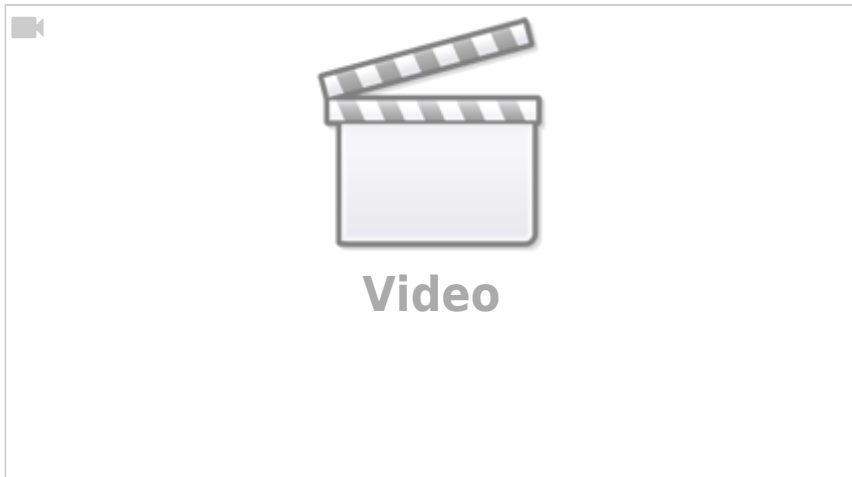
Anschlüsse. Physikalisch ist die Ursache in der gestrahlten Emission der Anschlüsse zu suchen. Ganz einfach gesagt haben die Emissionen schlichtweg keine Lust den Weg durch das Filter zu nehmen und entscheiden sich über den Luftweg. Wir können uns daher noch so anstrengen beim Filteraufbau wenn wir nicht für eine strikte Trennung zwischen den Ein- und Ausgängen sorgen, zum Beispiel durch genügend Abstand oder Schirmbleche.

6.8 Fazit

Die Auslegung von Filterschaltungen zur Bekämpfung leitungsgebundener Störgrößen, um z. B. die Grenzwerte nach CISPR25 oder DIN EN 61000-6-1 einzuhalten, ist selten handstreichartig erledigt. Falls Sie in der glücklichen Lage sind, dass genug Bauraum und Geld zur Verfügung stehen, können Sie den Leitungsfiler generell etwas großzügiger auslegen. Eine Filterstufe mehr oder weniger wird an dieser Stelle niemandem wehtun, und man kann getrost nach dem „Trial-and-Error“-Prinzip vorgehen, mit vorgefertigten Filterstufen (Kaufteilen). Ganz anders ist die Situation natürlich zu bewerten, falls gerade kein Platz mehr für zusätzliche Filter zur Verfügung steht oder eine derart hohe Stückzahl produziert werden soll, bei der jeder zusätzliche Cent schmerzt. Hier bietet sich eine systematische Filterentwicklung an, von der Störgrößenanalyse bis hin zur gezielten Reduzierung des dominanten Störmechanismus. Leitungsgebundene Störgrößen sind neben der gestrahlten Emission einer der beiden Hauptschwerpunkte der EMV-Arbeit. In keiner anderen Disziplin der EMV (Burst, Surge, ...) ist so viel Geduld gefordert. Obwohl die EMV so vielfältig ist, habe ich viele Jahre damit verbracht, dafür zu sorgen, dass in Fahrzeugen der Radioempfang in allen Frequenzbändern uneingeschränkt möglich ist. Der Schlüssel hierzu liegt meist nicht, wie vermutet, in der gestrahlten Emission, sondern im Ursprung bei leitungsgebundenen Störungen, welche sich über die Versorgungsleitungen ausbreiten.

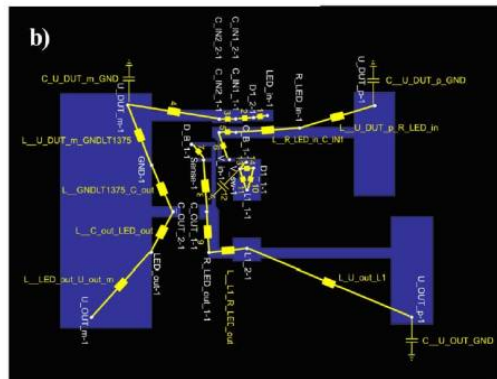
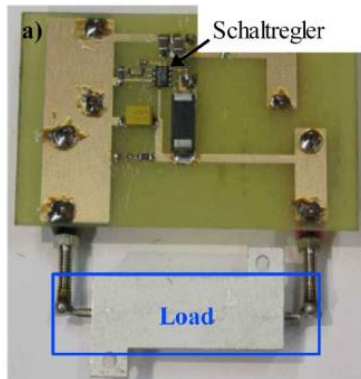
7.0 Simulation leitungsgebundener Störungen

Wie zuvor gesehen, lassen sich leitungsgebundene Emissionen wunderbar in einer Simulation darstellen. Die Frage, die sich stellt, muss natürlich lauten: Wie realitätsnah ist die Simulation? Können wir die Simulationen mit der Messung vergleichen?



Für eine Simulation der leitungsgebundenen Störgrößen müssen neben den Eigenschaften der Bauelemente auch die parasitären Eigenschaften des Aufbaus berücksichtigt werden. Die Induktivität der Leiterverbindungen auf dem PCB beeinflusst dabei die Störgrößen in gleicher Weise wie die kapazitive Kopplung der Elemente zur Referenzmasse bzw. zum Gehäuse. Zur Simulation der Störgrößen und deren Ausbreitungspfade stehen unterschiedliche Methoden wie die **TLM** (Transmission Line Method) oder die **PEEC-Methode** (Partial Element Equivalent Circuit) zur Verfügung. Die TLM berechnet mithilfe der Leitungsgleichungen die Eigenschaften einzelner Leitungsstücke des PCB, basierend auf den Leitungsbelägen [Paul, 1994]. Die PEEC-Methode hingegen ist dadurch gekennzeichnet, dass jedem elektrisch kurzen und homogenen Teil der Schaltung partielle Induktivitäten und Kapazitäten sowie Längs- und Querwiderstände zugeordnet werden.

Als Beispiel für eine Simulation leitungsgebundener Störungen wird eine Tiefsetzsteller-Schaltung (Buck-Converter) beschrieben, wie sie häufig zur Generierung verschiedener Spannungsebenen verwendet wird. Der Abwärtswandler reduziert z.B. die Bordnetzspannung in einem Kraftfahrzeug von 12–15 V auf eine konstante Ausgangsspannung von 5 V, bei einem maximalen Laststrom von 1,5 A. Die Schaltung beinhaltet neben der Spannungsregelung den Leistungstransistor, der mit einer Schaltfrequenz von 500 kHz angesteuert wird. Als Last wird ein einfacher Widerstand zur Aufnahme der abgegebenen Leistung verwendet. Nachfolgende Abbildung zeigt den Aufbau der Schaltung mit dem verwendeten Schaltregler im SO-8 SMD-Gehäuse (links, a) zusammen mit der Charakterisierung aller Leiterbahnen mit Ersatzimpedanzen auf der rechten Seite (b).

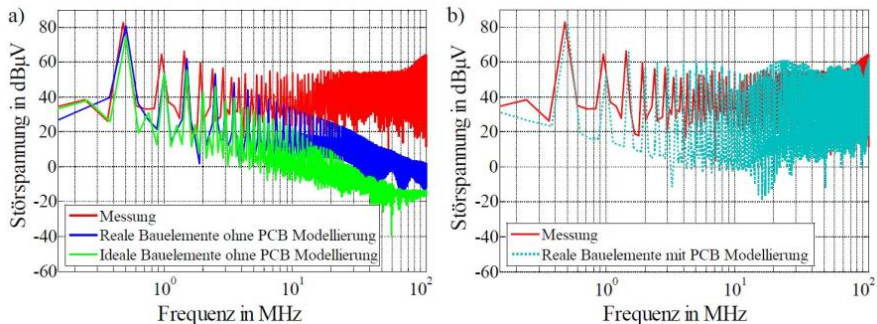


Die Simulation geht von gleichen Randbedingungen entsprechend einer Messung der leitungsgebundenen Störungen nach CISPR25 aus. Das bedeutet, dass die Leiterplatte 5 cm über der Referenzmasse angeordnet ist und die Störungen mit einer Netznachbildung ausgekoppelt werden. Die Gesamtsimulation wird als Netzwerksimulation in einer SPICE-Umgebung durchgeführt. Im ersten Schritt wird eine Simulation aufgebaut, die den funktionalen Teil ohne parasitäre Effekte berücksichtigt. Danach werden alle passiven Bauteile durch verlustbehaftete Ersatzschaltbilder der realen Elemente ersetzt.

Am aufwendigsten gestaltet sich die Nachbildung der auf der Leiterplatte verwendeten Verbindungen zwischen den einzelnen Elementen. Neben der Berechnung der Induktivität zwischen zwei Knoten wird die Kapazität zur Referenzmasse ermittelt. Das durch die PEEC-Methode generierte Modell erlaubt es, je Teilstück eine konzentrierte Induktivität und Kapazität zu definieren. Dieses Vorgehen hat den Vorteil einer deutlichen Vereinfachung der Gesamtsimulation, da die hohe Anzahl an in der PEEC-Nachbildung erzeugten Knoten nicht in der Simulation berücksichtigt werden muss.

Sehr dünne Leiterbahnen, die weder am aktiven Teil der Schaltung noch im Leistungszweig Verwendung finden, werden bei der Modellierung vernachlässigt. In Bild 2.9 b) sind die einzelnen Teilelemente dargestellt, die für eine ausreichende Modellierung des PCB nötig sind. Dauert die funktionale Simulation der Schaltung mit idealen Elementen weniger als eine Minute, so erstreckt sich eine Simulation mit der Abbildung der realen Elemente sowie der Eigenschaften des PCB auf einer modernen Workstation bereits über 180 min. Verlängert wird die Simulation zusätzlich durch die Nachbildung der zur Bewertung der Störspannungen notwendigen Netznachbildung. Die Verwendung „realer“ Bauelemente in der Simulation bezieht sich hierbei auf die Berücksichtigung der gemessenen HF-Ersatzschaltbilder aus einer Impedanzmessung der einzelnen Bauelemente.

Als Ergebnis ergeben sich die Störgrößen auf den Versorgungsleitungen im Zeitbereich. Die äquivalente Messung der Störspannung erfolgt bei gleichem Aufbau mit dem Oszilloskop. Dadurch lassen sich im ersten Schritt die Zeitsignale der Störgrößen in Messung und Simulation vergleichen und letztendlich mithilfe der FFT eine Aussage über das Störspektrum im Frequenzbereich berechnen.



Die Abbildung zeigt den Vergleich zwischen Messung und Simulation der in den Frequenzbereich transformierten Störspannung. Dabei kommen unterschiedliche Modellierungsmöglichkeiten zum Einsatz. Auf der linken Seite erfolgt die Simulation mit idealen und realen Bauelementen, wobei die Eigenschaften der Platine nicht berücksichtigt sind. Es zeigt sich, dass der Störpegel unter Annahme idealer Bauelemente lediglich für die Grundfrequenz des Abwärtswandlers richtig wiedergegeben wird. Bereits ab ca. 2 MHz weicht der simulierte Wert von der Messung um mehr als 20 dB ab. **Das funktionale Grundmodell ist somit nicht in der Lage, hochfrequente Störungen größer 2 MHz abzubilden.**

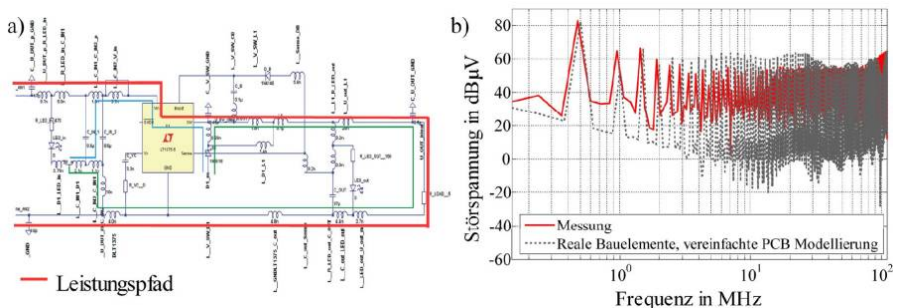


Funktionale Simulation ohne die Berücksichtigung parasitärer Elemente können lediglich die Störgröße der Grundfrequenz korrekt wiedergeben und sind somit für eine EMV-Bewertung unzureichend!

Das Ergebnis verbessert sich deutlich, wenn die realen HF-Eigenschaften der eingesetzten Bauelemente berücksichtigt werden. Die berechnete Störspannung stimmt nun mit annehmbaren Abweichungen bis ca. 5 MHz mit der Vergleichsmessung überein. Für eine Bewertung über dem geforderten Frequenzbereich ist dieses erweiterte Modell allerdings ebenfalls noch zu

ungenau. Erst für den Fall, dass die Geometrie der Leiterverbindungen mit berücksichtigt wird, lassen sich Frequenzanteile bis 108 MHz berechnen. In der zuvor gezeigten Abbildung auf der rechten Seite ist das Ergebnis der Gesamtsimulation mit Berücksichtigung der Leitungsinduktivitäten der Leiterplatte und der kapazitiven Kopplungen zur Referenzmasse im Vergleich zur Messung dargestellt. Die Abweichungen zur Kontrollmessung bleiben über dem gesamten Frequenzbereich < 10 dB. Im Vergleich der einzelnen Ergebnisse zeigt sich damit deutlich der Einfluss der Leiterverbindungen auf dem PCB. Ohne Berücksichtigung der parasitären Effekte auf der Leiterplatte ist somit eine genaue Berechnung der Störgrößen nicht möglich. Die Genauigkeit der Simulation wird jedoch durch einen signifikanten Anstieg der Simulationsdauer erkauft, da zur Berechnung der hochfrequenten Signale die Zeitschritte in der Netzwerksimulation stark verringert werden müssen.

Zur Reduktion der Simulationsdauer wird im Folgenden die Anzahl der verwendeten Ersatzelemente zur Beschreibung des PCB verringert. Dazu werden nur Pfade charakterisiert, welche sich im Leistungskreis befinden bzw. den Laststrom führen. Dabei kann sogar der Freilaufpfad vernachlässigt werden. Nachfolgende Abbildung zeigt den für die Vereinfachung charakterisierten Leistungspfad im Schaltplan. Die verbleibenden PCB-Traces werden als verlustfreie Verbindungen angenommen.



Die Simulationsdauer kann mit dieser Vereinfachung mehr als halbiert werden. Das Ergebnis des reduzierten Modells zeigt nach wie vor eine hohe Übereinstimmung mit den Werten aus der Kontrollmessung. Für eine erfolgreiche Simulation der leitungsgebundenen Störgrößen bis 108 MHz ist es somit ausreichend, den Leistungskreis der Schaltung zu modellieren bzw. die Strompfade mit der höchsten Änderungsgeschwindigkeit (di/dt) der Störströme zu berücksichtigen.

Eine gute Zusammenfassung mit vielen weiteren Infos als Veröffentlichung zum Download:

simulation_conducted_emission_buck-converter.pdf

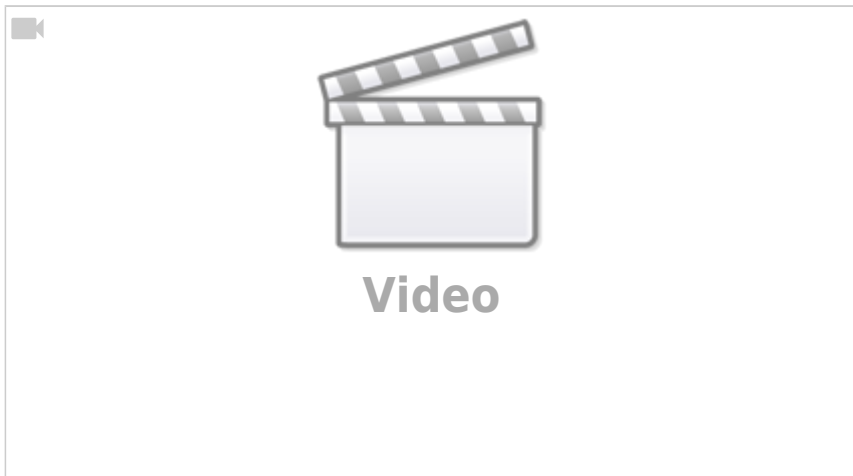
Das gezeigte Beispiel soll die Möglichkeit zur Simulation leitungsgebundener Störungen einer gängigen Schaltung aufzeigen. Nicht alle Teilprobleme der EMV lassen sich allerdings elegant im Netzwerksimulator beschreiben.

Teilsimulationen zum Download:

simulationen.zip

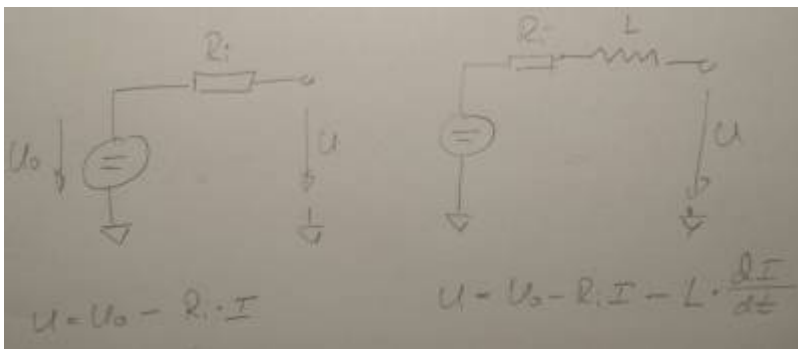
8.0 Zwischenkreiskapazität

Beim Öffnen eines elektronischen Gerätes fallen einem meist zuerst die Speicherkondensatoren auf. Aufgrund ihrer Größe auch nicht zu übersehen. Dabei handelt es sich je nach Anwendung typischerweise um Elektrolytkondensatoren mit mehreren mF Kapazität. Besonders große Zwischenkreiskapazitäten (engl. Bulk Capacitor) sind eher unerwünscht, da hinsichtlich der Lebensdauer die Elektrolyte stets die Schwachstelle vieler Geräte darstellen. Als elektronische Baugruppen noch repariert wurden, haben erfahrene Elektroniker diese Bauteile zuerst unter die Lupe genommen. Die häufigste Frage in diesem Zusammenhang lautet natürlich: Woher weiß ich als Schaltungsentwickler, welcher Kapazitätswert benötigt wird?



Diese Frage ist nicht einfach zu beantworten und benötigt viele Informationen rund um die Funktion der Komponente als auch deren Einbaulage. Aber gehen wir zuerst der Notwendigkeit für Zwischenkreise auf die Spur. Überlegen wir uns kurz, wie die Spannungsversorgung elektronischer Geräte hergestellt wird.

Typischerweise benötigen wir eine DC-Versorgung von 5V, 12V, 15V usw., welche natürlich möglichst konstant sein soll. Leider gibt es keine idealen Spannungsquellen, welche die Spannung an den Klemmen konstant hält, komme, was wolle. Wir müssen somit mit den realen Spannungsquellen in Form von Labornetzteilen, Batterien oder Gleichrichterschaltungen vorliebnehmen. Je besser die Spannungsquelle, desto unabhängiger wird die Ausgangsspannung vom Laststrom. Dieser Zusammenhang wird durch den Innenwiderstand der Quelle beschrieben, an welchem unter Last eine Spannung abfällt. Schlimmer noch als der Innenwiderstand oder der Leitungswiderstand einer entfernt angeschlossenen Klemme ist die sich durch die Zuleitungen bildende Induktivität. Jede noch so kurze Leitung kommt immer mit einer Induktivität, welche mit der Faustformel $1\mu\text{H}/\text{m}$ abgeschätzt werden kann.



Damit erhalten wir jetzt ein frequenzabhängiges Verhalten unserer Spannungsquelle. Je höher die Frequenz, desto größer wird die Impedanz der Leitung und bremst sozusagen unseren Nutzstrom aus. In der Praxis bedeutet dies für alle getakteten Anwendungen (Tiefsetzsteller, Motoransteuerungen, ...), dass, obwohl ein Stromfluss angefordert wird, im ersten Moment (je nach Induktivitätswert) nichts passiert. Vielleicht kennen Sie das Sprichwort: An Induktivitäten die Ströme sich verspäten – der Strom in der Zuleitungsinduktivität muss zuerst aufgebaut werden. Ein typisches Anzeichen, dass Ihr Zwischenkreis zu klein ist, zeigt sich beim Funktionstest, falls Sie es komischerweise nicht schaffen, die volle Leistung aufzunehmen (klar, die Spannung am Eingang bricht dauernd ein) oder die EMV-Ergebnisse für leitungsgebundene Störungen im sehr tiefen Frequenzbereich viel zu hoch sind. Wichtig ist in diesem Zusammenhang, dass Funktionstests stets mit realistischen Leitungslängen durchgeführt werden, ansonsten ist später im Feld mit merkwürdigen Ergebnissen zu rechnen. Das ist die eine Seite: Die Schaltung schafft es nicht, den Stromfluss aufzubauen. Genauso problematisch ist es, den Stromkreis wieder zu unterbrechen. Eines der schlimmsten Dinge, die man in der Elektrotechnik machen kann, ist eine

Kapazität kurz zu schließen oder eben den Stromfluss durch eine Induktivität zu unterbrechen. Fließt beim kurzgeschlossenen Kondensator ein sehr hoher Strom, führt die Unterbrechung des Stromflusses an der Induktivität zu einer Spannungsüberhöhung bis hin zu Lichtbögen. Ursache für beide Phänomene ist die Energieerhaltung. Die in der Spule gespeicherte magnetische Energie muss kontinuierlich sein und kann nicht springen. Sollten Sie es wagen, den Stromkreis zu unterbrechen, reagiert die Induktivität damit, die Spannung zu erhöhen, damit der Strom weiter fließen kann – zur Not über die Luft.

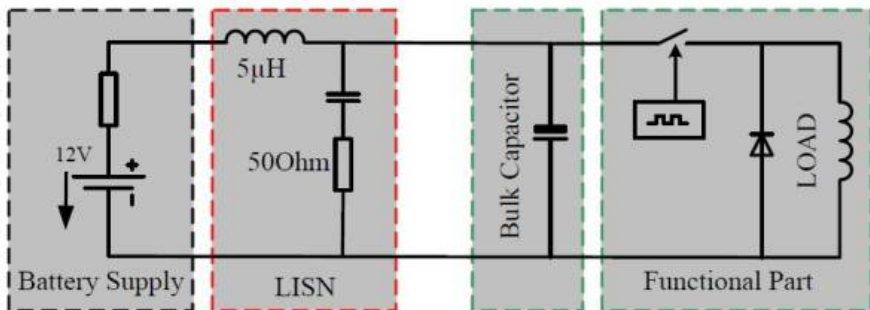
Zwischenkreiskapazitäten haben neben der Spannungsstabilisierung eine noch wichtigere Aufgabe – überschüssige Energie aus der Leitungsinduktivität muss beim Abschalten des Stromflusses aufgenommen werden können. Ist der Zwischenkreis zu klein ausgelegt, erzeugen beide Vorgänge (Ein- und Ausschalten) EMV-Störungen im Frequenzbereich bis ca. 1MHz. Daher kann der Zwischenkreis neben der funktionalen Auslegung auch der EMV zugeschrieben werden. In den meisten Fällen werden die Schaltungsentwickler etwas größere Kapazitäten vorsehen, als funktional zwingend notwendig.

Wir wissen jetzt, wozu die Zwischenkreiskapazität gut ist, allerdings müssen wir uns noch Gedanken darüber machen, über deren Größe. Prinzipiell gibt es wie so oft drei Möglichkeiten: Berechnen, Simulieren oder Trial and Error. Im vorherigen Kapitel haben wir gesehen, dass sich mit Hilfe einer funktionalen Simulation, selbst ohne genaue Kenntnis über die parasitären Eigenschaften der Bauelemente und des PCBs, im Frequenzbereich bis einige MHz sehr gute Ergebnisse erzielen lassen. Daher bietet es sich an, den Zwischenkreis über eine funktionale Simulation zu optimieren. Zu beachten sind allerdings auf jeden Fall die parasitären Eigenschaften der großen Becher-Elektrolytkondensatoren. Diese sind aufgrund der großen Abmessungen keineswegs zu vernachlässigen. In den nachfolgenden Unterkapiteln schauen wir uns die Notwendigkeit des Zwischenkreises und die Berechnung der auftretenden Störspannung am Beispiel eines einfachen automotive Tiefsetzstellers genauer an.

8.1 Beispiel Tiefsetzsteller

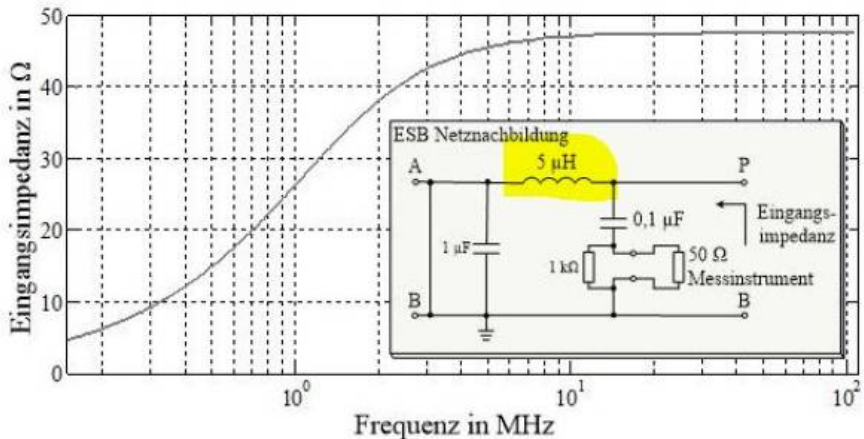
Elektronische Fahrzeugkomponenten werden je nach Hersteller an verschiedenen Positionen im Fahrzeug angeordnet. Dadurch variiert die benötigte Anschlusslänge der Komponente zur Batterieversorgung. Zur Nachbildung der Fahrzeugumgebung wird in der Komponentenmessung eine Netznachbildung verwendet, die in erster Näherung das induktive Verhalten der Versorgungsleitungen berücksichtigt. Werden Ströme durch eine Pulsweitenmodulation (PWM) zyklisch an- und abgeschaltet, wird das EMV-Verhalten wesentlich durch die Speicherkapazität bis in den MHz-Bereich bestimmt. Dabei spielt besonders der Kapazitätswert eine bedeutende Rolle. Ist

er zu gering, hängt das EMV-Verhalten der Komponente von der Leitungslänge ab. Ist der Kapazitätswert zu hoch, wird unnötiger Bauraum verbraucht. Speicherkondensatoren (Bulk-Kondensatoren) dienen hauptsächlich dazu, die Lastschwankungen, die während der Schaltperiode entstehen, auszugleichen. Betrachtet werden im Folgenden die leitungsgebundenen Störspannungen beim Abschalten von Lastströmen in der Umgebung einer CISPR25-Messung bzw. einer Abschätzung im Fahrzeug.



Die Abbildung zeigt das vereinfachte Ersatzschaltbild des Messaufbaus zur Messung der leitungsgebundenen Störspannung entsprechend CISPR25. Als DUT (Device Under Test) wird eine einfache PWM-Schaltung mit ohmsch/induktiver Last betrachtet. Der funktionale Teil der Schaltung (PWM-Generierung, Schaltvorgang) wird im Folgenden als idealer Vorgang angenommen. Von Interesse ist die Auslegung der Bulk-Kondensatoren (Speicherkondensatoren) sowie die auftretenden Störspannungen in Abhängigkeit der Induktivität der LISN (Netznachbildung) und somit der zum Anschluss der Komponente vorhandenen Leitungsinduktivität.

Die zur Auskopplung der HF-Ströme verwendete automotive Netznachbildung enthält eine $5\ \mu\text{H}$ Seriendrossel, ausgeführt als Luftspule. Durch die Ausführung als Luftspule behält die Drossel ihre Eigenschaften unabhängig von der Strombelastung. Spulen mit permeablem Kern sind stets abhängig vom Stromfluss und damit von der magnetischen Feldstärke. Die Ursache dafür ist, dass wir beim Übergang von der magnetischen Feldstärke zur magnetischen Induktion (später Flussdichte zusammen mit der durchsetzten Fläche) eine Linearisierung der **B-H Kurve** durchführen. Für ein Messmittel bietet es sich an, diese Fehlerquelle auszuschließen.



Die verwendete Spule in Serienschaltung hat zwei wesentliche Aufgaben:

- Unterbrechung der Störströme über die Batterie → Störstrom muss durch den Messempfänger
- Nachbildung der Impedanzverhältnisse im Fahrzeug

Allerdings muss auch die Eigenschaft der Drossel als Energiespeicher beachtet werden. Der in die Drossel eingeprägte Strom stellt die energietragende Größe dar, deren Verlauf stetig sein muss. Eine schnelle Änderung der Stromstärke (Abschalten des PWM-Impulses) hat entsprechend dem Induktionsgesetz einen Spannungsabfall über der Spule zur Folge, welcher bei der leitungsgebundenen Messung erfasst wird.

8.2 Einfluss der Leiterlänge / Induktivität auf die Störspannung

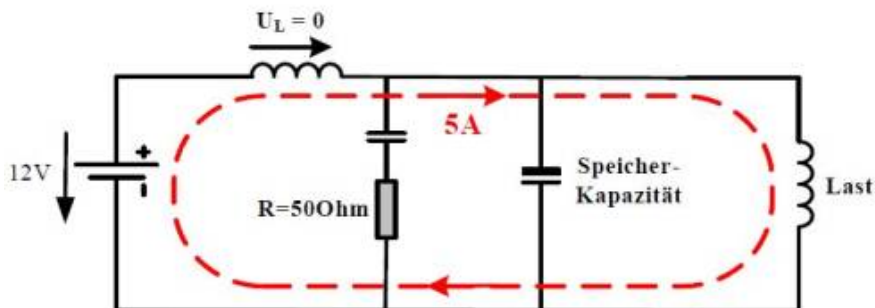
Im ersten Schritt wird der stationäre Zustand bei geschlossenem (Leistungs-)Schalter betrachtet. Die Spannung über dem Speicherkondensator stellt sich dabei auf die Batteriespannung ein. Der differenzielle Strom ergibt sich aus der Lastimpedanz, wobei er ungehindert im Lastkreis fließen kann.

Für das Berechnungsbeispiel wird von folgenden Randbedingungen ausgegangen:

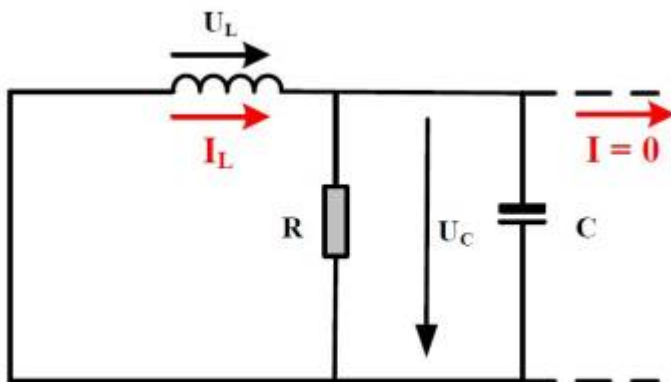
- Innenwiderstand der Batterie wird vernachlässigt
- Stationärer Laststrom $I = 5\text{ A}$

- Stationäre Spannung am Tiefsetzsteller $U = 12V$

Nachfolgende Abbildung zeigt den Stromfluss im stationären Zustand für das vereinfachte Ersatzschaltbild des Tiefsetzstellers.



Im nächsten Schritt wird der Leistungsschalter geöffnet und damit der Stromfluss von der Versorgung zur Last unterbrochen. Im Ausgangskreis des Tiefsetzstellers kommutiert der Strom nun auf die Freilaufdiode. Im Eingangs- oder Versorgungskreis ergibt sich das folgende Ersatzschaltbild. Für den aktuell fließenden Strom von $I_L = 5A$ in der Spule besteht keine Möglichkeit, weiter zu fließen. Das bedeutet, die in der Spule vorhandene Energie muss durch den Zwischenkreis aufgenommen werden.



Aus dem vereinfachten Ersatzschaltbild ergibt sich eine lineare

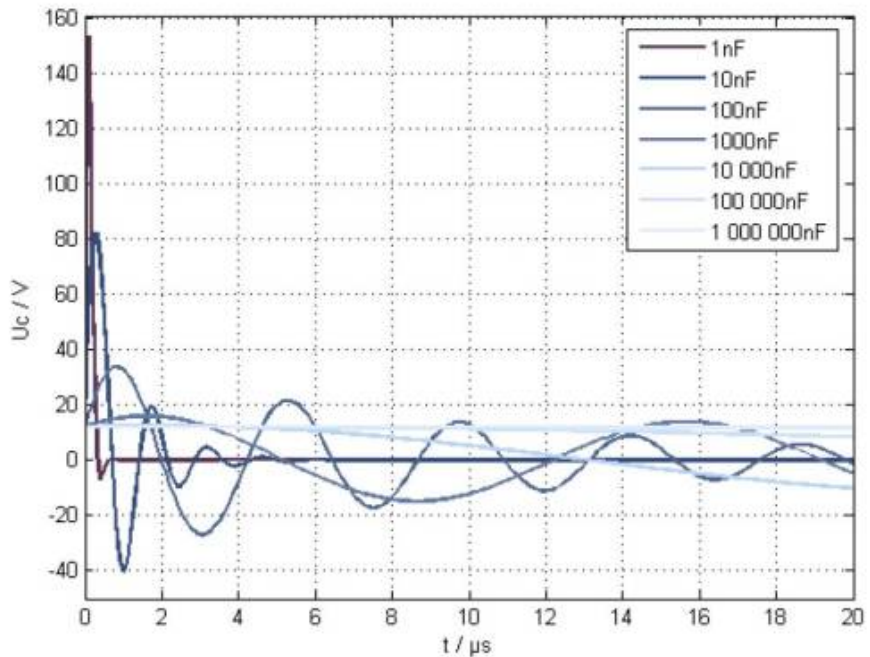
Differenzialgleichung zweiter Ordnung, die den Spulenstrom über der Zeit beschreibt. Das System wird zum Zeitpunkt $t = 0$ mit den Anfangsbedingungen für den Spulenstrom $I_L = 5A$ und der Spannung über dem Kondensator $U_C = 12V$ angeregt.

$$\frac{d^2}{dt^2} i_L(t) + \frac{\frac{d}{dt} i_L(t)}{R C} + \frac{i_L(t)}{C L} = 0$$

Der zeitliche Verlauf des Spannungsabfalls über der Spule in der Netznachbildung ergibt sich als Ableitung des Spulenstroms:

$$u_L = L \left(\frac{d}{dt} i(t) \right)$$

Die Spannung über der Spule entspricht ungefähr der Spannung am Messempfänger. Nachfolgende Grafik zeigt die berechnete Störspannung am Messempfänger, welche durch das Abschalten des Stromes durch die Spule entsteht, für verschiedene Zwischenkreiskapazitäten.



Betrachtet man den zeitlichen Verlauf der Störspannung, zeigt sich eine deutliche Spannungsüberhöhung, resultierend aus dem Stromgradienten in der Drossel. Kann der in der Spule verbleibende Strom nicht vom Kondensator aufgenommen werden, erhöht sich die Spannung über der Drossel, welche über dem Messkreis (50 Ω Messimpedanz) abgebaut wird. Wird nur ein Ausschaltzyklus bei der Berechnung beachtet, ergibt sich der Vorgang unabhängig von der Schaltfrequenz bzw. Wiederholfrequenz der PWM-Ansteuerung, womit sich das Frequenzspektrum berechnen lässt.

Die Störspannung ergibt sich aus der Differenzialgleichung im Zeitbereich zu:

$$u_C(t) = A_1 e^{\lambda_1 t} + A_2 e^{\lambda_2 t}$$

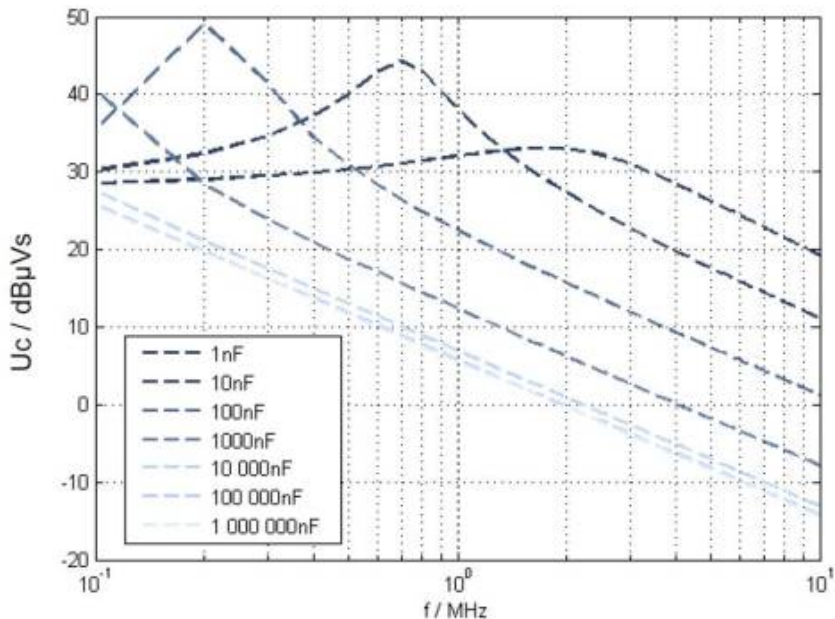
Mit den Zeitkonstanten λ_1 und λ_2 sowie den allgemeinen Konstanten A_1 und A_2 .

Über das Fourier-Integral gelangt man zu der Störspannung über der Frequenz mit der allgemeinen Lösung:

$$U_C(f) = \int_0^{\infty} u_C \cdot e^{-j\omega t} dt$$

$$U_C(f) = \frac{A_1}{j\omega - \lambda_1} + \frac{A_2}{j\omega - \lambda_2}$$

Nachfolgende Abbildung zeigt das Spektrum der Störspannung entsprechend des zuvor hergeleiteten Dichtespektrums für verschiedene Speicherkapazitäten. Wie erwartet steigt mit sinkendem Kapazitätswert das Spektrum der Störspannung deutlich an. Je nach Kapazitätswert bildet sich neben der breitbandigen Anhebung des Spektrums eine Resonanz mit der Spule der Netznachbildung im betrachteten Frequenzbereich aus.



Wichtig: Bitte beachten, dass wir bisher stets von idealen Bauelementen

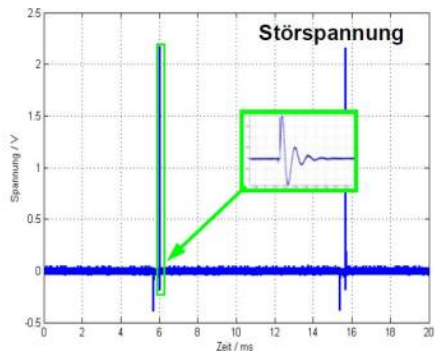
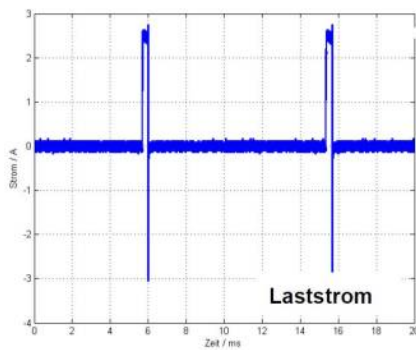
ausgegangen sind. Allerdings besitzen im Besonderen Elektrolytkondensatoren sehr hohe parasitäre Elemente, welche nicht vernachlässigt werden dürfen. Zusätzliche Elemente lassen die analytische Gleichung in ihrem Umfang explodieren, welche dann nur noch schwer geschlossen lösbar ist.

8.3 Schaltungsbeispiel



Als Schaltungsbeispiel schauen wir uns die Messergebnisse eines realen, wenn auch sehr einfachen, Tiefsetzstellers an. Der einzige Unterschied zum obigen ESB besteht darin, dass während der Messungen für jeden Leiter eine Netznachbildung

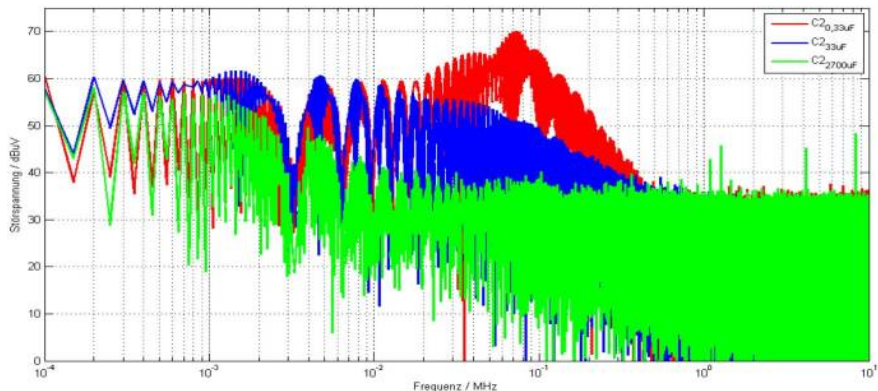
verwendet wurde. Die Ansteuerung des MOS-FET erfolgt über zwei Timer NE555, mit denen sich die Frequenz und das Puls-Pausen-Verhältnis bestimmen lassen. Als Last wird ein Widerstand verwendet. Eine induktive Last ist nicht notwendig, da der Lastzweig die Störspannung in diesem Fall nur bedingt beeinflusst. Der MOS-FET verbindet die Last mit der Spannungsversorgung entsprechend dem eingestellten Puls-Pausen-Verhältnis.



Die PWM-Frequenz wird mit Hilfe der Timer auf ca. 100 Hz eingestellt. Mit dem Lastwiderstand wird die maximale Amplitude definiert. Sie beträgt während der Schaltzeit ca. 2,8 A. Obige Abbildung zeigt den zeitlichen Verlauf des Laststroms sowie die Spannung an der Netznachbildung. Deutlich zu sehen ist die Spannungsüberhöhung, die während des Ausschaltvorganges auftritt. Das System aus Induktivität der Netznachbildung und Speicherkapazität wird durch den Schaltvorgang zum Schwingen angeregt. Nachfolgend werden drei unterschiedliche Speicherkondensatoren als Test für die

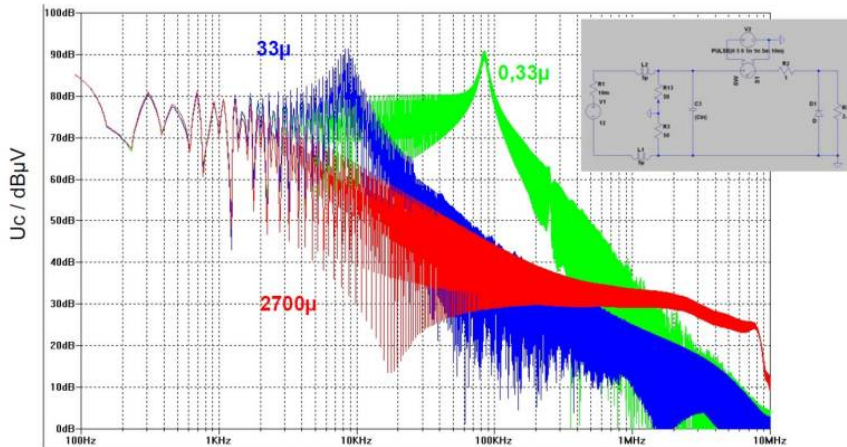
Spannungsüberhöhung beim Abschalten untersucht.

1. 0,33 μF Folienkondensator
2. 33 μF Elektrolyt
3. 2700 μF Elektrolyt



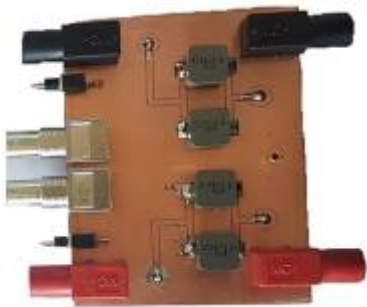
Obige Abbildung zeigt die mit Hilfe der Fourier-Transformation ermittelten Amplituden der Störspannung an der Netznachbildung. Am deutlichsten bildet sich die Resonanz für Variante 1 bei ca. 85 kHz aus. Aufgrund des geringen Kapazitätswerts entsteht wie erwartet die höchste Störspannung im Vergleich zu den Varianten 2 und 3. Deutlich zu sehen ist, wie die Störspannung mit steigendem Kapazitätswert abnimmt. Zur Reduktion der Störungen, bzw. soll der Prüfling möglichst unabhängig von der Zuleitungsinduktivität werden, darf der Kapazitätswert nicht zu gering gewählt werden.

Zum Abschluss ein Beispiel, wie die zuvor angestellten Berechnungen und Messungen in einer Simulation abgebildet werden können. Bereits eine sehr vereinfachte SPICE-Simulation kann den Einfluss der Speicherkapazität auf das Störspektrum sehr gut nachbilden. Richtig wiedergegeben wird die Resonanzfrequenz für die Variante 1 (0,33 μF), allerdings mit einer maximalen Amplitude, die mit 90 dB μV deutlich größer ist als die gemessene Spannung. Dies liegt an den nicht berücksichtigten Leitungsinduktivitäten / Leitungsbelägen des PCB sowie der deutlich geringeren Zeitbereichsauflösung des Oszilloskops im Vergleich zur Simulation.



8.4 Fazit

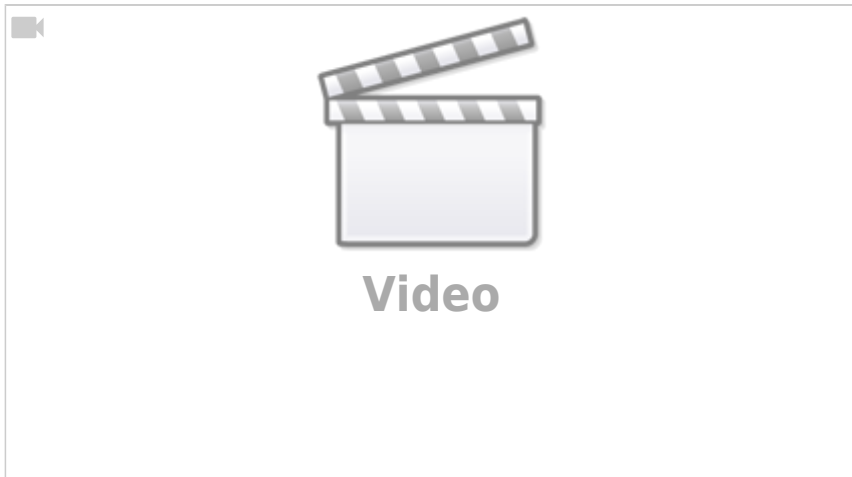
Bei der Modulation von Lastströmen sollte ein besonderes Augenmerk auf die Stromkommutierung in der Zuleitungsinduktivität bzw. Induktivität der Netznachbildung gelegt werden. Die beim Stromabriss entstehenden Überspannungen müssen durch die Speicherkondensatoren (Bulk-Kondensatoren) aufgenommen werden, da sie sich ansonsten über dem Innenwiderstand des Messempfängers entladen. Leider ist es oft nicht eindeutig möglich, die Störungen dem Spannungsabfall über der Drossel zuzuordnen, da der auftretende Störstrom ebenfalls aus differentiellen HF-Strömen im DUT resultieren kann. Um einen Einfluss durch die Zuleitung sicher auszuschließen, sind Messungen mit unterschiedlicher Leitungsinduktivität notwendig. Ergibt sich dabei keine Änderung der auftretenden Störspannung, ist davon auszugehen, dass auch im Fahrzeug unterschiedliche Leitungslängen das Störspektrum des DUT nicht beeinflussen. Funktionale Simulationen sind ein geeignetes Instrument, die Zwischenkreiskapazität so zu dimensionieren, dass die Zuleitungsinduktivität sowohl für die EMV als auch den Spannungseinbruch keine Rolle mehr spielt.



Auf jeden Fall sinnvoll und mit nur wenig Aufwand verbunden ist der Einsatz von Netznachbildungen bereits während der funktionalen Tests bzw. Inbetriebnahme der Geräte. Dadurch zeigt sich bereits in einer frühen Phase, ob die Zwischenkreiskapazität ggf. zu gering gewählt wurde. Dazu sind nicht immer teure und kalibrierte Netznachbildungen notwendig. Besonders für kleinere Ströme ($<20\text{A}$) ist es ausreichend, mit Hilfe einer externen Platine zwei Netznachbildungen diskret aufzubauen und in das Setup mit einzubinden. Einzig bei der Auswahl der Induktivität ist darauf zu achten, dass die Spule nicht in Sättigung gerät. Der Schritt, mit dem Oszilloskop die Emissionen zu messen, ist dann entsprechend einfach umzusetzen.

Die Abbildung zeigt eine Platine, welche beide Netznachbildungen nach CISPR25 (automotive) enthält. Über die BNC-Buchsen kann die Störspannung abgegriffen werden. Die beiden Schalter ermöglichen das Zuschalten eines 50Ω Abschlusswiderstands.

9.0 Kampf der gestrahlten Emission



Nachdem wir mit Hilfe von verschiedenen Filtern die leitungsgebundenen Emissionen erfolgreich unterhalb der Grenzwerte gedrückt haben, müssen wir uns jetzt mit den gestrahlten Emissionen beschäftigen. Um es vorweg zu nehmen: Die sorglosen Zeiten, bei denen wir für die leitungsgebundenen Emissionen die Zusammenhänge logisch erklären konnten, sind vorbei. Wer sich schon einmal mit der Antennentheorie beschäftigt hat, wird wissen, wovon ich spreche. Und genau hierher kommt das der EMV schon so lange anhaftende Voodoo-Image. Das ist aber auch kein Wunder. Bewegt man zwischen zwei Messungen die Versorgungsleitungen oder Kommunikationsverbindungen, kann das eine Änderung der Ergebnisse um mehrere dB bedeuten. Bin ich mit meinen Emissionen nahe dem Grenzwert, kann es sein, dass der Messaufbau über den weiteren Tagesablauf entscheidet ...



Gemessen werden die gestrahlten Emissionen mit dem zuvor beschriebenen Messverfahren (Antennenmessverfahren). Ich vermute, Sie kennen alle die Bilder von Absorberhallen, welche die Firmen gerne verwenden, um ihre Ausstattung zu präsentieren. Klar, eine Absorberhalle sieht man nicht alle Tage und ist vor allem keine

billige Angelegenheit. Die Abbildung zeigt die Messhalle der Porsche Engineering, in der sowohl Einzelkomponenten als auch gesamte Fahrzeuge untersucht werden können. Aber wozu der ganze Aufwand? Die Absorberkegel und Kacheln an den Wänden?

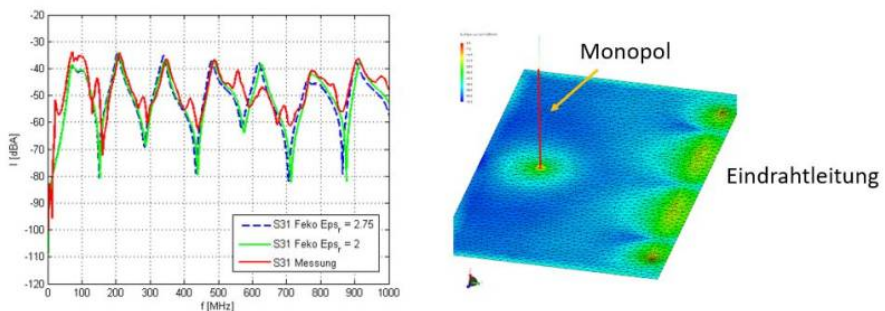
Wie immer geht es darum, reproduzierbar und frei von äußeren Einflüssen die

von einer Komponente abgestrahlten Emissionen zu erfassen. Es spricht nichts dagegen, die Emissionen auf der grünen Wiese zu messen, sogenannten Freifeldmessplätzen. Dazu sollte sich der Messplatz in einer „elektromagnetisch ruhigen Umgebung“ befinden bzw. es muss möglich sein, sicher zwischen externen Störgrößen und den vom Prüfling generierten Störungen zu unterscheiden. Freifeldmessplätze sind aufgrund der allgegenwärtigen elektronischen Systeme nur noch selten anzutreffen. Besser ist es natürlich, sich in eine elektrisch komplett geschlossene Kabine zu setzen, einen [Faradayschen Käfig](#), welcher von außen keine Störeinflüsse zulässt. Allerdings handelt man sich damit im Vergleich zum Freifeldmessplatz ein Problem ein: Die vom Prüfling emittierten Störgrößen werden an der Wand reflektiert und kommen zurück. Im schlimmsten Fall messen wir dann an der Antenne doppelte Störpegel oder es werden aufgrund von [Interferenzen](#) einzelne Frequenzen eliminiert. Es ist sogar möglich, dass sich durch die Addition der Wellen neue Frequenzanteile (Schwebungen) ausbilden, die in Wirklichkeit nicht vorhanden sind. Ein klarer Nachteil gegenüber Freifeldmessplätzen, bei denen die Emissionen einfach in der Umgebung verschwinden. Die gute Nachricht: Auch hierfür gibt es eine technische Lösung. Sogenannte Absorber, welche an den Wänden ringsum angebracht werden, verhindern, dass die elektromagnetischen Wellen zurückgeworfen werden. Der Name ist Programm. Die Wellen werden absorbiert und in Wärme umgesetzt. Je nach Frequenzbereich arbeiten die Absorber mal besser, mal schlechter. Deshalb werden meist verschiedene Materialien und Techniken eingesetzt, zum Beispiel Kegelabsorber oder Ferritkacheln. Weiterführende Infos zum Thema Absorber finden Sie auf der Homepage der Firma [Frankonia](#), einem Komplettanbieter für EMV-Messkabinen.

Jetzt könnte man ja auf die Idee kommen und fragen, warum wir Emissionsprobleme nicht einfach mit Feldberechnungsprogrammen simulieren und vorhersagen. Die mindestens Mittdreißiger unter uns wissen noch, wie ein Mobiltelefon mit einer Antenne aussieht. Der Entfall globiger Antennenstrukturen bei gleichbleibend hoher Übertragungsqualität ist nur gelungen über Geometrieoptimierungen mit Hilfe von Berechnungsprogrammen. Das funktioniert bei einem Handy deshalb so gut, da die Geometrie der Gesamtanordnung bekannt ist und sich eine exakte Untersuchung aufgrund der hohen Stückzahlen schnell bezahlt macht. Für die EMV-Messung ist es zwar möglich, Simulationen durchzuführen, allerdings nicht oder nur selten mit Absolutwerten. Die Ursache liegt im meist komplexen Aufbau und den unbekannten Materialien. Wer kennt schon die frequenzabhängige [Permittivität](#) der Leiterkarte oder die ebenfalls frequenzabhängige [Permeabilität](#) der Gehäusestruktur? Wir werden im Verlauf des Kapitels sehen, dass je größer der von uns betrachtete Frequenzbereich ist, wir die Strukturen umso genauer beschreiben müssen. Die Ursache liegt in der Wellenlänge unserer

Emissionssignale.

Ich möchte Ihnen ein einfaches Beispiel vorstellen, welches das Dilemma unbekannter Parameter beschreibt. Das Beispiel simuliert die Feldemission einer einfachen Leitung (Eindrahtleitung: ein einzelner Draht über einer Massefläche). Die Emissionen werden indirekt gemessen über eine Monopolantenne in einem Abstand von 1 m zur Leitung. Indirekt deshalb, da vereinfachend der Strom im sogenannten Fußpunkt der Antenne gemessen wird. Die Antenne wird im Fußpunkt (also am unteren Ende) mit $50\ \Omega$ gegen die Referenzmasse abgeschlossen. Genau in diesem Abschlusswiderstand wird der Strom gemessen. Da wir uns in einer geschirmten Kabine befinden und auch sonst keine elektromagnetischen Quellen vorhanden sind, können wir direkt darauf schließen, dass der Strom auf Feldemissionen der Eindrahtleitung zurückzuführen ist. Bei bekannter Übertragungsfunktion des Monopols (Zusammenhang Feld zu Fußpunktspannung) ist es natürlich möglich, auf die elektrische und magnetische Feldstärke zurückzurechnen.



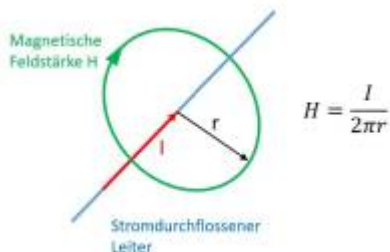
Das Bild zeigt auf der rechten Seite den generellen Aufbau der Gesamtanordnung in der Simulation. Zu sehen ist, wie das auf der **MoM-Methode** basierende Simulationsprogramm die Oberfläche in einzelne Teilprobleme unterteilt, das sogenannte Mesh (Gitternetz). Diese Methode wird besonders gerne für große (groß bezogen auf die Wellenlänge) Systeme angewendet wie Flugzeuge, Schiffe usw. Für jedes Netzsegment werden über alle betrachteten Frequenzen die Oberflächenströme entlang der Kanten berechnet. Damit lassen sich, wie im Bild gezeigt, die Amplituden (Spitzenwerte) über den Ort darstellen. Deutlich zu sehen sind die sich auf der Leitung ausbildenden stehenden Wellen (zu erkennen an der Stromverteilung unterhalb der Leitung).

Natürlich lassen sich damit nicht nur bunte Bilder erzeugen, sondern wie angekündigt die Emissionen der Leitung über der Frequenz darstellen. Die linke Seite zeigt den Strom im Antennenfußpunkt im Vergleich Messung vs. Simulation.

Jetzt stecken wir in einer weiteren Zwickmühle: Wie würden Sie die Kurve Messung zu Simulation interpretieren? Man könnte argumentieren: Wow, sehr gutes Ergebnis, die Amplituden stimmen oft überein, die Resonanzstellen werden getroffen, natürlich gibt es immer leichte Abweichungen zwischen Messung und Simulation. Falls Sie eher auf der pessimistischen Seite des Lebens unterwegs sind, würden Sie sagen: Da stimmt ja gar nichts. Die Amplituden stimmen so gut wie nie, die Resonanzstellen werden an den falschen Frequenzen berechnet. Wird die Grafik linear gezeichnet, wird das Ergebnis noch schlimmer bis unansehnlich. Wie so oft steckt die Wahrheit irgendwo in der Mitte. Man könnte auch vereinfachend sagen: Der Messaufbau entspricht nicht dem Simulationsaufbau. Und genau darin liegt das Problem. Die Simulation zeigt, dass bereits geringe Unsicherheiten in der angenommenen Permittivität der Leiterisolation die Resonanzstellen im höheren Frequenzbereich verschieben. Der Messaufbau benötigt natürlich Kontakte, um die Messgeräte anzuschließen, die Leitung hängt vielleicht etwas durch, alle geometrischen Messungen sind fehlerbehaftet und werden in der Simulation natürlich nicht berücksichtigt. Sie sehen, bereits kleine Änderungen im Aufbau oder die elektrischen Eigenschaften der eingesetzten Materialien haben einen Einfluss auf das Ergebnis der gestrahlten Emissionen. Nur durch einen definierten Aufbau (wo liegen die Leitungen) ist es möglich, annähernd reproduzierbare Ergebnisse zu erhalten.

9.1 Feldemissionen

Wir wollen hier noch einmal genau nachschauen, an welcher Stelle die Emissionen generiert werden. Alle Emissionsarten lassen sich aus den Maxwell'schen Gleichungen mit Hilfe der Vereinfachungen für stationäre und quasistationäre Felder herleiten. Quasistationäre Felder beschreiben langsam veränderliche Felder, welche wir später noch im Detail anschauen. Am einfachsten sind die **stationären magnetischen und elektrischen Felder** zu erklären.



Man stelle sich dazu einen beliebigen stromdurchflossenen Leiter vor. Das erzeugte magnetische Feld ist direkt proportional zum Strom und sinkt mit $1/r$ bei zunehmendem Abstand zum Leiter. Beispiel: Sie wollen mit einem Hall-Sensor das Erdmagnetfeld (Angabe meist in magnetischer Flussdichte B mit dem Zusammenhang $B = \mu_0 H$) von ca. $48 \mu\text{T}$ (Europa) messen. Mit dem einfachen Emissionsmodell ergibt sich, dass ein stromdurchflossener Leiter mit 20 A in ca. 10 cm Abstand bereits ein Magnetfeld in der Größe des Erdmagnetfeldes generiert. Im Abstand von 80 cm zum Leiter beeinflussen Sie das Messergebnis immerhin noch um ca. 10% .

Die Konsequenz daraus ist natürlich, dass wir mit jeglichem Stromfluss magnetische Felder generieren, die ggf. außerhalb unserer Geräte messbar sind. Auch hier gelten die im Abschnitt „Woher kommen die Störungen“ hergeleiteten Zusammenhänge periodischer Signale. Somit erzeugen rechteckförmige oder häufiger anzutreffende dreiecksförmige Stromverläufe eine unendliche Anzahl an Oberschwingungen und damit hochfrequente magnetische Emissionen.

Stationäre **elektrische Felder** sind immer dann anzutreffen, sobald eine Potenzialdifferenz vorhanden ist, worüber wir eine elektrische Spannung definieren. Theoretisch sind dazu nicht einmal leitfähige Medien im klassischen Sinne notwendig (Kupferleiter, Alu etc.), wie der berühmte Luftballonversuch, der uns die Haare zu Berge stehen lässt, verdeutlicht. Dadurch ist es auch möglich, elektronische Schaltungen durch reines Auflegen der Hand oder Berühren zu beeinflussen, was meist zu großem Erstaunen führt. Für den einfachsten Fall eines Plattenkondensators mit zwei unendlich ausgedehnten Platten ergibt sich die elektrische Feldstärke aus der bekannten Formel als Quotient der Spannung und des Abstands. Der Kapazitätswert ergibt sich als Funktion der aufgespannten Fläche und des Plattenabstands.

Das Gegenteil von statischen Feldern stellen Felder dar, deren Amplitude (Max.-Min.-Werte) über der Zeit variiert. Diese können in Form von periodischen Abläufen auftreten (der zeitliche Verlauf wiederholt sich mit der Periodendauer T) oder einmaligen langsamen Schalthandlungen. Rein stochastische Signale wie das natürliche Umgebungsrauschen werden hier nicht betrachtet. Übertragen wir dieses sich über die Zeit ändernde Signal mit Hilfe einer Leitung oder auf unserer Leiterkarte von einem Halbleiter zum nächsten, gehen wir selbstverständlich davon aus, dass entlang der Leitung zu jedem Zeitpunkt die identische Spannung anliegt. Sie ändert sich zwar mit der Zeit, allerdings nicht mit dem Ort der Messung. Der EMV- oder HF-Experte spricht dann von **quasistationären** Vorgängen, also zeitlich veränderbar, aber nicht ortsabhängig. Maxwell hat uns gezeigt, dass diese Grundannahme der Ortsunabhängigkeit elektrischer Signale nicht stimmt! Elektrische Signale besitzen nur eine endliche Ausbreitungsgeschwindigkeit, welche bei den folgenden Randbedingungen

berücksichtigt werden muss:

- sehr lange Übertragungswege
- sehr kurze Impulse

oder eine Kombination aus beiden Gesichtspunkten.

Bei diesen schnell veränderlichen Vorgängen ist es sogar möglich, dass sich von der elektrischen Struktur eine elektromagnetische Welle löst, welche sich in alle Raumrichtungen ausbreitet. Typischerweise handelt es sich bei diesen elektrischen Strukturen um Antennen oder Leitungsverbindungen. Prinzipiell ist es allerdings möglich, dass sich von jeglichem leitfähigen Gebilde (z.B. leitfähige Gehäuse, Leiterkarten, Steckverbinder usw.) eine elektromagnetische Welle ablösen kann. Die sich ausbreitende elektromagnetische Welle kann gemessen werden und unterliegt über einen weiten Frequenzbereich Grenzwerten, die eingehalten werden müssen. Mit Hilfe der Maxwellschen Gleichungen ist es möglich, die Ausbreitung der Welle auch mathematisch zu beschreiben – händisch allerdings nur begrenzt für symmetrische und mathematisch beschreibbare Strukturen. Komplexe Strukturen lassen sich über Methoden wie die Finite-Elemente-Methode oder die Method of Moments berechnen. Aufwändige Feldberechnungen kommen allerdings nur in sehr speziellen Fällen zum Einsatz, bei denen Systeme optimiert werden. Dass elektromagnetische Wellen nur bei sehr schnell veränderlichen Vorgängen, zum Beispiel bei hohen Frequenzen oder schnellen Einzelimpulsen (Blitz- oder ESD-Entladungen), auftreten, ist nur die halbe Wahrheit. Ob sich eine elektromagnetische Welle ausbilden kann, hängt von einem weiteren wichtigen Faktor ab, der

Wellenlänge.

Die Wellenlänge ist definiert als das Verhältnis der Ausbreitungsgeschwindigkeit zur Frequenz ($\lambda = c_0/f$). Allgemeiner wäre es, anstatt der Lichtgeschwindigkeit die tatsächliche Ausbreitungsgeschwindigkeit der Welle zu verwenden. Allerdings breiten sich unsere Wellen überwiegend in Luft aus, mit einer Geschwindigkeit nahe ($\approx 97\%$) der Lichtgeschwindigkeit.

9.1.1 Wann löst sich denn nun eine elektromagnetische Welle von meiner Struktur ab?

Diese Frage ist leider ebenfalls nicht eindeutig zu beantworten. Wie gut bzw. ab welcher Frequenz eine Struktur in der Lage ist, eine Welle zu emittieren, hängt von deren Geometrie, Materialien und dem Abstand zur Massefläche ab. Als Daumenregel kann der Faktor $\lambda/10$ herangezogen werden. Das bedeutet: Ist die betrachtete Struktur (Leitungslänge, PCB-Trace, Gehäuse usw.) größer als $\lambda/10$

der höchsten betrachteten Frequenz, muss von einer elektromagnetischen Emission bei dieser Frequenz ausgegangen werden. Natürlich nur, falls jemand diese Frequenz anregt (z.B. durch ein Rechtecksignal).

Beispiel 1: Die kleinste von Ihnen zu messende Frequenz beträgt 30 MHz. Bei einer Ausbreitung der Welle mit c_0 ergibt sich eine Wellenlänge von 10 m. Mit der Faustformel $\lambda/10$ erhalten wir eine kritische Länge von 1 m. Das bedeutet für uns, dass wir in unserem Blockschaltbild für alle Leitungsverbindungen > 1 m Maßnahmen zur Reduktion der gestrahlten Emission vorsehen müssen.

Beispiel 2: Wir starten in umgekehrter Reihenfolge und betrachten die geometrischen Abmessungen unseres Systems. Dabei stellen wir z.B. fest, dass die längste Leitungsverbindung 10 cm beträgt. Daraus folgt für uns, dass Wellenlängen von ca. 1 m, also Frequenzen ≥ 300 MHz, emittiert werden können.

Vielleicht haben Sie sich auch schon einmal darüber geärgert, dass günstige Produkte meist nur sehr kurze Anschlussleitungen (Netzkabel) besitzen. Neben der Materialeinsparung bringt die Verkürzung der Leitungen auch einen Vorteil bei der Betrachtung der gestrahlten Emissionen.

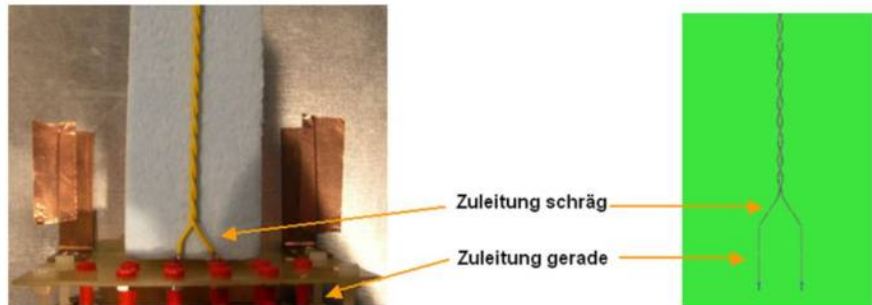
9.2 Antennenströme oder schon wieder Gleichtaktströme

Auch für die gestrahlte Emission macht es Sinn, die Störgrößen in Gleich- und Gegentaktsignale zu zerlegen. Die Ursache dafür ist ganz einfach. Schauen Sie noch einmal nach oben zur Entstehung magnetischer Felder entlang einer Leitung.

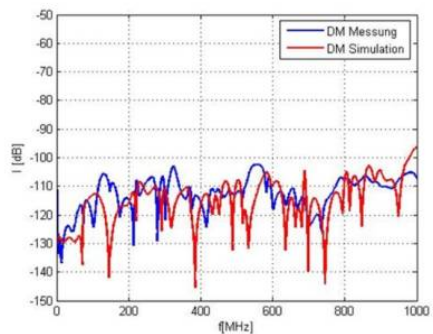
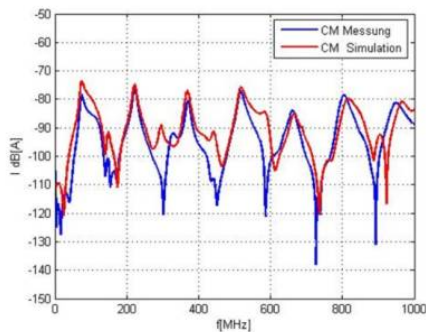
Im für uns idealen Fall liegt neben dieser „Hinleitung“ auch eine „Rückleitung“, welche den Rückstrom führt. Genau so haben wir das Gegentaktsignal definiert – unser Nutzsignal! Liegen diese beiden Leitungen nahe genug beisammen, löschen sich die entstehenden Magnetfelder nahezu aus. Je näher sie geometrisch beisammen sind, desto besser funktioniert die Auslöschung. Glücklicherweise funktioniert das nicht nur mit Magnetfeldern, sondern auch für elektromagnetische Wellen. Durch einen Trick lassen sich die Leitungen im Mittelwert geometrisch exakt an den identischen Ort bringen. Sicher sind Ihnen auch schon verdrehte Leitungen aufgefallen, typischerweise für Signalleitungen (z.B. CAN, FlexRay, ...).

Ich möchte Ihnen den Effekt verdrehter Leitungen mit Hilfe eines Experiments verdeutlichen. Wir nehmen zwei 2 m lange Leitungen, verdrehen diese genau 140-mal und bewerten erneut die Abstrahlung mit Hilfe eines Monopols (vgl. oben). Die Bewertung erfolgt erneut über den im Monopol messbaren Strom im

Antennenfußpunkt (bei 50 Ω Abschluss der Antenne). Durch die Verdrillung reduziert sich die Gesamtlänge auf 180 cm. Ursprünglich war das Experiment dazu gedacht zu zeigen, wie gut sich verdrehte Leitungen im Vergleich zur realen Messung simulieren lassen. Lassen Sie sich daher durch die beiden Messkurven nicht verwirren.



Die Ergebnisse sprechen für sich und es wird deutlich, warum Gleichtaktsignale häufig als Antennensignale bezeichnet werden.



Vergleichen Sie die beiden Bilder zuerst hinsichtlich der Amplituden. Deutlich zu sehen ist, wie die Amplituden auf der linken Seite über einen weiten Bereich um ca. 30 dB über der rechten Seite liegen (Achtung: negative dB-Angaben. Je negativer der Wert, desto kleiner die absolute Amplitude). Angeregt werden beide Konfigurationen mit identischen Signalamplituden. Lediglich die Phasenlage zwischen den beiden Leitern wird geändert.

Gleichtaktsignale $\rightarrow$ Phase zwischen den Leitern = 0° . Gegentaktsignale $\rightarrow$ Phase zwischen den Leitern = 180° .

Bei Gegentaktanregung (Differential Mode, DM) löschen sich die Emissionen nahezu vollständig aus. Die im Antennenfußpunkt messbaren Ströme liegen nur knapp über der Rauschgrenze. Daraus ergibt sich folgende goldene Regel:



Sorgen Sie dafür, dass zu jedem Stromfluss (Nutzstrom) möglichst geometrisch nah ein Rückstrompfad existiert. Wir treffen diese Regel später erneut unter dem Stichwort Kopplung.

Verdrillte Leitungen oder die Vermeidung von „aufgespannten Flächen“ durch nahe Hin- und Rückleiter helfen uns nicht nur, die Emissionen zu reduzieren – das Prinzip ist reversibel. Das bedeutet, Störgrößen entlang der Leitung können sich bei einer Störfestigkeitsprüfung (externe Felder wirken auf die Anordnung ein) nicht ausbilden bzw. löschen sich ebenfalls gegenseitig aus.

10.0 Vorhersage / Berechnung der gestrahlten Emission

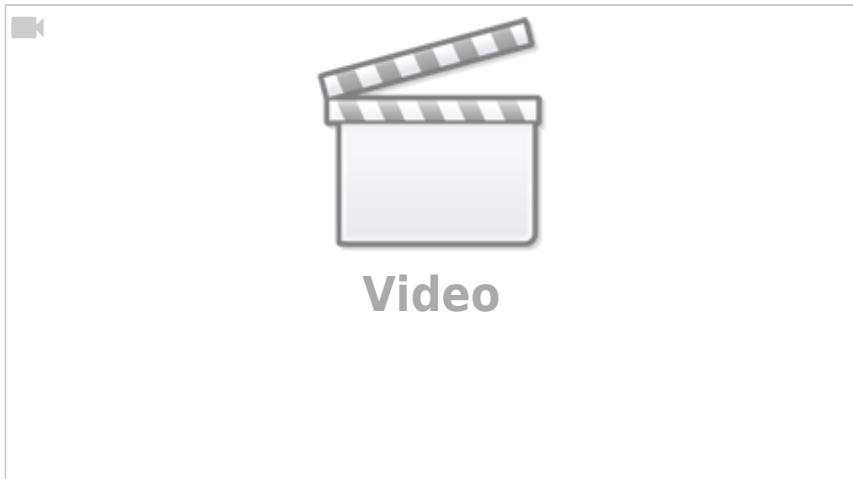


Sie können sich sicher vorstellen, dass die Vorhersage der Emissionen keine leichte Aufgabe darstellt. Wir haben ja bereits gesehen, dass selbst geringe Abweichungen im Simulationsaufbau das Ergebnis beeinflussen. Somit geht es in einer Simulation mehr darum, den Einfluss einzelner Maßnahmen und Resonanzstellen zu erkennen, als das Ergebnis exakt vorherzusagen. Im nachfolgenden Kapitel geht es darum, wie wir die gestrahlten Emissionen einer Anordnung modellieren können. Zugegeben, diesen Schritt werden Sie nur

gehen, falls Sie wissenschaftliches Interesse an dem Thema haben, aufgrund sehr hoher Stückzahlen das System optimieren oder weil Sie verzweifelt nach Lösungen für Ihr Emissionsproblem suchen. Trifft keiner der Punkte auf Sie zu, können Sie getrost weiterspringen zum Kapitel Maßnahmen.

Zur Nachbildung des gesamten Messaufbaus für die EMV-Simulation, entsprechend einer Komponentenmessung, ist es notwendig, jedes Teilsystem zu qualifizieren und in die Modulkette einzubeziehen. Es ist bekannt, dass der Kabelbaum bzw. die elektrische Verbindung zweier Teilsysteme einen signifikanten Einfluss auf die abgestrahlte Feldstärke hat und somit gesondert betrachtet werden muss. Das Ziel ist es, die Gesamtsimulation aus den einzelnen Teilsystemen modular aufzubauen, womit im Vorfeld die einzelnen Systeme getrennt voneinander bewertet und validiert werden.

Die internen Leitungsparameter bzw. die elektrischen Eigenschaften von Kabelbäumen können über die [Telegrafengleichungen](#) analytisch beschrieben werden oder mit Hilfe einer 3D-/2D-Simulation auf Basis gängiger Z-, Y- oder S-Parameter ausgedrückt werden. Für die Bewertung der vom Kabelbaum ausgehenden Abstrahlung ist es jedoch nur bedingt möglich, eine geschlossene Lösung zu finden. Eine genaue Aussage ist lediglich über eine 3D-Feldsimulation zu ermitteln, welche jedoch eine sehr genaue Modellierung des Gesamtaufbaus voraussetzt. Als Gesamtergebnis für das Teilsystem Kabelbaum wird im Folgenden versucht, Kabelbaummodelle zu erstellen, die sowohl die Leitungsparameter berücksichtigen als auch die vom Kabelbaum **ausgehende Abstrahlung**. Im ersten Schritt wird die Abstrahlung zu einem einfachen Simulationsmonopol berechnet und mit Hilfe von Messungen verifiziert. Die Modelle können direkt im Netzwerksimulator angewendet werden und stehen als SPICE-Modul in Netzlistenform zur Verfügung bzw. können dann direkt in LTspice verwendet werden. Um die reale Messumgebung nachzubilden, wird im zweiten Schritt versucht, die vereinfachte Simulation zum Monopol auf das reale Setup zu extrapolieren. Dieser Schritt erfolgt durch eine Kalibrationsmessung, welche die Umgebung sowie die verwendeten Messantennen berücksichtigt. Kommen Kabelbäume mit mehr als fünf Adern zum Einsatz, ergibt sich einerseits eine sehr hohe Unsicherheit im Vergleich zwischen Messung und Simulation und andererseits verschiedene Möglichkeiten der Anregung. Es zeigt sich, dass je nach Kabelbaumlage und interner Verdrillung ein Adernpaar mit maximaler sowie minimaler Abstrahlung existiert. Zur Beschreibung des statistischen Systems ist es somit notwendig, den Worst Case zu finden und in der Simulation anzuwenden.



10.1 Möglichkeiten zur Berechnung der Abstrahlung

Die Berechnung der Abstrahlung von Leiteranordnungen soll anhand zweier Verfahren erläutert werden. Verfahren I berechnet die Abstrahlung der Leiteranordnungen über die Hilfsgröße der Gleichtaktströme. Verfahren II bedient sich einer 3D-Feldsimulation der Gesamtanordnung auf eine Hilfsantenne. Die Berechnung mit Hilfe des Summenstromes geht vereinfacht davon aus, dass die Gesamtabstrahlung aus der Gleichtaktabstrahlung herrührt. Jegliche Differential-Mode-Abstrahlung wird vollständig vernachlässigt. Allerdings bedarf es keiner aufwändigen 3D-Simulation. Die einfachen leitungsgebundenen Modelle sind dafür ausreichend und dementsprechend schneller in der Anwendung. Dieses Verfahren eignet sich auch hervorragend für entwicklungsbegleitende Messungen, bei denen mithilfe einer Stromzangenmessung auf die gestrahlte Emission geschlossen werden kann. Als Nachteil ist zu sehen, dass natürlich nur die vom Kabelbaum ausgehenden Emissionen berücksichtigt werden können. Direktabstrahlungen, zum Beispiel vom Gehäuse, werden ignoriert.

10.1.1 Verfahren I, Abschätzung der Abstrahlung mit Hilfe des Summenstromes



Der Summen- oder Gleichtaktstrom entlang eines Kabelbaums lässt sich messtechnisch sehr einfach erfassen. Mit Hilfe einer Stromzange und eines Messempfängers/Spektrumanalysators lassen sich die Amplituden über der Frequenz ermitteln. In der Simulation können die Einzelströme entlang eines Kabelbaums einfach aufsummiert werden. Zu beachten ist jedoch, dass die Ströme phasenrichtig, also als komplexe Größen, behandelt werden. Vorhandene Gegentaktsignale löschen sich dann entsprechend der 180°-Phasenverschiebung bei der Summenbildung aus.

Die Berechnung der aus dem Gleichtaktstrom resultierenden Abstrahlung erfolgt über die Betrachtung eines [Hertzschen Dipols](#). Hertzsche Dipole sind infinitesimal kleine Antennen, welche in alle Raumrichtungen die identische Abstrahlung aufweisen. Clayton R. Paul leitet in seinem Buch [Introduction to EMC](#) im Kapitel Antennen sehr anschaulich den Zusammenhang zwischen den Gleichtaktströmen und der resultierenden Feldstärke her. Die Abstrahlung in V/m (sprich Volt pro Meter) ist demnach abhängig vom Gleichtaktstrom I_c , der Leiterlänge l und direkt proportional zur betrachteten Frequenz f . Mit zunehmendem Abstand d zum Kabelbaum sinkt die Feldstärke mit $1/r$.

$$\left| \hat{E}_{c,max} \right| = 6.28 \times 10^{-7} \frac{\left| \hat{I}_c \right| f l}{d}$$

Aus der Gleichung lassen sich wichtige Erkenntnisse zur Reduktion der gestrahlten Emission ableiten:

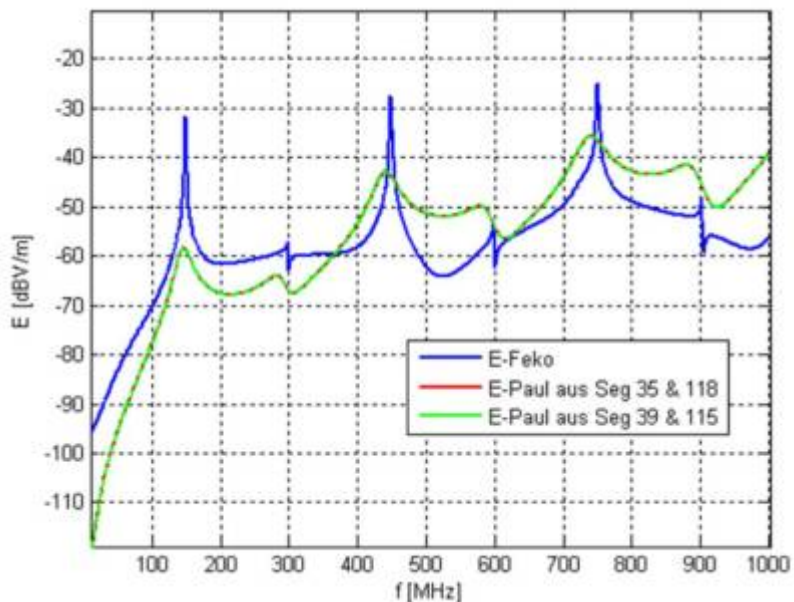


Zur Reduktion der gestrahlten Emission:



- Leiterlängen so kurz wie möglich
- Mit Hilfe von Filtern keine Störströme auf den Leitungen zulassen
- Hohe Frequenzen stets mit Hilfe von Tiefpassfiltern reduzieren!

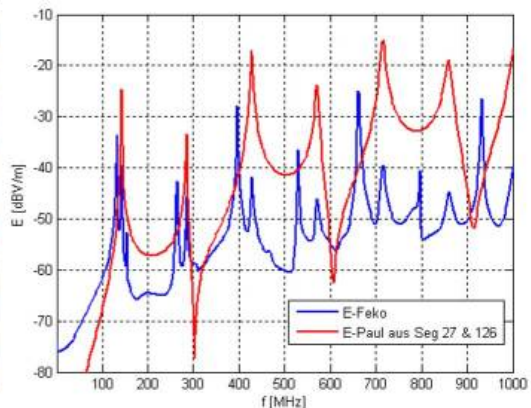
Das Verfahren kann dann sinnvoll angewendet werden, wenn keine leitfähige Massefläche vorhanden ist, entsprechend CE- bzw. FCC-Messverfahren. Zur Verifikation der gemachten Herleitung wird die Abstrahlung einer Leiterschleife mit Feko als Nahfeldberechnung bestimmt. Mit der Gleichung zur Feldberechnung wird aus den Segmentströmen die Common-Mode-Abstrahlung bestimmt. Entsprechend einer Messung nach EN 61000-6-3 wird die Leiterschleife 1 m über einer unendlich ausgedehnten Massefläche zur Nachbildung des Hallenbodens angeordnet. Bewertet wird die elektrische Feldstärke in einem Abstand von 3 m.



Die Abbildung zeigt die simulierte Nahfeldstärke aus Feko im Vergleich zur berechneten Feldstärke entsprechend unserer Gleichung. Zur Berechnung der abgestrahlten Feldstärke werden jeweils zwei benachbarte Segmentströme der Zweidrahtleitung phasenrichtig zu einem Summenstrom addiert. Die vereinfachte Berechnung erfasst den Grundpegel der abgestrahlten Feldstärke, hervorgerufen

durch den Strom in der Leiterschleife. Es zeigt sich auch, dass die Berechnung unabhängig von den verwendeten Segmenten (also an welcher Stelle entlang des Kabelbaums der Strom erfasst wird) funktioniert. Die absoluten Abweichungen sind besonders an den Resonanzstellen des Kabelbaums natürlich nicht zu vernachlässigen. Allerdings eignet sich das Verfahren somit für eine sehr schnelle Bewertung der Abstrahlung. Ein weiteres Einsatzgebiet besteht darin, bei vorhandenen Musterständen die Abstrahlung aus einer Summenstrommessung mit einer Stromzange zu extrapolieren.

Wird die entsprechend dem CISPR-Aufbau vorgegebene leitfähige Tischfläche mit in das Modell aufgenommen, zeigen sich deutlich größere Abweichungen zwischen simulierter und berechneter Abstrahlung der Leiteranordnung.

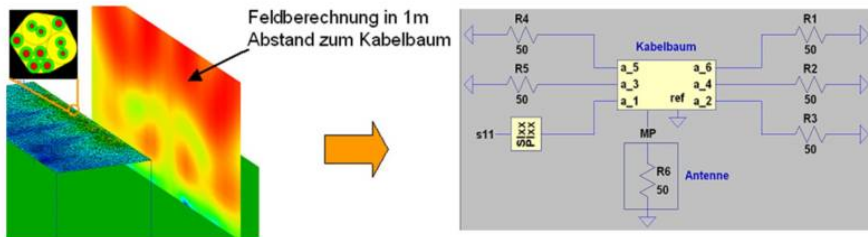


Eine genügende Übereinstimmung ist lediglich im Bereich bis ca. 400 MHz feststellbar. Oberhalb dieser Grenze ist das Ergebnis stark durch die Tischfläche geprägt, womit sich sehr hohe Abweichungen zur vereinfachten Berechnung ergeben. Die hohen Abweichungen waren zu erwarten, da die gesamte Berechnungsmethode auf Hertzschen Dipolen basiert, deren Felder sich ungehindert in alle Raumrichtungen ausbreiten können.

10.1.2 Verfahren II, Bewertung der Abstrahlung mit Hilfe eines Monopols als Hilfsmittel

Die Wunschvorstellung eines Anwenders der Simulationsmodelle besteht für den Kabelbaum aus einer Black Box, die neben der Funktionalität zusätzlich die

Abstrahlung berücksichtigt und als Bewertungskriterium ausgibt. Das Kabelbaummodell stellt als zusätzlichen Port die Abstrahlung entsprechend dem CISPR-Setup zur Verfügung.



Die generelle Vorgehensweise zur Erstellung der Kabelbaummodelle kann in drei Teilschritte untergliedert werden.

1. 3D-Feldsimulation der Leiteranordnung

Mit Hilfe einer Feldsimulation wird die Abstrahlung zu einem Hilfsmonopol berechnet. Die Bewertung erfolgt dabei über **Streuparameter**, welche die Beziehung zwischen den einzelnen Ports (Ports sind hier die Ein- und Ausgänge der Leitungsverbindungen und der Fußpunkt des Hilfsmonopols) beschreiben.

2. Überführung der Ergebnisse in die Netzwerksimulation

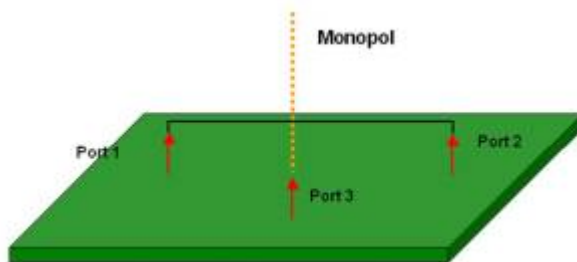
Die ermittelten Streuparameter aus Schritt 1 können in eine Netzliste umgewandelt werden. Im Folgenden wird die Software IdEM verwendet, welche mit Hilfe eines Vector-Fitting-Verfahrens die Streuparameter in äquivalente Netzlisten im SPICE-Format umwandelt. Die Schwierigkeit beim Vector Fitting besteht darin, sowohl stabile als auch passive Vielpole zu generieren.

3. Erweiterung der vereinfachten Simulation auf die reale Messumgebung

Mit Hilfe von Kalibriermessungen ist es möglich, die vereinfachte Simulation auf die reale Messumgebung zu extrapolieren. Dabei wird über einen Korrekturfaktor der Zusammenhang zwischen dem simulierten Monopol und der realen Empfangsantenne hergestellt. Diese Vorgehensweise wurde gewählt, um eine komplexe Simulation der Messantenne zu umgehen. Weiterhin ergibt sich bei der Vernachlässigung der Messumgebung in der Simulation eine enorme Reduktion der Simulationsdauer im Vergleich zur Gesamtsimulation mit CISPR-Tischfläche und Antennensetup.

Allerdings ist spätestens für eine Kalibriermessung der gesamte Messaufbau, analog der Messung nach Norm, notwendig. Dieser Schritt muss nur einmal durchgeführt werden, womit ggf. auch auf Referenzdaten zurückgegriffen werden kann.

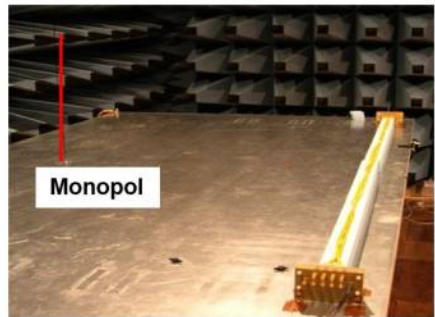
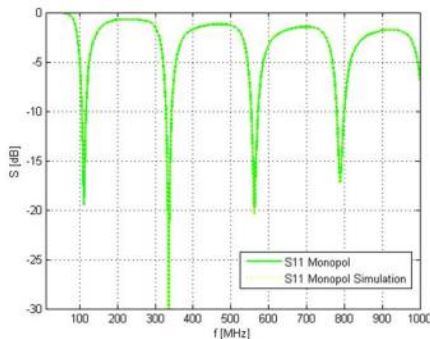
10.2 Schritt 1: Vorarbeiten



Zur Bewertung der Abstrahlung wird im Folgenden ein Hilfsmonopol in die Simulation mit eingefügt. Der Fußpunkt des Monopols wird dabei als zusätzlicher Port in die Gesamtmatrix der

Leiterparameter hinzugefügt. Grundlage der Berechnungen stellen die Streuparameter der Anordnung dar. Die Gesamtmatrix der Streuparameter erweitert sich mit betrachtetem Monopol und n Leitern auf $S = [(2n+1) \times (2n+1)]$. Die Fußpunktspannung am Monopol kann somit entsprechend einem beliebigen Port in der Netzwerksimulation einbezogen werden. Der Monopol wurde als bewertende Größe herangezogen, da er sowohl im Messaufbau als auch in der Simulation mit einfachen Mitteln generiert werden kann. Da es sich um einen einfachen Kupferleiter handelt, kann er selbst für hohe Frequenzen sehr gut nachgebildet werden. Als Nachteil ist zu sehen, dass die vom Monopol charakteristischen Resonanzen mit in das Gesamtergebnis eingehen. Die Abbildung zeigt den Gesamtaufbau der Anordnung mit Portbelegung. Im Aufbau ist der Monopol 1 m vom betrachteten Kabelbaum entfernt. Die Länge des betrachteten Kabelbaums beträgt 2 m. Zugegeben, im folgenden Beispiel wird die Messmethode nach CISPR herangezogen, welche stets von einer leitfähigen Tischfläche ausgeht. Für Aufbauten nach der CE- oder FCC-Kennzeichnung mit einer nicht leitenden Unterlage ist die Integration nicht unmöglich, aber deutlich schwieriger. Über eine in den Tisch eingelassene N-Buchse ([allg. HF-Steckverbinder](#)) haben wir nicht nur ein sehr gutes Massepotenzial, sondern auch eine stabile Anbindung.

Nachfolgende Abbildung zeigt den Gesamtaufbau zur messtechnischen Verifikation sowie die simulierte und gemessene Eingangsreflexion des Hilfsmonopols. Messung und Simulation stimmen deutlich überein, sobald die Materialparameter und Geometrie in der Simulation berücksichtigt werden.



Aus rein statischen Verhältnissen wurde der Monopol mit einer Gesamtlänge von 0,66 m und einem Radius von 1,1 mm gewählt. Über die N-Buchse kann der Monopol von unten kontaktiert werden. Die Eingangsreflexion des Monopols zeigt die charakteristischen Resonanzstellen entsprechend der Leiterlänge mit $\lambda/4$, $3\lambda/4$, $5\lambda/4$ und $7\lambda/4$.

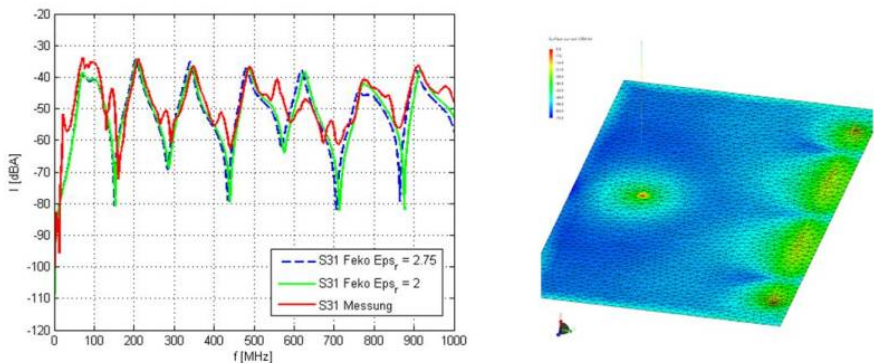
10.2.1 Validierung mit einfacher Eindrahtleitung und Doppelleitung

Zur Überprüfung des Verfahrens wird im ersten Schritt eine einfache Eindrahtleitung aus der Simulation in die Netzwerkkumgebung überführt. Im zweiten Schritt wird die Simulation einer Doppelleitung mit dem messtechnischen Aufbau verglichen. Das Ergebnis der Eindrahtleitung haben wir bereits kennengelernt. Wie gesehen gelingt es uns nur, die Simulation der Realität anzunähern, falls wir alle eingesetzten Materialparameter kennen. Kopfzerbrechen bereiten meist die frequenzabhängigen Werte des $\tan\delta$ und ϵ_r des Leitungsmantels.

Die in der Simulation und Messung generierten Streuparameter werden mit Hilfe von **IDEM** in ein äquivalentes Netzwerkmodell umgewandelt. IdEM ist ein auf Matlab basierendes Tool, welches uns aus den gemessenen Streuparametern Netzwerkmodelle für einen SPICE-Löser generiert. Das bedeutet, wir müssen nur einmal die S-Parameter erfassen und können dann alle weiteren Simulationen z.B. in LT-Spice durchführen. Da wir den Monopol als Port definiert haben, können wir so bei beliebiger Anregung der restlichen Ports die Fußpunktspannung am Monopol simulieren. Die Sache hat nur einen Haken: Die erstellten SPICE-Netzlisten sind voll mit aktiven Spannungsquellen. Das bedeutet, die Kunst liegt darin, passive Modelle zu generieren, die zu keinem Zeitpunkt Energie erzeugen. Das macht bei einem Kabelbaum durchaus Sinn. Leichter gesagt als getan, da besonders im Bereich sehr kleiner Pegel, also im Rauschen, stets undefinierte Pegel erfasst werden, welche IdEM dann als aktives Netzwerk interpretieren kann.

Die Entwickler von IdEM bieten Lösungen an, die Netzwerke aktiv in den passiven Bereich zu drücken. Bitte testen Sie trotzdem jedes generierte Netzwerk mit einer Sprungfunktion an einem der Ports sowie auf den DC-Arbeitspunkt.

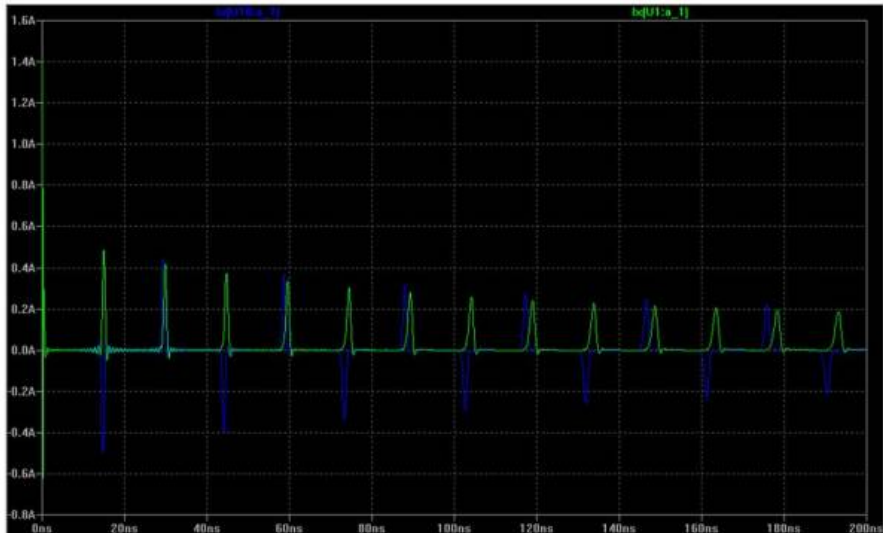
Die Auswertung erfolgt im Netzwerksimulator, wobei für die Abstrahlung die Fußpunktspannung am Monopol oder der entsprechende Fußpunktstrom bewertet wird.



Die Abbildung zeigt den Vergleich zwischen gemessenem und simuliertem Fußpunktstrom der Monopolantenne. Dabei wird ein Leiterende der Eindrahtleitung mit einer AC-Quelle mit 1 V angeregt, wobei das andere Leiterende mit 50 Ω abgeschlossen ist. Das Simulationsmodell als auch die Messdaten werden in die Netzwerksimulation überführt und verglichen. Es zeigt sich weiterhin der Einfluss des Materialparameters ϵ_r . Die Dielektrizität ist normalerweise nicht im Detail bekannt bzw. kann nur abgeschätzt werden, da sie sowohl von der Frequenz als auch von der Temperatur abhängig ist. Bereits die einfache Eindrahtleitung zeigt, dass eine Abweichung zwischen Messung und Simulation nicht zu vermeiden ist. Die maximale Abstrahlung der Leiteranordnung wird über dem gesamten Frequenzbereich relativ gut getroffen, mit nur einer wesentlichen Abweichung bei ca. 620 MHz mit 10 dB Abweichung. An den unteren Resonanzstellen ist eine deutlich größere Abweichung zu erkennen, die neben der allgemeinen Unsicherheit aus der Auflösung der Messergebnisse herrührt. Bei der Messung steht eine Resonanzauflösung von ca. 0,6 MHz zur Verfügung (maximale Frequenzauflösung des verwendeten Agilent ENA = 1601 Punkte).

Ob die generierten Modelle passiv sind, lässt sich am einfachsten in einer Zeitbereichssimulation überprüfen. Diese Überprüfung ist sehr wichtig, da aktive Modelle eine funktionale Simulation komplett verfälschen würden. Dazu regen wir unser Modell mit einem sehr breitbandigen Impuls an, idealerweise einem Dirac-Impuls, und überprüfen, ob sich Frequenzanteile ausbilden, die nicht mehr

abklingen. Echte Dirac-Impulse kann auch eine Simulation meist nicht berechnen, aber es ist möglich, sehr hohe Flankenanstiegs- und Abfallzeiten zu verwenden. Nachfolgende Abbildung zeigt das Ergebnis, wenn wir unser Leitungsmodell mit einem schnellen Impuls anregen, bei verschiedenen Abschlüssen am Leitungsende.



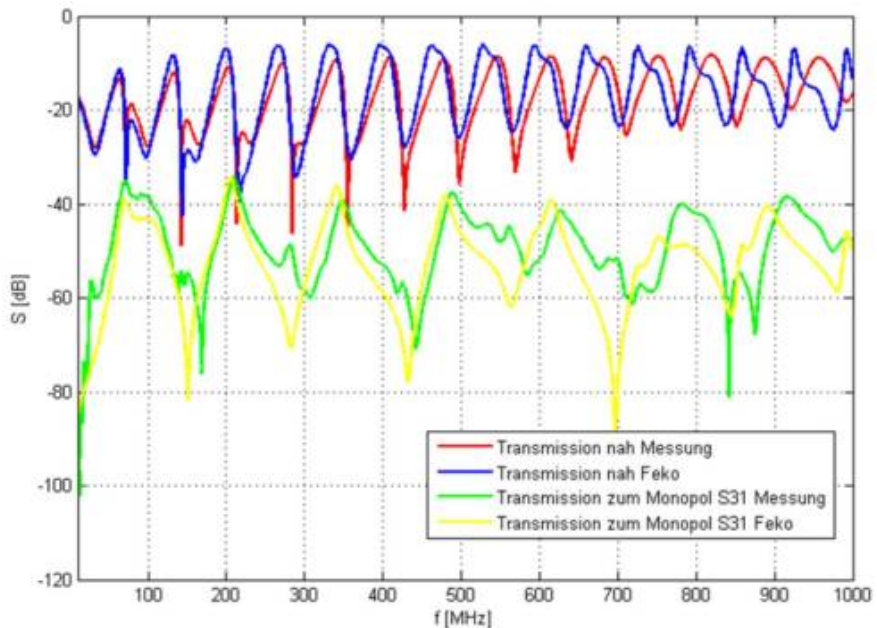
Dargestellt ist jeweils der Eingangsstrom in die Leitung. Bei der grünen Kurve ist das Leitungsende offen, bei der blauen Kurve mit einem in der Simulation als ideal angenommenen 2700 μF Kondensator gegen Masse abgeschlossen. Deutlich zu sehen sind die durch die Fehlanpassung hervorgerufenen Reflexionen auf der Leitung, wobei sich das System in keinem Fall aufschwingt. Das Ergebnis ist entsprechend positiv, das bedeutet, wir können die erzeugten Modelle für weitere funktionale Berechnungen verwenden.

10.2.2 Mess- / Simulationsaufbau mit einer Doppelleitung

Als Erweiterung der Eindrahtleitung wird eine Doppelleitung in Simulation und Messung miteinander verglichen. Die Messung der Streuparameter muss in zwei Schritten erfolgen, da aufgrund der vier Leitungsanschlüsse der Doppelleitung und des Monopols eine 5×5 -Matrix entsteht, typische Netzwerkanalysatoren allerdings nur mit vier Ports ausgestattet sind.

Wir validieren erneut zuerst unsere Simulation mit einer Messung im Labor. Dabei

greifen wir uns nur einzelne Streuparameter aus der Matrix mit 25 Einträgen heraus.



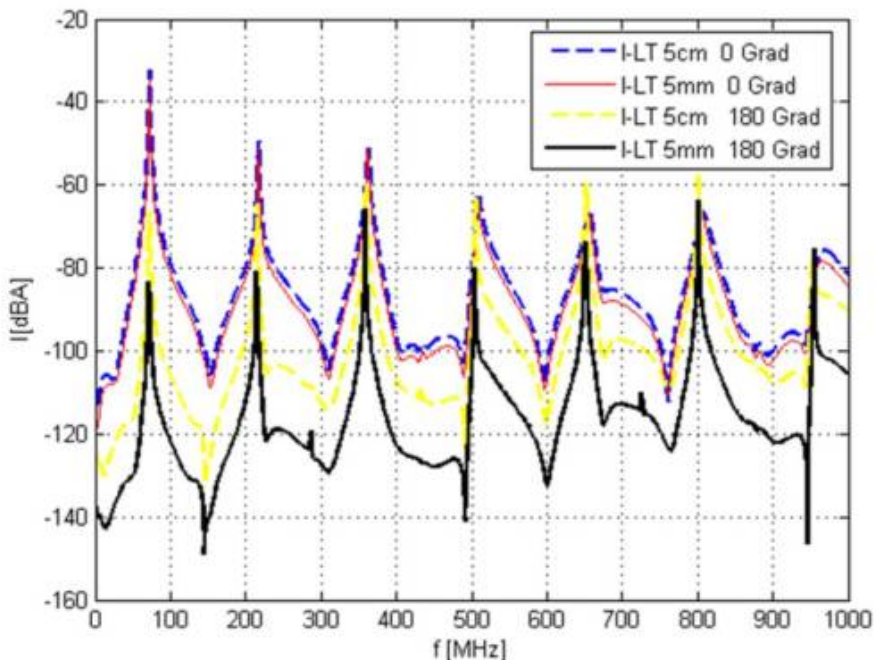
Die Grafik zeigt das Übersprechen der Eingangsports der beiden Adern sowie die Transmission eines Eingangsports zum Monopol. Der Grundpegel der Abstrahlung wird erfasst, wobei mit steigender Frequenz die Abweichung vom Simulationsmodell zur Messung zunimmt. Es ist deutlich zu sehen, dass die Abweichungen in der Abstrahlung auch in den leitungsgebundenen Parametern zu erkennen sind. Auch hier nimmt mit steigender Frequenz die Abweichung zu. Die Genauigkeit der Modelle zur Abstrahlung kann somit auch aus den Abweichungen der Leitungsmodelle zu den gemessenen Daten extrapoliert werden. Aus Kapitel 9.1 ist bekannt, dass die Abstrahlung der Leiteranordnungen in Zusammenhang mit den auf den Leitern fließenden Strömen steht. Stimmen die vom Modell zur Verfügung gestellten Eingangsimpedanzen nicht, ist es natürlich nicht möglich, die Abstrahlung exakt zu berechnen. Das gleiche Ergebnis resultiert aus der Betrachtung der Eingangsimpedanz bzw. Eingangsreflexion der Leiter in Messung und Simulation.

Bedeutung der Common- und Differential-Mode-Abstrahlung

Die Bewertung und Validierung der Modelle mittels Streuparameter beinhaltet

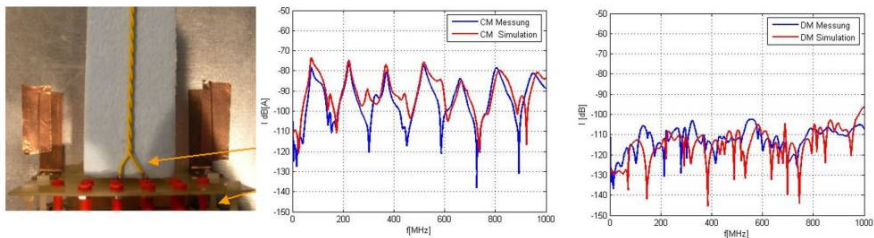
den Nachteil einer rein nodalen (Port gegen Bezugsmasse) Auswertung. Es wird somit lediglich der Common-Mode (Gleichtaktfall) abgedeckt. Die Gegentaktabstrahlung hingegen ist aufgrund der dabei auftretenden Feldauslöschungen deutlich schwieriger zu beschreiben. Zum einen sind die betrachteten Störpegel deutlich geringer, zum anderen darf die Phaseninformation der Parameter nicht vernachlässigt werden. Es ist bekannt, dass für Kabelbäume mit einer geringen Anzahl an Einzeladern die Gegentaktabstrahlung zur Gleichtaktabstrahlung nicht zu vernachlässigen ist. Für Kabelbäume mit einer hohen Anzahl an Einzeladern entsteht bereits im unteren Frequenzbereich (ab wenigen MHz) eine Modenwandlung des differentiellen Eingangssignals, sodass die Abstrahlung als Gleichtaktabstrahlung auftritt.

Zur Verdeutlichung der Gegentaktabstrahlung werden zwei Doppelleitungskonfigurationen miteinander verglichen. Dabei wird der Leiterabstand s in der Simulation von 5 cm auf 5 mm reduziert. Die Auswertung erfolgt im Netzwerksimulator, nachdem die Modelle aus den berechneten Streuparametern erstellt wurden. Durch die Reduktion des Leiterabstandes reduziert sich bei gegenphasiger Anregung der Leiter die Abstrahlung aufgrund der Feldauslöschung.



Betrachtet man die Gleichtaktanregung der beiden Leiterkonfigurationen, ist zu erkennen, dass sich der reduzierte Leiterabstand nicht auf die Abstrahlung auswirkt (blau vs. rot). Bewertet wird auch hier wieder der Fußpunktstrom des Monopols an 50 Ω . Die Gegentaktabstrahlung hingegen wird um bis zu 20 dB reduziert (schwarz vs. gelb). Generell ist zu beobachten, dass die Gegentaktabstrahlung unterhalb der Gleichtaktabstrahlung liegt bei identischer Anregung (blau vs. gelb, rot vs. schwarz).

Twisted-Pair-Leitungen werden häufig eingesetzt, um den Effekt der Feldauslöschung noch weiter zu verstärken. Eine Gesamtauslöschung der Abstrahlung für differenzielle Anregung würde sich ergeben, wenn beide Leiter exakt an der geometrisch gleichen Stelle verlaufen. Dies ist allerdings physikalisch nicht möglich, womit eine Annäherung über verdrehte Doppelleitungen erfolgt. Je höher der Verdrillungsgrad (Schlaglänge) der Leiter, desto besser wirkt die differenzielle Feldauslöschung. Im Versuchsaufbau zu Twisted-Pair-Leitungen werden zwei 2 m lange Leitungen mit insgesamt 140 Schlägen verdreht, was zu einer Twisted-Pair-Länge von 180 cm führt.



Bei der Modellierung der Leitungen ist zu beachten, dass die Zuleitungen einen erheblichen Einfluss auf das Gesamtergebnis haben. Sie dürfen in der Simulation nicht vernachlässigt werden. Die Auswertung erfolgt erneut im Netzwerksimulator (SPICE). Deutlich zu sehen ist die enorme Reduktion der Gegentaktabstrahlung aufgrund der Verdrillung beider Adernpaare. Wie bereits bei der Eindrahtleitung gesehen, stimmt die maximale Abstrahlung bei Gleichtaktanregung im Vergleich zwischen Messung und Simulation sehr gut überein. Abweichungen ergeben sich erneut aufgrund fehlerhafter Materialparameter (ϵ_r : Resonanzverschiebung) sowie an den Resonanzstellen. Bei der Gegentaktabstrahlung wird lediglich der Grundpegel erfasst. Eine Übereinstimmung zwischen Messung und Simulation ist nicht mehr gegeben. Der vorhandene Teil der Gegentaktabstrahlung wird hauptsächlich durch die Zuleitungen zum Messmodell hervorgerufen. Im Simulationsmodell wird zwar die Zuleitung berücksichtigt, allerdings ist es kaum möglich, diese (Winkel, exakte Länge) detailliert darzustellen. Eine weitere

natürliche Grenze in Bezug auf die Genauigkeit zwischen Messung und Simulation stellt das Rauschlevel des verwendeten VNA (Vector Network Analyzer) dar. Zusätzlich geht die Simulation von einer idealen, homogenen Verdrillung entlang der gesamten Leiterlänge aus. Im Messaufbau kann dies aufgrund der von Hand gefertigten Leitungen kaum erreicht werden, wodurch sich eine weitere Unsicherheit in Bezug auf die Gleichtaktabstrahlung ergibt.

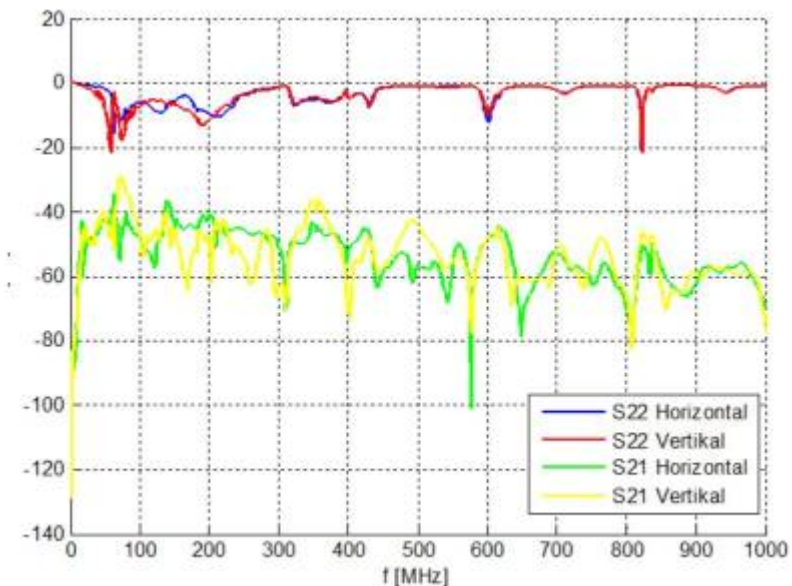
10.3 Schritt 2: Modellerstellung

Wir haben uns in Schritt 1 Gedanken darüber gemacht, wie wir den komplexen Aufbau aus Leiterverbindungen zu einer einfachen Monopolantenne in Simulation und Messung charakterisieren können. Jetzt geht es darum, die erfassten Daten in eine Netzwerksimulation zu überführen. Wir sind dann in der Lage, den Kabelbaum und die bewertende Antenne wie ein herkömmliches Bauteil in die Analogsimulation einzubinden. Wie bereits besprochen, benötigen wir dazu ein mathematisches Verfahren, das uns die ermittelten Streuparameter in ein äquivalentes Netzwerkmodell umwandelt. Gute Ergebnisse wurden bisher mit der Software IdEM erzielt, einem auf Matlab basierenden Vector-Fitting-Tool.

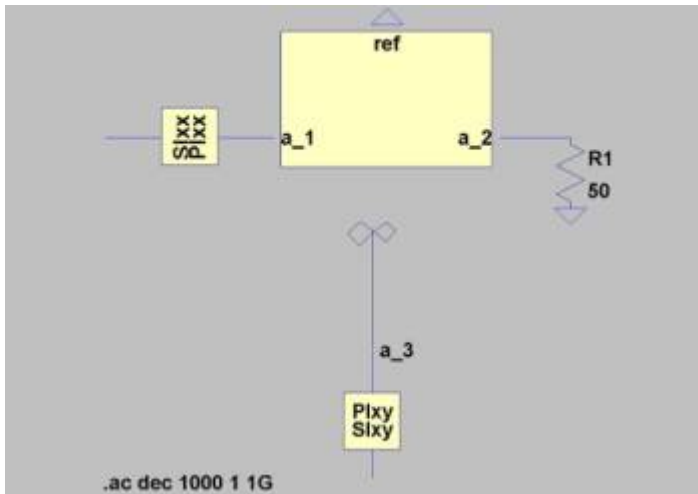
10.4 Schritt 3: Umrechnung auf die Laborantenne

Wie bereits erwähnt, ist es notwendig, dass die Modelle die Abstrahlung zu den vorhandenen Laborantennen berücksichtigen. Bisher wurde die Abstrahlung lediglich zu einer Hilfseinrichtung, dem Monopol, betrachtet. Als Messmittel deckt der Monopol jedoch nur einen geringen Teil des benötigten Frequenzbereichs ab und das auch nur für CISPR-Messungen in der automotive Welt, womit eine Übertragung auf bikonische oder logarithmisch-periodische Antennen notwendig wird. Als grundlegende Untersuchung wird im ersten Schritt versucht zu zeigen, dass es prinzipiell möglich ist, Modelle zu generieren, die das Verhalten zu den Laborantennen beschreiben. Dazu wird eine Eindrahtleitung in Verbindung mit der bikonischen Antenne untersucht. Der Eindraht wird auf dem CISPR-Tisch in der Absorberhalle aufgebaut. Der Abstand zur Tischkante beträgt 10 cm, wobei die Leitung 5 cm oberhalb des Tisches fixiert wird. Die Antenne wird entsprechend CISPR25 mit einem Abstand von einem Meter zur Tischkante positioniert. Die im Folgenden erstellten Modelle werden aus Messdaten generiert, wobei die Auswertung erneut im Netzwerksimulator (LTspice) erfolgt. Die Modellerstellung erfolgt mit IdEM. Zunächst werden die Eingangsreflexionsfaktoren der bikonischen Antenne und die Transmissionsfaktoren vom Eindraht auf die bikonische Antenne für beide

Polarisationsarten der Antenne miteinander verglichen. Die verwendete bikonische Antenne ist spezifiziert in einem Frequenzbereich von 20–230 MHz, welches sich auch in der Anpassung der Antenne erkennen lässt. Die Polarisation und damit die geometrischen Verhältnisse um die Antenne zeigen einen geringen Einfluss auf die Anpassung der Antenne. Obwohl die Antenne nur bis 230 MHz spezifiziert ist, wird sie im folgenden Versuch bis 1 GHz verwendet. Es geht im ersten Schritt nicht darum, die gesamte Kette der Messtechnik nachzubilden, sondern darum, die Anwendbarkeit des Verfahrens aufzuzeigen.



Wir sehen, dass die Polarisation der Antenne teils deutlichen Einfluss auf die Transmission vom Eindraht zur Antenne zeigt (S21 Transmission Leitung → Antenne). Für jede Polarisation der Antenne ist es nun möglich, aus den Messdaten ein äquivalentes Netzwerkmodell zu erstellen und in der SPICE-Umgebung anzuwenden.

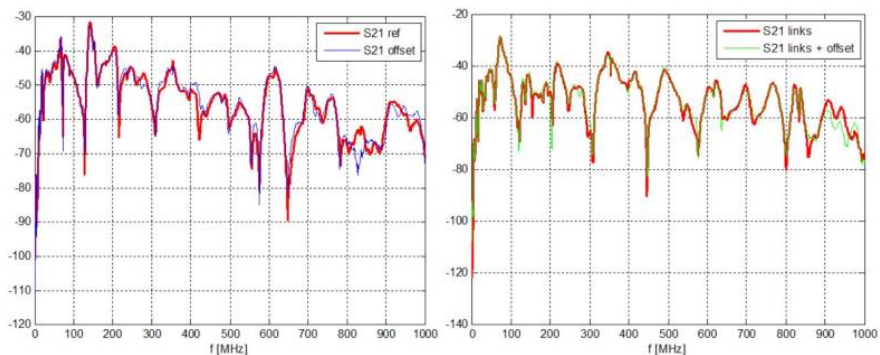


Der LT-Spice-Screenshot zeigt beispielhaft das Modell für horizontale Polarisation. Dabei ist der zusätzliche Anschlusspunkt für die bikonische Antenne vorhanden. Der Frequenzbereich der Modelle ergibt sich aus den Einstellungen des Netzwerkanalysators mit 0,3–1000 MHz. Streuparameter lassen sich in einer Netzwerksimulation komfortabel berechnen. Dazu sind allerdings zwei unterschiedliche Modelle notwendig. Das anregende Modell beinhaltet eine über der Frequenz variable Spannungsquelle mit konstanter Amplitude, welche die zu untersuchende Schaltung über einen 50-Ω-Innenwiderstand anregt. Die zweite Aufgabe des anregenden Modells besteht darin, die Eingangsreflexion am anregenden Port zu ermitteln. Das Transmissionsmodell (passiver Subcircuit) ermittelt jeweils den Sxy-Parameter, also die Transmission vom anregenden Port zum betrachteten Port. Wird LT-Spice verwendet, müssen die Unternetzwerke (Subcircuits) im gleichen Dateipfad vorhanden sein wie die Simulationsdatei. Im Zip-Download finden Sie die Netzwerke und ein Beispiel zur Transmissionsberechnung. Die Streuparameterberechnung findet immer im Frequenzbereich statt und gibt das Ergebnis der Parameter im Wertebereich zwischen 0 und 1. Bitte achten Sie auf die richtige Anschlussbezeichnung (Anschlussrichtung) der Ports.

test_reflection_transmission.zip

10.4.1 Reproduzierbarkeit

Ein wichtiger Punkt bei der Beurteilung der Abstrahlung ist die Reproduzierbarkeit der vorgenommenen Messungen. Als Voruntersuchung wird die Abweichung in der Transmission zur Antenne untersucht, wobei eine Verschiebung der bikonischen Antenne um 5 cm in Richtung des Kabelbaums erfolgt. Es ist somit mit einer leichten Erhöhung bzw. höheren Kopplung der beiden Systeme zu rechnen.



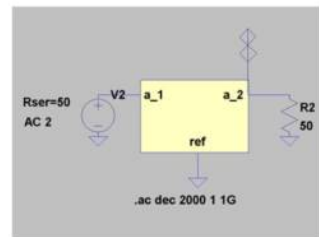
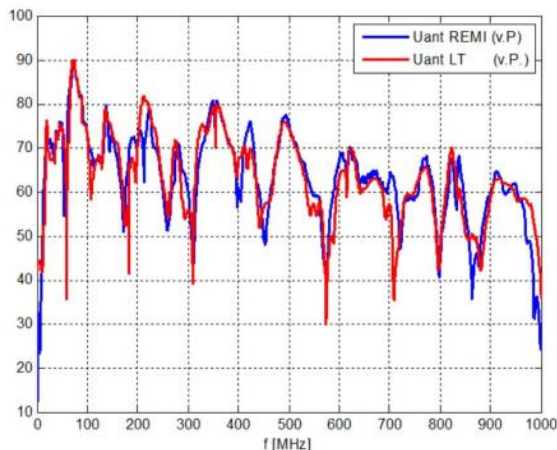
Die Abbildung zeigt die Transmission des Kabelbaums zur bikonischen Antenne (Eingang Kabelbaum → Antenne), rechts für eine vertikale Polarisation der Antenne, links für horizontale Polarisation. Wir sehen, eine Abweichung von 5 cm zeigt kaum Einfluss auf die Transmission zur bikonischen Antenne. Die Ergebnisse sind nahezu identisch. Es ist somit von einer sehr hohen Reproduzierbarkeit auszugehen, selbst bei einer leicht fehlerhaften Positionierung der Empfangsantenne.

10.4.2 Spielen I

Suchen wir uns unsere erste Anwendung für die gefundene Lösung, eine gestrahlte Emission komfortabel in LT-Spice zu berechnen. Dazu nehmen wir einen Frequenzgenerator, oder früher auch als **Wobbelgenerator** bezeichnet, der in der Lage ist, über einen weiten Frequenzbereich ein konstantes Ausgangssignal in der Amplitude auszugeben. Der Generator durchläuft die Frequenzen periodisch in einem einstellbaren Bereich. Beispiel: Starte bei 1 kHz, erhöhe jede 100 ms um 100 kHz bis 50 MHz, danach von vorne. Diese Funktion können natürlich heute auch Netzwerkanalysatoren übernehmen, das Bild in

Wikipedia zeigt ein sehr altes Gerät. Der Wobbelgenerator regt unseren Kabelbaum an, ok, bisher ist es nur eine Ader, wobei wir gleichzeitig mit dem Messempfänger an der Antenne die erzeugten Emissionen messen, analog einer Abnahmemessung.

Mit dem zuvor in LT-Spice erstellten Modell machen wir jetzt genau das Gleiche. Wir regen die Leitung mit der identischen Spannung wie der Wobbelgenerator an, 2 V, und berechnen nun die an der Antenne auftretende Spannung. Um mit Hilfe des Antennenfaktors dann noch die Feldstärke auszugeben, wäre nur noch ein letzter mathematischer Schritt notwendig, welchen sowohl der Messempfänger als auch die Simulation machen müssten. Deshalb verzichte ich an dieser Stelle darauf, da wir keine weitere Erkenntnis gewinnen würden. Es ändert sich lediglich die Achsenbeschriftung.



Der Vergleich beider Kennlinien zeigt eine gute bis sehr gute Übereinstimmung der beiden Spannungen im Antennenfußpunkt. Das bedeutet, für dieses einfache Beispiel sind wir in der Lage, die Emission der Leitung sehr genau vorherzusagen.

10.4.3 Spielen II

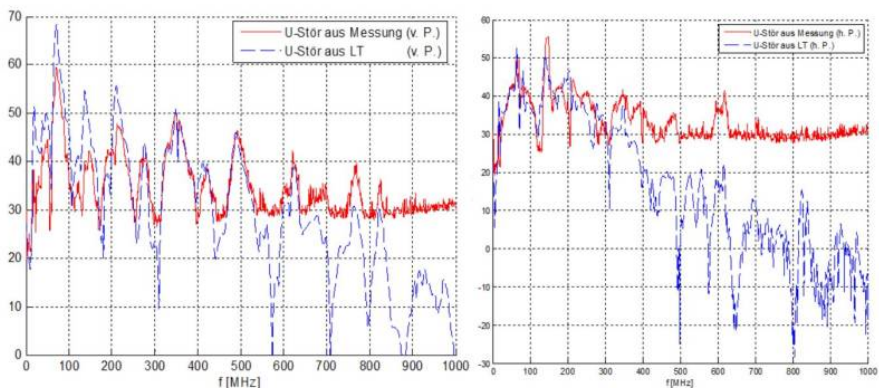
Es lässt sich zu Recht anmerken, dass sich im vorangegangenen Beispiel sowohl die Quelle als auch die Empfangseinheit im komfortablen Frequenzbereich aufhalten bzw. wir nur mit rein sinusförmigen Größen arbeiten. Deshalb gehen wir einen Schritt weiter und regen die Leitung jetzt mit einer hochfrequenten Impulsquelle an, dem IGUF 2910 von der Firma Schwarzbeck. Die

Fußpunktspannung wird dann über der Frequenz an der bikonischen Antenne mit dem Messempfänger aufgezeichnet.

Um die Daten aus der Simulation mit der Messung vergleichen zu können, ist hier etwas mehr Aufwand notwendig. Der schnelle Impuls des Generators muss erst einmal nachgebildet werden, um ihn in LT-Spice einzubinden. Der gesamte Messablauf läuft jetzt wie folgt ab:

1. Beschreibung der Pulsquelle im Zeitbereich mittels Oszilloskop
2. Import der Zeitbereichsdaten aus 1.) in Spice und Anregung des Modells
3. FFT der Fußpunktspannung an der Antenne in Spice
4. Messung der Abstrahlung mit angeschlossener Leitung
5. Vergleich zwischen Messung und Simulation der Funkstörspannung

Um ein Modell der anregenden Quelle zu erhalten, wird das Signal im Zeitbereich unter Schritt 1. mit dem Oszilloskop aufgezeichnet. Die als ASCII-Datei zur Verfügung stehenden Daten werden direkt in Spice als Spannungsquelle eingebunden. Für die Messung wird der Pulsgeber an die Leitung angeschlossen und die Abstrahlung auf die bikonische Antenne mit dem Messempfänger gemessen, Schritt 4. Bei der Durchführung der Simulation in Schritt 2. und 3. wird das bekannte, mit dem VNA gemessene Modell der Leitung mit der bikonischen Antenne verwendet. Die am Antennenfußpunkt simulierte Spannung wird mittels FFT in den Frequenzbereich überführt und in Schritt 5. mit den gemessenen Werten aus Schritt 4. verglichen. Damit erhalten wir erneut die komplette Nachbildung des Versuchsaufbaus im Netzwerksimulator. Jetzt allerdings arbeiten wir wie später vermutlich auch bei funktionalen Simulationen zuerst im Zeitbereich und transferieren das Ergebnis dann in den Frequenzbereich.



Gemessene vs. simulierte Abstrahlung bei Pulsanregung. Links: horizontale Polarisation, rechts: vertikale Polarisation

Die Abbildung zeigt den Vergleich zwischen Messung und Simulation für die Pulsanregung. Für die vertikale Polarisation ergibt sich eine deutlich höhere maximale Abstrahlung als für die horizontale Polarisation. Die Rauschgrenze des Messempfängers liegt hier aufgrund der verwendeten Bandbreite bei ca. 30 dB μ V. Im Bereich messbarer Pegel zeigen sich jetzt deutlich höhere Abweichungen in Messung und Simulation, teilweise für absolute Aussagen nicht mehr akzeptabel, für vergleichende Bewertungen im annehmbaren Bereich. Die Gründe für die Abweichungen liegen in:

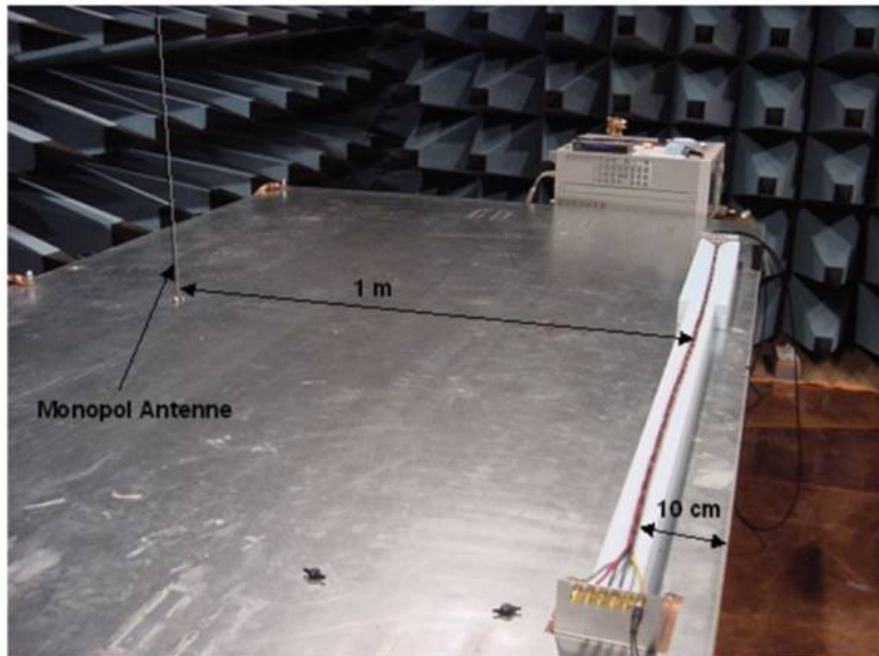
- Die FFT-Funktion in LT-Spice kann einen Messempfänger für breitbandige Signale nur hinreichend genau nachbilden
- Die begrenzte Dynamik und Auflösung in Schritt 1. bei der

Charakterisierung mit dem Oszilloskop, 8 Bit Auflösung



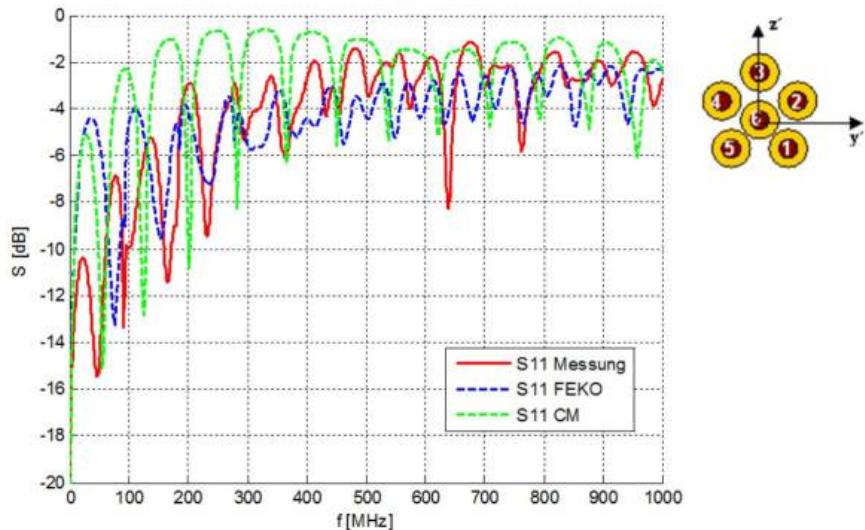
10.4.4 Genug gespielt

Die Betrachtung der Eindrahrtleitung ist für uns wichtig, um ein generelles Verständnis der Abstrahlung zu erhalten. Der Zweidrahtkabelbaum ist bereits in der Lage, Kommunikationsleitungen abzubilden, aber für viele Anwendungen natürlich noch nicht repräsentativ. Wir machen daher weiter mit einer sechsadrigen Leitung. Damit sind wir bereits z.B. in der Lage, einen BLDC-Motor anzusteuern mit 2x Versorgungsleitung, 3x Hallsignalen und einer digitalen Masse. Nachfolgende Abbildung zeigt den Versuchsaufbau zur Charakterisierung des Kabelbaums zusammen mit dem Hilfsmonopol.



Zur Mehrportmessung müssen die Ports entsprechend der Gesamtmatrix aufgeteilt werden. Der Netzwerkanalysator stellt vier Ports zur Verfügung, womit sich für zwölf Ports minimal 15 Messungen ergeben. Beim vorliegenden Kabelbaum wird versucht, die Position der Adern über der gesamten Länge konstant zu halten. Dazu sind die einzelnen Adern mit Klebeband an ihrer relativen Position fixiert. Im weiteren Verlauf der Messung wird gezeigt, wie sich ein Knick im Leiter bzw. eine kleine Geometrieänderung in den Streuparametern bemerkbar macht. Weiterhin wird der Einfluss des Abstandes der Leiteranordnung zur Monopolantenne messtechnisch untersucht. Im Vergleich zwischen Messung und Simulation werden einzelne Streuparameter herausgegriffen und miteinander verglichen. Eine Auswertung der gesamten 12×12 -Matrix bringt keine neuen Erkenntnisse. Die Simulation erfolgt für dieses Beispiel mit zwei verschiedenen Simulationsverfahren. Zum einen kommt die bekannte Feldberechnung mit Hilfe der MoM zum Einsatz, zum zweiten ein Tool, CM, welches nur den Kabelbaum über die [Transmission-Line-Methode](#) charakterisiert. Dabei werden zuerst die Kapazitäts- und Induktivitätsbeläge der Einzeladern berechnet, welche sich aufgrund der Verlegung innerhalb eines Bündels ergeben. Je nach Position innerhalb des Kabelbaums können diese Werte natürlich über den Ort variieren. Da es sich um eine Nachbildung des Kabelbaums ähnlich einer Netzwerkanalyse mit diskreten Komponenten handelt, R, L, C, bietet

dieses Verfahren natürlich enorme Vorteile hinsichtlich der Berechnungsdauer im Vergleich zu einem Feldlöser. Clayton R. Paul hat übrigens auch zu diesem Thema ein Grundlagenbuch verfasst: [Paul, Multiconductor Transmission Lines](#)



Das Ergebnis ist erschreckend. Wir haben drei Verfahren, zwei Simulationsverfahren und eine Messung, mit drei Ergebnissen. Die einzige Erkenntnis, die wir daraus ziehen können, ist, dass wie zu erwarten war Resonanzen auftreten, deren Amplitude und Frequenz nur messtechnisch bestimmbar sind. Eine ähnlich gute Übereinstimmung wie mit den Ein- und Zweidrahtleitungen ist leider nicht zu finden. Ebenso fällt es schwer, eine generelle Aussage über die Abweichung zwischen Messung und Simulation zu treffen. Das lässt nur einen Schluss zu, den erfahrene Laboringenieure schon lange kennen:



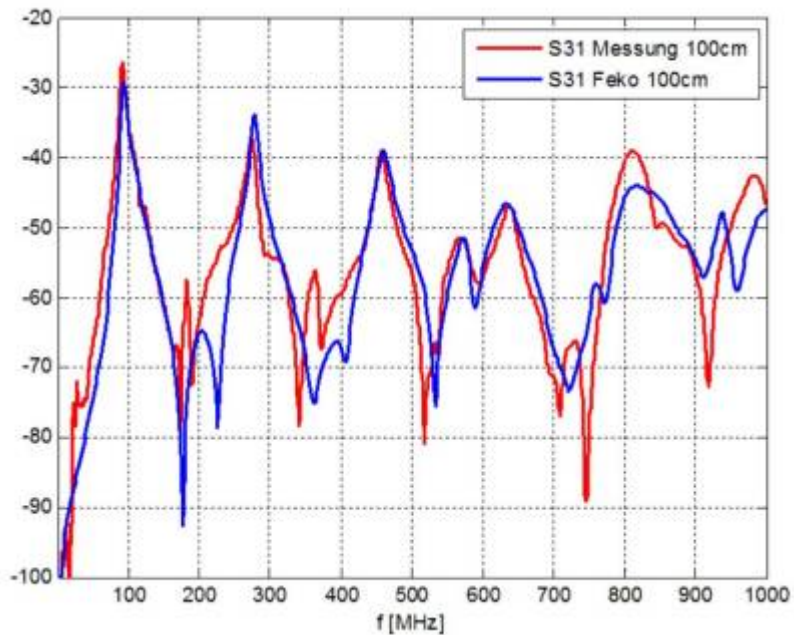
Das Verhalten von Einzeladern innerhalb eines Kabelbaums ist hinsichtlich Resonanzstellen und Emission nicht vorhersagbar.

Daher ist es nicht verwunderlich, dass sich EMV-Ingenieure für die Kabelbäume oder generell alle Leitungsverbindungen Referenzaufbauten herstellen mit

definierter Leiterverlegung. Über Kabelbinder oder Klebebänder, verändert das nicht die Resonanzstellen? Siehe Einfluss ϵ_r , werden die Einzeladern in Position gehalten und sorgen so dafür, dass der Einfluss unterschiedlicher Verlegung auch laborübergreifend reduziert werden kann.

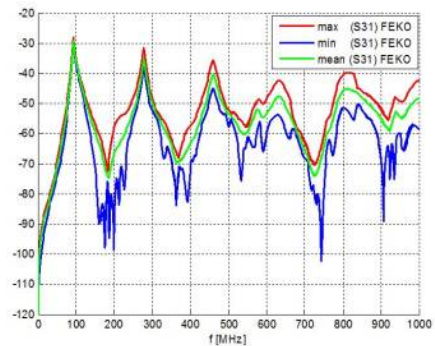
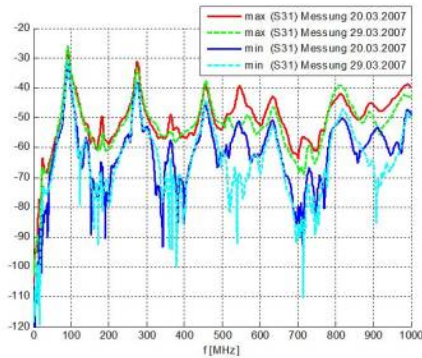
Zurück zum Diagramm Messung vs. Simulation. Es zeigt sich, dass die Daten aus Feko generell dichter an der Messung liegen als das berechnete Ergebnis mit der Transmission-Line-Methode, welche ein sehr stark resonantes Verhalten zeigt, das so in der Realität nicht vorhanden ist, Resonanzstellen werden überbewertet. Im Vergleich zwischen minimaler und maximaler Eingangsreflexion und Transmission zeigt es sich, dass selbst bei der geringen Anzahl an Leitern eine hohe Varianz bzw. Abweichung innerhalb der Leiterparameter besteht. Das Ergebnis aus Feko zeigt im Grundpegel hinreichende Übereinstimmung mit den Messergebnissen. Resonanzverläufe bzw. Resonanzstellen werden auch hier nur bedingt wiedergegeben.

Im nächsten Schritt wird die Abstrahlung zum Monopol betrachtet. Die Portbelegung hat sich aufgrund einer zusätzlich durchgeführten Messung allerdings geändert. Die Feko-Berechnung gibt als Ergebnis die Abstrahlung zum Monopol direkt aus. Wir betrachten nun die Transmission vom Leiterport (3) auf den Hilfsmonopol (1).



Der Vergleich zwischen Messung und Simulation für die Transmission zum Monopol zeigt eine hohe Übereinstimmung über dem gesamten Frequenzbereich. Erneut ist die zusätzliche Resonanz, hervorgerufen durch die Tischresonanz, im Bereich um 20 MHz in der Messung zu sehen. Der prinzipielle Verlauf wird durch die Simulation gut getroffen. Die in der Legende beschriebenen 100 cm weisen auf einen Abstand vom Kabelbaum zum Monopol von einem Meter hin. Eine Reduktion des Abstands zum Monopol in Messung und Simulation zeigt gleiche bzw. ähnliche Ergebnisse und Übereinstimmung und wird nicht weiter betrachtet.

Es soll nun untersucht werden, wie reproduzierbar die Messungen zur Abstrahlung sind. Dabei wird das gesamte Setup neu aufgebaut und die Transmission zum Monopol erneut betrachtet.

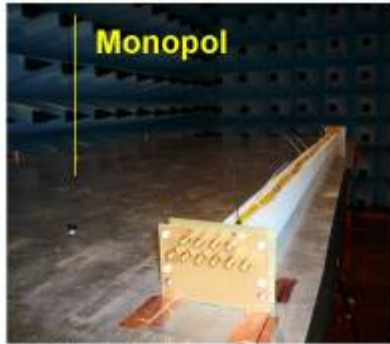


Links: Min. / Max. aus unterschiedlichen Messungen, rechts: Min. / Max.-Werte aus Simulation über den Adern

Die obige Abbildung zeigt, dass die minimal und maximal auftretenden Transmissionen zum Monopol in der Wiederholungsmessung bis 450 MHz nahezu identisch sind. Das bedeutet, dass wir uns bis zu dieser Frequenz keine Gedanken über Abweichungen aus dem Aufbau machen müssen. Ab 500 MHz gibt es an einigen Resonanzstellen deutliche Abweichungen bzw. gegenläufige Amplituden in der Transmission. Dazu muss jedoch erwähnt werden, dass der Kabelbaum zwischen den Messungen in der Mitte um 90° abgewinkelt wurde, um danach wieder in die Ausgangsform zurückzukehren. Bei der Abwinkelung wurden sicher einige Leiter von ihrem vormaligen Platz verdrängt bzw. an eine andere Stelle geschoben. Von weiteren Reparaturmaßnahmen wurde dabei abgesehen, kompletter Neuaufbau der Leiter etc.

Bis ca. 300 MHz zeigt sich kaum eine Abweichung zwischen minimaler und maximaler Abstrahlung. Das bedeutet, dass alle Leiter in etwa denselben Anteil zur maximalen Abstrahlung besitzen. Je größer der Abstand zwischen minimaler und maximaler Kurve ist, desto unterschiedlicher ist die Abstrahlung der Einzelleiter. Es existiert somit immer ein oder mehrere Leiter mit maximaler Abstrahlung. Dass verständlicherweise nicht alle Leiter dieselbe Abstrahlung aufweisen, zeigt sich auch in der Feko-Simulation. In der Abbildung rechts ist die minimal und maximal mögliche Transmission zum Monopol aus der Simulation dargestellt. Eine Mittelwertberechnung bei komplexen Zahlen, welche die Streuparameter ja darstellen, macht nur bedingt Sinn, veranschaulicht allerdings gut die im Mittel vorhandenen Übertragungswerte. Zur Berechnung wurde die Phaseninformation ignoriert und nur das Betragssignal verwendet.

10.4.5 Abstrahlungsverhalten von Mehradernkabelbäumen



10 Adern Aufbau

Im vorhergehenden Abschnitt haben wir das Verhalten eines Kabelbaums betrachtet, dessen Einzelleitungen über die Länge fixiert wurden.

Folgende Fragen blieben allerdings bisher unbeantwortet:

- Was passiert bei statistisch verteilten Leitern?
- Welchen Einfluss zeigt die zufällige Verlegung der Adern über die Leiterlänge?

Um den Einfluss der statistischen Verteilung besser zu verstehen, wurde ein Kabelbaum mit zehn Adern untersucht. Die Einzelleitungen sind nicht verdreht, sondern werden vom Ein- bis zum Ausgang als unabhängige Drahtleitungen geführt. Alle Leiter werden als äquivalent betrachtet, womit keine Nummerierung oder Farbcodierung verwendet wird. Als Bewertungskriterium wird die Transmission zum Versuchsmonopol betrachtet und ausgewertet. Dabei wird die mittlere Abstrahlung aller Adern betrachtet sowie der Anteil zur maximalen Abstrahlung einer Ader.

Entsprechend der Portdefinition ergeben sich für den Versuchsaufbau 21 Ports. Die Leiter haben eine Gesamtlänge von 1,97 m, der Abstand zum Monopol beträgt einen Meter. Die Leiter sind 5 cm über der leitenden Tischfläche fixiert. Bevor eine Bewertung der Abstrahlung erfolgt, werden zuerst die internen Leitungsparameter verglichen.

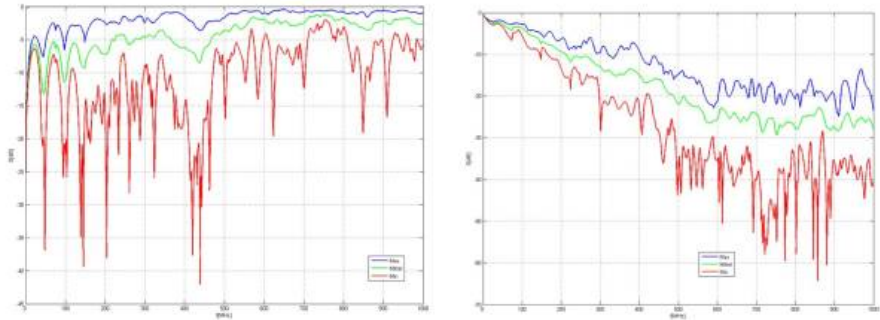


Abbildung: Links: Maximal-, Minimal- und Mittelwert der Eingangsreflexion, rechts: Maximal-, Minimal- und Mittelwert der Transmission (fern)

Die Abbildung zeigt die Mittelwerte mit den dazugehörigen Minimal- und Maximalwerten der Eingangsreflexion für die Ports S_{ii} , $i = 1 \dots 20$, sowie die Werte zur Transmission (fern) S_{ij} , i, j entsprechend einer Leitung, oder allgemein $i \in \{1 \dots 10\}$, $j \in \{11 \dots 20\}$. Das Ergebnis muss man erst einmal sacken lassen Mit einer derartig hohen Abweichung innerhalb des Bündels war im Vorfeld nicht zu rechnen. Bis ca. 200 MHz verhält sich das Kabelbündel entsprechend einer Eindrahtleitung mit den typisch ausgeprägten Resonanzstellen, die sich aus der Leiterlänge ergeben. Ab 200 MHz überwiegt die Kopplung der benachbarten Adern bzw. die Wirkung des Kabelbaums als Gesamtsystem. Die Transmission entlang einer Einzelader stimmt für kleine Frequenzen für alle Adern überein. Allerdings nehmen die Unterschiede mit steigender Frequenz zu bis zu einem maximalen Abstand, Worst Case, von 20 dB Unterschied.

Es bleibt noch die Frage, welcher Unterschied sich für die Transmission zum Monopol ergibt. Die großen Abweichungen innerhalb des Kabelbaums legen nahe, dass es Kombinationen von Adernpaaren bzw. einzelne Leitungen gibt, die einen hohen bzw. erhöhten Anteil zur Abstrahlung beitragen im Vergleich zu anderen Leitungen oder Adernpaaren. Um diese Annahme zu verifizieren, werden die gemessenen Parameter in ein SPICE-Modell überführt und auf verschiedene Anregungsarten sowie Kombinationen hin untersucht.

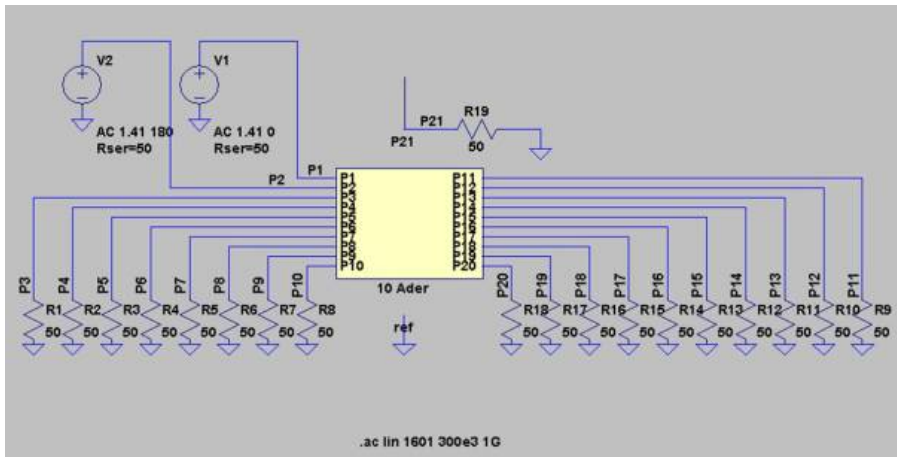


Abbildung: Beispiel zur Analyse in SPICE, dargestellt: Anregung Leiter 1 und Leiter 2 mit Gegentakt

Es werden jeweils zwei Eingangsports herausgegriffen und nacheinander mit einer Gleichtakt- sowie Gegentaktanregung beaufschlagt. Bewertet wird dabei die Fußpunktspannung am Monopol, welche direkt mit der vom Kabelbaum abgestrahlten Feldstärke gleichzusetzen ist. Aus der Anzahl der Eingangsports an einem Ende des Kabelbaums ergeben sich für die Anregung mit zwei Quellen 45 Kombinationen.

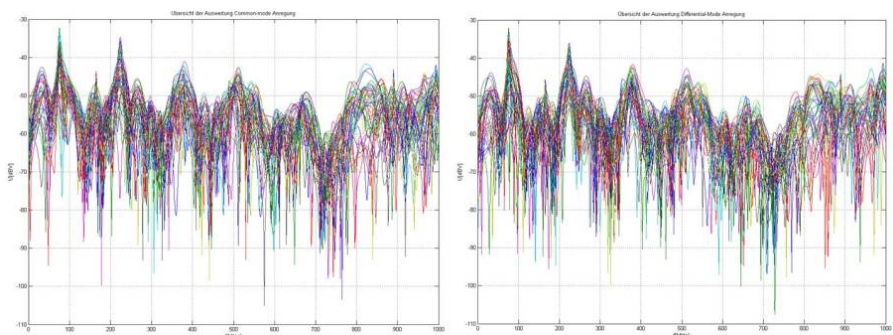


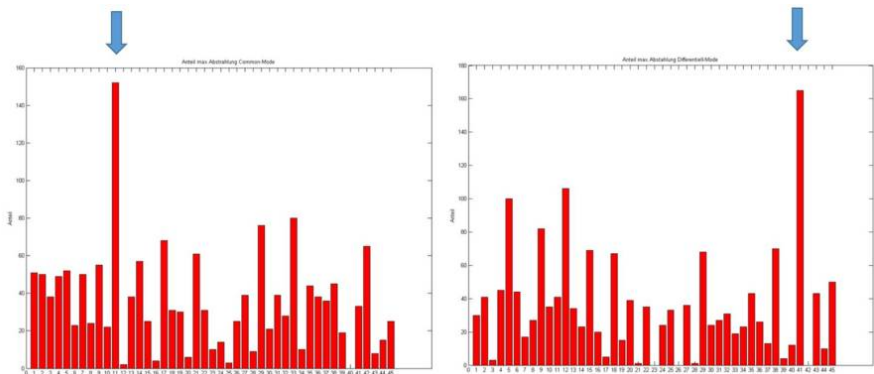
Abbildung: Spannung am Monopol für die verschiedenen Anregungskombinationen. Links: Gleich-, rechts: Gegentaktanregung

„Wir sehen den Wald vor lauter Bäumen“

Die in der obigen Abbildung dargestellte Abstrahlung, bewertet als Fußpunktspannung am Monopol in $[dB\mu V]$, zeigt, dass aus der Summe der

Einzelabstrahlungen kaum Information zu entnehmen ist. Wie erwartet weicht die Abstrahlung einzelner Aderkombinationen, wie bereits bei den internen Leitungsparametern erkannt, um bis zu 20 dB voneinander ab. **Es zeigt sich, dass obwohl die Leiter sehr dicht aneinander liegen keine signifikante Feldauslöschung bei Gegentaktanregung auftritt.** Dies bedeutet, dass der in zwei Adern eingespeiste Gegentaktanteil sich bereits im unteren Frequenzbereich als reine Gleichtaktstörung ausbreitet und somit keine Feldauslöschung auftreten kann. Die Vorteile einer Zweidrahtleitung zur Reduktion der Emission sind innerhalb eines Leitungsbündels damit nicht mehr gegeben.

Leider gibt es keinen Leiter oder Leiterpaar, das mit Abstand die höchste Abstrahlung hervorruft. Oft ist es nur ein geringer Frequenzbereich, in dem ein Aderpaar signifikant höhere Abstrahlung aufweist. Um dennoch eine Aussage über die Verteilung der Abstrahlung machen zu können, wird der Anteil maximaler Abstrahlung gewichtet. Das bedeutet, dass an jedem diskret vorhandenen Frequenzpunkt nach dem Aderpaar mit maximaler und minimaler Abstrahlung gesucht wird für Gleich- und Gegentaktanregung über dem gesamten Frequenzbereich.



Anzahl an der maximalen Emission über dem Frequenzbereich für einzelne Leiterpaare.

Deutlich größer als erwartet zeigen sich die Unterschiede innerhalb der Abstrahlung der einzelnen Aderpaare. Kombination 11 mit den Leitungen 2 und 5 zeigt für die Gleichtaktanregung zweier Leiter die maximale Abstrahlung bzw. über dem gesamten Frequenzbereich den Anteil höchster Abstrahlung. Für die Gegentaktanregung ist Kombination 41 mit den Leitern 5 und 10 verantwortlich für die im Mittel höchste Abstrahlung. Auffällig ist natürlich, dass in beiden Fällen Leiter 5 an der maximalen Abstrahlung beteiligt ist.

10.4.6 Übertragung der Simulationsergebnisse auf die reale Messumgebung

In den bisherigen Kapiteln wird versucht, die Abstrahlung der Kabelanordnungen auf die Hilfseinrichtung „Monopol“ zu berechnen und zu bewerten. Der Monopol kann in diesem Zusammenhang jedoch nur als Hilfsmittel betrachtet werden, da er mit einer realen Messung im Labor nichts zu tun hat. Entsprechend den CISPR-Anforderungen wird der Frequenzbereich bis 1000 MHz in die Arbeitsbereiche der verschiedenen Antennen unterteilt mit:

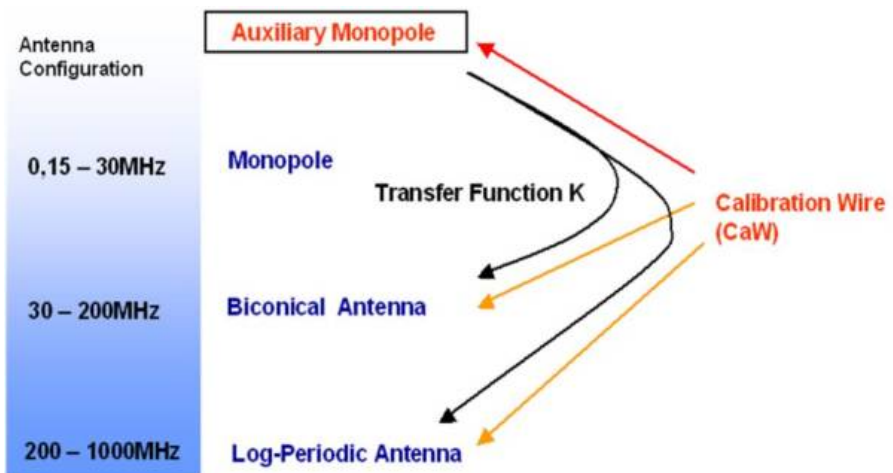
- 0,15–30 MHz Monopolantenne
- 30 MHz–200 MHz bikonische Antenne
- 200 MHz–1000 MHz logarithmisch-periodische Antenne

Je nach Laborausstattung kann der Frequenzbereich ab 30 MHz auch durch eine Kombination, BiLog, aus bikonischer und logperiodischer Antenne abgedeckt werden. Soll das gesamte Komponentensystem in der EMV-Messung nachgebildet werden, müssen auch die Antennen nachgebildet werden bzw. in das Gesamtmodell mit aufgenommen werden. Die Modellierung der Antennen ist jedoch nicht trivial und erzwingt in Zusammenhang mit der benötigten Tischmodellierung eine sehr hohe Anzahl an Meshzellen für die Berechnung mit der MoM-Methode. Allein durch die Vernetzung der Tischfläche werden ca. 40.000–50.000 Elemente, Mesh-Triangles, erzeugt, die bei der Berechnung mit Hilfe eines herkömmlichen PCs zu einer Simulationsdauer von mehreren Tagen führen. Soll weiterhin die Antenne in das System eingebracht werden, erhöhen sich die Meshzellen um weitere 10.000–20.000 für jedes Antennensystem. Für eine schnelle Bewertung sind Rechenzeiten, die über einer Woche liegen, natürlich nicht annehmbar, besonders falls verschiedene Varianten des Kabelaufbaus betrachtet werden müssen. Abhilfe bringt hier die Verknüpfung zwischen Messung und Simulation. Bindeglied zwischen den beiden Systemen stellt der Simulationsmonopol dar. Der Monopol wurde deshalb gewählt, da er in der Simulation sehr einfach zu realisieren ist sowie einfach in das Messequipment integrierbar ist. Der in den vorhergehenden Abschnitten gezeigte Vergleich zwischen Messung und Simulation für den Monopol zeigt eine sehr hohe Übereinstimmung der Parameter. Das Ziel ist es nun, das gesamte Messequipment sowie die Umgebung in einer einzigen Kalibriermessung zu erfassen. Die Kalibriermessung gibt dabei den Zusammenhang zwischen der Fußpunktspannung am Hilfsmonopol sowie der Fußpunktspannung an der Antenne wieder. Zur Berechnung der Kalibrierung, oder im Folgenden Korrektur genannt, wird eine Kalibrierleitung verwendet, die am gleichen Ort wie der spätere Kabelbaum aufgebaut wird. Die vom Kabelbaum erzeugte Feldstärke an

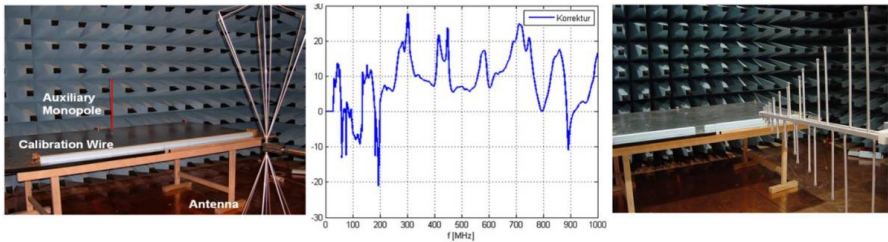
den Antennen wird über die Fußpunktspannung miteinander in Relation gesetzt und als Korrekturkurve ermittelt. Die Korrekturkurve enthält somit Informationen über:

- Antennentyp
- Antennenpolarisation
- Umgebungseigenschaften wie Tischfläche, Masseanbindung der Tischfläche, Absorber etc.

Aus der Simulation ist die Transmission der Leiteranordnungen zum Monopol sehr einfach bestimmbar. Sind die Umgebungsparameter bekannt, lässt sich nun mit Hilfe der Korrektur die Abstrahlung zur Laborantenne ermitteln.



Die Korrekturkurven bestimmen somit das Verhältnis der Störung am einfach zu simulierenden Hilfsmonopol zur Laborantenne. Die Messung erfolgt im Labor mit den dort zur Verfügung stehenden Messmitteln und einem Netzwerkanalysator. Dabei wird mit einem Port des Netzwerkanalysators eine einfache Eindrahtleitung angeregt, in der oberen Abbildung als Calibration Wire bezeichnet. Jeweils ein weiterer Port wird an den Hilfsmonopol sowie die Laborantenne angeschlossen. Mit der gemessenen 3x3-Matrix lässt sich die gesuchte Übertragungsfunktion oder der Korrekturfaktor K ermitteln. Nachfolgende Abbildung zeigt die Antennenkorrektur für eine horizontal polarisierte bikonische und logarithmisch-periodische Antenne.



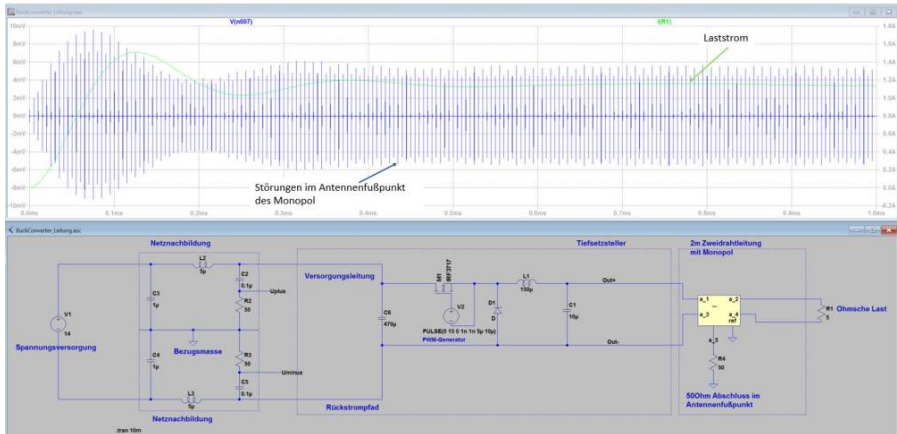
Ich hoffe, Sie konnten meinen Gedanken bis hierher folgen. Zur Verdeutlichung ein Beispiel mit dem bekannten Tiefsetzsteller.

Aufgabe:

Mit Hilfe einer Simulation möchten wir herausfinden, welche Emissionen generiert werden, falls wir die Last eines Tiefsetzstellers über eine zwei Meter lange Leitung anbinden.

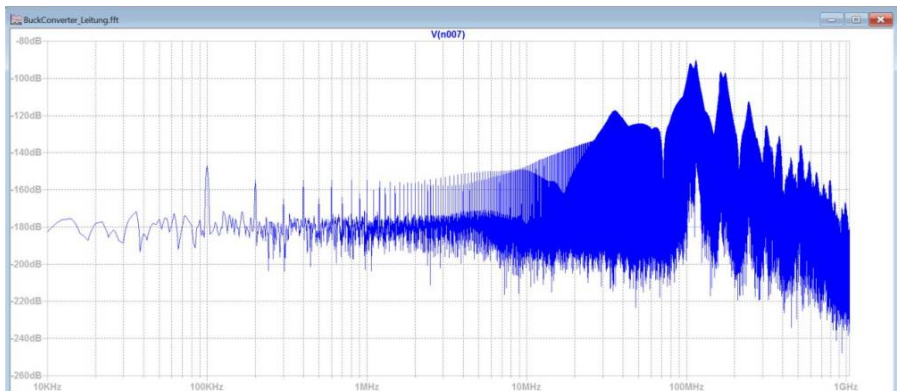
Eine typische Anwendung wäre die Versorgung einer entfernten Last aus einer zentralen Spannungsversorgung. Die funktionale Simulation ist identisch mit der Simulation zur Bestimmung der leitungsgebundenen Emission. Neu hinzugekommen ist ein Modell der Übertragungsleitung. Nachfolgende Abbildung zeigt die um die Zweidrahtleitung erweiterte Simulation des Tiefsetzstellers mit dem berechneten Ausgangsstrom und der Spannung im Antennenfußpunkt, Monopol. Deutlich zu sehen ist, wie im Einschwingvorgang erhöhte Emissionen auftreten. Damit wir vergleichbare Ergebnisse zur Messung erhalten, müssen wir diesen Bereich in den weiteren Betrachtungen ausschließen.

Das Kabelmodell mit Hilfsmonopol wurde aus den simulierten Streuparametern mit Hilfe des Vector Fittings in ein äquivalentes SPICE-Modell übersetzt. Dazu sind eine Vielzahl an passiven Bauelementen, aber auch Strom- und Spannungsquellen notwendig, welche sich negativ auf die Simulationsdauer auswirken. Ist der Tiefsetzsteller ohne die Leitung in wenigen Sekunden berechnet, benötigt mein Notebook jetzt bereits mehrere Minuten.



Die Übertragung der Fußpunktspannung am Monopol bringt das gesuchte Ergebnis für die vom Kabelbaum ausgehende Emission. Schön zu sehen sind im Frequenzbereich > 100 MHz die Resonanzstellen, welche sich aufgrund der Ausbreitung der Störsignale entlang der Leitung ausbilden. Da wir zuvor das gesamte Setup bis 1 GHz qualifiziert haben, dürfen wir uns auch nur diesen Frequenzbereich ansehen.

buckconverter_leitung.zip



Wir kennen jetzt das Frequenzspektrum der Fußpunktspannung des Monopols, von LT-Spice berechnet in dBV, sprich dBVolt. Wir suchen jedoch die emittierten Störungen als Feldstärke in dBV/m, sprich dBMikroVolt pro Meter. Die Umrechnung zur Messantenne haben wir im vorhergehenden Kapitel besprochen. Das bedeutet, wir müssen lediglich zu unserem Ergebnis in dBV die

Transferfunktion addieren und erhalten damit das Spektrum der Fußpunktspannung der Messantenne. Jetzt der finale Schritt vom Fußpunkt der Messantenne zur Feldstärke. Dazu verwendet man den vom Antennenhersteller ermittelten Korrekturfaktor, oder auch Antennenfaktor genannt, für die jeweilige Antenne und Polarisisation. Die finale Feldstärke, welche wir mit den Grenzwerten vergleichen dürfen, erhält man somit aus:

$$E \text{ [dB}\mu\text{V/m]} = U_{\text{Monopol}} \text{ [dBV]} + 120 + \text{Transferfunktion K [dB]} + \text{Antennenfaktor [dB/m]}$$

Gesuchte Feldstärke

Ergebnis aus LTSpice

Umrechnung von dBV in dBμV

Korrektur vom Monopol zur Messantenne

Umrechnung von der Fußpunktspannung der Messantenne zur Feldstärke

Der letzte Rechenschritt kann, soweit ich weiß, nicht in LTSpice durchgeführt werden, da im Frequenzbereich keine externen Daten geladen werden können. Somit bleibt nur, eine Auswertung in Matlab oder Excel durchzuführen. Obwohl wir so viel Aufwand in die Erstellung der Transferfunktion und Simulation gesteckt haben, werden wir diesen letzten Rechenschritt nur selten durchführen. Ursache dafür ist, dass wir es nur mit sehr hohen Anstrengungen schaffen, die von der Schaltung generierten Störgrößen nachzubilden, siehe Kapitel PCB-Simulation. Sinnvoller ist es daher, die Simulation mit Hilfsmonopol zu nutzen, um z.B. Filtermaßnahmen oder das Schaltverhalten der Leistungshalbleiter zu optimieren. Mit Hilfe von Parametersweeps lassen sich Hardwaremaßnahmen, besonders Filterkomponenten, komfortabel bewerten, bevor der finale Hardwaretest erfolgt.

10.5 Fazit zur Simulation

Wir haben gesehen, dass sich sowohl die Störungen auf eine Platine als auch die gestrahlten Störungen simulieren lassen. Relativvergleiche lassen sich dabei sehr einfach und mit überschaubarem Aufwand anstellen. Schwieriger wird es, falls wir herausfinden wollen wo der absolute Pegel sich befindet. Lassen wir dazu die einzelnen Schritte noch einmal Revue passieren und wir stellen uns vor dass wir die Emissionen einer elektronischen Komponente sowohl für leitungsgebundene als auch gestrahlte Emissionen mit absolut richtigem Pegel simulieren möchten.

Arbeitsschritte:

1. Festlegen des zu untersuchenden Frequenzbereich

2. Charakterisierung aller passiven Bauelemente (HF- Ersatzschaltbilder z.B. mit Networkanalysator oder Impedanzanalysator)
3. Geeignete funktionale Modelle der verwendeten Halbleiter (z.B. Mosfets, Treiber, Dioden) erstellen oder suchen. Meist bieten die Hersteller hier viele Modell in Spice-Syntax an.
4. Abschätzen oder berechnen der parasitären Kapazitäten und Induktivitäten auf der Leiterkarte.
5. Erstellen der Simulationsmodelle für die Anschlussleitungen, mindestens Impedanznachbildung.
6. Ab jetzt können die leitungsgebundenen Emissionen berechnet werden
7. Erweitertes Kabelbaummodell erstellen zur Bewertung der Abstrahlung (3D Feldberechnung mit Hilfsmonopol + Korrekturkurve zur realen Antenne)
8. Simulation der gestrahlten Emission

Sie sehen, eine EMV Simulation ist nichts was man kurz aus dem Ärmel schütteln kann. Dazu bedarf es lange Vorbereitung und Erfahrung mit den verwendeten Bauelementen. Nicht ohne Grund besitzen viele große Firmen eigene Abteilungen welche sich nur mit der Funktionsbeschreibung und Modellerstellung der verwendeten Halbleiterkomponenten beschäftigen. Für Geräte mit geringer Stückzahl lohnt sich der Aufwand meist nicht und man arbeitet kostengünstiger in dem man sich auf Relativvergleich, ohne den Anspruch auf absolute Pegelmaße, beschränkt.

11 Was tun gegen gestrahlte Emissionen

Wir haben viele Möglichkeiten gesehen, wie wir messtechnisch und mit einer Simulation die Pegelmaße ermitteln können. Spannender ist natürlich die Frage: „Was können wir dagegen tun?“. Die wichtigste Regel dazu ist identisch mit den Regeln zur Bekämpfung der leitungsgebundenen Emission:



Für alle Ein- und Ausgänge der Funktionsblöcke im Blockschaltbild muss eine EMV-Maßnahme vorgesehen werden!

Auch hier ist unser schärfstes Schwert die Bfilterung. Man könnte den obigen Merksatz auch umschreiben und sagen, alle Ein- und Ausgänge müssen mit

einem Filter ausgestattet sein. Im ersten Moment denkt man vielleicht an klobig große Filterschaltungen, allerdings reicht in vielen Fällen ein einfacher Kondensator aus. Anders als bei der leitungsgelassenen Emission müssen allerdings jetzt alle Strukturen, welche sich in der Größenordnung der betrachteten Wellenlänge befinden, bewertet werden. Hierzu zählen meist die bereits ausführlich betrachteten Kabelbäume, aber auch ein vorhandenes (leitfähiges) Gehäuse, welches als Antenne wirken kann.

Bei den Gehäusen werden meist die Gerätedeckel und die Öffnungen, z.B. zur Kühlung, unterschätzt. Deckel oder angeschraubte Gehäuseteile müssen stets niederohmigen ($< 1\text{m}\Omega$) Kontakt zum Gehäuse haben. Andernfalls bildet sich entlang der Oberfläche eine hochfrequente Spannung aus, welche als Antenne wirken kann. Um eine niederohmige Anbindung zu gewährleisten, gibt es sogenannte EMV-Dichtungen, welche zwischen die Gehäuseteile eingelegt und dann fest verschraubt werden. Vielleicht sind Ihnen auch schon Geräte aufgefallen, die merkwürdig viele Schrauben verwenden, um einen Gehäusedeckel zu befestigen. Das Ziel ist hier, wie bereits beschrieben, die Anbindung so niederohmig wie möglich zu gestalten.

Öffnungen in den Gehäusen lassen sich nur selten vermeiden. Vielleicht möchten Sie ein Display integrieren oder für eine Konvektion zur Kühlung sorgen. Da solche Öffnungen meist nicht sehr groß sind, stellen sie eher ein Einfallstor für sehr hochfrequente Signale, typischerweise $>500\text{MHz}$, dar. Für tiefere Frequenzen ist die Wellenlänge zu groß, man könnte auch sagen, sie passen



schlicht nicht durch die Öffnung

Auch für diese Problemstellung gibt es bereits vorgefertigte Lösungen in Form von Gittern oder Schirmblechen (Display), um das Gehäuse abzudichten. Sie merken, es gibt zahlreiche Lösungen, um die EMV einzuhalten, es gilt sie nur zu kennen und anzuwenden!

Bitte schauen Sie sich auf den Homepages der angegebenen Firmen um und Sie werden feststellen, es gibt eine große Vielfalt an Lösungen rund um die Schirmung von Gehäusen oder einzelnen Schaltungsteilen Ihrer Platine:

- MTC
- Würth Elektronik

Bitte werden Sie stets vorsichtig und misstrauisch, falls Ihnen jemand eine EMV-Tapete oder ein Moskitonetz mit integriertem EMV-Schutz anbietet. Die EMV befasst sich mit sehr vielen Themenfeldern, die Wechselwirkung magnetischer und elektrischer Felder mit dem menschlichen Organismus gehört jedoch nicht dazu. Die Sensibilität einzelner Menschen hinsichtlich hochfrequenter Felder möchte ich nicht zur Diskussion stellen. Leider bieten unseriöse Händler

Lösungen zu Wucherpreisen an, unter dem Deckmantel der EMV bzw. pseudowissenschaftlicher Aussagen.

12.0 EMV im Schaltplan und PCB

Wie wir bereits wissen, beginnt die EMV-Arbeit mit der Erstellung des Blockschaltbilds, in dem wir erste EMV-Maßnahmen definieren. Sobald wir einen der funktionalen Blöcke mit Leben füllen, müssen wir auch dort wieder EMV-Aspekte berücksichtigen. Das Spiel wiederholt sich so lange, bis wir auf der Leiterkarte angekommen sind. Hier gilt es dann, die kritischen Pfade zu erkennen und mit Maßnahmen zu hinterlegen. Erst dann haben wir ein durchgängiges Gesamtkonzept von der Idee bis zur Bauteilebene.

12.1 Kritische Pfade

Um EMV-Maßnahmen im Schaltplan und auf dem PCB umzusetzen müssen wir natürlich im ersten Schritt klären wo liegt das Problem oder besser, nach was suchen wir denn? Dazu erinnern wir uns noch einmal an das Kapitel 3 „Woher kommen die Störgrößen“ in dem wir geklärt haben wer denn nun hochfrequente Signale generiert. Und genau nach diesen Kandidaten suchen wir jetzt:

- Getaktete Signale
- Taktgeneratoren, Quarze, Schwingkreise allgemein
- Kommunikationsleitungen
-

Oder man könnte schlichtweg auch sagen: Alle Signal-Traces auf denen **funktional** mehr als nur DC-Signale unterwegs sind. Funktional bedeutet hier ohne hochfrequenten Anteile, keine Funktion!

Diese Signale markieren wir im Schaltplan und hinterlegen stets ein EMV-Maßnahme. Maßnahmen können sein:

- Lokale Filterung
- Platzierung der Filterschaltung möglichst nahe an der Störquelle
- Rückstrompfade möglichst kurz halten
- Getrennte Masseführung (gezielte Rückstrompfade)

Das nachfolgende Beispiel zeigt einen Schaltplanausschnitt einer

[illegible]

Es verbleibt in diesem Beispiel eine Schleife in welcher wir von einem dreieckförmigen Stromfluss ausgehen. Dieser Strom ergibt sich aus dem funktionalen Strom unseres Tiefsetzstellers, vgl. dazu die Simulation in Kapitel

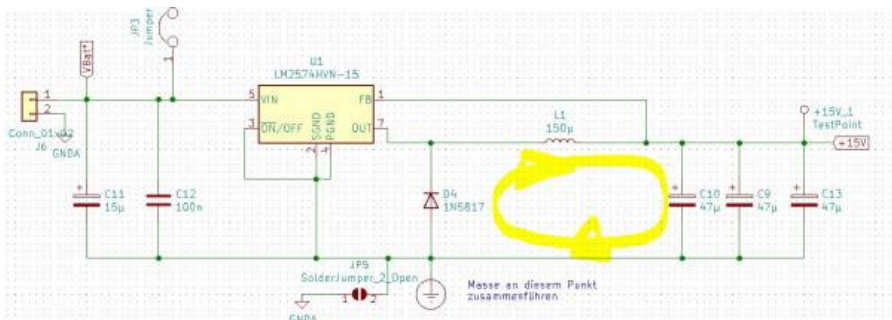
2.3.1. Im Schaltplan sind diese Strompfade grün gekennzeichnet. Um genau zu sein, ergibt sich ein dreieckförmiger Strom nur in der Spule und somit im Ausgangskondensator. Durch die Diode und das IC jeweils der rechte bzw. linke Teil des Dreiecks, entsprechend der nachfolgenden Abbildung. Wir sprechen in diesem Fall davon, dass der Strom kommutiert, was nichts anderes bedeutet als dass er auf einen anderen Knoten übergeht.



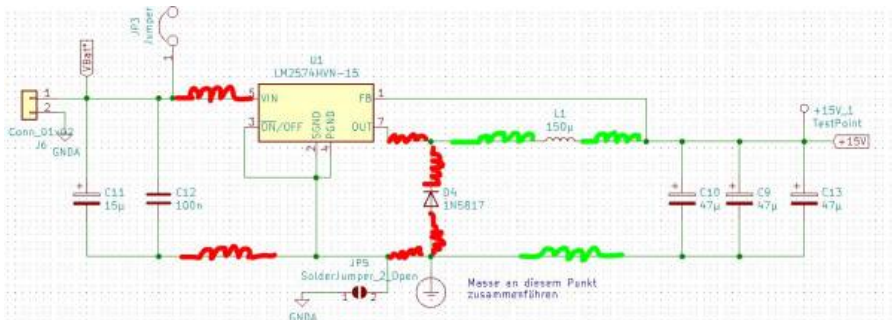
Diese kommutierenden Ströme bringen gleich mehrere Probleme mit sich. Einerseits generiert dieser Stromverlauf weit höhere Emissionen als der rein dreieckförmige Verlauf (rein aus der Fourier-Reihe betrachtet), zum anderen erhalten wir aufgrund der springenden Ströme Spannungsüberhöhungen an parasitären Kapazitäten.

12.2 Gute und schlechte Induktivitäten

Wir wissen bereits, dass wir den Strom in der Hauptinduktivität (Spule) nicht einfach unterbrechen dürfen. Sie erinnern sich: Der Energiefluss muss stets stetig sein! Daher bietet wird dem Strom einen alternativen Pfad über die Diode an sobald der Mosfet deaktiviert wird. Daher nennen wir die Diode auch Freilaufdiode. Nachfolgende Abbildung markiert den Freilaufpfad der Schaltung.



Wir wissen ja bereits, dass sämtliche Leitungsverbindungen und damit auch einzelne PCB-Leitungsabschnitte ebenfalls Induktivitäten darstellen. Auf unserer Leiterkarten können wir als groben Richtwert ca. 10nH/cm annehmen. Das bedeutet aber auch, dass es Induktivitäten in unserer Schaltung gibt, welchen wir keinen Freilaufpfad anbieten können.



Im gezeigten Schaltplanausschnitt sind in Rot die kritischen parasitären Induktivitäten eingezeichnet, für die kein Freilaufpfad zur Verfügung steht. Die grünen Leitungsinduktivitäten bilden mit der Hauptinduktivität eine Serienschaltung und dürfen zu dieser addiert werden womit sie den Freilaufpfad über die Diode verwenden. Wir können somit in unserer Schaltung „gute und schlechte“ parasitäre Induktivitäten unterscheiden.

Im Layout versuchen wir dann bei der Bauteilplatzierung die ungewollten roten Induktivitäten so klein wie möglich zu halten.

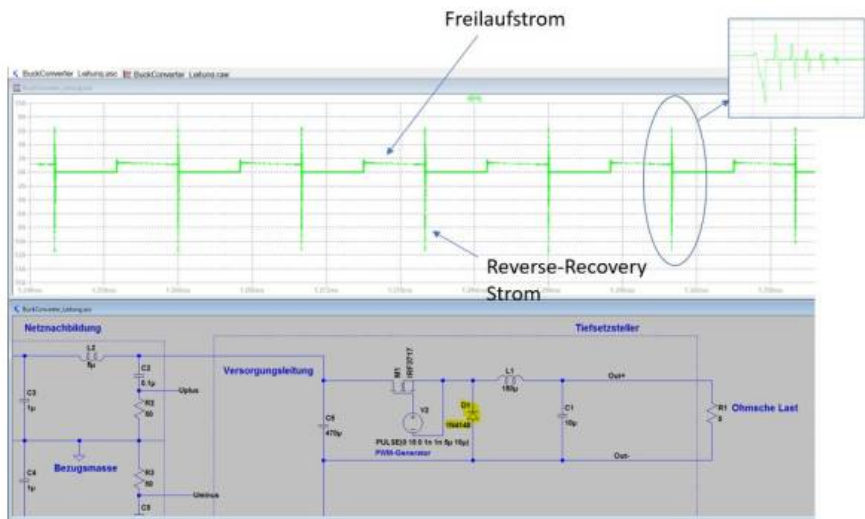


Pfade mit diskontinuierlichem Stromfluss müssen hinsichtlich deren parasitären Induktivitäten bewertet werden.

12.3 Reverse Recovery

In manchen Fällen kämpfen wir auch direkt gegen parasitäre Eigenschaften einzelner Bauelemente. Solch ein typischer Effekt stellt der Reverse Recovery dar. Dieser Effekt ist an p-n Übergängen von Halbleitern, typischerweise bei Dioden, also auch unserer Freilaufdiode zu beobachten. Eine Diode sperrt den Strom in Rückwärtsrichtung nicht sofort, sondern es kommt zuerst zu einem

rückwärtigen Strom durch die Diode. Man kann sich dabei vorstellen, dass die freien Ladungsträger, welche noch aus der leitfähigen Phase zur Verfügung stehen, zuerst ausgeräumt werden müssen. Dieser sehr kurze Strompeak liegt meist in der Größenordnung des zuvor geführten Strom in Vorwärtsrichtung und erzeugt durch seine sehr schnellen Flanken hochfrequente Anteile. Spice-Modelle realer Dioden berücksichtigen diesen Effekt womit er in einer Simulation darstellbar wird. Dazu muss in der zuvor betrachteten Simulation des einfachen Tiefsetzstellers die ideale Diode ersetzt werden. Zum Beispiel durch die bekannte Diode 1N4148.



Schön zu sehen ist zum einen der eigentliche Reverse-Recovery Impuls und im Zoom wie durch diesen Impuls die Schaltung zu einer hochfrequenten Schwingung angeregt wird. Ganz verhindern lässt sich dieser Effekt leider nicht. Allerdings können wir ihn minimieren indem schnelle Dioden eingesetzt werden mit einer möglichst geringen Sperrverzugszeit, im Datenblatt unter Q_{rr} zu finden.

12.2 Bauteilplatzierung

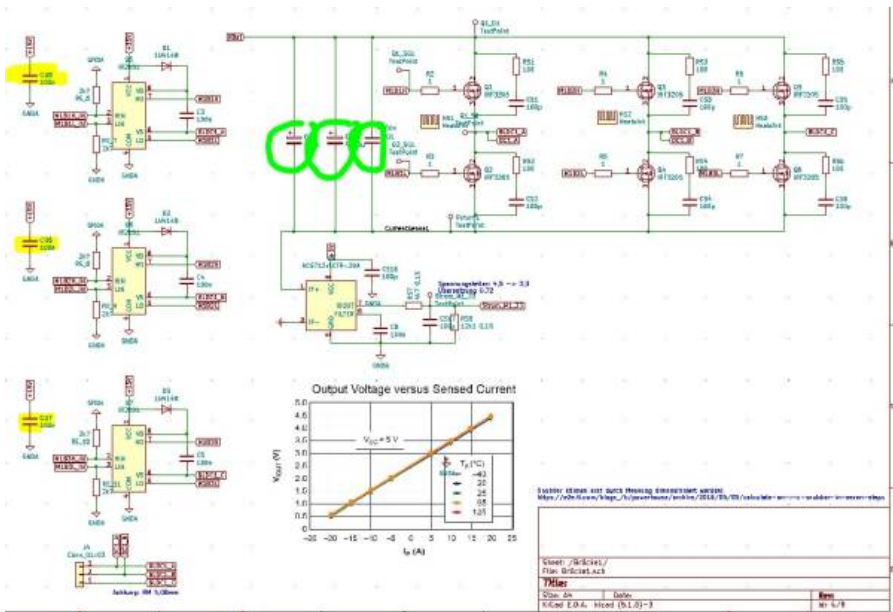
Bei der Platzierung der Bauteile auf der Platine können viele Fehler gemacht werden, da sozusagen hier der Ort ist an dem die Störungen entstehen. Die gute Nachricht ist, dass meist wenige Regeln ausreichen für einen EMV gerechtes

Layout.

Entkopplung der Netze

Eine wichtige Regel lautet, dass wir im Schaltplan nicht mehr davon ausgehen, dass es überhaupt identische Potentiale gibt. Das bedeutet wir berücksichtigen stets, dass es aufgrund parasitärer Induktivitäten entlang der PCB-Traces zu Spannungsabfällen kommen kann. Daher müssen wir die oben eingeführten Blockkondensatoren im Schaltplan jeweils an der zugehörigen Komponente berücksichtigen und dann natürlich im PCB auch so übernehmen.

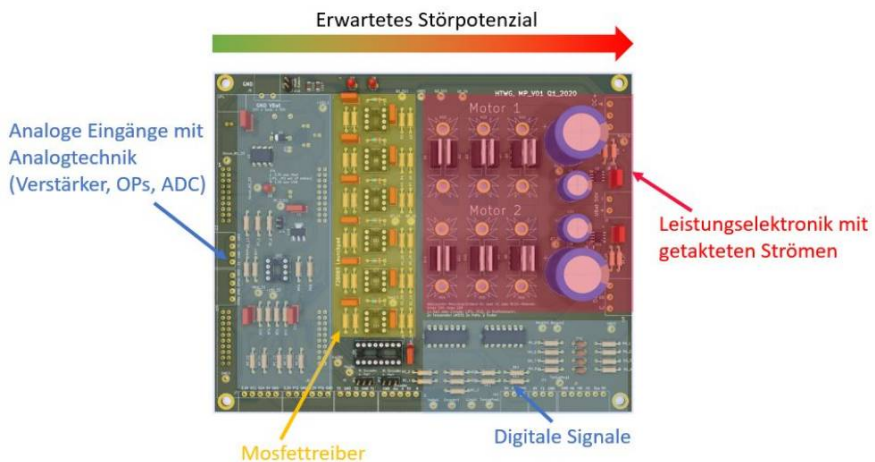
Wichtig ist es auf diesen Punkt hinzuweisen falls der Schaltplan und dann das folgende PCB Layout durch unterschiedliche Mitarbeiter bearbeitet werden.



Elektrotechnisch richtig wäre es natürlich alle Kondensatoren mit gleichem Potenzial (hier die grün und gelb markierten Bauteile) auf einem separaten Arbeitsblatt zu zeichnen. Allerdings fehlt uns dann später bei der PCB Erstellung jeglicher Bezug zur eigentlichen Aufgabe der Bauteile: Lokale Filter und Glättung der Versorgungsspannung.

Gruppierung der Bauteile

Durch Kopplungen auf der Leiterplatte (galvanische, kapazitive oder induktive Kopplung) breiten sich hochfrequente Störungen natürlich auch innerhalb der Leiterkarte aus. Zum Schutz sensibler Schaltungsteile, wie analoger Schaltungen oder HF-Systeme mit sehr geringen Signalpegeln, müssen wir dafür sorgen, dass diese Bereiche möglichst separiert werden. Oder anders gesagt, wir unterteilen die Schaltung in Gruppen entsprechend dem erwarteten Störpotenzial. Die Zonen können auch größeren funktionalen Gruppen aus dem Blockschaltbild entsprechen.



Das Beispiel zeigt eine Platine mit einem zweifachen Motortreiber für BLDC Motoren mit Zonen unterschiedlicher Störpotenziale. Das höchste Störpotenzial geht natürlich von den zwei Leistungsendstufen aus. Daher sollten dieser Schaltungsteile möglichst so platziert werden, dass seine Ein- und Ausgangssignale nicht mit weiteren Schaltungsteilen in Berührung kommen. Oft ist es hilfreich sich für die einzelnen Schaltungsteile eine Tabelle mit möglichen Störaussendung- und Störfestigkeitsproblemen zu erstellen. Wie immer das Henne - Ei Spiel. Nur wer die Probleme kennt kann sie angehen.

Schaltungsteil	Störpotential	Störaussendung	Störfestigkeit
----------------	---------------	----------------	----------------

Leistungselektronik	Sehr hoch	- Getaktete Ströme - Spannungsabfall durch hohe Ströme - Reverse Recovery - Falsche Totzeiteinstellung - Kühlkörperflächen führen zu hohen Gleichtaktströmen - Steckeranschluss und Leitungen zum Aktor (Motor) - Schaltüberspannung (Überschwingen beim Schalten) - Anregung resonanter Strukturen beim Schalten (engl. ringing)	- Ungewolltes Schalten von Leistungshalbleitern - Zerstörung der Halbleiter durch Überschwingen beim Schalten
Treiber	Hoch	- Hohes di/dt durch Kapazitätsumladung - Digitale Signale hoher Spannungsamplitude	- Ungewolltes Schalten von Leistungshalbleitern
Digitalteil	Mittel	- Je nach Übertragungsrate hohes du/dt - Kommunikationsleitungen mit variierendem Spektrum	- Einkopplung führt zu einer Fehlinterpretation der Signale
Digitalteil μC	Hoch	- Takt- und PWM-Signale mit sehr hohen Frequenzen - Quarz-Oszillator	- Ungewollte Signalzustände z.B. am Resetpin - Unterspannung aufgrund galvanischer Kopplung (eng. brownout)
Analogteil	Gering	- Meist keine direkte Abstrahlung sondern aufgrund von internen Kopplungen benachbarter Schaltungsteile	- Direkte Einkopplung in den Analogen Signalpfad - Indirekte Störung über Versorgungsspannung (galv. Kopplung) - Zu geringe Signal- zu Rauschabstand - Sehr geringe Nutzsignale $< 1mV$

Natürlich ist es möglich, dass innerhalb der genannten Flächen sich weitere Insel verschiedener Schutzklassen befinden welche einer besonderen Behandlung benötigen.

Nachdem die Tabelle erstellt ist gilt es zu allen gefundenen Punkte geeignete Maßnahmen zu definieren, bzw. Maßnahmen vorzusehen von denen wir ausgehen sie helfen uns in irgendeiner Weise bezüglich der EMV. Ob wir alle Punkte bedacht haben wissen wir spätestens nach den EMV Prüfungen.

Beispielmaßnahmen:

Risiko	Maßnahme
Getaktete Stöme	Lässt sich nicht vermeiden, da funktionaler Bestandteil der Schaltung. Definierter Rückstrompfad schaffen, aufgespannte Flächen so klein wie möglich.
Spannungsabfall	Großflächige Leitungsführung, Leiterlänge so gering wie möglich halten (priorisiertes Routing)
Reverse Recovery	Bauteilvergleich mit Simulation verschiedener Bauteile. Ggf. andere Technologie wählen.
...	...

Masseführung

Beim Thema **Masse** sollten wir zuerst den Unterschied zur **Erde** klären. Der Begriff Masse finde ich persönlich eher unglücklich und verwirrt eher, besser wäre der Begriff Rückleiter. Zugegeben, ich verwende meist ebenfalls meist den Begriff Masse.

Die Erde hingegen ist unser Schutzpotenzial und der PE- Leiter zu Hause sollte daher auch nur einen Strom im absoluten Notfall als im Schutzfall führen. Unsere FI- Schutzschalter erkennen dann in diesem Falls einen Fehlerstrom. Das Erdpotenzial kann mit Masse verbunden werden, ist aber nicht zwingend erforderlich. Die Situation, dass wir uns in einem Dreileitersystem finden wir in ganz vielen Anwendungen.

Beispiel Hausnetz:

Spannungsführender Leiter: L

Rückleiter (Masse): N

Erde = Schutzleiter: PE

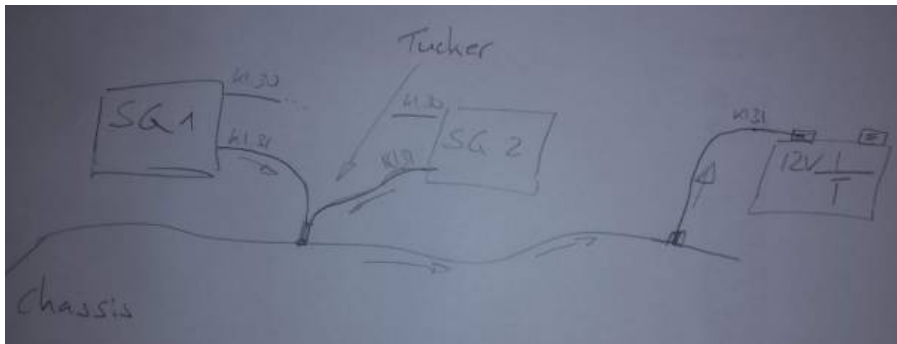
Der Schutzleiter PE und der Rückleiter N besitzen im stromlosen Zustand identisches Potential. Die Auftrennung erfolgt im Keller direkt am Ringerder der Betonplatte.

Beispiel Fahrzeugumgebung:

Spannungsführender Leiter: Kl. (Klemme) 30

Rückleiter (Masse): Kl. 31

Erde = Chassis



Im Unterschied zum Hausnetz ist es im Fahrzeug, abgesehen von HV-Systemen, nicht notwendig nach Fehlerströmen zu suchen. Der Rückstrompfad erfolgt direkt am Steuergerät über die Kl. 31, wird dann aber lokal auf das Chassis geführt. Auch hier herrscht im unbestromten Zustand an jeder Stelle identisches Potential. Erst durch einen Spannungsabfall aufgrund des ohmschen Widerstandes der Leitung oder aufgrund von Wechselströmen in Zusammenspiel mit der Leitungsinduktivität ergeben sich unterschiedliche Potenziale für Kl. 31 und das Chassis.

Es drängt sich die Frage auf warum die Steuergeräte den Rückstrom nicht direkt lokal auf das Chassis führen. Die Frage lässt sich damit beantworten, dass zum einen über das Steuergerätegehäuse und die Halterung keine niederohmige Verbindung gewährleistet werden kann und zum anderen, dass die Anzahl an (unlackierten) Bolzen, meist Tuckerbolzen, begrenzt ist.

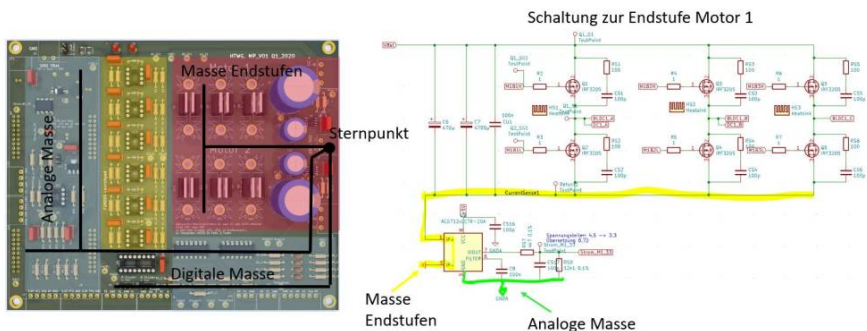
Bei genauerer Betrachtung der obigen Abbildung zu den Stromverläufen fällt auf, dass es sich eigentlich um genau die Situation der galvanischen Kopplung handelt, welche wir eben gerade nicht haben wollen. In den meisten Fällen stellt das Chassis des Gesamtfahrzeug eine derart niederohmige und niederinduktive Verbindung dar, dass der entstehende Spannungsabfall nicht ins Gewicht fällt.

Im Dreileitersystem können wir die Ströme in Gleich- und Gegentaktströme unterteilen. Zu Reduktion von Gleichtaktströmen haben wir verschiedene Filtertypen kennengelernt. Leider lassen sich die sogenannten Y-Kondensatoren nur einsetzen falls die Störströme lokal auf das Chassis bzw. den PE-Leiter abgeführt werden können. Nicht möglich ist dies im Fahrzeug bei z.B. nicht

leitenden Gehäusen und im Hausnetz allgemein bei Geräten ohne PE- Anschluss. Zur Unterdrückung von Gleichtaktstörungen bleibt uns somit nur die Möglichkeit über Gleichtaktdrosseln.

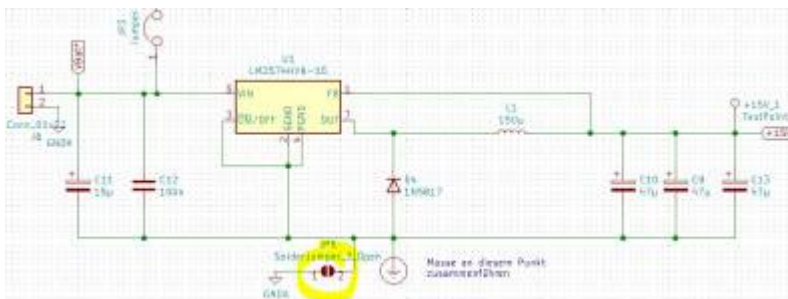
Auf der Platine arbeiten wir normalerweise stets im Zweileitersystem. Das Chassispotential oder PE-Potential liegt typischerweise nur noch am Stecker an zur Anbindung der von Y-Kondensatoren und wird auf dem PCB nicht weiter geführt. PE und N- Anschluss dürfen wir keinesfalls zusammenführen, da dies zu einer sofortigen Auslösung des FI-Schutzschalters führen würde. Im Automotiv Bereich obliegt es den Regularien des Fahrzeugherstellers ob wird KI. 31 und eine ggf. vorhandene Anbindung über das Gehäuse zur Masse zusammenführen dürfen.

Auf der Platine stellt unser größtes Problem eine unsaubere Referenz bzw. Masse dar. Vielleicht ist ihnen auch schon einmal aufgefallen, dass Messergebnisse mit dem Oszilloskop von der Position der Anbindung der Tastkopfreferenz abhängig sind. In so einem Fall liegt eine „hüpfend“ Masse vor, ein klassisches Problem der galvanischen Kopplung. Daher ist es sinnvoll sich eine Ruhepunkt zu definieren welches den Referenz Null-Knoten darstellt. In einer Simulation wäre das der Spice Null-Knoten. Von diesem Referenzpunkt gehen wir dann sternförmig zu den einzelnen Zonen in unserem Schaltbild und sorgen so dafür, dass ein Spannungsabfall entlang eines Rückstrompfades sich nicht auf benachbarte Systeme auswirken kann.



Obige Abbildung zeigt die sternförmige Masseverteilung unserer BLDC Motorendstufe. Direkt am Stecker der Leistungselektronik wird der Sternpunkt definiert von welchem aus die Masse in die einzelnen Zonen verteilt wird. Rechts im Bild ist die Umsetzung anhand eines Schaltplanausschnitt dargestellt. Die meisten Layoutprogramme unterstützen mehr als als eine Darstellung der Masse. Unterschieden wird meist zwischen analoger und digitalen Masse sowie dem Erdsymbol. In der Simulation ist dies jedoch meist nicht möglich, da es nur eine

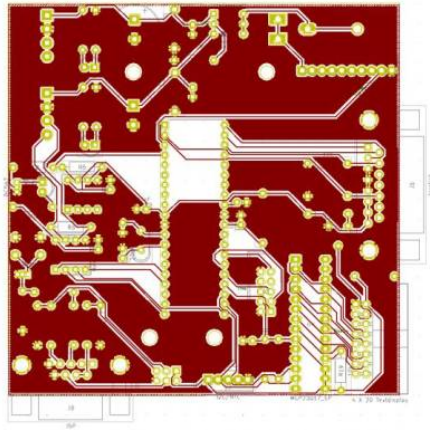
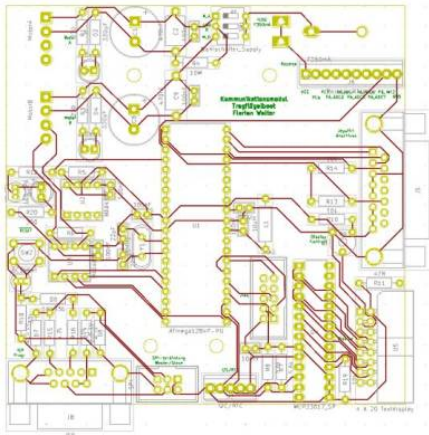
Hilfreich ist es die Knoten der unterschiedlichen Massepunkte erst zum Ende des Layoutprozess zu verbinden, da das Programm andernfalls von identischen Potenzialen ausgeht womit keine Unterscheidung mehr möglich ist (siehe unten, lokale Masseführung)



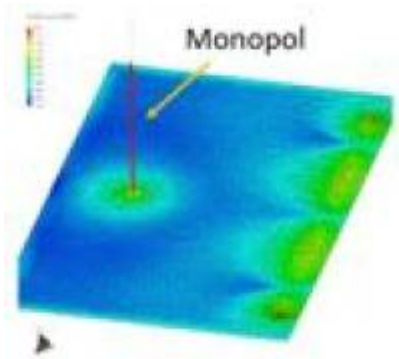
hatte war dies anders. Je mehr Kupfer aufgelöst werden sollte, desto schneller musste das Ätzbad erneuert werden. Man hatte also ein natürliches Interesse daran möglichst wenig Kupfer zu entfernen.
Heute hat sich hinsichtlich der EMV die Erkenntnis durchgesetzt sämtliche freien



Flächen mit der Masse (besser dem Rückleiter) zu füllen. Frei nach dem Motto: Viel hilft viel. Oder anders ausgedrückt, die eigentliche Leiterbahn wird aus der Massefläche freigeschnitten. Da wir an der Hochschule eine Fräse verwenden haben wir jetzt ein ähnliches Problem wie zuvor mit dem Ätzbad. Je mehr Kupfer entfernt werden muss, desto schneller verschleißt unser Fräser. Nachfolgende Abbildung zeigt eine typische Platine im Vergleich mit und ohne Masseflächen.



Ich empfehle allerdings stets den Rückleiter zuerst mit Traces auszubilden und sich nicht auf die folgende Massefläche zu verlassen. Um eine möglichst niederohmige Massefläche zu erhalten sollten bei industriell gefertigten Leiterplatten möglichst viele Verbindungen (Vias) zwischen den Masseflächen auf der Ober- und Unterseite hergestellt werden. Es ist nicht schlimm falls die Platine dann wie ein Schweizer Käse aussieht.



Das Ziel der großflächigen Masse basiert darauf den Rückstrom geometrisch möglichst nahe am Leiter zu führen. Dabei hilft uns, dass die Rückströme gewillt sind von sich aus möglichst nahe am Leiter zu verlaufen. Wir müssen ihm sozusagen nur die Chance dazu zu geben.....

Für Platinen mit mehr als zwei leitfähigen Lagen (Multilayer) wird meist eine komplette Lage der Masse gewidmet und zusätzlich frei Flächen mit Masse aufgefüllt in allen Lagen. Man erhält damit eine Art Schirmung, zu mindestens für die innenliegenden Leiterbahnen. Schön zu sehen ist dieses Verhalten in der Simulation zur Eindrahtleitung für Frequenzen $> 100\text{MHz}$. Die Flächenfarben zeigen den Stromfluss unterhalb der Leitung welche 5cm über einer leitfähigen Tischoberfläche angebracht wird.

Es wäre theoretisch auch möglich anstatt mit Masse, mit dem Versorgungspotenzial zu füllen. Allerdings existieren in den Schaltungen natürlich meist mehrere unterschiedliche Potenziale und es wäre hinsichtlich auftretender Kurzschlüsse ziemlich unhandlich.

Fragen Sie einen EMV-Ingenieur nach einer speziellen Maßnahmen für das gegebene Problem wird er vermutlich sagen: Kommt darauf an Das bedeutet alle hier aufgestellten Regeln sind Hinweise welche die EMV in Summe verbessern aber alleine meist nicht der Weisheit letzten Schluss darstellen. Die Massefläche ist ein schönes Beispiel: Im Allgemeine wie beschrieben eine gute Idee. Allerdings auch hier stets kritisch hinterfragen. Bei offenen Spulen ([Stabkerndrosseln](#)) möchten wir natürlich nicht, dass das Magnetfeld in die unterliegende Massefläche einkoppelt. Hier ist es besonders wichtig nicht mit gestückelten Masseflächen Schleifen unter der Spule aufzuspannen.

Abbiegende Traces

Um abbiegende oder abknickende Leiterbahnen ranken sich viele Mythen. Der Einfluss auf die EMV wird meist überschätzt. Fakt ist, die Zeiten halbkreisförmiger bzw. bogenförmiger Abbiegungen sind vorbei. Diese stammen noch aus der Zeit als Leiterbahnen von Hand aufgeklebt wurden. Bei der Erstellung der nachfolgenden Beispiele konnte ich in KiCAD die Einstellung dazu auf die Schnelle



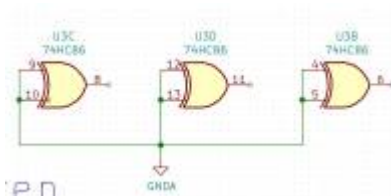
gar nicht finden. Es verbleiben somit 90° und 2x45° Winkel. Durchgesetzt haben sich die 45° Winkel. Dabei ergeben sich folgende Vorteile. Liegt die Signallaufzeit der Signale in der Größenordnung der Laufzeiten auf der Leiterkarte können sich natürlich auch hier Reflexionen ausbilden. Das bedeutet wir

müssen auch die Leiterbahnen mit dem Wellenwiderstand abschließen. Die allermeisten Layoutprogramme haben mittlerweile kleine Berechnungsprogramme integriert welche für uns den Wellenwiderstand der Traces berechnen. Basis dazu sind die Maxwellschen Gleichungen welche sich für unsere hier zweidimensionalen Probleme auf der Leiterkarte geschlossen lösen lassen. Bei jeder Abbiegung der Leiterbahn erzeugen wir jedoch eine Inhomogenität des Wellenwiderstandes welche sich mit Hilfe der 45° Bögen reduzieren lassen. Vermutlich wäre für diesen Fall ein runder Bogen die besser Wahl. Ich hoffe Sie erinnern sich, 300MHz entspricht einer Wellenlänge von 1m. Das bedeutet unsere Signale müssen schon sehr schnell sein (>1GHz) damit wir uns darüber Gedanken machen müssen.

Es hält sich auch die Überlieferung, dass sich in der PCB Fertigung weniger Probleme im Ätzbad und bei der Lackierung ergeben. Allerdings kenne ich dazu keine gesicherte Quelle bzw. Aussage eines Herstellers. In den Regeln (Design Rules) der Hersteller ist mir bisher dazu nichts bekannt.

Somit verbleibt nur noch vermutlich der am häufigsten verwendet Grund für 45° Winkel: Es sieht einfach professioneller aus!

Ungenutzte Pins



In manchen Fällen kann es vorkommen, dass integrierte Schaltkreise mehr Funktionalität bieten als notwendig. Typischerweise bei binären Schaltkreisen oder Operationsverstärkern. Obwohl nicht verwendet, sollten die Eingänge dieser Schaltkreise auf ein definiertes Potential

gebracht werden. Damit stellen wir sicher, dass es keine „in der Luft hängende

(floating)“ Potentiale gibt welche dann über eine interne Kopplung hochfrequente Signale emittieren oder von außen Störungen einfangen.

Software

Es hört sich vielleicht merkwürdig an, aber auch Software kann einen kleinen Teil zur EMV beitragen. Über die Software wird in vielen Fällen die Taktfrequenz leistungselektronischer Systeme oder Übertragungsraten eingestellt. Ok, die Vorgabe der Frequenz erfolgt in der funktionalen Entwicklung, allerdings können sie den Softwareentwickler bitten

- verschiedene Taktfrequenzen für die EMV-Prüfung zu implementieren mit einer einfachen Umschaltung oder über verschiedene Softwarestände
- falls funktional möglich ein Spread-Spektrum Verfahren zu implementieren, siehe [spread-spectrum-frequency-modulation-reduces-emi.pdf](#)
- Bei ungenutzten Pins die internen Pull-Up Widerstände des μC aktivieren

13.0 ESD und Blitzentladungen

Wir alle kennen sowohl diese kleinen lästigen Entladungen, welchen wir ständig ausgesetzt sind, wie natürlich auch die meist imposanten Blitzentladungen. Beide haben gemein, dass lokal die Spannungsfestigkeit der Luft, die sogenannte [Durchschlagfestigkeit](#), überschritten wird. Der exakter Wert dieser Festigkeit ist von mehreren Faktoren wie die Luftfeuchtigkeit, dem Luftdruck und die Zusammensetzung der Luft abhängig. Als Daumenwert kann für die Durchschlagfestigkeit in Luft 30kV pro cm (Einheit entspricht einer elektrischen Feldstärke) angenommen werden. Man spricht daher auch von der Durchbruchfeldstärke.

Bei der ESD-Prüfung versuchen wir jetzt exakt diese Impulse nachzubilden und beaufschlagen entsprechend alle Stellen unseres Prüflings welche der zukünftige Anwender berühren kann. Das kann eine ganz schön langwierige und ggf. auch langweilige Aufgabe darstellen. Stellen sie sich vor wie viele Punkte ein Kunde im Fahrzeug anfassen kann....

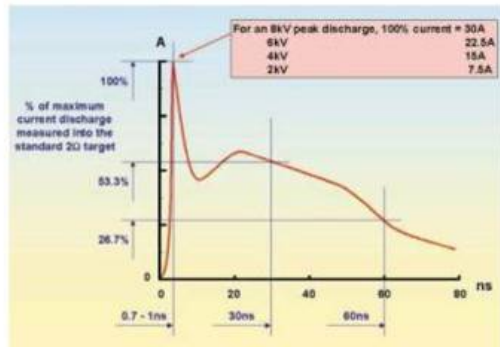
Nachfolgende Abbildung zeigt eine ESD-Pistole von [Teseq](#) und den typischen

Impulsverlauf einer ESD-Entladung. Beachten sie bitte die Zeitachse: ns!



ESD-Pistole

Bild: Teseq.com



ESD-Testimpuls

Bild: compliance-club.com

Bitte überlegen sie sich für die Vorlesung (Prüfung) folgende Punkte:

- Wie würden sie eine ESD Prüfung definieren? Wann ist die Prüfung bestanden?
- Wie würden sie herausfinden welche Prüfspannungen sie verwenden müssen, welchen Prüfaufbau?
- Wie lange soll die Prüfung dauern?

Das bedeutet wir tun so als würden wir eine eigene ESD Prüfnorm entwerfen. Ganz abwegig ist diese Situation im Berufsleben nicht. Besonders Fahrzeughersteller vertrauen bei Festigkeitsprüfungen auf ihre eigenen Erfahrungen. Es hilft ja nichts wenn ihr Produkt reihenweise Ausfälle zeigt und sie den Kunden auf eine eingehaltene EMV-Norm hinweisen. Da die ESD Belastung bei der Ladungstrennung hauptsächlich von den eingesetzten Materialien abhängt, ist es naheliegend die jeweilige Situation zu untersuchen. Zum Beispiel:

- Haben verschiedenartige, nichtleitende Materialien Kontakt zueinander
- Beschreiben diese Teile eine Relativbewegung (im Fahrzeug z.B: Motorfest, Karosseriefest)
- Kann es zu einer Ladungstrennung kommen

Der Klassiker bei der Aufladung sind stets Riementriebe oder Transportrutschen. In manchen Fällen ist es auch wichtig den gesamten Lebenszyklus der Produkte

hinsichtlich ESD zu bewerten. Bei der Montage von Fahrzeugsteuergeräten bei der Fahrzeugmontage werden sehr viele Maßnahmen integriert, damit die sensible Elektronik im Fahrzeug durch ESD keinen Schaden nimmt. Allerdings können diese Maßnahmen nicht dauerhaft garantiert werden. Denken sie an die Reparatur einzelner Steuergeräte in einer beliebigen Werkstatt. Auch hier muss der Tausch ohne ESD Schaden über die Bühne gehen, was bedeutet, dass der ESD Schutz jetzt das Gerät selber übernehmen muss.

Überlegen und recherchieren sie bitte jetzt verschiedene Möglichkeiten wie wir uns und elektronische Bauelemente in der Entwicklung und Produktion schützen können. Vielleicht kennen sie ja auch schon verschiedene Maßnahmen aus dem Praktikum oder ihrer Ausbildung.

ESD in der Entwicklung

Das ESD Thema wurde bis in die 90er Jahre in den Entwicklungsabteilung ziemlich lasch gehandhabt. Von intensiven ESD - Schutzmaßnahmen der Halbleiter bzw. elektronischen Geräte im Allgemeinen war noch keine Spur zu sehen. Das hat sich geändert, als die ESD bedingten Ausfälle in der Entwicklung und damit die Entwicklungskosten angestiegen sind. Durch die erhöhte Packungsdichte integrierter Schaltkreise wurde die ESD Anfälligkeit zusätzlich erhöht.

Obwohl wir an der Hochschule hinsichtlich ESD noch kein gutes Bild abgeben kann ich ihnen für die Abschlussarbeit nur dringend raten die oben recherchierten ESD- Schutzmaßnahmen bei einer Hardwarearbeit anzuwenden. Ansonsten kann der für uns in der Elektrotechnik schlimmste Fall auftreten. Die Schaltung funktioniert nur ein bisschen oder halb! Wie immer wäre es besser die Schaltung würde gar nicht funktionieren, da wir Fehler meist zuerst bei uns im Aufbau bzw. Schaltplan suchen. Kein Problem falls sie die Schaltung zum zehnten mal aufbauen und das Verhalten des integrierten Bausteins sehr gut kennen. In der Entwicklung ist es ja oft eine Erstinbetriebnahme und bis sie den Fehler gefunden haben sind mehrere Tage ins Land gegangen. Typische Anzeichen für eine ESD-Vorschädigung eines integrierten Schaltkreises:

- Der Ausgangspuffer einzelner Pins am μC lässt sich nicht ansteuern
- Der μC funktioniert, aber einzelne Pins zeigen merkwürdigen Verhalten, z.B. halber Logikpegel, falsche ADC Werte
- Der Operationsverstärker verhält sich nicht wie erwartet
- Pins die hochohmig sein sollten sind plötzlich niederohmig und umgekehrt
- Sind mehrere gleiche Funktionen in einem Gehäuse funktionieren manche nicht wie erwartet
- ...

Während der Abschlussarbeit verlieren sie mit etwas Glück nur wenige Tage bis sie den Fehler gefunden haben. Die Kosten, falls sie einen Dauerlauffest wiederholen müssen aufgrund einer ESD Vorschädigung der seit mehreren Monaten im Klimaschrank läuft, belaufen sich meist auf sechsstelligen Eurobeträge, abgesehen dem daraus resultierenden Terminverzug. Das Gemeine an der ESD Entladung ist zusätzlich, dass wir nur einen Bruchteil der auftretenden Entladungen wahrnehmen. Viele Entladungen geringerer Spannung können von uns auf die Halbleiter übergehen ohne, dass wir davon etwas mitbekommen. Daher kommt man heutzutage bei den meisten Firmen nur noch in Entwicklungs- und Produktionsstätten mit entsprechender ESD-Schutzausrüstung sowie einer Leitfähigkeitsprüfung am Eingang mit entsprechender Schleuse.

Blitzentladungen

Blitzentladungen sind uns in ihrem Auftreten und Entstehung vermutlich allen bekannt. Bei der Beeinflussung unterscheiden wir, wie auch bei der ESD-Entladung, drei verschiedene Szenarien.

- Direkte Beeinflussung / Beaufschlagung
- Indirekte Beeinflussung, z.B. Überspannung (abgeschwächter Impuls)
- Beeinflussung durch entstehenden elektromagnetischen Impuls

Blitzentladungen lassen sich nach deren Polarität und Ladung charakterisieren. Die größte Anzahl an Blitzentladungen bemerken wir nur indirekt über deren elektromagnetisches Feld, da es sich um Wolke zu Wolke Entladungen handelt. Auf diesem Prinzip arbeiten Ortungssysteme zur Detektion von Ort und Intensität der Entladungen. Eine direkte Blitzentladung in ein elektronisches Gerät ist sehr selten und nur für elektrische Anlagen im Freien von Bedeutung (z.B. Sende-Empfangsantennen, Wetterbalons,). Bei Entladungsströmen von mehreren Kiloampere hilft uns auch der beste Überspannungsschutz nicht mehr weiter. Am Einfachsten stellt man sich eine Blitzentladung als Stromquelle vor. Das bedeutet es gibt überschüssige Elektronen welche den Ladungsaustausch suchen. Überschreitet die Potenzialdifferenz Wolke-Wolke oder Wolke-Erde die Durchbruchfeldstärke in der Luft kommt es zur Entladung. Bei einer Entladung über einen Blitzableiter ergibt sich die resultierende Spannungsüberhöhung aus dem Widerstand welcher der Blitzableiter dem Stromfluss entgegenstellt. Das bedeutet, dass die Gefährdung bei einem direkten Blitzeinschlag in ein Gebäude mit Blitzableiter darin besteht, dass unser Referenzpotenzial, also der PE-Schutzleiter und damit auch der Neutralleiter, angehoben wird. Die Überspannung kommt als in diesem Fall sozusagen durch die Hintertür (über PEN). Daher ist es bei bestehenden Blitzableitern wichtig stets für eine möglichst

niederohmige Verbindung von der Fangleitung bis ins Erdreich zu sorgen. Ohne Blitzableiter sucht sich der Blitzstrom seinen Weg durch das Gebäude. Wir wissen bei einem gegebenen Stromfluss hängt die eingetragene Energiemenge vom Strom (quadratisch) und vom Widerstand ab. Je größer der Widerstand, desto größer die eingetragene Energie und damit die Brandgefahr.

Häufiger, aber für unserer zu schützenden Geräte nicht weniger gefährlich, sind auftretende Überspannungen aufgrund einer entfernten Blitzentladung in eine Versorgungsleitung. Der Überspannungsimpuls breitet sich entlang der Leitung aus bis die Potenzialanhebung auch auf unser Gerät trifft. Der typische Überspannungsimpuls lässt sich mit einer Anstiegszeit von $1,2\mu\text{s}$ und einem Rücken von $50\mu\text{s}$ beschreiben. Wir sprechen dabei von einem **Surge** Impuls. Auf einer Freileitung benötigt der Impuls ca. 5ms. Das bedeutet der Impuls ist deutlich schneller als dessen Laufzeit entlang der Leitung. Trifft der Impuls auf das Leitungsende wird er aufgrund des für ihn offenen Endes noch einmal verdoppelt und läuft zurück. Wir haben dann die Situation einer klassischen Reflexion entlang einer Leitung. Das ist ein Grund dafür, dass Aussiedlerhöfe als letzter Teilnehmer an der Versorgungsleitung oft häufiger von Überspannungsschäden betroffen sind als angeschlossene Häuser entlang der Leitung. Aufgrund des ohmschen Widerstandes der Leitung erfährt der Impuls eine Dämpfung entlang seines Ausbreitungswegs womit die Spannungsamplitude von der Entfernung zur Einschlagsstelle abhängig ist.

Verdeutlichen wir uns noch einmal die Größenordnungen unserer beiden Überspannungsimpulse:

- ESD: Anstiegszeit innerhalb weniger ns (10^{-9}s)
- Surge: Anstiegszeit innerhalb weniger μs (10^{-6}s)



Bei einer Entladung wird sich somit auch eine elektromagnetische Welle lösen, welche sich kreisförmig um die Entladung ausbreitet. Bei Blitzentladungen kennen wir das typische Störgeräusch im Radioempfang während einer Entladung. Aufgrund der meist hohen Entfernung zur Blitzentladung nehmen wir hier nur ein kurzes Knacken im Radioempfang wahr, was bedeutet, dass hauptsächlich Funkdienste beeinträchtigt sind. ESD Entladungen hingegen finden lokal und damit sehr nahe an unseren Geräten statt. Die Entladungsdauer ist noch einmal eine Größenordnung geringer als die der Blitzentladung womit das emittierte

Frequenzspektrum bis hoch in den GHz Bereich reicht.

Bei der ESD-Prüfung wird daher neben der direkten Beaufschlagung das

Verhalten hinsichtlich einer indirekten Beeinflussung ermittelt. Der ESD Impuls wird dabei auf eine Koppelplatte (Bild rechts) entladen. Die Platte ist eine einfache leitfähige Platte (0,5m x 0,5m) aus Kupfer oder Aluminium welche über hochohmige Widerstände an die Referenzmasse der ESD-Pistole angebunden wird. Der Prüfling befindet sich während der Entladung in 10cm Abstand zur Koppelplatte und damit im direkten feldgebundenen Einflussbereich der Entladung.

Weitere Infos zur ESD Prüfung finden sich unter DIN EN 61000-4-2, Prüfung der Störfestigkeit gegen die Entladung statischer Elektrizität.

Schutz vor Überspannungen

Es gibt mehrere Möglichkeiten sich gegenüber Überspannungen zu schützen. Meist ist eine Maßnahme alleine nicht ausreichend, sondern es bedarf mehrere Schutzelemente um die notwendige Spannungsreduktion zu erreichen. Wir sprechen dann von einem gestaffelten Schutz. Alles Schutzmaßnahmen vereinen, dass die Überspannung gegenüber dem Referenzpotenzial abgebaut wird. Das bedeutet, die Schutzelemente werden stets parallel zum Verbraucher positioniert. Vielleicht kennen Sie von zu Hause Steckdosenleisten mit integriertem Überspannungsschutz. Je nach Güte sind eine oder mehrere der folgenden Element darin verbaut.

Gasableiter

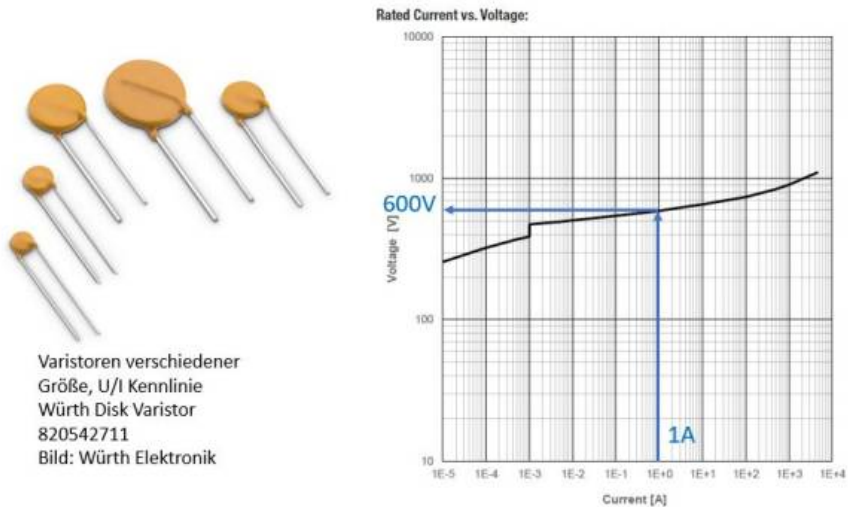


Gasableiter sind sozusagen unsere erste Speerspitze zur Reduktion der Überspannung, also ein Fall fürs Grobe Der Aufbau entspricht im Wesentlichen einem kleinen Plattenkondensator bei dem ab dem Überschreiten der Durchbruchspannung eine gezielte Entladung stattfindet. Die Ansprechspannung ab wann sich eine Entladung ausbildet kann über das eingesetzte Gas (z.B. Stickstoff oder Edelgase wie Neon, Argon) und den Druck im Inneren der Elemente reguliert werden. Beim Einsatz ist zu beachten, dass der Strom

im Entladungsfall nur durch die Lichtbogenspannung begrenzt wird. Nach dem Abklingen der Überspannung erlischt der Lichtbogen im nächsten Nulldurchgang der Netzspannung. Bei Gleichspannung ist dies nicht der Fall und wir müssen über eine externe Maßnahme für eine Strombegrenzung oder Stromabriss sorgen, zum Beispiel mit einer Sicherung.

Varistoren

Varistoren sind spannungsabhängige, nichtlineare Widerstände aus Metalloxid. Im Gegensatz zum Gasableiter ist das Ansprechverhalten unabhängig von der Anstiegszeit der Spannung (du/dt). Nachfolgende Abbildung zeigt Varistoren verschiedener Größe und eine typische U/I Kennlinie in logarithmischer Darstellung.



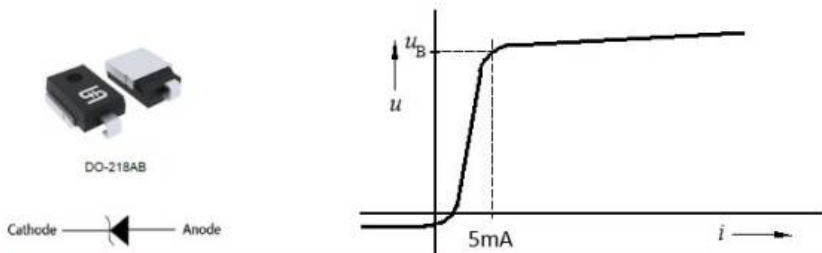
Die Kennlinie gehört zu einem Varistor der typischerweise zum Schutz einer Anwendung am 230V Versorgungsnetz eingesetzt wird. Zu beachten ist, dass bereits bei der Netzspannung ein geringer Ableitstrom fließen kann. Mit steigender Spannung nimmt der Widerstand des Ableiters schlagartig ab und ermöglicht den Kurzschluss des Überspannungsimpuls. In vielen Fällen wird die Spannung angegeben bei der 1A fließen, in unserem Fall also knapp 600V. Würth Elektronik gibt für seine Varistoren die 5A Spannung an. Das Diagramm zeigt, dass die Spannung ein vielfaches der Nennspannung betragen muss, damit ein signifikanter Ableitstrom fließen kann.

Das bedeutet, wir müssen uns nach unserem Grobschutz mit Funkenstrecke oder Varistor jetzt noch um Schutzelemente kümmern welche die Spannung noch weiter reduzieren können.

Dioden

Dioden kennen wir natürlich sehr gut aus der funktionalen

Schaltungsentwicklung. Als Überspannungsschutz eignen sich nur bedingt bzw. bis zu ihrer Durchlassspannung. Zum Schutz werden daher vielmehr Zener-Dioden eingesetzt. Diese Dioden werden in Sperrichtung verbaut und wir nutzen den Effekt einer schlagartig eintretenden Leitfähigkeit sobald die Zenerspannung überschritten wird. Supressordioden oder TVS (Transient Voltage Suppressor) sind den Zenerdioden sehr ähnlich, allerdings mit höheren Durchbruchspannungen (analog der Zenerspannung) bzw. mit größerer Stromtragfähigkeit verfügbar. Alle Diodentypen sind in ihrer Grundform unidirektional womit wir zum Schutz von positiven als auch negativen Überspannungen jeweils zwei Dioden in Reihe schalten müssen „back to back“.



ELECTRICAL SPECIFICATIONS ($T_A = 25^\circ\text{C}$ unless otherwise noted)								
Part number	Marking code	Breakdown voltage V_{BR} at I_T (V) (Note 1)		Test current I_T (mA)	Working stand-off voltage V_{WM} (V)	Maximum blocking leakage current I_R at V_{WM} (μA) (Note 1)	Maximum peak impulse current I_{PPM} (A) $t_p = 10/1000 \mu\text{s}$	Maximum clamping voltage V_C at I_{PPM} (V)
		Min.	Max.					
TLD5S10AH	TLD5S10A	11.1	12.3	5.0	10.0	15	212	17.0

Die Abbildung zeigt ein Ausschnitt aus dem Datenblatt der TVS Diode DO-218. Die Durchbruchspannung ist hier als der Wert angegeben ab dem ein Strom von 5mA fließt. Bitte beachten sie den schlagartigen Übergang in die Leitfähigkeit:

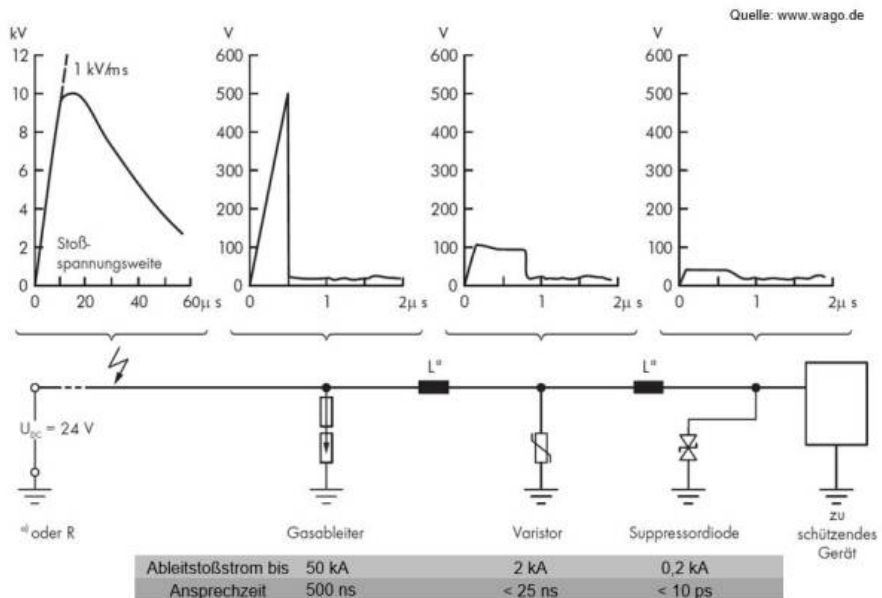
- 10,0V → 15 μA
- 11,1V - 12,3V → 5mA

Zu beachten ist, dass TVS Diode eine relativ hohe Kapazität aufweisen welche bis in den nF Bereich reichen kann. Das bedeutet wir können die Dioden ggf. nicht zum Schutz hochfrequenter Signale einsetzen (Audio, Funkanwendungen usw.).

Gestaffelter Schutz

Ein gestaffelter oder mehrstufiger Schutz enthält mehrere Bauelemente und baut

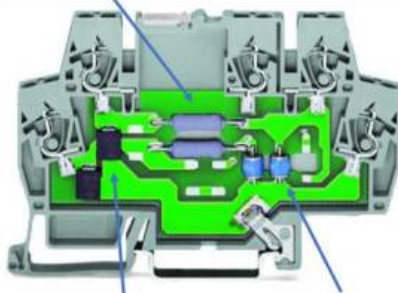
die auftretende Überspannung nach und nach ab. Nachfolgend ein Beispiel für einen gestaffelten Schutz für eine 24V Versorgungsschiene.



Ein Problem bleibt allerdings. Liegen alle Schutzelemente am identischen Potenzial (also gleicher Anschluss) würde der Überspannungsimpuls unsere TVS Diode zerstören bevor der Gasableiter überhaupt angesprochen hat. Daher müssen wir die einzelnen Elemente gegeneinander Entkoppeln bzw. wir müssen davor sorgen, dass der gedämpfte Impuls verzögert an den weiteren Schutzelementen ankommt. Ströme bremsen bzw. verzögern können wir mit Hilfe einer Induktivität welche zwischen den jeweiligen Stufen angebracht werden.

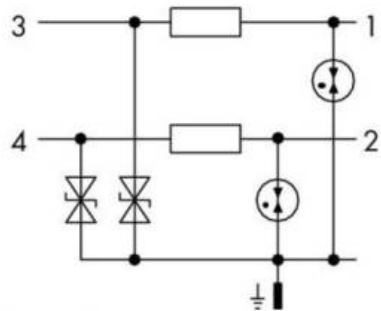
Abschließend ein Beispiel der Firma Wago welche die gesamte Schutzschaltungen in Reihenklemmen zur Hutschienenmontage integriert.

Entkopplungsinduktivität



Bidirektionale TVS Dioden

Gasableiter



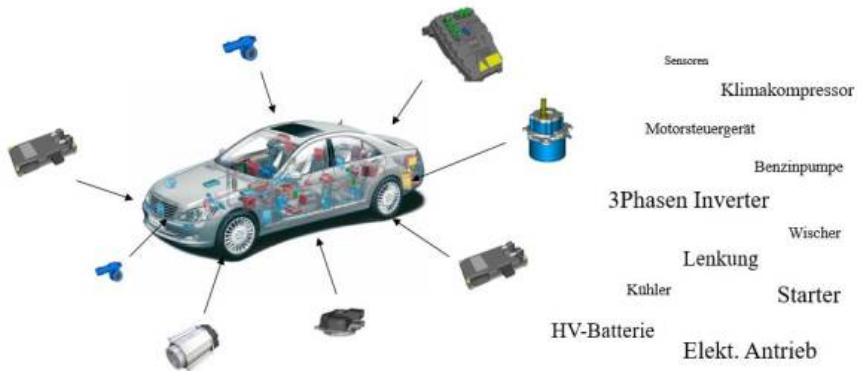
Bilder: Datenblatt [Wago 792-800](#)
 Überspannungsschutzmodul für Signaltechnik

Die Schutzbeschaltung ist ausgelegt zum Schutz von 24V DC Versorgungsmodulen gegenüber Überspannungen nach DIN EN 61643-21 (Überspannungsschutzgeräte für Niederspannung).

Achten Sie doch einfach beim nächsten Einkauf im Baumarkt oder Aldi darauf welche Schutzelemente in Mehrfachsteckdosen mit integriertem Überspannungsschutz verbaut sind. Vielleicht ist ein Hinweis auf dem Typenschild zu finden.

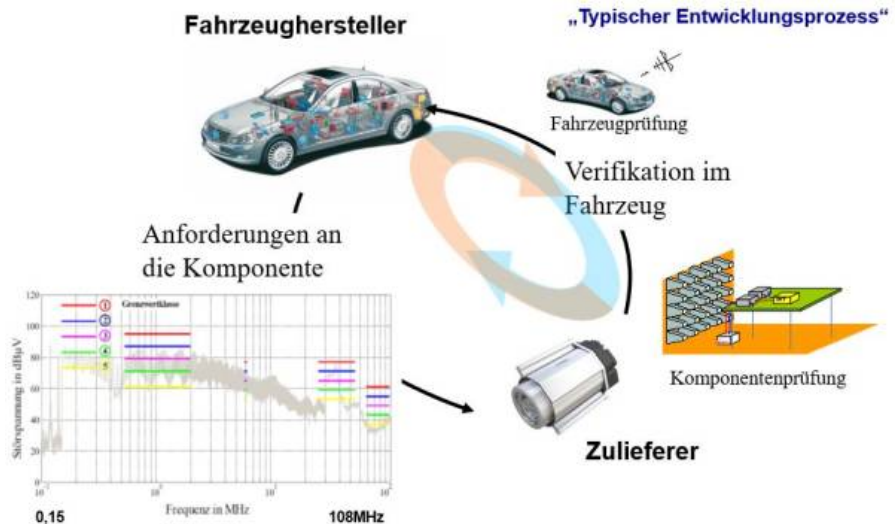
14.0 Automotive EMV

Auch Fahrzeuge benötigen eine Art CE-Kennzeichen. Allerdings sprechen wir dann von der Straßenzulassung oder Homologation. Obwohl EMV-Grenzwerte vorgegeben sind setzen die Fahrzeughersteller meist deutlich strengere Anforderungen an die Komponenten als auch an das Gesamtfahrzeug an. Wie bereits zu Beginn erwähnt, EMV ist auch ein Qualitätsmerkmal. Aufgrund der hohen Elektronikdichte im Fahrzeug natürlich auch nicht weiter verwunderlich, insbesondere da wir dem Kunden alle Radio und Funkdienste zur Verfügung stellen wollen.

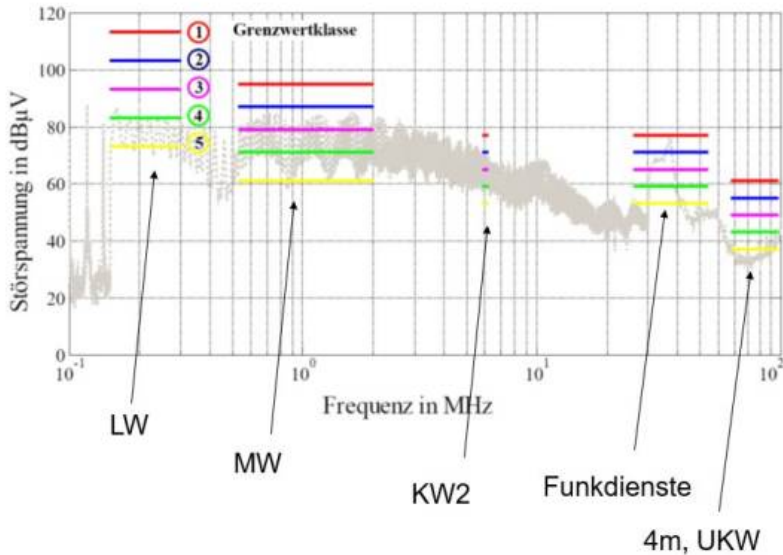


Während des Entwicklungsprozess einer neuen Fahrzeugkomponenten werden die ersten EMV-Messungen bereits mit frühen A-Mustern durchgeführt oder idealerweise zusätzlich zu funktionalen Tests am Brettmuster. Spätestens bevor die Komponenten in Fahrzeugprototypen oder der Baustufe verwendet werden erfolgt eine reduzierte EMV-Messung der wichtigsten Prüfungen (leitungsgebundene und gestrahlte Emission, Störfestigkeit). In den weiteren Musterständen werden die EMV-Eigenschaften der Komponenten dann bis zur Serienreife entwickelt.

Ob die Komponente allen Anforderungen genügt zeigt sich allerdings erst bei der EMV-Prüfung im Fahrzeug. Hier wird explizit danach gesucht ob die Emissionen der Komponente Funksignale beeinflussen können. Dazu wird zum Beispiel die Störspannung am Ausgang der verschiedenen Antennenverstärker bei aktiver Komponente gemessen. Wie wir im späteren Verlauf des Kapitels sehen werden kann es leider vorkommen, dass eine Komponente im Fahrzeug auffällig werden kann obwohl im Labor alle Grenzwerte eingehalten werden. Der Unterschied liegt natürlich zum einen an der Leitungsimpedanz im Fahrzeug, welche wir während der Komponentenmessung über die Netznachbildungen annähern, sowie der verschiedenen Masseverhältnisse (Anbindung zum Chassis).



Nachfolgende Abbildung zeigt eine Messung der leitungsgebundener Emissionen mit den Grenzwerten nach CISPR25. Man unterscheidet verschiedene Grenzwertklassen. Der Fahrzeughersteller wählt die Grenzwertklasse für das jeweilige Steuergerät je nach Gefährdungspotenzial und Lage im Fahrzeug. Je näher sich die Komponente an Antennenstrukturen oder Verstärkern befindet, desto strenger muss der Grenzwert gewählt werden.



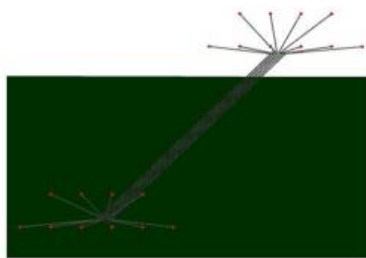
Die Grenzwerte sind nicht für den gesamten Frequenzbereich definiert, lediglich im Bereich der Arbeitsfrequenz der verschiedenen Funksignale. Da in Europa die Lang- Mittel- und Kurzwelle kaum noch nachgefragt werden, kann es sein, dass der Hersteller die Grenzwerte für diesen Bereich ignoriert.

Hier wird nur beispielhaft auf die leitungsgebundene Emission eingegangen. Für die gestrahlte Emission wird in gleicher Weise vorgegangen.

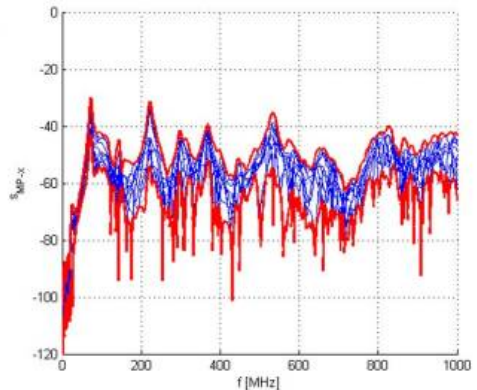
Leitungsverbindungen im Fahrzeug

Wie bereits in den vorangegangenen Kapiteln gesehen spielen Leiterverbindungen in der EMV eine große Rolle. Sie verbreiten die Emissionen nicht nur wie eine Antenne sondern führen leitungsgebundene Störströme direkt zu benachbarten Geräten. Die Störausbreitung muss nicht nur auf der Versorgungsleitung erfolgen, sondern kann sich aufgrund von Überkopplungen zu benachbarten Leitern im Kabelbaum auf beliebigen Adern verteilen. Nachfolgende Abbildung zeigt für das Beispiel eines zehnnadrigen Kabelbaum über einer leitfähigen Tischfläche die Emissionen zu einer Monopolantenne in einem Meter Abstand. Bewertet werden jeweils die Transmission (Bild rechts).

10 Adern Kabelbaum, Transmission zum Mon

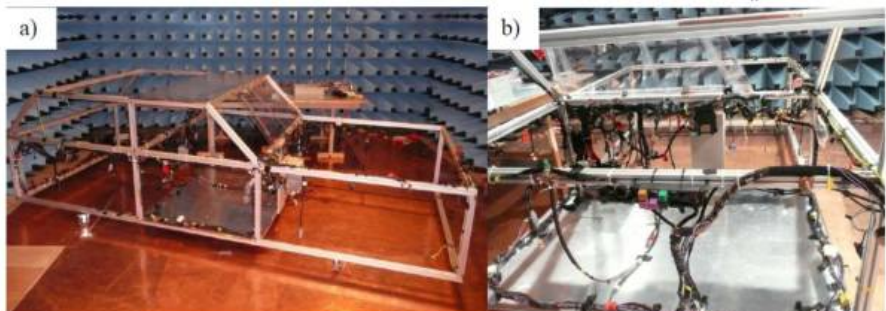


Simulationsmodell



Mir geht es nicht darum den exakten Verlauf mit Ihnen zu besprechen, sondern einfach noch einmal die Varianz aufzuzeigen welche sich mit dem Einsatz eines Kabelbaums ergibt. Die roten Linien zeigen die minimal und maximal mögliche Transmission zu einer Monopolantenne. Je nach verwendeter Ader ergeben sich darin Unterschiede von bis zu 20dB. Das bedeutet, je nach gewählter Ader innerhalb des Bündels können die Messwerte zur gestrahlten Emission entweder 10dB über dem Grenzwert oder 10dB unter dem Grenzwert liegen. Zehn Adern innerhalb eines Kabelbaum sind eher die Regel als die Ausnahme. Teilweise verlaufen armdicke Leitungen im Inneren der Fahrzeugstruktur. Um die Kopplung und Impedanzeigenschaften der Leitungen genauer zu untersuchen schauen wir uns nun das Verhalten realer Fahrzeugkabelbäume genauer an.

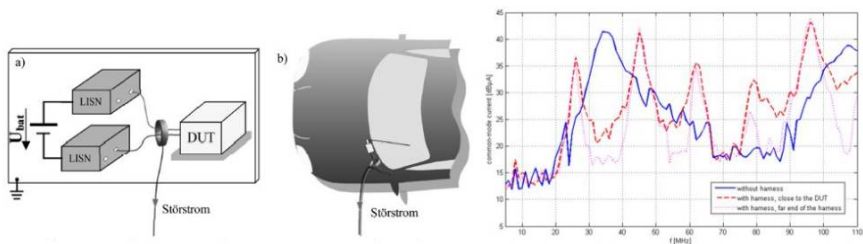
Dazu wird der Fahrzeugkabelbaum eines gesamten Fahrzeugs an einem Hilfsrahmen befestigt und die Impedanzen an verschiedenen Stellen vermessen. Kabelwege lassen sich dadurch einfach nachverfolgen sowie bei der Emissionsmessung die Ausbreitung von Störströmen im Gesamtaufbau verfolgen. Nachfolgende Abbildung zeigt den Laboraufbau zur Untersuchung der Kabelbäume.



Führen wir jetzt unseren ersten Versuch mit dem Aufbau durch und integrieren den uns sehr gut bekannten Tiefsetzsteller. In unserem Testaufbau bewerten wird dann die ausgehenden Gleichtaktströme. Nachfolgendes Beispiel zeigt den Aufbau in Anlehnung zur Untersuchung eines Scheibenwischermotors mit dazugehörigem Steuergerät.

Wir vergleichen dabei die Messwerte für den Gleichtaktstrom in dB μ A:

- Aufbau analog der Komponentenmessung nach CISPR25 (Bild links a, rechts without harness)
- Aufbau im Fahrzeug, Störstrommessung direkt an der Komponente (close to the DUT)
- Aufbau im Fahrzeug, Störstrommessung am Verbraucher (Bild links b, far end of the harness)

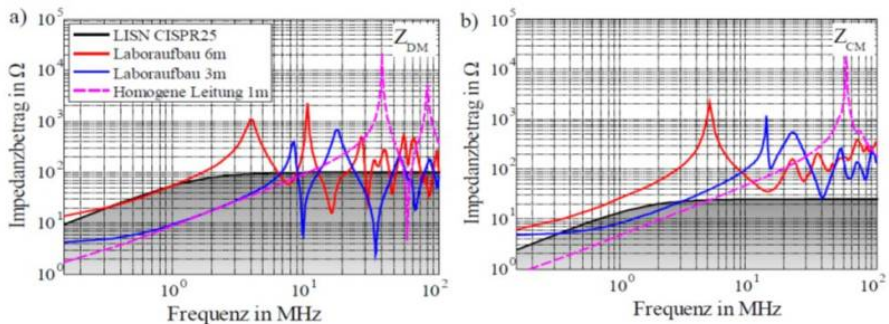


Bereits auf den ersten Blick zeigen sich deutliche Unterschiede zwischen der Messung im klassischen Laboraufbau mit der Netznachbildung und dem Fahrzeugkabelbaum. Die Leitungsresonanzen stechen deutlich hervor und heben jeweils den Störstrom entlang der Leitung deutlich mit bis zu 10dB an. Messbar, aber nur an wenigen Frequenzpunkten wirklich signifikant fällt der Unterschied an verschiedenen Positionen entlang des Kabelbaums aus.

Diese erste einfache Messung erklärt sehr anschaulich, dass eingehalten Grenzwerte nicht immer eine Garantie dafür sind erfolgreiche Fahrzeugmessungen zu erhalten. Es macht daher Sinn, dass wir uns die Impedanz der Leitungsverbindungen noch einmal genauer anschauen. Das identische Problem existiert natürlich auch bei der klassischen AC Netznachbildung. Auch hier kann es zu einer Beeinflussung kommen obwohl alle EMV Anforderungen erfüllt werden. Aufgrund der meist geringeren Vernetzung zu Hause oder in der Industrieumgebung fallen solche Probleme allerdings deutlich seltener auf.

Nachfolgende Abbildung zeigt die Gleich- und Gegentaktimpedanz verschiedener Leiter (Zweidrahtleitungen) im Vergleich zur Messung mit Netznachbildungen. Die

3m und 6m Leitung befinden sich dabei innerhalb eines Kabelbündels mit weiteren Leitern bei unbekannter relativer Lage zueinander, die Eindrahtleitung ist gleichförmig auf einer leitfähigen Tischfläche aufgebaut.



Die schwarze Kennlinie zeigt die standardisierte Eingangsimpedanz bei der Verwendung von zwei Netznachbildungen. Im differentiellen Pfad ergibt sich eine Serienschaltung mit 100Ω , für die Gleichtaktimpedanz eine Parallelschaltung und damit 25Ω . Schön zu sehen sind die aufgrund der Leitungsresonanzen entstehenden Minima und Maxima der Impedanzen in Abhängigkeit der Leitungslänge. Bei allen Messungen zeigen sich folgende Effekte:



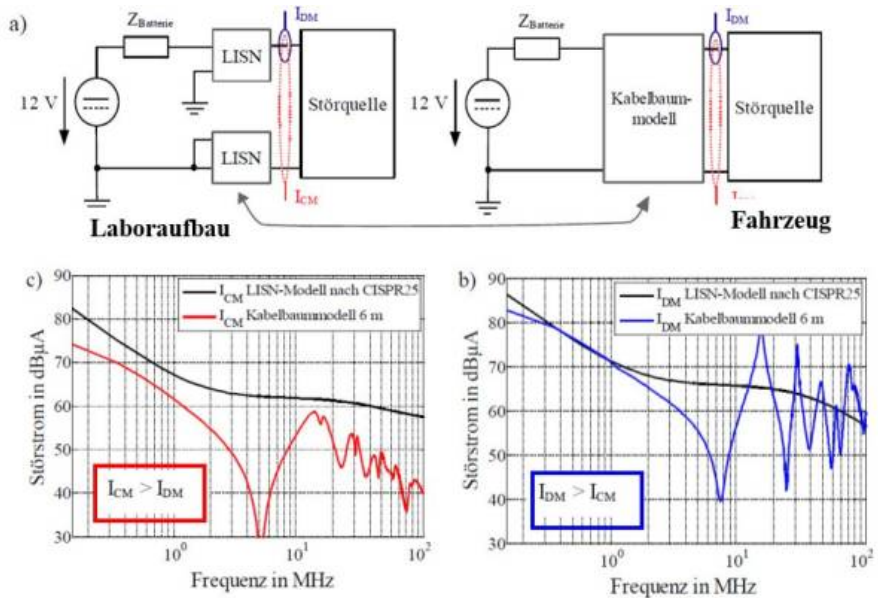
- Je je mehr Adern sich im Kabelbaum befinden, desto weniger ausgeprägt sind die Resonanzstellen!
- Gleichtaktimpedanzen sind von diesem Effekt deutlicher betroffen als die Gegentaktimpedanz.

Der erste Effekt ist deutlich zu sehen wenn Sie die 1m lange homogene Leitung (homogen bedeutet hier geradlinig parallel verlegtes Leiterpaar) mit der 6m langen Leitung innerhalb eines Kabelbaums vergleichen. Die Resonanzstellen der 1m Leitung zeigen eine sehr hohe Güte, wobei die Resonanzen der 6m Leitung eher verschliffen wirken und ungleichförmig. Da der differentielle Strompfad den Hin- und Rückleiter berücksichtigt, ein Gleichtaktstrom jedoch nur eine Richtung kennt, treten die Gegentaktesonanzen entsprechend bei tieferen Frequenzen auf. Typischerweise für frequenzunabhängigen Materialien (Isolation) bei der halben Gleichtaktfrequenz.

Für uns hauptsächlich von Interesse sind Resonanzstellen mit geringen Impedanzwerten im Vergleich zur Komponentenmessung, da sich hier ein hoher Störstrom ausbilden kann. Genau diese Überhöhungen der Störströme konnten wir zuvor in unserem einführenden Messbeispiel anhand des Scheibenwischermotors erkennen. Das bedeutet, wir sind der Ursache für die angestiegenen Emissionen dicht auf dem Fersen ...

Im direkten Vergleich der Gleich- und Gegentaktimpedanz sticht die Gegentaktimpedanz mit ausgeprägten Resonanzstellen unterhalb der Gegentaktimpedanzen der Netznachbildungen hervor, womit wir den zweiten allgemeinen Effekt beschreiben: Gleichtaktimpedanzen zeigen innerhalb eines Kabelbündels bereits bei wenigen Adern nahezu keine ausgeprägten Resonanzstellen mehr. Die Gleichtaktimpedanz ist in diesem Beispiel stets größer als die Gegentaktimpedanz der Netznachbildungen. Die Ursache liegt darin, dass Gleichtaktströme jede Möglichkeit nutzen sich zur Masse zu schließen. Je mehr Adern zu Verfügung stehen, desto größer sind die kapazitiven Kopplungen zu benachbarten Leitern welche wiederum kapazitiv zur Masse gekoppelt sind. Für Gegentaktströme gibt es deutlich weniger Ausbreitungspfade, sie müssen entlang des Kabelbaums zurückfließen.

Versuchen wir diese These jetzt zu verallgemeinern. Nachfolgendes Beispiel zeigt dazu das Setup für den Laboraufbau und die Fahrzeugmessung. Mit Hilfe einer Stromzange wird jeweils der Gleich- und Gegentaktstörstrom für zwei verschiedene Varianten gemessen und gegenübergestellt. Unterschieden wird jetzt eine Störquelle deren dominante Emission entweder eine Gleichtaktstörung (Bild c) oder eine Gegentaktstörung (Bild b) darstellt.



Im Teilbild c) zeigt sich deutlich, dass der Gleichtaktstrom in der Fahrzeugsituation über dem gesamten Frequenzbereich stets kleiner ist als die Referenzmessung im Labor. Dies deckt sich mit der Erkenntnis, dass die Gleichtaktimpedanz stets größer ist als die Laborimpedanz. Ganz anders die Situation bei der Gegentaktemission. Bei einer dominierenden Gegentaktstörung fallen die Leitungsresonanzen jetzt deutlich ins Gewicht und wir erhalten eine Überhöhung der Störströme.

$$l \leq \frac{c_0 / \sqrt{\epsilon_r}}{2 \cdot f_{max}}$$

Ob wir also bei der Fahrzeugmessung im Vergleich zur Labormessung Probleme bekommen können hängt somit von mehreren Parametern ab. Entscheidend ist ob die Resonanzstellen unserer Anschlussleitungen größer oder kleiner sind als die Impedanz der im Labor verwendeten

Netznachbildungen. Geringere Impedanzwerte generieren meist höhere Störströme entlang der Leitungen. Ganz verallgemeinern lässt sich die Aussage leider nicht, da die Störquelle an den entsprechenden Frequenzpunkten in der Lage sein muss den Störstrom auch zu treiben (zur Verfügung zu stellen, abhängig vom frequenzabhängigen Innenwiderstand der Störquelle). Die Resonanzstellen einer Leitungsverbindung lassen sich mit der Ausbreitungsgeschwindigkeit und der Permittivität des Kabelmantels abschätzen. Ok, wer kennt schon die Dielektrizitätskonstante seiner eingesetzten Kabel welche meist auch noch frequenzabhängig ist ... verwenden Sie einfach für die

Ausbreitungsgeschwindigkeit 2/3 der Lichtgeschwindigkeit! In die Berechnung der Resonanzstellen geht die Leiterlänge mit ein, womit wir durch Umstellen der Gleichung eine kritische Länge ermitteln können ab der Leitungsresonanzen auftreten können (Gleichung links).

Nachfolgende Tabelle fasst die zuvor hergeleiteten Erkenntnisse noch einmal zusammen.

Kabelbaumart	Einfluss auf Gleich- und Gegentaktstörungen	
	$I_{CM} > I_{DM}$	$I_{CM} < I_{DM}$
Kurze, homogen verlegte Leitung	Überhöhung der Störgrößen falls $l > l_{max}$	Überhöhung der Störgrößen falls $l > l_{max}$
Inhomogen verlegte lange Leitung	Dämpfung der Störungen	Überhöhung der Störgrößen

In vielen Fällen sind die Emissionen im Fahrzeug kleiner als die ermittelten Werte im Labor. Dies liegt daran, dass oft längere Kabelbäume mit einer hohen Adernanzahl zum Einsatz kommen und dadurch die Resonanzpunkte nicht mehr ausgeprägt sind. Besonders die Resonanzen der Gleichtaktimpedanz werden innerhalb der Kabelbündels stark gedämpft. Damit verbleibt das Problem für diese Art Kabelbäume hauptsächlich dann, wenn die Gegentaktstörungen die Gleichtaktstörungen dominieren. In den vorhergehenden Kapiteln haben wir gesehen, dass in vielen Fällen jedoch die Gleichtaktemission dominant ist, womit sich das Problem noch weiter reduziert. Kritisch zu bewerten sind demnach generell kurze Leitungen (1-2m) welche nicht innerhalb eines Bündels laufen. Hier kann es zur Anhebung sowohl von Gleich- als auch Gegentaktströmen kommen. Die Eingrenzung der Fälle spiegelt sich auch mit den Erfahrungen aus dem EMV-Labor: Die Problematik höherer Emissionen im Fahrzeug als in der Komponentenmessung kann vorkommen, allerdings nur in einer geringen Anzahl an untersuchten Systemen. Mit den hergeleiteten Ergebnissen sind wir in der Lage eine Aussage darüber zu treffen ob unsere Komponente betroffen sein wird oder nicht.