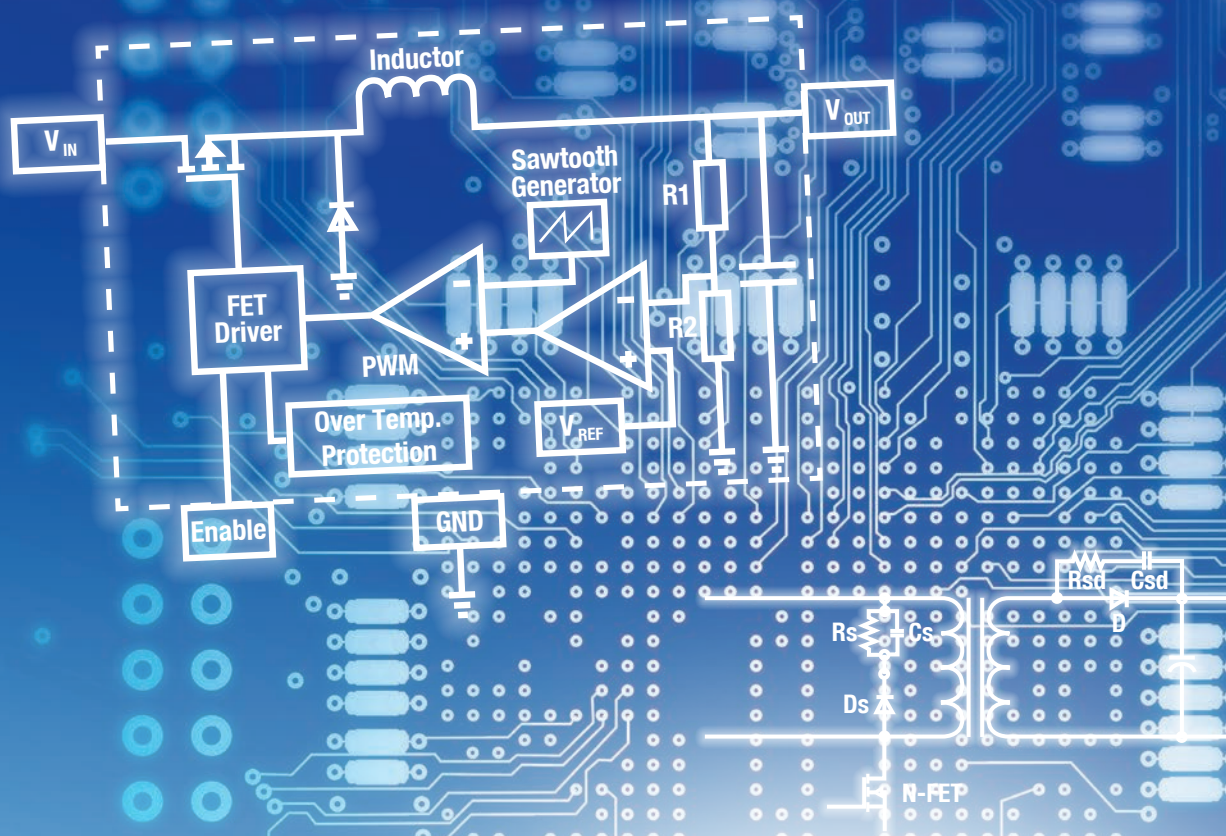




DC/DC

BOOK OF KNOWLEDGE

Praktische Tipps für Anwender

von Steve Roberts M.Sc. B.Sc.
Technischer Direktor, RECOM

RECOM

DC/DC BOOK OF KNOWLEDGE

Praktische Tips für Anwender

Von Steve Roberts, M.Sc. B.SC
Technical Director, RECOM

Erstausgabe

© 2014 2014 Alle Rechte vorbehalten. RECOM Engineering GmbH & Co KG,
Österreich (nachstehend RECOM)

Die Inhalte dieses Buchs dürfen ohne schriftliche Genehmigung von RECOM in keiner Form, weder vollständig noch in Auszügen, reproduziert, vervielfältigt oder verbreitet werden. Die in diesem Buch vermittelten Informationen sind nach bestem Wissen des Autors korrekt. Für sachliche Fehler, Auslassungen oder typographische Fehler wird keinerlei Haftung übernommen. Die Diagramme veranschaulichen typische Anwendungen und erheben keinen Anspruch auf Vollständigkeit.

Vorwort der Recom Geschäftsleitung

Als wir vor 25 Jahren unsere ersten DC/DC-Wandler entwickelten, gab es kaum Fachliteratur zum Thema, ebenso wenig wie internationale Standards. Wir mussten unseren Kunden jedoch dringend in irgendeiner Weise praktische Anwendungsinformationen zugänglich machen, was uns dazu veranlasste, unseren ersten Produktkatalog durch einen Anhang mit praktischen Anwendungshinweisen zu ergänzen. Im Laufe der Jahre und mit wachsender Erfahrung ist der Inhalt dieses Anhangs immer umfangreicher geworden. Obwohl die Hinweise noch immer in eher rudimentärer Form vorliegen, wird das mittlerweile 70 Seiten umfassende Hinweispaket, das zum Download auf unserer Webseite bereitsteht, von unserer Kundenbasis gerne angenommen.

Durch den Fortschritt der Halbleitertechnologie und die Verschiebung hin zur hochintegrierten, digitalen Elektronik ist das Wissen über analoge Technologien in vielen Entwicklungslaboren, Universitäten und technischen Hochschulen immer weiter auf dem Rückzug. Oft stellen wir fest, dass es an praktischem Know-how über die analoge Schaltkreisentwicklung fehlt, gerade im Hinblick auf angewandte Techniken, Überprüfungs- und Messmethoden sowie im Hinblick auf Störimpuls-Filterung und Rauschunterdrückung. Darum sahen wir, als Experten auf diesem Gebiet, die Notwendigkeit ein umfassendes technisches Handbuch zu erstellen, das von Hardware-Designern und Studenten als Referenzquelle genutzt werden kann.

Anfang 2014 begann unser Technischer Direktor Steve Roberts in seiner Freizeit damit, das umfangreiche Wissen der RECOM Gruppe im Bereich Entwicklung, Prüfung und Anwendung von DC/DC-Wandlern schriftlich niederzulegen. Trotz seiner anspruchsvollen Aufgabe bei RECOM und neben der Entwicklung neuer Produkte sowie der technischen Planung für unsere neuen Forschungs- und Testlabore hat er es geschafft, das Werk pünktlich zur Electronica 2014 fertigzustellen.

Steve hat ein Handbuch vorgelegt, das unserer Überzeugung nach der Engineering-Community und allen an DC/DC-Wandlern und ihren Einsatzmöglichkeiten Interessierten erheblichen Nutzen bringen wird. Das Handbuch ist in der Erstausgabe als gedrucktes Buch und PDF-Dokument in englischer, deutscher und chinesischer Sprache erhältlich.

Die Geschäftsführung Gmunden, 2014

RECOM Group

Vorwort des Autors

Ein AC/DC- oder DC/DC-Wandlermodul erfüllt eine oder mehrere der folgenden Aufgaben:

- i. Anpassung der sekundären Stromversorgung an die Primärstromversorgung
- ii. Trennung von Primär- und Sekundärkreisen
- iii. Schutz vor Fehlfunktionen, Kurzschlüssen und Überhitzung
- iv. Vereinfachung der Einhaltung von Sicherheits- und Leistungsstandards sowie der Vorgaben der EMV-Gesetzgebung

Um das zu erreichen, stehen mehrere Technologien zur Verfügung, beginnend beim einfachen Linearregler bis hin zur mehrstufigen, digital gesteuerten Stromversorgung. Dieses Buch erklärt die verschiedenen verfügbaren DC/DC-Schaltkreise und Topologien, sodass Nutzer die Vorteile, Einschränkungen und operativen Grenzen solcher Lösungen besser einschätzen können. Die Sprache dieses Handbuchs ist zwangsläufig technischer Natur, wobei versucht wurde, die Sachverhalte so einfach wie möglich darzustellen, ohne die zugrunde-liegenden Technologien zu trivialisieren.

Der Autor verfügt über viele Jahre Erfahrung in der Beantwortung typischer Kundenfragen, hat an vielen Design-Ins mitgearbeitet, Seminarvorträge gehalten, Fachartikel verfasst und sogar an der Erstellung von YouTube-Videos mitgewirkt. Trotz seines geballten Wissens ist er der Meinung, dass es jeden Tag wieder etwas Neues zu diesem vielfältigen und weitreichenden Thema zu lernen gibt. Dieses Buch trägt den Untertitel „Praktische Tipps für Anwender“, denn seine Intention ist es, das Thema Energieumwandlung zu „entmystifizieren“, obwohl es ebenso viele Lösungen wie Anwendungen gibt. Wenn es gelingt, Ihnen ein wenig von unserem Wissen zu vermitteln, hat es sein Ziel erfüllt.

Die Informationen in diesem Buch wurden nach bestem Wissen zusammengestellt und auf Richtigkeit überprüft. Sollten Sie als Leser dennoch Fehler, Auslassungen oder Ungenauigkeiten entdecken, dürfen Sie es mich gerne wissen lassen.

Steve Roberts Gmunden, 2014

Technischer Direktor
s.roberts@recom-power.com

RECOM

Inhaltsverzeichnis

1. Einführung	11
1.1 Lineare Spannungsregler	11
1.1.1 Wirkungsgrad des Linearspannungsreglers	13
1.1.2 Sonstige Eigenschaften des Linearspannungsreglers	14
1.1.3 LDO-Regler	15
1.2 Schaltregler	16
1.2.1 Schaltfrequenz und Größe der Induktivität	17
1.2.2 Schaltregler-Topologien	18
1.2.2.1 Nichtisolierter DC/DC-Wandler	18
1.2.2.1.1 Schalttransistoren	19
1.2.2.1.2 Buck-Wandler	20
1.2.2.1.3 Anwendungen des Buck-Wandlers	22
1.2.2.1.4 Boost-Wandler	23
1.2.2.1.5 Anwendungen des Boost-Wandlers	25
1.2.2.1.6 Buck-Boost- (invertierender) Wandler	25
1.2.2.1.7 Buck-Boost- diskontinuierlicher (DCM- discontinuous mode) und kontinuierlicher Betrieb (CM- continuous mode)	27
1.2.2.1.8 Synchrone und asynchrone Umwandlung	28
1.2.2.1.9 Zweistufen-Boost-Buck (Ćuk-Wandler)	30
1.2.2.1.10 Zweistufen-Boost-Buck-SEPIC-Wandler	32
1.2.2.1.11 Zweistufen-Boost-Buck-ZETA-Wandler	34
1.2.2.1.12 Mehrphasige DC/DC-Wandler	35
1.2.2.2 Isolierter-DC/DC-Wandler	37
1.2.2.2.1 Flyback DC/DC-Wandler	37
1.2.2.2.2 Forward DC/DC Converter	38
1.2.2.2.3 Aktiver Klemmvorwärtswandler	41
1.2.2.2.4 Push-Pull-Wandler	42
1.2.2.2.5 Halb- und Vollbrückenwandler	45
1.2.2.2.6 Bus- oder ratiometrischer Wandler	46
1.2.2.2.7 Ungeregelter Push-Pull-Wandler	47
1.2.3 Parasitäre Elemente und deren Effekte	50
1.2.3.1 Quasiresonanz-Wandler	53
1.2.3.2 Resonanzmoduswandler	54
1.2.4 Wirkungsgrad von DC/DC-Wandler	56
1.2.5 PWM-Regelungsverfahren	57
1.2.6 Regelung durch DC/DC-Wandler	59
1.2.6.1 Regelung von Mehrfachausgängen	59
1.2.6.2 Remote Current Sense Regulation	62
1.2.7 Beschränkung im Hinblick auf den Eingangsspannungsbereiches	65

1.2.8 Synchrongleichrichtung	66
1.2.9 Planartransformatoren	67
1.2.10 Gehäusetypen von DC/DC-Wandlern	69
2. Feedbackschleifen	71
2.1 Einleitung	71
2.2 Konstruktion mit offener Schleife	71
2.3 Geschlossene Schleifen	72
2.4 Kompensation von Rückkopplungsschleifen	75
2.4.1 Instabilität der rechten Halbebene	77
2.5 Anstiegsausgleich oder slope compensation	77
2.6 Analyse der Schleifenstabilität in analogen und digitalen Feedbacksystemen	78
2.6.1 Experimentelle Ermittlung der analogen Schleifenstabilität	78
2.6.2 Ermittlung der analogen Schleifenstabilität unter Verwendung der Laplace-Transformation	79
2.6.3 Ermittlung der digitalen Schleifenstabilität mit Hilfer der bilinearen Transformation	81
2.6.4 Digitale Feedback-Schleife	83
3. Datenblatt-Parameter richtig verstehen	85
3.1 Messverfahren - DC-Charakteristik	85
3.2 Messverfahren - AC-Charakteristik	87
3.2.1 Messung des minimalen und maximalen Impuls-Pausen-Verhältnisses	88
3.2.2 Ausgangsspannungsgenauigkeit	89
3.2.3 Temperaturkoeffizient der Ausgangsspannung	89
3.2.4 Lastregelung	90
3.2.5 Kreuzregelung	91
3.2.6 Netzregelung (line regulation)	91
3.2.7 Ausgangsspannungsgenauigkeit im ungünstigsten Fall	92
3.2.8 Berechnung des Wirkungsgrades	92
3.2.9 Eingangsspannungsbereich	93
3.2.10 Eingangsstrom	93
3.2.11 Kurzschluss und Überlastung	95
3.2.12 Ein/Aus-Steuerung (Remote ON/OFF)	97
3.2.13 Isolationsspannung	98
3.2.14 Isolationswiderstand und -kapazität	101
3.2.15 Dynamische Belastungskennlinie	102
3.2.16 Ausgangsrestwelligkeit	103
3.3 Zum Verständnis thermischer Parameter	104

3.3.1 Einleitung	104
3.3.2 Thermische Impedanz	105
3.3.3 Temperatur-Derating	107
3.3.4 Zwangskühlung	109
3.3.5 Leitungs- und Strahlungskühlung	110
4. DC/DC-Wandler-Schutzmaßnahmen	111
4.1 Einleitung	111
4.2 Verpolungsschutz	111
4.2.1 Verpolungsschutz durch Seriendioden	112
4.2.2 Verpolungsschutz durch Überbrückungs- oder Shunt diode	113
4.2.3 Verpolungsschutz mit P-FET	113
4.3 Eingangssicherung	114
4.4 Ausgangs-Überspannungsschutz	115
4.5 Eingangsüberspannungsschutz	115
4.5.1 SCR-Crowbar-Schutz	116
4.5.2 Clamping-Elemente	117
4.5.2.1 Varistor	117
4.5.2.2 Suppressordiode	119
4.5.3 OVP unter Verwendung mehrerer Bauteile	120
4.5.4 OVP-Standards	121
4.5.5 OVP durch Abschaltung	122
4.6 Spannungseinbruch und -unterbrechung	123
4.7 Einschaltspitzenstrom-Begrenzung	124
4.8 Lastbegrenzung	127
4.9 Unterspannungsabschaltung	129
5. Befilterung von Wandlereingang und- ausgang	130
5.1 Einleitung	130
5.2 Rückwelligkeitsstrom (back ripple current)	131
5.2.1 Messung des Rückwelligkeitsstroms	131
5.2.2 Gegenmaßnahmen zur Reduktion des Rückwelligkeitsstroms	133
5.2.3 Auswahl des Eingangskondensators	135
5.2.4 Eingangsstrom von parallel geschalteten DC/DC-Wandlern	136
5.3 Ausgangsbefilterung	137
5.3.1 Differenzmodus- Ausgangsbefilterung	138
5.3.2 Gleichtakt-Ausgangsbefilterung	140
5.3.3 Gleichtakt drosseln	141
5.4 Vollfilterung	144
5.4.1 Filter-PCB-Layout	145

6. Sicherheit	147
6.1 Stromschläge	149
6.1.1 Isolationsklasse	149
6.1.2 Stromgrenzwerte für den menschlichen Körper	150
6.1.3 Schutz vor Stromschlägen	152
6.1.4 Schutzerdung	155
6.2 Gefährliche Energielevels	157
6.2.1 Schmelzsicherungen	158
6.2.1.1 Ansprechzeiten von Schmelzsicherungen und Einschaltspitzströmen	160
6.2.2 Rückstellbare Überstromschutzschalter (circuit breaker)	161
6.3 Inhärente Sicherheit	163
6.4 Eigensicherheit	164
6.4.1 Brennbare Materialien	165
6.4.2 Rauch	166
6.5 Verletzungsrisiko	167
6.5.1 Heiße Oberflächen	168
6.5.2 Scharfe Kanten	168
6.6 Auf Sicherheit ausgerichtete Konstruktionsweisen	169
6.6.1 FMEA	172
6.7 Medizintechnische Sicherheit	173
7. Betriebszuverlässigkeit	176
7.1 Vorhersage der Betriebszuverlässigkeit	176
7.2 Umweltbelastungsfaktoren	179
7.3 Verwendung von MTBF-Angaben	180
7.4 Demonstrierte MTBF	181
7.5 MTBF und Temperatur	182
7.6 Auf Betriebszuverlässigkeit ausgerichtete Konstruktionsweise	183
7.7 Empfehlung zum PCB-Layoutentwurf zum Erzielen einer hohen Zuverlässigkeit	184
7.8 Kondensatorzuverlässigkeit	188
7.8.1 MLCC	188
7.8.2 Tantal- und Elektrolytkondensatoren	192
7.9 Zuverlässigkeit von Halbleitern	195
7.10 ESD	197
7.11 Induktivitäten	198
8. LED-Charakteristiken	201
8.1 Betrieben von LEDs mit Dauerstrom	203
8.2 Einige DC-Konstantstromquellen	203

8.3 Verbindung von LEDs in Ketten	205
8.4 Parallelschaltung von LED-Ketten	206
8.5 Symmetrierung des LED-Stroms in parallelgeschalteten Ketten	207
8.6 Parallele Ketten oder Grid-Array - was ist besser?	209
8.7 Dimmen von LEDs	211
8.7.1 Analoges Dimmen gegenüber PWM-Dimmen	211
8.7.2 Wahrgenommene Helligkeit	214
8.7.3 Abschlusswort zum Dimmen	215
8.8 Thermische Überlegung	216
8.9 Temperature Derating	216
8.9.1 Hinzufügen von automatischem Temperatur-Derating zu einem LED-Treiber	217
8.9.2 Übertemperaturschutz unter Verwendung von analogem Temperatursensor-IC	218
8.9.3 Over-temperature Protection using an Analogue Temperature Sensor IC	218
8.9.4 Übertemperaturschutz unter Verwendung eines Microcontrollers	220
8.10 Korrektur der Helligkeit	221
8.11 Einige Schaltungskonzepte unter Verwendung eines RCD-Treibers	223
9. DC/DC-Anwendungsideen	230
9.1 Einleitung	230
9.2 Polaritätswechslung	230
9.3 Spannungsverdoppler	231
9.4 Kombination von Schaltreglern und DC/DC-Wandlern	232
9.5 Schaltung von Wandlern in Serie	235
9.6 Erhöhung der Isolation	236
9.7 5-V-Bereinigung	237
9.8 Verwendung des CTRL-Steuerungspins	238
9.9 Verwendung des V_{ADJ} -Pin	239
Quellen	242
Über RECOM	243
Danksagung	244

1. Einführung

Moderne AC/DC- und DC/DC-Wandler sind so ausgerichtet, dass sie eine effiziente Spannungsumsetzung gewährleisten und eine geregelte, sichere und konstante Gleichstromspeisung für verschiedene elektronische Instrumente, Geräte und Systeme bereitstellen. Vor nicht allzu langer Zeit waren Transformatoren, Gleichrichter und lineare Spannungsregler die Haupttechnologie in der Spannungsumsetzung; aber so wie LEDs langsam die Glühlampe ersetzen, verdrängt auch der DC/DC-Wandler allmählich den Linearspannungsregler, und primärgetaktete Netzteile ersetzen den einfachen 50-Hz-Netztransformator. Im letzten Jahrzehnt war die Entwicklung der Schaltregler durch einen enormen technischen Fortschritt gekennzeichnet, was die Nutzung der Vorteile neuer Schaltkreise, Komponenten und Materialien, die früher einfach nicht vorhanden waren, möglich machte. Dieser Fortschritt führte zu erhöhter Leistung und verbessertem Temperaturverhalten sowie zu einer wesentlichen Verringerung der Größe, des Gewichts und der Kosten der Stromversorgungsgeräte. Infolgedessen werden Schaltnetzteile heute in großen Mengen verwendet und sind Standardtechnologie sowohl in der DC/DC- als auch in der AC/DC-Leistungsumsetzung.

1.1 Lineare Spannungsregler

Lineare Spannungsregler liefern eine stabile Ausgangsspannung von einer mehr oder weniger stabilen Eingangsspannungsquelle. Selbst wenn die Eingangsspannung schnell schwankt, bleibt die Ausgangsspannung im Normalbetrieb stabil. Das bedeutet, dass sie die Eingangsspannungswelligkeit nicht nur bei der Grundfrequenz, sondern auch bis zur fünften oder zehnten Oberwelle sehr wirksam ausfiltern können. Die Beschränkung liegt nur in der Reaktionsgeschwindigkeit des Feedback-Kreises des Fehlerverstärkers.

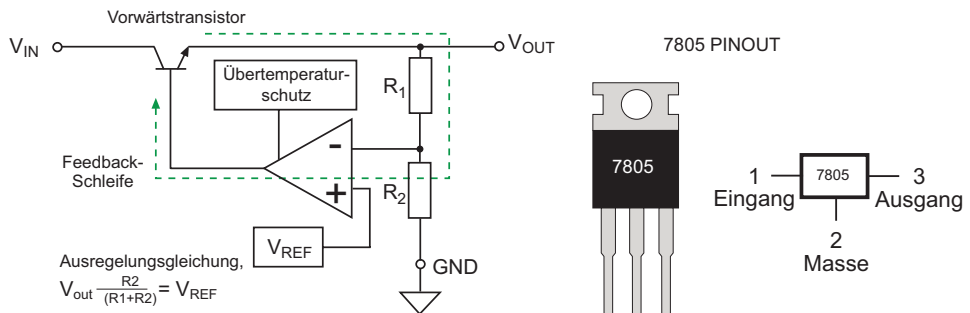


Abb. 1.1 Blockdiagramm des Linearspannungsreglers mit 3 Pins und Pinout

Die meisten Linearregler stellen einen geschlossenen Regelkreis dar. Abb. 1.1 veranschaulicht diesen Typus der Spannungsregelung. Der Vorwärtstransistor ist das Regelement, faktisch ein Varistor, der den vom Eingang zum Ausgang fließenden Strom beschränkt. Der Spannungsteiler R_1/R_2 ist so gewählt, dass an der gewünschten Ausgangsspannung die nach unten geteilte Spannung am invertierten Eingang des Fehlerverstärkers dieselbe ist wie die Spannung V_{REF} am nichtinvertierten Eingang. Der Fehlerverstärker steuert den Ausgang so, dass die Potentialdifferenz zwischen den Eingängen von selbigen stets gleich Null ist.

Wenn die Ausgangsspannung infolge von Lastabnahme oder erhöhter Eingangsspannung steigt, wird die Spannung am invertierten Eingang des Fehlerverstärkers höher als die V_{REF} -Spannung, und der Ausgang des Fehlerverstärkers senkt den Basisstrom vom Transistor und in Folge die Ausgangsspannung des Linearreglers bis der Sollwert wieder erreicht ist. Im umgekehrten Fall, wenn die Last zunimmt oder die Eingangsspannung abfällt, sinkt die Spannung am invertierten Eingang unterhalb der V_{REF} -Spannung, und der Basisstrom zum Transistor steigt, um die Ausgangsspannung zum Ausgleich zu erhöhen. So regelt dieselbe Feedback-Schleife sowohl Schwankungen der Eingangsspannung (line regulation) als auch der Lastwechsel (load regulation). Es versteht sich von selbst, dass die Vergleichsspannung sehr stabil sein und einen ausgezeichneten Temperaturkoeffizienten haben sollte, damit eine stabile und genaue Ausgangsspannung erzielt wird. Mit einem guten PCB-Layout ist jedoch eine Ausgangswelligkeit kleiner als $50\mu V_p$ leicht zu erreichen.

Das vereinfachte Blockschaltbild des Reglers mit 3 Pins in Abb. 1.1 zeigt keinen Kurzschlusschutz. Ist der Ausgang kurzgeschlossen würde über den Transistor ein sehr hoher Strom fließen, was eine Überlastung und Zerstörung zur Folge haben könnte, sodass eine interne Zusatzbeschaltung notwendig ist, um den Ausgangsstrom zu beschränken (Abb. 1.2). Der Strombegrenzer verwendet den Spannungsabfall über den Messwiderstand R_S , um den Ausgangsstrom zu überwachen. Wenn der Strom so hoch ist, dass die Spannung $0,7\text{ V}$ überschreitet, beginnt Q_2 zu leiten und den Strom von Q_1 abzuziehen, sodass der Basisstrom verringert und der Ausgangsstrom beschränkt wird. Daraus folgt: $I_{LIMIT} = 0,7V/R_S$.

Die Strombegrenzung muss erheblich über dem maximalen Strom, der während des Normalbetriebs fließen würde, liegen. Typischerweise ist dieser um 150 % bis 200 % höher als der Nennstrom. Da der Regler im Kurzschlussfall nicht abgeschaltet wird, ist er bei kurzgeschlossenem Ausgang konstant überlastet.

Einige preiswerte Linearregler vertrauen einfach auf den Übertemperaturschutz und schalten den Transistor ab, bevor er als "Kurzschlusschutz" durchgebrannt ist. Dies kann zwar als Schutz des Linearspannungsreglers fungieren, jedoch kann sich die primäre Stromversorgung überhitzen und ausfallen, wenn diese nicht entsprechend dimensioniert ist, um dem Kurzschlussstrom so lange standzuhalten, bis sich der Regler abschaltet.

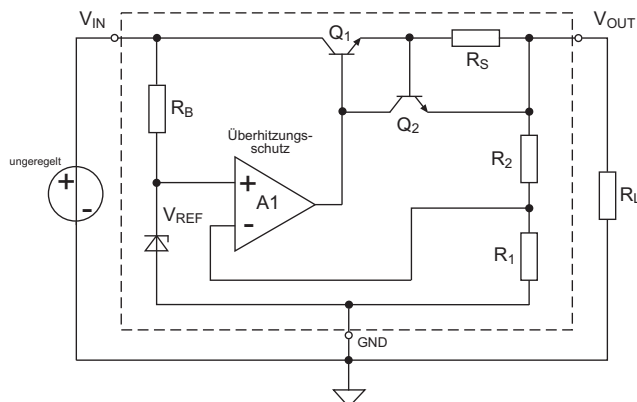


Abb. 1.2: Linearregler mit Strombegrenzung ("Kurzschlusschutz")

Praktischer Hinweis

Die Differenz zwischen der Eingangs- und Ausgangsspannung fällt am Transistor ab. Wenn die Eingangsspannung beispielsweise 12V (z. B. Akku) und die geregelte Ausgangsspannung 5V beträgt, soll der Transistor 7V abfangen. Dies bedeutet, dass mehr Leistung im Regler in Wärme umgesetzt werden muss, als zur Last übergeben wird (siehe auch die Erörterung der Berechnungen des Wirkungsgrades im nächsten Abschnitt). Daher benötigen die meisten Linearregler einen Kühlkörper. Logischerweise kann der Linearregler nicht kompensieren, wenn die Eingangsspannung unter die Ausgangsspannung absinkt, weshalb die Ausgangsspannung der Eingangsspannung abwärts folgt. Wenn die Eingangsspannung jedoch zu stark abfällt, ist die interne Stromversorgung zum Fehlerverstärker und V_{REF} gefährdet, und der Ausgang kann instabil werden oder anfangen zu schwingen.

Linearregler liefern auch im Bereitschaftsbetrieb (Stand-by) schlechte Ergebnisse. Selbst wenn keine Last anliegt, benötigt ein typischer Regler der 78xx-Serie immer noch ca. 5mA, um den Fehlerverstärker sowie die Referenzspannungsquelle zu versorgen. Bei einer Eingangsspannung von 24V verursacht dieser Ruhestrom eine Leerlaufleistung von 120mW.

Praktischer Hinweis

Vorteile von Linearreglern: preiswert, gute Regelcharakteristik, niedriges Rauschen, geringes Störverhalten und ausgezeichnetes Lastwechselverhalten. Nachteile: hoher Verbrauch im Ruhezustand, nur einfache Ausgänge und extrem niedriger Wirkungsgrad.

1.1.1 Wirkungsgrad des Linearspannungsreglers

Der Wirkungsgrad η eines Linearreglers ist durch das Verhältnis der übergebenen Ausgangsleistung P_{OUT} zum Leistungsverbrauch P_{IN} bestimmt.

$$\eta = \frac{P_{OUT}}{P_{IN}} \quad \begin{aligned} P_{OUT} &= V_{OUT} I_{OUT} \\ P_{IN} &= V_{IN} I_{IN} \\ I_{IN} &= I_{OUT} + I_Q \end{aligned}$$

Gleichung 1.1: Wirkungsgrad eines Linearspannungsreglers

I_Q ist der Ruhestrom durch den Regler im linearen Leerlaufbetrieb. Die Gleichung kann wie folgt umgeformt werden:

$$\eta = \frac{(V_{OUT} I_{OUT})}{V_{IN} (I_{OUT} + I_Q)}$$

Gleichung 1.2: Erweiterte Gleichung für den Wirkungsgrad des Linearspannungsreglers

Das folgende Beispiel gilt für einen typischen 5V 3-Pin-Spannungsregler mit einer Eingangsspannung von 10VDC, einem Ausgangsstrom von 1A und einem Ruhestrom von 5mA.

Der Wirkungsgrad wird wie folgt berechnet:

$$\eta = \frac{5V \times 1A}{10V \times 1.005A} = 0.49$$

So beträgt der Gesamtwirkungsgrad 49 %, und die Verlustleistung im Wandler überschreitet sogar die zur Last übergeben 5W. Wenn die Eingangsspannung bis zum Minimum von 7VDC sinkt, erhöht sich der Wirkungsgrad auf 70 %; dies stellt jedoch den praktisch maximalen Wirkungsgrad dar, da für eine ordnungsgemäße Regelung ca. 2V „Spielraum“ notwendig sind. Es ist sofort erkennbar, dass der Wirkungsgrad eines Reglers dieses Typs stark von der Eingangsspannung und der Last abhängt und nicht konstant ist. Dies bedeutet auch, dass der Spannungsregler mit einem ausreichend großen Kühlkörper ausgestattet sein sollte, um einen sicheren Betrieb auch im ungünstigsten Fall, bei maximaler Eingangsspannung und maximalem Ausgangsstrom, zu ermöglichen.

1.1.2 Sonstige Eigenschaften des Linearspannungsreglers

Einerseits haben Linearregler viele Vorteile, andererseits aber auch einige gravierende Nachteile, die besondere Sorgfalt bei Anwendung und Einsatz des Reglers erfordern.

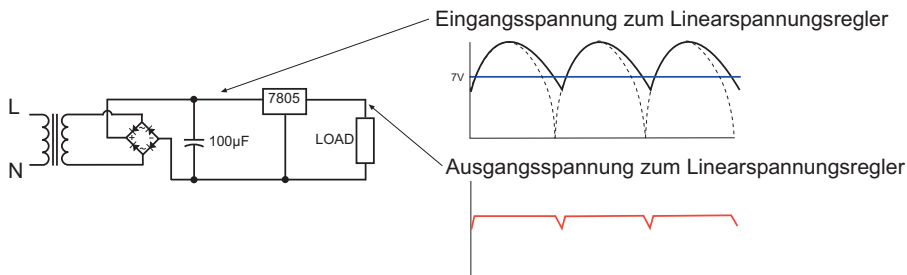


Abb. 1.3: Drop-Out-Problem beim Linearspannungsregler

Praktischer Hinweis

Wenn die Potentialdifferenz zwischen Eingang und Ausgang, wie oben erwähnt, unterhalb des erforderlichen „Spielraums“ (üblicherweise 2 V) liegt, kann der Regelkreis nicht mehr korrekt funktionieren. Ein verbreitetes Anwendungsproblem besteht, wenn der gleichgerichtete Wechselstromeingang eine hohe Spannungswelligkeit besitzt, weil der Glättungskondensator viel zu klein ist (Abb. 1.3). Wenn die Eingangsspannung bei jeder Halbperiode unter die Abfallspannung (engl.: drop out voltage) abfällt, zeigt der geregelte Ausgang periodische Abfälle mit doppelter Netzfrequenz. Diese kurzzeitigen Abfälle werden nicht vom Multimeter angezeigt, das nur die Durchschnitts-Ausgangsspannung misst, können aber doch „nicht erklärbare“ Schaltkreis-Probleme in der zu versorgenden Schaltung hervorrufen. Dieser Effekt kann beseitigt werden, indem man entweder größere Glättungskondensatoren verwendet oder das Windungsverhältnis des Transformators erhöht – beides teure Optionen.

1.1.3 LDO-Regler

Der im Standard-Linearspannungsregler verwendete bipolare Durchlasstransistor wird als Stromverstärker verwendet. Der Erregungsstrom vom Ausgang des Fehlerverstärkers wird mit der geringen Signalstromverstärkung des Transistors (H_{FE}) multipliziert, um den Laststrom zu versorgen. Die H_{FE} des Leistungstransistors ist sehr niedrig, typischerweise 20-50, weshalb häufig eine Darlingtonschaltung mit mehreren Transistoren verwendet wird, um die effektive Stromverstärkung zu erhöhen und den vom Ausgangsstrom des Fehlerverstärkers geringer halten zu können. Der Nachteil des Darlington-Transistors besteht darin, dass die Rücksetzspannung durch V_{BE} für jede Stufe zunimmt; so wird die typische Abfallspannung für einen Standard-Linearspannungsregler, der einen PNP-Transistor verwendet, um einen NPN-Darlington-Transistor zu erregen, zu:

$$V_{Dropout} = 2 V_{BE} + V_{CE} \approx 2V \text{ (bei Zimmertemperatur)}$$

Bei niedrigen Umgebungstemperaturen sinkt H_{FE} , weshalb für eine zuverlässige Regelung aller Betriebsbedingungen ein „Spielraum“ von 2,5 bis 3V erforderlich werden kann.

Durch Ersetzen des bipolaren Transistors durch einen P-Kanal-FET arbeiten Low-Drop-Out- (LDO- = Niedrigabfallspannung-)Linearregler mit einer Rücksetzspannung von nur wenigen hundert Millivolt. Die Abfallspannung beträgt dann einfach die Durchlassspannung am FET. Diese ergibt sich aus dem Widerstand R_{DS} multipliziert mit dem Laststrom I_{LOAD} . Da R_{DS} typischerweise sehr niedrig ist, ist die Abfallspannung ebenfalls niedrig.

FETs werden selten in ihrem ohmschen Bereich verwendet, da die Stromverstärkung einer komplizierten Relation folgt und sowohl von der Temperatur als auch der Last abhängig ist (siehe Abb. 1.4). In jedem Fall kompensiert der Fehlerverstärker jedoch jegliche Drift und Nichtlinearität in der $V_{GS}-V_{TH}$ -Kurve, da er lediglich die Ausgangsspannung mit der Vergleichsspannung vergleicht und seinen Ausgang entsprechend abstimmt.

Der Nachteil von LDO-Reglern besteht darin, dass die $V_{GS}-V_{TH}$ -Kurve bei hohen Gate-Spannungen sehr steil und bei niedrigen Gate-Erregungsspannungen sehr flach ist, sodass der Fehlerverstärker einerseits eine sehr niedrige Ausgangs-Dämpfung haben sollte und dennoch schnell auf Last- oder Eingangsspannungsänderungen reagieren können muss (leicht gedämpft sein sollte). Das Ergebnis ist ein notwendiger Kompromiss zwischen den beiden Betriebsextremen, was entweder bei hochinduktiven oder bei hochkapazitiven Belastungen zu Problemen führen kann.

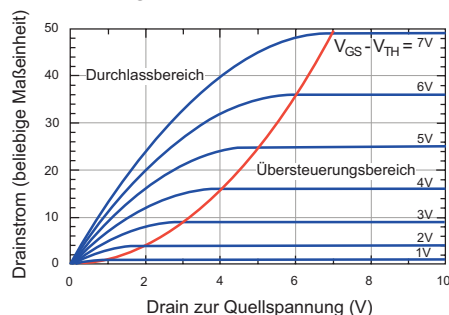


Abb. 1.4: FET-Charakteristik

LDO-Linearregler sind deshalb anfälliger für Überspannungsschäden und erfordern bessere Filterung und umfangreicherer Maßnahmen zur Transientenvermeidung. Auch haben sie einen stärker eingeschränkten Eingangsspannungsbereich (Input Voltage Range).

Sowohl Standard- wie LDO-Linearregler sind für Ausfälle anfällig. Falls der Vorwärtstransistor ausfällt, ist die Ursache in der Regel ein Kurzschlussfehler zwischen Kollektor und Emitter. Dies bedeutet, dass der Ausgang ohne jegliche Regelung direkt an den Eingang durchgeschaltet wird, was zur Zerstörung des Anwendungsschaltkreises führen kann. Abb. 1.5 zeigt eine mögliche Schutzschaltung, in der eine leistungsstarke Zener-Klemmdiode verwendet wird, die im Falle eines Regelungsfehlers die Sicherung verschmort.

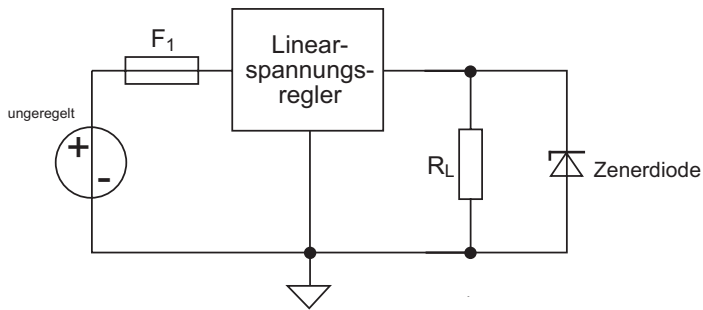


Abb. 1.5: Einfacher Ausgangs-Überspannungsschutz

1.2 Schaltregler

Im Unterschied zu Linearreglern, die die Verlustleistung in Wärme umsetzen, um die Ausgangsspannung zu beschränken, verwenden Schaltregler Energiespeicherungseigenschaften von induktiven und kapazitiven Komponenten, um Leistung in diskreten Energiepaketen zu übertragen. Diese Energiepakete sind entweder im Magnetfeld einer Induktivität oder im elektrischen Feld eines Kondensators gespeichert. Der Schaltregler stellt sicher, dass nur die Energie, die für die Last aktuell erforderlich ist, in jedes Paket übertragen wird, weshalb diese Topologie sehr effizient ist. Abb. 1.6 zeigt die vereinfachte Struktur eines Schaltreglers.

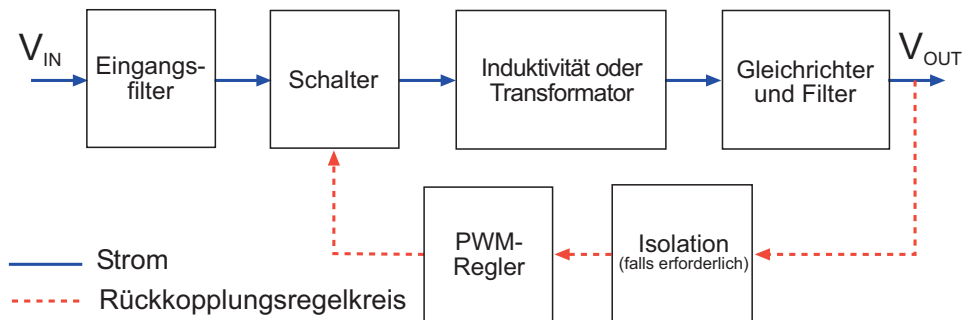


Abb. 1.6: Blockdiagramm eines Schaltregler

Um Energie in kontrollierten Mengen vom Eingang zum Ausgang zu übertragen, ist ein komplizierteres Regelungsverfahren notwendig als beim Linearspannungsregler. Der meistverbreitete Typ der Regelung ist PWM (Pulse Width Modulation = Pulsweitenmodulation), bei der die vom Eingang zum Ausgang übertragene Energiemenge durch einen in seiner Breite variablen Impuls in einem fixen Zeitintervall moduliert wird. Die relative Einschaltdauer der PWM, δ , ist das Verhältnis von t_{ON} (Zeitraum, in dem Energie aus der Quelle fließt) zur Periode T (Umkehrgröße der Schaltfrequenz f_{OSC}).

$$\delta = \frac{t_{ON}}{T}, \text{ wobei } T = \frac{1}{f_{OSC}}$$

Gleichung 1.3: Definition der relativen Einschaltdauer

Bei Schaltreglern ist die geregelte Ausgangsspannung direkt proportional zur relativen Einschaltdauer der PWM. Der Regelkreis verwendet das „Großsignal“-Impuls-Pausen-Verhältnis, um das Leistungsschaltbauelement zu steuern. Hingegen verwendet der Linearspannungsregler die „Kleinsignal“-Regelschleife, um den Strom durch den Vorwärtstransistor zu beschränken. Die PWM-Regelung ist mit wesentlich besseren Wirkungsgraden möglich als die lineare Regelung, da Hauptverluste nur während jeder Änderung der Schalterstellung anfallen und nicht dauerhaft Verlustleistung „verbraten“ werden muss. Anders gesagt: Da die FETs, entweder komplett an- oder abgeschaltet sind, fällt eine ungleich geringere Verlustleistung an.

1.2.1 Schaltfrequenz und Größe der Induktivität

Die Größe der Schalt- und Speicherelemente des Schaltreglers sind umgekehrt proportional zur verwendeten Schaltfrequenz. Die Energie, die in einer Induktivität gespeichert werden kann, beträgt:

$$P(L) = \frac{L I^2 f}{2}$$

Gleichung 1.4: Gespeicherte Energie in einer Induktivität

Die in einer Induktivität gespeicherte Energiemenge ist proportional zur Frequenz. Um eine konstante Energiemenge zu speichern, kann z. B. die Induktivität L halbiert werden, wenn die Frequenz verdoppelt wird.

Für kapazitive Elemente lautet die Gleichung für die gespeicherte Leistung:

$$P(C) = \frac{C V^2 f}{2}$$

Gleichung 1.5: Gespeicherte Energie in einem Kondensator

Hier wiederum, kann die Größe des Kondensators durch Erhöhung der Frequenz verringert werden, ohne Kompromisse bei der gespeicherten Energiemenge eingehen zu müssen. Diese Verkleinerungen der materiellen Größe sind sowohl für den Hersteller, als auch für den Kunden, ausschlaggebend, weil so geringeres Gewicht erzielt werden kann und mitunter weniger Verpackungsaufwand erforderlich wird. Jedoch sind mit höherer Schaltfrequenz auch EMV-Probleme zu erwarten.

Daher existiert ein EMV-Kompromiss, der die höchste praktische Schaltfrequenz auf ca. 500kHz beschränkt (einige sehr kleine Konstruktionen können bei 1MHz oder höher takten, erfordern jedoch sehr sorgfältiges PCB-Layout und beste EMV-Abschirmung).

1.2.2 Schaltregler-Topologien

Der Begriff Topologie verweist auf verschiedene Arten von Schaltvorgängen und Kombinationen von Energiespeicherelementen, die für die Übertragung, Steuerung und Regelung der Ausgangsspannung oder des Ausgangsstroms aus einer Eingangsspannungsquelle möglich sind.

Die verschiedenen Topologien für Schaltregler können in zwei Hauptgruppen unterteilt werden:

- a) Nichtisolierte Wandler, in welchen die Eingangsquelle und die Ausgangslast während des Betriebs einen gemeinsamen Stromweg nutzen.
- b) Isolierte Wandler, in welchen Energie durch transformatorgekoppelte magnetische Bauelemente (Transformatoren) übertragen wird, wobei die Kopplung zwischen der Versorgung und der Last ausschließlich durch ein Magnetfeld erreicht wird, was eine galvanische Trennung zwischen dem Eingang und Ausgang erlaubt.

1.2.2.1 Nichtisolierter DC/DC-Wandler

Die Auswahl aus der Vielfalt der verfügbaren Topologien basiert auf Überlegungen wie Kosten, Performance und Regelcharakteristik, die durch die Hauptbetriebscharakteristiken bestimmt werden. Keine Topologie ist besser oder schlechter als die andere. Jede Topologie hat sowohl Vorteile als auch Nachteile, und so ist die Wahl eine Frage der Bedürfnisse des Nutzers und der Bedarfsanforderungen an die Spannungsversorgung.

Für nichtisolierte DC/DC-Wandler gibt es fünf transformatorlose Grundtopologien:

- i. Buck- oder Abwärtswandler
- ii. Boost- oder Aufwärtswandler
- iii. Buck-Boost- oder Auf-Abwärtswandler
- iv. Zweistufiger invertierender Buck-Boost-Wandler (Ćuk-Wandler)
- v. Zweistufiger nicht-invertierender Buck-Boost-Wandler (Sepic-Wandler, ZETA-Wandler)

Die nachfolgenden Erklärungen setzen voraus, dass das PWM-Regelschema einen Regelkreis (nicht dargestellt) hat und für die gewünschte Ausgangsspannung das korrekte Impuls-Pausen-Verhältnis gewählt wurde. Angenommen sind außerdem ideale Schalter (Schalttransistoren oder -dioden) sowie ideale Kondensatoren und Induktivitäten, um die Eigenschaften der Übertragung jeder Topologie besser zu demonstrieren. Bevor wir uns jedoch die Topologien anschauen, erscheint es opportun, einige Worte über die Ansteuerung von Schalttransistoren zu sagen.

1.2.2.1.1 Schalttransistoren

FETs werden bei Verwendung als Schalttransistoren üblicherweise übersteuert betrieben, wo der Drain-Quellen-Widerstand am Minimum ist und Leistungsverluste im Schalter minimal sind. Solange die Gate-Spannung V_{GS} wesentlich größer ist als die Schwellenspannung V_{TH} , befindet sich der FET über den ganzen Lastbereich in Übersteuerung (Sättigung) (siehe Abb. 1.7).

Beim Betrachten des vereinfachten asynchronen Buck-Wandlerkreises in Abb. 1.7 ist erkennbar, dass zwei FETs vorhanden sind, wobei der eine auf Masse (Niederspannungsseite) und der andere auf V_{IN+} (Hochspannungsseite) schaltet.

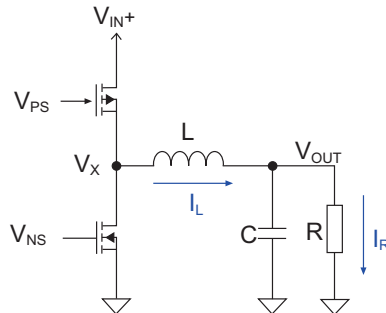


Abb. 1.7: Vereinfachter asynchroner Impuls-Abwärtsregler

Der niederspannungsseitige FET (N-Kanal) schaltet durch, wenn $V_{NS} \gg V_{TH}$, und schaltet ab, wenn $V_{NS} < V_{TH}$.

Der hochspannungsseitige FET (P-Kanal) schaltet durch, wenn $V_{PS} \ll (V_{IN} - V_{TH})$, und schaltet ab, wenn $V_{PS} > (V_{IN} - V_{TH})$. P-Kanal-FETs weisen jedoch typischerweise die dreifache Verlustleistung eines gleichgroßen N-Kanal-FETs auf und sind zudem teurer. In vielen Leistungsanwendungen ist dies nicht zulässig, weshalb als hochspannungsseitiges Schaltelement ein N-Kanal-FET bevorzugt wird. Dies bedeutet jedoch, dass der hochspannungsseitige Treiber imstande sein soll, eine Ausgangsspannung zu generieren, die höher als die Eingangsspannung V_{IN} ist.

Eine allgemein gebräuchliche Lösung für einen hochspannungsseitigen N-Kanal-Treiber besteht darin, das Rechtecksignal bei V_X zu verwenden, um die Versorgungsspannung des hochspannungsseitigen Treibers durch einen Bootstrapkondensator und Diode D_1 zu erhöhen.

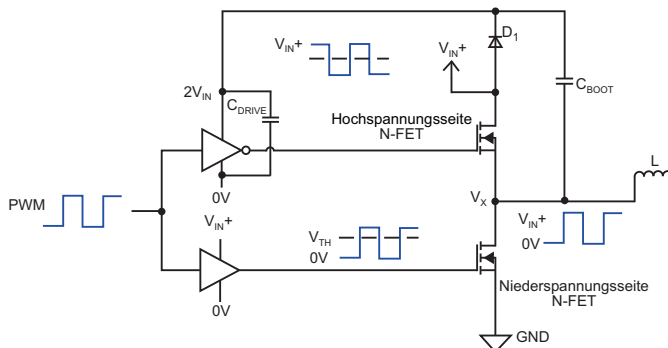


Abb. 1.8: Bootstrapschaltkreis für hochspannungsseitigen Treiber

Der Kondensator C_{BOOT} ist immer dann auf $V_{IN}+$ durch D1 geladen, wenn $V_X = GND$ ist, und entlädt immer dann $2 \times V_{IN}+$ in den hochspannungsseitigen Treiberkondensator C_{DRIVER} , wenn $V_X = V_{IN}+$ ist. So kann der hochspannungsseitige Treiber das Gate des hochspannungsseitigen N-FETs über die Eingangsspannung hinaus treiben.

Der Nachteil dieses einfachen Bootstrapschaltkreises besteht darin, dass der Bootstrapkondensator in PWM-Betriebszyklen mit entsprechendem Impuls-Pausen-Verhältnis nicht ausreichend Zeit hat, um den Kondensator C_{Drive} aufzuladen. So ist ein Betrieb bei einem Impuls-Pausen-Verhältnis in der Nähe von 100 % nicht möglich. Das schränkt die Eingangsspannung und den Lastbereich des Wandlers ein.

Eine Lösung dieses Problems besteht darin, eine separate Ladungspumpe (engl.: charge pump oscillator) zu verwenden, um C_{Drive} über den gesamten Impuls-Pausen-Verhältnissbereich oberhalb von $V_{IN}+$ aufgeladen zu halten. Solche Ladungspumpschaltungen sind oft in den Regler oder hochspannungsseitigen Treiber IC integriert (siehe unten stehendes Beispiel des hochspannungsseitigen Treibers MAX1614 mit integrierter Ladungspumpe).

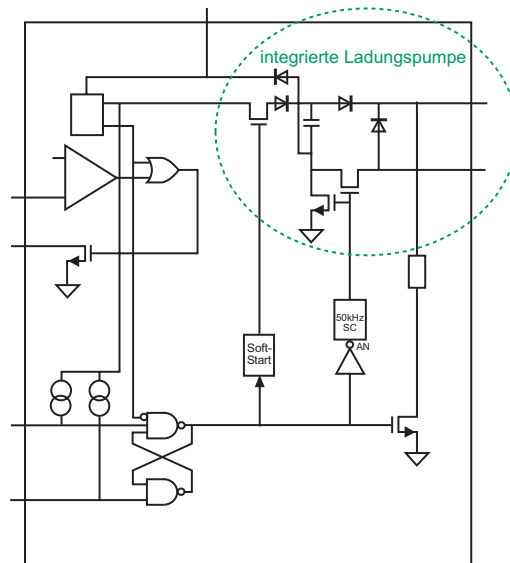


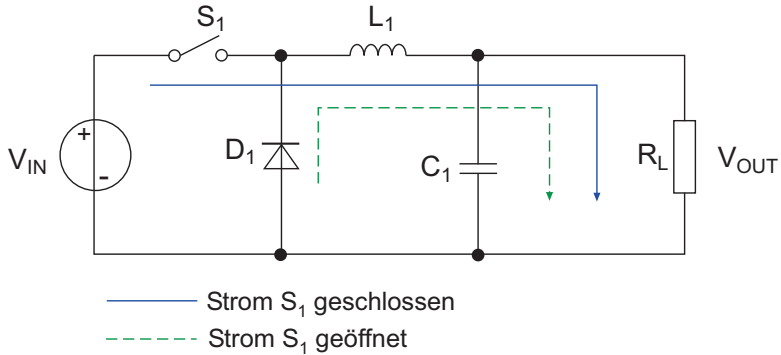
Abb. 1.9: Blockschaltbild des hochspannungsseitigen Treibers MAX1614

1.2.2.1.2 Buck-Wandler

Wie der Name sagt, wandelt der Abwärts- oder Buck-Wandler eine hohe Eingangsspannung in eine stabilisierte niedrigere Ausgangsspannung um. Ein vereinfachtes Schaltbild und die Hauptstrom- und -spannungswellenformen sind der Abb. 1.10 zu entnehmen.

Am leichtesten zu verstehen ist dieser Schaltkreis, wenn man sich vorstellt, dass L_1 und C_1 einen Tiefpassfilter bilden. Wenn der Schalter S_1 geschlossen ist, erhöht sich die Spannung an der Last langsam linear, während der Kondensator C_1 durch L_1 aufgeladen wird.

Wenn S_1 dann geöffnet wird, wird die im Magnetfeld der Induktivität gespeicherte Energie am Schalterende der Induktivität durch die Diode D_1 bei 0V fixiert. Somit hat diese Energie keine andere Wahl, als sich über den Kondensator und die Last zu entladen und in Folge wird die Spannung an der Last langsam absinken. Die Durchschnittsausgangsspannung ist dann das Impuls-Pausen-Verhältnis des PWM-Regel-signals multipliziert mit der Eingangsspannung.



$$V_{OUT} = V_{IN} \frac{t_{ON}}{T} = \delta V_{IN}, \text{ gilt, wenn } V_{IN} > V_{OUT}$$

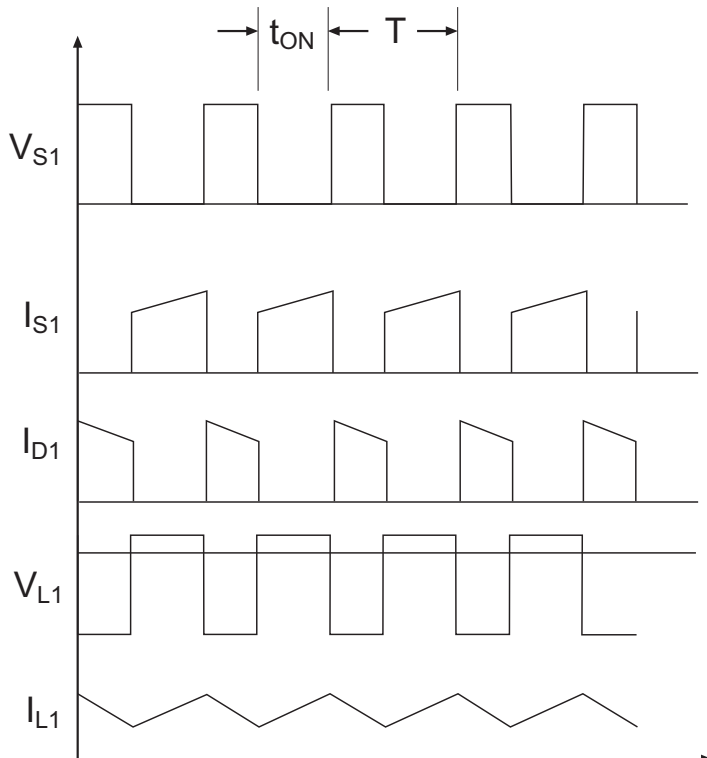


Fig. 1.10: Buck Regulator Simplified Schematic and Characteristics

Die Übertragungsfunktion kann abgeleitet werden, indem man die Spannung-Zeit-Produkte der Induktivität in den Zuständen „AUS“ und „AN“ gleichsetzt. Aufgrund des Energieerhaltungsprinzips müssen die beiden Produkte gleich groß sein.

Für den AN-Zustand: $\text{Energie}_{\text{IN}} = (V_{\text{IN}} - V_{\text{OUT}}) t_{\text{ON}}$

Für den AUS-Zustand: $\text{Energie}_{\text{OUT}} = V_{\text{OUT}} t_{\text{OFF}}$, wo $t_{\text{OFF}} = T - t_{\text{ON}}$ und $\delta = t_{\text{ON}} / T$

Das Ersetzen ergibt: $(V_{\text{IN}} - V_{\text{OUT}}) t_{\text{ON}} = V_{\text{OUT}} (T - t_{\text{ON}})$

$$V_{\text{IN}} t_{\text{ON}} = V_{\text{OUT}} T$$

$$V_{\text{OUT}} = V_{\text{IN}} (t_{\text{ON}} / T)$$

$$V_{\text{OUT}} / V_{\text{IN}} = \delta$$

Gleichung 1.6: Übertragungsfunktion des Buck-Wandlers

1.2.2.1.3 Anwendungen des Buck-Wandlers

Die Vorteile des Buck-Wandlers bestehen darin, dass die Verluste sehr niedrig sind (da Wirkungsgrade von > 97% leicht erreichbar sind, besonders als synchrone Ausführung [siehe Abschnitt 1.2.2.1.8]); die Ausgangsspannung kann beliebig im Bereich von V_{REF} bis V_{IN} gewählt werden, und die Differenz zwischen V_{IN} und V_{OUT} darf sehr groß sein. Außerdem darf die Schaltfrequenz einige hundert kHz betragen, um eine sehr kompakte Konstruktion mit kleinen Induktivitäten und einem dynamischen Übergangsverhalten zu erzielen. Wenn der Schalt-FET gänzlich ausgeschaltet ist, ist die Leerlaufstromaufnahme der gesamten Schaltung vernachlässigbar klein. Aus diesen Gründen stellt der Abwärtsregler in vielen Anwendungen eine sehr attraktive Alternative zum Linearspannungsregler dar.

Praktischer Hinweis

Die Serie RECOM R-78xx ist eine pin-kompatible Alternative zur 78xx-Serie. R-78xx ist ein komplettes Buck-Wandler-Abwärtsreglermodul, für dessen Normalbetrieb keine externe Beschaltung erforderlich ist. Er bietet ein Wirkungsgrad von 97%, eine Eingangsspannung von bis zu 72VDC und Leerlaufstromaufnahme (engl.: quiescent current) von lediglich 20nA.

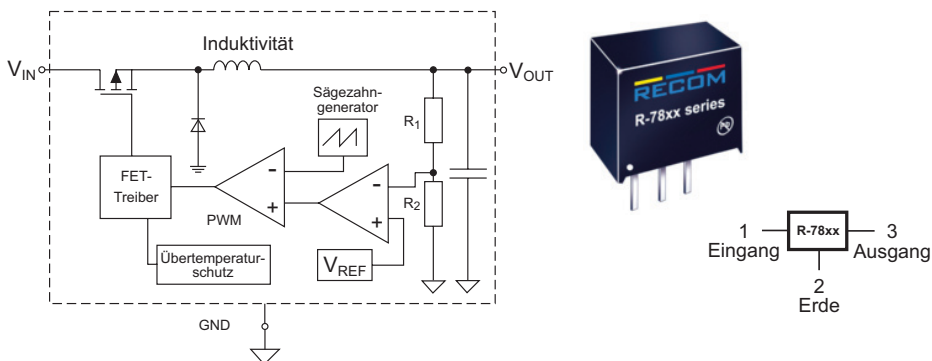


Abb. 1.11: Schaltregler Buck-Wandler und Pinout

Praktischer Hinweis

Ein Nachteil des Buck-Wandlers liegt darin, dass der Feedbackkreis des PWM-Reglers, um korrekt zu regeln, eine minimale Spannungswelligkeit am Ausgang erfordert, da die Regelung typischerweise Zyklus für Zyklus erfolgt. Die Spannungswelligkeit am Ausgang hängt vom Impuls-Pausen-Verhältnis ab und erreicht ihr Maximum bei 50% Tastverhältnis. Deshalb ist es nicht möglich diese Welligkeit (engl.: ripple/noise), auf den durch Linearregler erreichbaren μV -Pegel herabzusetzen. Wenn eine sehr „saubere“ Versorgung notwendig ist, kann nach dem Abwärtsregler ein Linearspannungsregler eingesetzt werden, um von den Vorteilen beider Technologien zu profitieren. Im unten stehenden Beispiel werden die unregelmäßigen 24VDC durch den Schaltregler mit einem Wirkungsgrad von 95% auf 15V gesenkt. Der Linearspannungsregler liefert dann einen „sauberen“ 12-V-Ausgang mit $< 5\mu\text{V}$ Restwelligkeit und Rauschen. Der gesamte Wirkungsgrad des Systems beträgt ca. 76%, im Vergleich zu weniger als 50%, wäre nur der Linearspannungsregler eingesetzt.

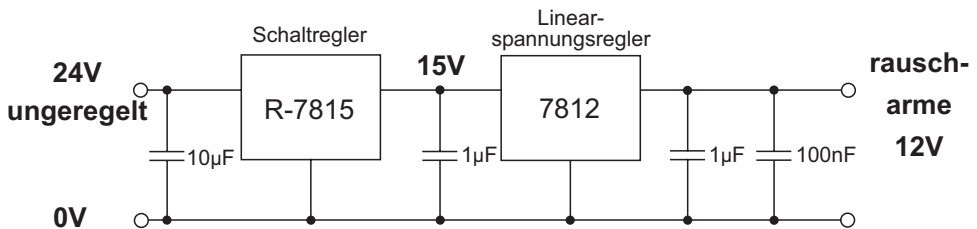
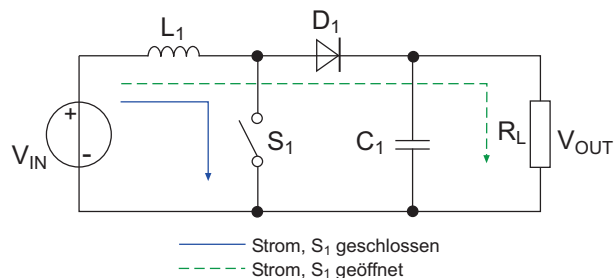


Abb. 1.12: Kombination von Impuls-Abwärtsregler und Linearregler

1.2.2.1.4 Boost-Wandler

Wie der Name schon sagt, wandeln Aufwärts- oder Boost-Wandler eine niedrige Eingangsspannung, in eine stabilisierte, höhere Ausgangsspannung. Ein vereinfachtes Schaltbild, sowie die Hauptstrom- und Spannungsformen sind in Abb. 1.13 gezeigt.



$$V_{OUT} = V_{IN} \frac{1}{1 - \delta}, \text{ gilt, wenn } V_{IN} < V_{OUT}$$

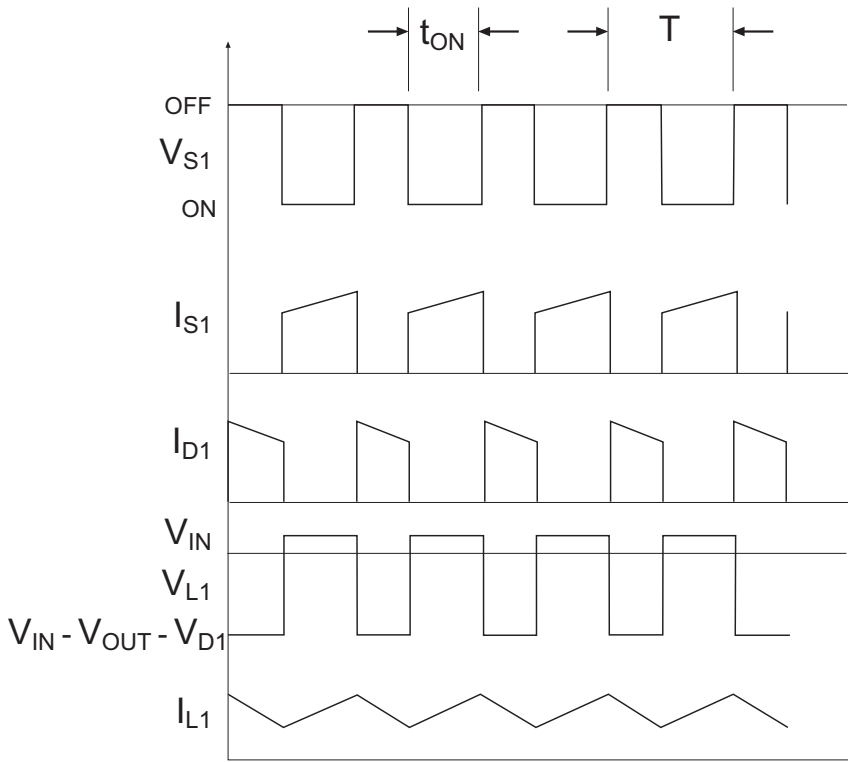


Abb. 1.13: Vereinfachtes Schema und Charakteristik des Boost-Wandlers

Wenn S_1 geschlossen ist, fließt der Strom durch die Induktivität L_1 , der linear im Verhältnis V_{IN}/L_1 steigt. Während dieser Periode, wird der Laststrom von der in C_1 gespeicherten Energie versorgt. Wird nun der Schalter wieder geöffnet, führt die gespeicherte Energie in der Induktivität dazu, dass die Ausgangsspannung der Eingangsspannung überlagert wird. Der resultierende Strom fließt durch die Freilaufdiode D_1 , versorgt die Last und lädt C_1 wieder auf. Der Strom durch die Induktivität nimmt linear und proportional zu $(V_{OUT} - V_{IN})/L_1$ ab. Die Ableitung der Übertragungsfunktion ähnelt der im vorherigen Abschnitt, nur sind die Grundgleichungen umgeordnet:

Für den AN-Zustand: $Energy_{IN} = V_{IN} t_{ON}$

Für den AUS-Zustand: $Energy_{OUT} = (V_{OUT} - V_{IN}) t_{OFF}$

$$\frac{V_{OUT}}{V_{IN}} = \frac{1}{1 - \delta}$$

Gleichung 1.7: Übertragungsfunktion des Boost-Wandlers

1.2.2.1.5 Anwendungen des Boost-Wandlers

Der Vorteil des Boost-Wandlers besteht darin, dass die Ausgangsspannung, abhängig vom Impuls-Pausen-Verhältnis des PWM-Signals, gleich oder größer als V_{IN} ist. Daher eignet er sich z.B. besonders gut für akkubetriebene Anwendungen, bei denen die benötigte Versorgungsspannung, der zu versorgenden Anwendung, um den sogenannten Boost-Faktor höher sein soll, als die Batteriespannung. In der Praxis jedoch sind Boost-Faktoren von x_2 , x_3 oder mehr nur schwer zu realisieren, da die Höhe der Eingangsstromimpulse zum Boost-Faktor proportional zunimmt. D.h., es zieht ein Wandler, der die Eingangsspannung verdreifachen soll, den dreifachen Eingangsstrom. Dieser pulsierende Eingangsstrom kann hohe EMV- und auch Spannungsabfallprobleme in den Eingangszuleitungen hervorrufen, die, wenn überhaupt, nur mit sehr hohem Aufwand in Griff zu bekommen sind.

Ein Nachteil besteht beim Boost-Wandler auch darin, dass der Ausgang ohne einen zweiten Schalter in Serie zum Eingang nicht ausgeschaltet werden kann, da ein Abschalten des PWM-Reglers die Last nicht vom Eingang trennt.

Praktischer Hinweis

Es muss unbedingt darauf geachtet werden, dass die Eingangsspannung nicht über die gewünschte Ausgangsspannung ansteigen kann. Der PWM-Regler würde dann s_1 stets geöffnet halten, und Eingang und Ausgang würden ohne Regelung durch L_1 und D_1 direkt verbunden werden. So könnten hohe Ströme fließen, die sowohl den Wandler als auch die Last sehr schnell zerstören würden. Um dieses Problem zu vermeiden, ist eine Topologie erforderlich, die sowohl den Buck- als auch den Boost-Betrieb zulässt.

1.2.2.1.6 Buck-Boost- (invertierender) Wandler

Der invertierende Flyback-Wandler, auch Buck-Boost-Wandler genannt, wandelt eine Eingangsspannung in eine geregelte negative Ausgangsspannung um, die höher oder niedriger als der absolute Wert der Eingangsspannung sein kann. Das vereinfachte Schema in Abb. 1.14 zeigt den Hauptschaltplan und die damit verbundenen Wellenformen.

Wenn S_1 geschlossen ist, fließt der Strom I_{L1} , der direkt proportional zu V_{IN}/L_1 ansteigt, in diesem Schaltkreis durch L_1 . Diode D_1 blockiert währenddessen den Strom in die Last. Während dieser Zeit wird der Laststrom vom Ausgangskondensator C_1 gespeist. Wenn der Schalter S_1 geöffnet ist, bewirkt die in L_1 gespeicherte Energie, dass das dem Schalter zugewandte Ende der Induktivität negativ wird (das andere Ende der Induktivität ist geerdet). Der daraus resultierende negative Strom fließt jetzt durch D_1 in die Last, bestehend aus C_1 und R_L . Dieser Strom nimmt proportional zu V_{OUT}/L_1 ab. Aufgrund der Stromflussrichtung ist die Ausgangsspannung in Bezug auf das Massepotential negativ. Deshalb eignet sich diese Topologie nur zur Erzeugung negativer Spannungen.

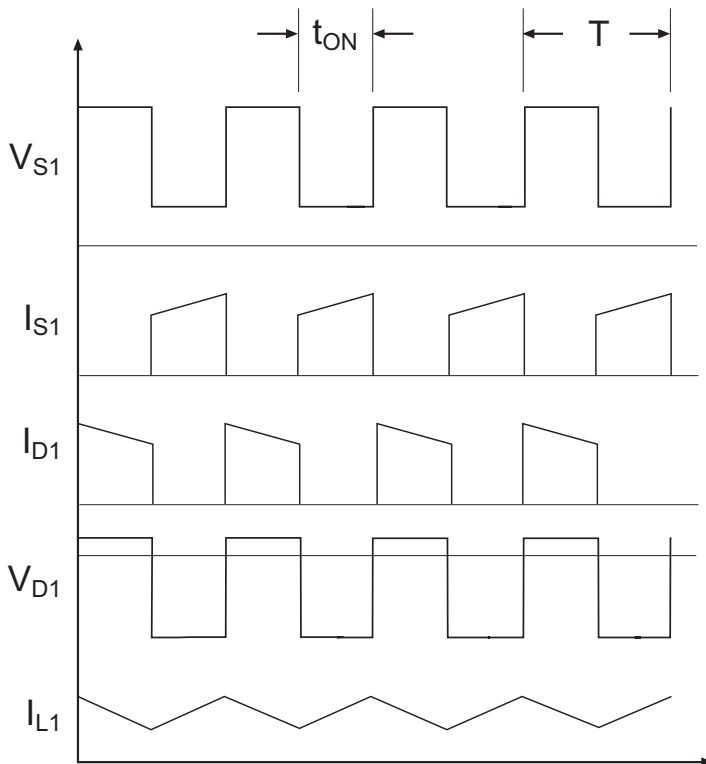
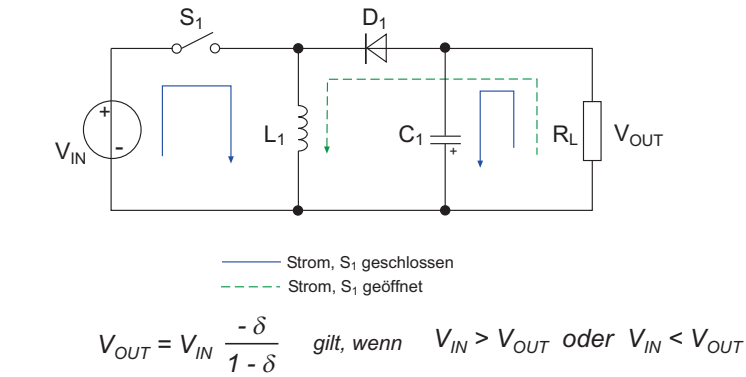


Abb. 1.14: Vereinfachtes Schema und Charakteristik des Buck-Boost-Wandlers

Die Ableitung der Übertragungsfunktion ähnelt der in den vorhergehenden Kapiteln, jedoch mit folgenden Grundgleichungen:

Für den AN-Zustand: $Energy_{IN} = V_{IN} t_{ON}$

Für den AUS-Zustand: $Energy_{OUT} = -V_{OUT} t_{OFF}$

$$\frac{V_{OUT}}{V_{IN}} = \frac{-\delta}{1 - \delta}$$

Gleichung 1.8: Übertragungsfunktion des invertierenden Buck-Boost-Wandlers

Praktischer Hinweis

Der Vorteil des Buck-Boost-Wandlers besteht darin, dass die Eingangsspannung höher oder niedriger sein kann, als die gewünschte geregelte Ausgangsspannung. Das kann z. B. besonders in solchen Anwendungen nützlich sein, die stabilisierte 12V am Ausgang aus einer 12-V-Blei-Säure-Batterie erfordern, die eine Klemmenspannung zwischen 9V – wenn entladen – und 14V – wenn voll aufgeladen – haben kann.

Praktischer Hinweis

Buck-Boost-Wandler sind auch zur Stabilisierung von Solarzellenversorgungen hilfreich. Eine Solarzelle liefert relativ hohe Spannung und Strom bei hellem Sonnenschein, jedoch niedrige Spannung und Strom, wenn die Sonne durch Wolken überschattet wird. Da sich das Spannung-Strom-Verhältnis ändert, kann der Buck-Boost für das Maximum Power Point Tracking (MPPT) (engl.: Tracking des Punktes maximaler Leistung) verwendet werden, da das Übersetzungsverhältnis durch eine entsprechende Regelung kontinuierlich angeglichen werden kann.

Der größte Nachteil ist die invertierte Ausgangsspannung. Bei Verwendung von Akkus ist die Inversion der Ausgangsspannung jedoch irrelevant, weil die Batterie in Bezug auf Masse schwebend belassen und $-V_{OUT}$ dann mit Masse verbunden werden kann, um eine positive Ausgangsspannung zu erzeugen. Ein anderer Nachteil ist, dass der Schalter S_1 keine Verbindung mit Masse hat. Dies bedeutet, dass ein Pegelumsetzer im PWM-Schaltkreis notwendig ist, was weitere Kosten und zu einer höheren Komplexität des Schaltungsdesigns führt.

1.2.2.1.7 Buck-Boost - diskontinuierlicher (DCM – discontinuous mode) und kontinuierlicher Betrieb (CM – continuous mode)

Mit den Abwärts-(Buck-, step-down-) oder Aufwärts-(Boost-, step-up-) Topologien wird die während jedes EIN-Impulses des PWM-Signals übertragene Energie teilweise durch die Last bestimmt. Wenn die Last verringert wird, wird das Impuls-Pausen-Verhältnis zum Ausgleich verkürzt. Bei der Buck-Boost-Topologie wird das Impuls-Pausen-Verhältnis verwendet, um die Eingangs-Ausgangsspannungsübersetzung zu verändern; dies steht nicht unmittelbar mit der Last in Verbindung. Was passiert nun bei Laständerungen?

Wenn die Last am Wandler hoch ist, verläuft der Strom in der Induktivität I_{L1} wie in Abb. 1.14 dargestellt, in einer dreieckigen Wellenform, die niemals auf null fällt. Der Stromfluss ist ununterbrochen (kontinuierlich). Wenn die Last am Buck-Boost-Wandler jedoch sehr niedrig ist, reicht die Energie in jedem EIN-Impuls aus, um die gewünschte Spannung am Ausgangskondensator schon vor dem Schaltvorgang vollständig zu erreichen. Der Strom I_{L1} fällt dann für den Rest der EIN-Zeit des PWM-Signals auf null. In diesem Fall wird der Strom durch die Speicherinduktivität intermittierend oder diskontinuierlich genannt.

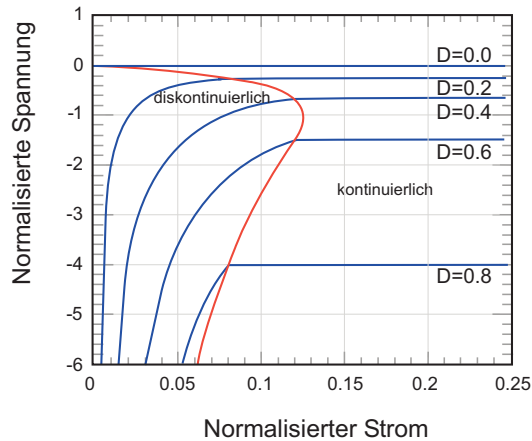


Abb. 1.15: CM- und DCM-Übergang

Der Betrieb im diskontinuierlichen Modus ist mit zusätzlichen Einflüssen auf die Übertragungsfunktion, wie z. B. der Induktivität, der Eingangsspannung und des Ausgangsstroms, verbunden. Daher wird die in Gleichung 1.7 angegebene einfache Übertragungsfunktion komplexer:

$$\frac{V_{OUT}}{V_{IN}} = \frac{V_{IN} \delta^2 T}{2 L_1 I_{OUT}}, \text{ wobei } T = t_{ON} + t_{OFF}$$

Gleichung 1.9: Übertragungsfunktion für den Boost-Wandler im diskontinuierlichen Betrieb

Der Effekt des Übergangs vom kontinuierlichen zum diskontinuierlichen Betrieb ist eine Änderung im Eingangs-Ausgangsspannungsübersetzungsverhältnis bei niedrigen Lasten (Abb. 1.15). Die meisten Buck-Boost-Regler erhöhen deshalb ihre Taktfrequenz bei niedrigen Lasten, um innerhalb der Grenzen des kontinuierlichen Betriebs zu bleiben. Dies erhält das einfache Übertragungsfunktionsverhältnis auf Kosten einer komplizierteren EMV-Filterung, um einen breiteren Bereich von Taktfrequenzen abzudecken. Leider sind Induktivitäten, Kondensatoren und Widerstände im realen Leben nicht ideal, weshalb eine Änderung der Taktfrequenz wegen Nichtlinearität, parasitärer Effekte und unerwünschten Kopplungseffekten häufig auch andere Fehler verursacht.

1.2.2.1.8 Synchrone und asynchrone Umwandlung

In den zuvor vorgestellten Topologien wird eine Freilaufdiode in allen Konstruktionen verwendet. Eine Alternative besteht darin, die Diode durch einen FET zu ersetzen, der mit einem out-of-phase-Signal zum PWM-Signal eingeschaltet wird und die Funktion der Diode übernimmt. Ein Schaltkreis, der einen FET verwendet, plus Diode nennt man asynchron und einen Schaltkreis mit zwei FETs synchron. Abb. 1.16 zeigt zwei alternative Schaltkreise für den Buck-Wandler.

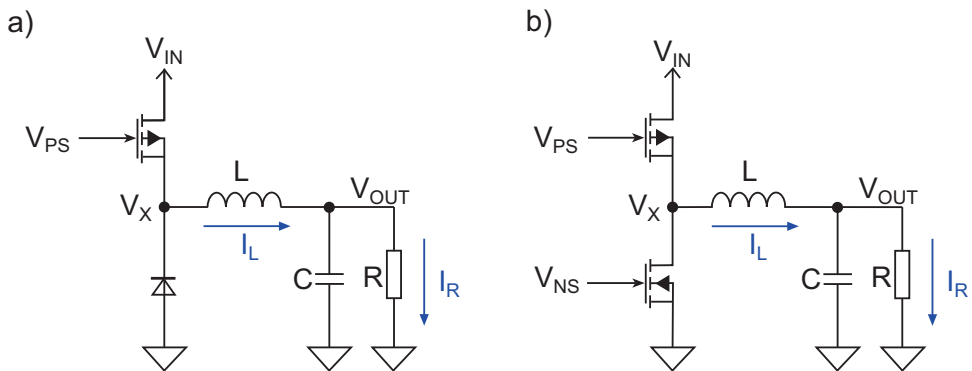


Abb. 1.16: Asynchroner (a) und synchroner (b) Buck-Wandler

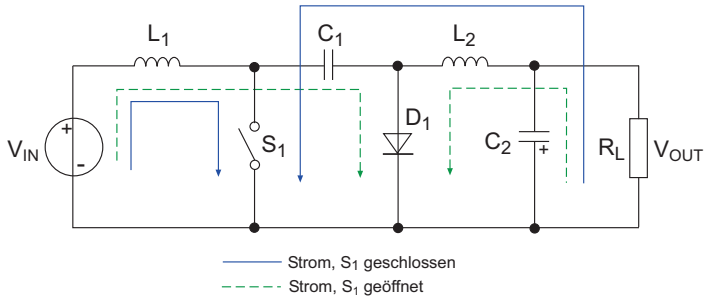
Das Ersetzen der Freilaufdiode durch einen FET hat mehrere Vorteile. Der $R_{DS,ON}$ eines FETs ist sehr niedrig und verursacht – im Gegensatz zur Diode – keinen so hohen Spannungsabfall, sodass die synchrone Ausführung, sowohl bei hohen Eingangsströmen, als auch bei niedrigen Ausgangsspannungen, geringere Verluste und somit besseren Wirkungsgrad zeigt. Die Effizienzerhöhung kann im Volllastbetrieb sehr signifikant sein, da die Verlustleistung der Freilaufdiode ca. das 4-fache der Verlustleistung eines typischen synchronen Wandler mit mittlerer Leistung von 15W betragen kann. Ein anderer Vorteil besteht darin, dass ein Hochstrom-FET von der Bauform her normalerweise viel kleiner ist als eine Leistungsdiode und man so erheblich Platz auf dem PCB sparen kann.

Der Nachteil des synchronen, gegenüber dem asynchronen, Schaltkreis besteht darin, dass die Komponentenkosten nicht nur für den zusätzlichen FET und dessen Ansteuerung höher sind, sondern auch für die Schaltung, die eine gleichzeitige Ansteuerung der beiden FETs verhindert. Ein anderer Nachteil ist, dass die synchrone Ausführung bei sehr niedrigen Lasten ($< 10\%$ der Volllast) in Wirklichkeit weniger effizient sein kann als die asynchrone. Ein Grund dafür sind die zusätzlichen Verluste im zweiten FET-Schaltkreis, der durch Auf- und Entladung der niederspannungsseitigen FET-Gate-Kapazität ebenfalls Verluste produziert. Ein anderer Faktor ist, dass der Strom durch die Induktivität in der asynchronen Ausführung durch die Diode daran gehindert wird, in die entgegengesetzte Richtung zu fließen, in der synchronen Ausführung aber sowohl positive als auch negative Ströme fließen können. Jeder negativer Strom stellt einen zusätzlichen Leistungsverlust dar, den die asynchrone Ausführung vom Prinzip her schon vermeidet.

Controller-ICs, die alle für den synchronen Betrieb notwendigen Signalpegel und das entsprechende Timing generieren, sind leicht erhältlich und haben häufig entweder sowohl hochspannungs- als auch niederspannungsseitige FETs oder nur nieder-spannungsseitige FETs bereits integriert. Zusätzliche Timing-Schaltungen sind oft in den Regel-ICs integriert, um den Wirkungsgrad im Niedriglastbetrieb durch Puls-Skipping (weniger häufiges Einschalten der FETs zur Verringerung von Schaltverlusten) oder durch Reduzierung der Taktfrequenz abhängig von der Last zu erhöhen.

1.2.2.1.9 Zweistufen-Boost-Buck (Ćuk-Wandler)

Auch der Ćuk-Boost-Buck-Regler (ausgesprochen: Tschuk) wandelt eine Eingangsspannung in eine geregelte, invertierte Ausgangsspannung – höher oder niedriger als die Eingangsspannung, je nach Impuls-Pausen-Verhältnis – um. Das vereinfachte Schema in Abb. 1.17 zeigt das Aufbauprinzip und die damit verbundenen Wellenformen. Meistens ist es ein Boost-Wandler, der kapazitiv mit einem invertierenden Buck-Wandler gekoppelt wird.



$$V_{OUT} = V_{IN} \frac{-\delta}{1-\delta}$$

$$V_{IN} > V_{OUT}, \delta < 0.5$$

$$V_{IN} < V_{OUT}, \delta > 0.5$$

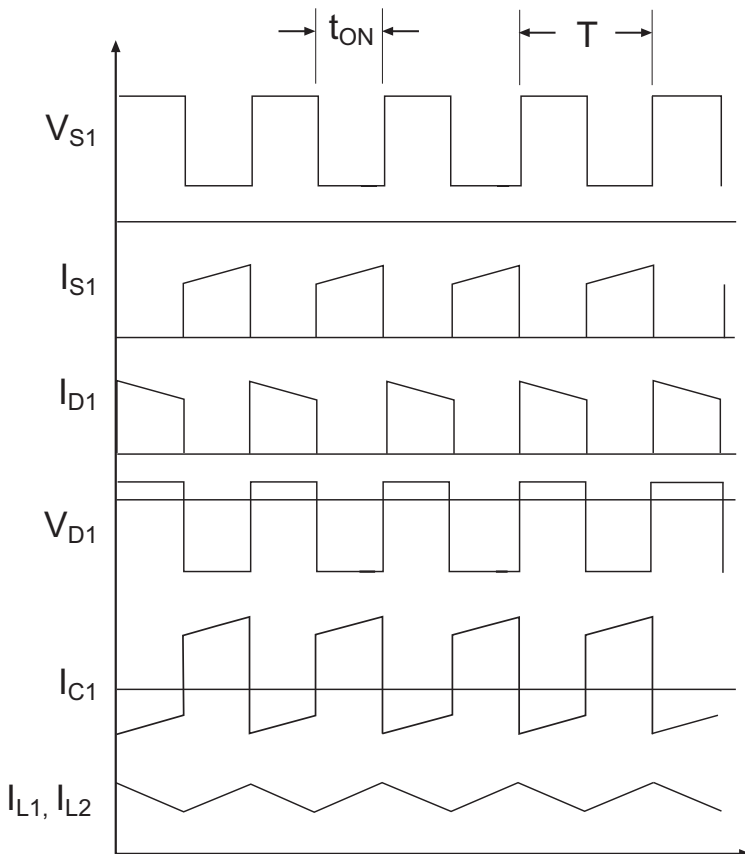


Abb. 1.17: Vereinfachtes Schema und Charakteristik des Ćuk-Wandlers

Im Vergleich zu den oben vorgestellten Topologien ist leicht erkennbar, dass diese Topologie zwei Induktivitäten (welche aufgrund des gemeinsamen Stromflusses auch einen gemeinsamen Kern nutzen können) und zwei Kondensatoren beinhaltet. Wenn Schalter S_1 geschlossen ist, fließt der Strom I_{L1} durch L_1 mit der Anstiegsgeschwindigkeit V_{IN}/L_1 . Gleichzeitig wird die positive Klemme von C_1 auf Masse gezogen, was C_1 veranlasst, mittels eines negativen Stroms durch L_2 den Kondensator C_2 aufzuladen und einen negativen Strom durch die Last R_L zu treiben. Der Strom durch L_2 steigt gemäß dem Verhältnis von $(V_{C1} + V_{OUT})/L_2$. Wenn S_1 geöffnet ist, löst die Energie – gespeichert im Magnetfeld von L_1 – eine Spannungsüberhöhung aus, die dann in Folge C_1 durch D_1 wiederauflädt. Der Strom durch L_1 fällt gemäß dem Verhältnis $(V_{C1}-V_{IN})/L_1$. Gleichzeitig entlädt sich der Kondensator C_2 durch L_2 und Diode D_1 , was einen kontinuierlich abnehmenden Strom L_2 entsprechend V_{OUT}/L_2 bewirkt. Der Kondensator C_1 spielt hier eine besondere Rolle, da er für den gesamten Energiefluss, vom Eingang bis zum Ausgang, verantwortlich ist. Der Wert von C_1 wird so gewählt, dass die Spannung im stabilen Zustand unbedingt konstant bleibt.

Aufgrund der sich wie oben beschrieben einstellenden Stromflussrichtung ist die Ausgangsspannung bezüglich des Massepotentials negativ. Deshalb eignet sich diese Topologie nur für die Erzeugung von negativen Spannungen. Für die Prüfung der Übertragungsfunktion für diese Topologie muss der Einfluss der beiden Induktivitäten geprüft werden.

Auf L_1 anwendbare Gleichungen:

Für den AN-Zustand: $Energy_{IN}(L_1) = V_{IN} t_{ON}$

Für den AUS-Zustand: $Energy_{OUT}(L_1) = (V_{C1} - V_{IN}) t_{OFF}$

Auf L_2 anwendbare Gleichungen:

Für den AN-Zustand: $Energy_{IN}(L_2) = (V_{C1} + V_{OUT}) t_{ON}$

Für den AUS-Zustand: $Energy_{OUT}(L_2) = -V_{OUT} t_{OFF}$

Substitution ergibt zwei Gleichungen für die $C1$ -Kondensatorspannung:

$$V_{C1} = V_{IN} \frac{1}{1 - \delta} \quad \text{und} \quad V_{C1} = \frac{-V_{OUT}}{\delta}$$

wobei die Lösungen dasselbe Ergebnis haben wie für den einstufigen Buck-Boost-Wandler:

$$\frac{V_{OUT}}{V_{IN}} = \frac{-\delta}{1 - \delta}$$

Gleichung 1.10: Übertragungsfunktion des Ćuk-Wandlers

Praktischer Hinweis

Der Vorteil des Ćuk-Wandlers gegenüber dem einstufigen Buck-Boost-Wandler besteht darin, dass die Ströme, die in L_1 und L_2 fließen, dieselben und ununterbrochen sind. Die Eingangs- und Ausgangsströme sind beide effizient LC-gefiltert, was eine Optimierung der EMV-Performance deutlich einfacher macht, da nur sehr geringe Hochfrequenzstörungen entstehen. Da die Ströme in beiden Induktivitäten gleich sind, können sie einen gemeinsamen Kern teilen, was den gesamten Aufbau deutlich vereinfacht und auch dazu beiträgt die Restwelligkeit am Ausgang zu verringern.

Die Konstruktion ist auch deshalb sehr effizient, weil durch das Auf- und Entladen der Kondensatoren über Induktivitäten hohe Stromspitzen mit den dazugehörigen ohmschen Verlusten vermieden werden. Außerdem ermöglicht der auf Masse bezogene Schalter S_1 den Einsatz von FETs mit niedrigen Verlusten und einfacher Ansteuerungsschaltung.

Der größte Nachteil des Ćuk-Wandlers besteht in der starken Abhängigkeit von C_1 . Alle vom Eingang zum Ausgang fließenden Ströme gehen durch diesen Kondensator, der unipolar sein muss, da die Spannungsrichtung beim Durchgang mit jeder Halbperiode wechselt. Die starke Welligkeit des Stroms führt intern zu Wärmeentwicklung, die den Temperaturbereich für sicheren Betrieb einschränkt. In der Praxis heißt das, dass relativ sperrige und meist teure Polypropylen-Kondensatoren zur Anwendung kommen müssen. Darüber hinaus muss der PWM-Regelkreis sehr sorgfältig entwickelt werden, um einen stabilen Betrieb zu gewährleisten. Bei vier reaktiven Komponenten (zwei Induktivitäten und zwei Kondensatoren) muss außerdem die Applikation mit großer Sorgfalt dimensioniert werden um Resonanzen im Steuerstromkreis tunlichst zu vermeiden.

1.2.2.1.10 Zweistufiger Boost-Buck-SEPIC-Wandler

Einer der Nachteile von Buck-Boost-Wandlern ist die invertierte Ausgangsspannung. Dieses Problem kann durch eine zweistufige Konstruktion beseitigt werden, genannt Single Ended Primary Inductor Converter (SEPIC).

Im Wesentlichen ähnelt die Konstruktion der des Ćuk (zwei Stufen: Boost-Wandler gefolgt vom Buck-Wandler), mit Ausnahme der SEPIC-Topologie, bei der die Induktivität L_2 und Diode D_1 die Anordnung tauschen. Dies führt im Endeffekt zur selben Polarität von Ausgangs- und Eingangsspannung.

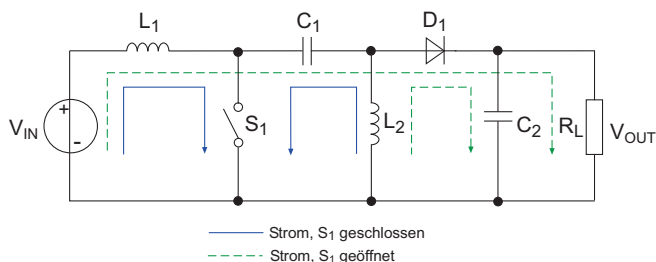


Abb. 1.18: Vereinfachtes Schema der SEPIC-Topologie

Die Energieübertragung ist ähnlich wie im Ćuk-Wandler, weshalb sich folgende Übertragungsfunktion ergibt:

$$\frac{V_{OUT}}{V_{IN}} = \frac{\delta}{1 - \delta}$$

Gleichung 1.11: Übertragungsfunktion des SEPEC-Wandlers

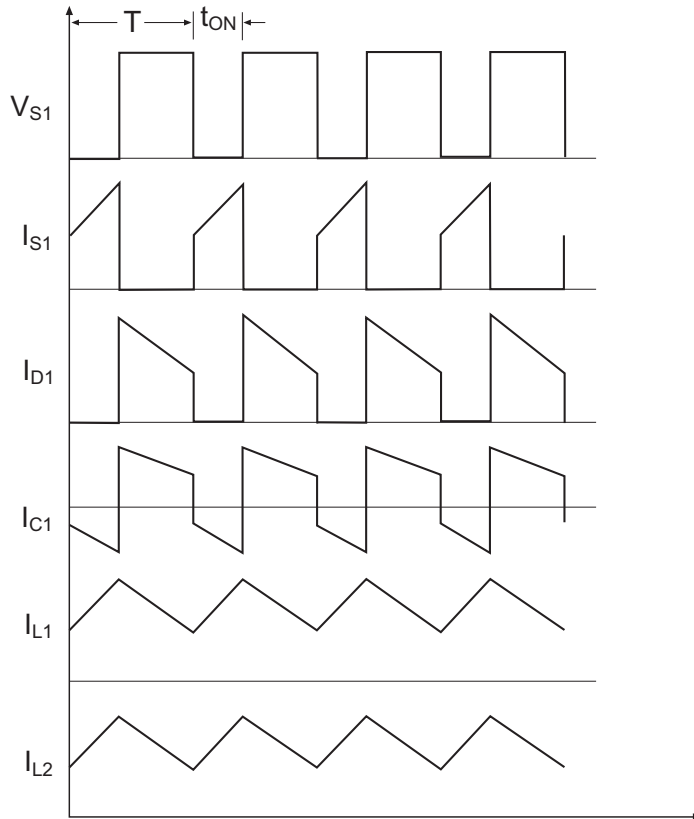


Abb. 1.19: Charakteristik des SEPIC-Wandlers

Die Tatsache, dass die Polarität der Ausgangsspannung der Eingangsspannung gleicht, macht den SEPIC-Schaltkreis für batteriegespeiste Anwendungen, mit wiederaufladbaren Zellen einsetzen, sehr nützlich. Das Batterieladegerät kann dann sowohl zum Wiederaufladen des Akkus, als auch zur gleichzeitigen Versorgung der Anwendung dienen, weil beide eine gemeinsame Masse teilen. Ähnlich wie der Ćuk-Wandler hat SEPIC eine kontinuierliche Eingangsstromwellenform, die die EMV-Filterung erheblich erleichtert.

Praktischer Hinweis

SEPICs werden auch für LED-Beleuchtungsanwendungen eingesetzt, da der Kondensator C_1 einen inhärenten Ausgangskurzschlussschutz darstellt, die Feedbackschleife leicht auf die Konstantstrom- anstelle der Konstantspannungsregelung modifiziert

werden kann und die gemeinsame V-Schiene die EMV-Filterung erleichtert (LED-Beleuchtungsanwendungen, müssen strenge Normvorgaben erfüllen was harmonische Oberwellenstörungen an der Eingangsseite angeht).

Die Nachteile sind, dass der SEPIC-Wandler eine pulsformige Wellenform des Ausgangsstroms, die der des konventionellen Einstufen-Buck-Boost-Wandlers ähnelt. Ebenso besitzt der SEPIC-Wandler, ähnlich wie der Ćuk-Wandler, eine komplexe vierpolige Rückführungsfunktion, was leicht zu Resonanzen führen kann.

1.2.2.1.11 Zweistufen-Boost-Buck-ZETA-Wandler

Eine weitere Variante der SEPIC-Topologie ist der ZETA- oder SEPIC-Gegenwandler. Anstelle einer Boost-Stufe gefolgt von einem Buck-Regler verwendet der ZETA-Wandler einen Buck-Wandler gefolgt von einer Boost-Stufe. Die umgeordnete Topologie besitzt den Vorteil einer SEPIC-Konstruktion, bei der sowohl Ausgangs- als auch Eingangspolarität positiv sind.

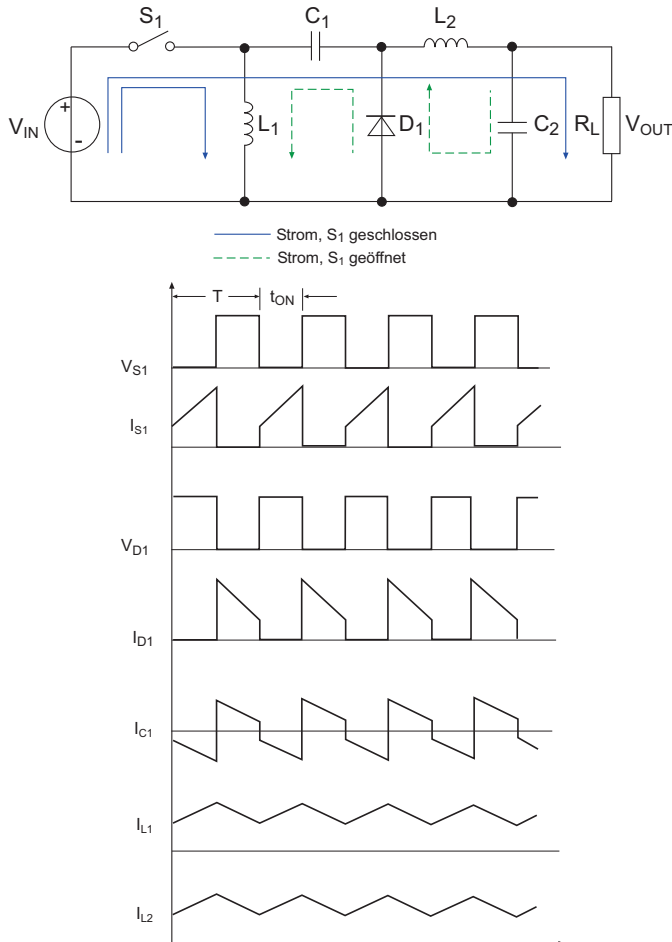


Abb. 1.20: Vereinfachtes Schema und Charakteristik des ZETA-Wandlers

Die Energieübertragung ähnelt der SEPIC-Topologie und ergibt somit dieselbe Übertragungsfunktion:

$$\frac{V_{OUT}}{V_{IN}} = \frac{\delta}{1 - \delta}$$

Gleichung 1.12: Übertragungsfunktion des ZETA-Wandlers

Der Vorteil der ZETA-Topologie gegenüber dem SEPIC-Wandler besteht darin, dass die Feedbackschleife stabiler ist, sodass sie einen weiteren Eingangsspannungsbereich (Input Voltage Range) abdecken kann und höhere Lastschwankungen zulässt, ohne Gefahr zu laufen in Resonanz zu gehen. Die Welligkeit der Ausgangsspannung ist ebenfalls wesentlich niedriger als bei einer vergleichbaren SEPIC-Konstruktion.

Der Nachteil besteht darin, dass die ZETA-Topologie eine höhere Welligkeit des Eingangsstroms aufweist, einen vergleichsweise größeren Kondensator C_1 für dieselbe Energieübertragung benötigt (die Zwischenspannung ist niedriger) und Schalter S_1 nicht massebezogen ist, also ein Pegelwandler erforderlich ist, um den P-Kanal-FET anzutreiben.

1.2.2.1.12 Mehrphasige DC/DC-Wandler

Mehrphasige DC/DC-Wandler sind ein gutes Beispiel für das Gleichgewichtsprinzip in der Elektronik. Dies bedeutet, dass für einen beliebigen erwünschten Vorteil ein Preis in Form eines ausgleichenden Nachteils bezahlt werden muss. Die Forderung nach immer schnelleren Schaltgeschwindigkeiten zur Erhöhung der Verarbeitungsleistung hat eine Herabsetzung der typischen Basisversorgungsspannung des Mikroprozessors von 5V auf 3,3V und weiter auf weniger als 1V verursacht, während die zunehmende Schaltkreiskomplexität zu einem Bedarf an immer höheren Versorgungsströmen geführt hat. Derartige Niederspannungs-/Hochstromversorgungen zu bauen, stößt mitunter auf teils unüberwindbare Schwierigkeiten.

Der Grund, dass mehrphasige DC/DC-Wandler immer mehr nachgefragt werden, liegt zum Teil in den Beschränkungen der Ausgangsfilterkomponenten. Um die Spannungswelligkeit am Ausgang bei höheren Lastströmen bis auf das gewünschte Niveau zu verringern, können die Werte sowohl aus technischen, als auch aus wirtschaftlichen Gründen, nicht beliebig groß gewählt werden. Außerdem bedeutet die Forderung nach immer kleineren Formfaktoren, dass Ausgangsinduktivitäten und -kondensatoren früher oder später durch Bauform und Baugrößeneinschränkungen limitiert sind. Daher ist eine neue Technologie notwendig. Um die Vorteile der mehrphasigen Technologie zu veranschaulichen, wird zuerst ein kurzer Blick auf die einphasige Form geworfen.

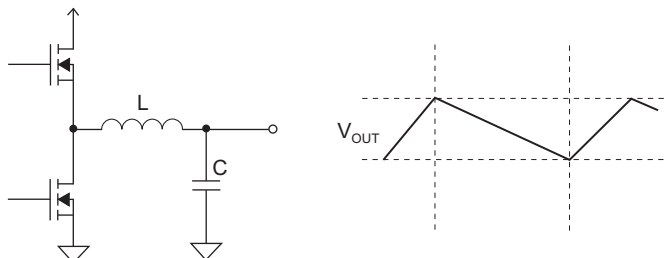


Abb. 1.21: Einphasiges DC/DC-Ausgangsmodell

Während der wiederholten Lade- und Entladungszyklen ändert sich die Ausgangsspannung um den Spitze-zu-Spitze-Wert der Restwelligkeit V_{RIPPLE} . Wenn der Laststrom steigt, steigt auch der Entladungsstrom, und der Ladestrom nimmt ebenso automatisch zu. Dies bedeutet, dass der Strom durch die FETs, Induktivität L und Kondensator C zunimmt. Um V_{RIPPLE} klein zu halten, müssen die Schaltfrequenz und/oder die Werte von L und C vergrößert werden. Um den Wirkungsgrad jedoch hoch zu halten, müssen FETs, Induktivität und Kondensator einen niedrigen Serienwiderstand haben, der zu sperrigeren Komponenten führt. EMV-Belange setzen der maximalen Schaltfrequenz eine Grenze.

Mehrphasige Wandler lösen dieses Problem, indem sie den Laststrom durch mehrere Komponenten gemeinsam nutzen. Abb. 1.22 zeigt das Prinzip unter Verwendung einer zweiphasigen Anordnung.

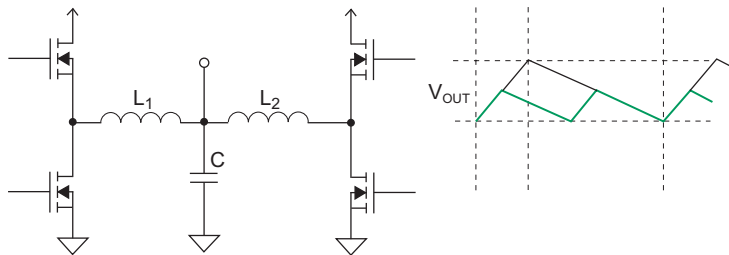
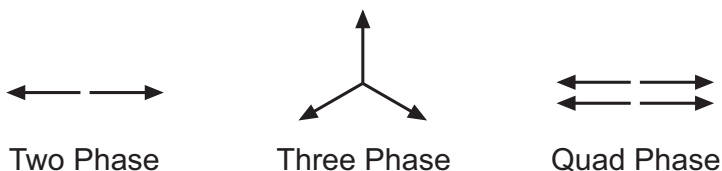


Abb. 1.22: Zweiphasiges DC/DC Ausgangsmodell

Ein Nachteil des mehrphasigen Ausgangs besteht in den höheren Kosten der Komponenten, da für jede zusätzliche Phase zwei zusätzliche FETs und eine weitere Induktivität notwendig sind. Außerdem muss der Ansteuerungsschaltkreis entsprechend konstruiert sein um phasenverschobene Mehrfachausgangssignale zu erzeugen. Wie jedoch zuvor erwähnt, können die Induktivitätswerte um einiges kleiner ausfallen, was zu einer wesentlich kompakteren Konstruktion führt. Der Kondensatorwert kann ebenfalls verringert werden. Die Vorteile gehen aber noch weiter. Vorausgesetzt, dass individuelle Ausgänge außer Phase angeschaltet werden, wird die maximale Amplitude der kombinierten Ausgangsspannung reduziert, der Strom fließt gleichmäßiger, und elektromagnetische Störungen werden somit deutlich reduziert. Das bedeutet, dass der Grad der Filterung am Eingang ebenfalls um einiges verringert werden kann. Schließlich wird die Reaktionszeit bei Laständerungen beschleunigt und die Einstellzeit verringert, da auch die Ausgangskondensatoren verkleinert werden können.

Bei zweiphasigen Ausgängen beträgt die Phasenverschiebung typischerweise 180° ; bei dreiphasigen Ausgängen 120° . Vierphasige Ausgänge werden jedoch typischerweise als zwei Paare, die in Gegenphase laufen, angeordnet. Der Grund liegt darin, dass die Entwicklung der Eingangs-EMV-Filter einfacher ist, wenn im Schaltkreis nicht zu viele phasenverschobene Eingangsrückströme fließen.



Kombinierte mehrphasige Ansteuer-ICs stehen zur Verfügung, die als Buck-, Boost- oder SEPIC-Konfigurationen konfiguriert werden können und beinhalten Schaltkreise zum Kurzschlussschutz, sowie auch zur Absicherung gegenüber unzulässig niedrigen Eingangsspannungen (under-voltage lockout).

1.2.2.2 Isolierter DC/DC-Wandler

In der Familie der isolierten DC/DC-Wandler gibt es eine Vielfalt von Topologien, aber nur drei davon kommen für moderne DC/DC-Wandler in Betracht. Dieser Abschnitt beschränkt sich auf die Behandlung von Flyback-, Vorwärts- sowie Push-Pull-Topologien des Wandlers. Bei diesen Typen isolierter Wandler wird die Energieübertragung vom Eingang zum Ausgang durch einen Transformator ausgeführt. Wie bei den nicht isolierten Wandlern wird die Regelung vom PWM-Controller ausgeführt, wobei die Kontrolle der Ausgangsspannung wiederum in der Feedbackschleife erfolgt. Wir gehen wiederum von Idealkomponenten aus.

Der weitere Unterschied zwischen transformatorbasierten isolierten Wandler-Topologien und den zuvor besprochenen nicht isolierten Topologien besteht darin, dass die Buck-, Boost- oder Buck-Boost-Funktion mit dem Wicklungsverhältnis des Transformators erreicht werden kann und der PWM-Treiber somit nur als einfacher Energiepaket-Controller fungieren muss, der, je nach Anforderung an die Eingangsspannung und die Last am Wandler, mehr oder weniger Energie vom Eingang zum Ausgang überträgt.

Der Nachteil beim Einsatz eines Transformators liegt darin, dass die Energieübertragung von der Primärwicklung zur Sekundärwicklung zusätzliche Verluste nach sich zieht. Während der nichtisolierte Buck-Wandler eine Energieumwandlungseffizienz von 97% erreichen kann, gelingt es isolierten, transformatorbasierten Wandlern kaum, 90% zu überschreiten.

1.2.2.2.1 Flyback-DC/DC-Wandler

Der Flyback-Wandler wandelt eine Eingangsspannung in eine geregelte Ausgangsspannung um, indem er Energie während der EIN-Zeit des PWM-Signals im Transformator Kern speichert und sie während der AUS-Zeit in die Sekundärwicklung überträgt. Abb. 1.23 zeigt den vereinfachten Schaltkreis und Abb. 1.24 die damit verbundenen Spannungs- und Stromwellenformen.

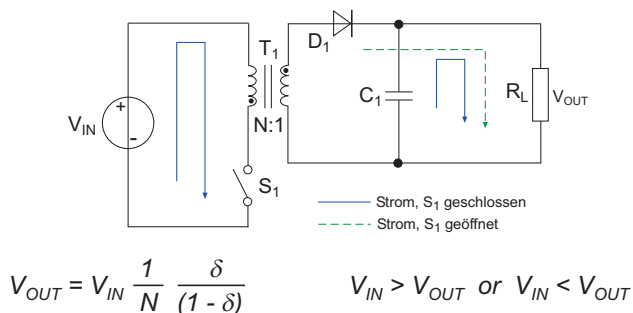


Abb. 1.23: Vereinfachtes Schema des Isolierten Flyback-Wandlers

Wenn Schalter S_1 geschlossen ist, fließt der Strom I_{S1} durch die Transformatorprimärwicklung T_1 mit einer Induktivität L_P und steigt mit einer Anstiegsgeschwindigkeit V_{IN}/L_P an. Während dieser Zeit fließt kein Strom durch die Sekundärwicklung L_S zur Last. Der Laststrom wird in dieser Zeit von Kondensator C_1 versorgt.

Wenn sich S_1 öffnet, bewirkt der Abbau des Magnetfelds im Transformator, dass sich die Polarität der Spannung an den Primär- und Sekundärwicklungen ändert. Die in der Primärwicklung gespeicherte Energie wird jetzt zur Sekundärwicklung übertragen. Die Spannung der Sekundärwicklung erhöht sich steil und ein sich mit der Geschwindigkeit V_{OUT}/L_S verringernder Stromimpuls fließt in die Last und C_1 . Diode D_1 fungiert als Spitzengleichrichter.

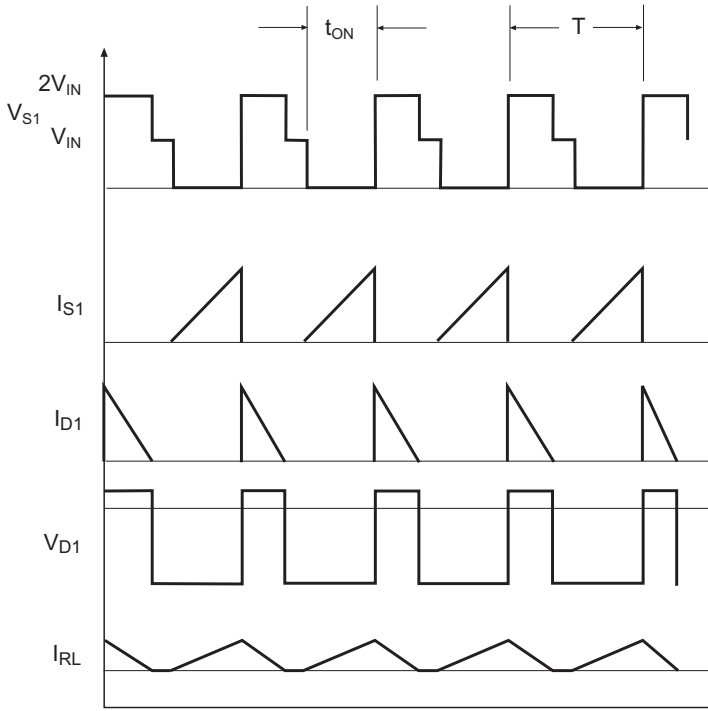


Abb. 1.24: Charakteristik des isolierten Flyback-Wandlers

Die entsprechenden Energiegleichungen sind:

Für den AN-Zustand: $Energy_{IN} = \frac{V_{IN} t_{ON}}{N}$, wobei N = Windungsverhältnis

Für den AUS-Zustand: $Energy_{OUT} = V_{OUT} t_{OFF}$

Substitution ergibt: $\frac{V_{IN} t_{ON}}{N} = V_{OUT} (T - t_{ON})$

nach Umordnung: $\frac{V_{OUT}}{V_{IN}} = \frac{1}{N} \frac{\delta}{1 - \delta}$

Gleichung 1.13: Übertragungsfunktion des isolierten Flyback-Wandlers

Praktischer Hinweis

Die Übertragungsfunktionen des Buck-Boost-Wandlers und des isolierten Flyback-Wandlers unterscheiden sich also nur durch den Transformator-Windungsverhältnisfaktor $1/N$. Der Vorteil der Konstruktion des Flyback-Wandlers besteht darin, dass eine sehr hohe Vervielfachung der Ausgangsspannung bei kurzen Impuls-Pausen-Verhältnissen zulässig ist, weshalb sich diese Topologie gut für hohe Ausgangsspannungen eignet. Ein weiterer Vorteil besteht darin, dass Mehrfachausgänge (ggf. mit verschiedenen Polaritäten) durch Hinzufügen von mehreren Sekundärwicklungen leicht realisiert werden können. Die Anzahl der Komponenten ist ebenfalls sehr niedrig, was diese Topologie für preiswerte Konstruktionen prädestiniert.

Durch Überwachung der Ausgangsspannung bzw. des Ausgangsstroms und einen isolierten Rückführkreis (typischerweise durch einen Optokoppler) lässt sich ein sehr stabiler geregelter Ausgang erzeugen. Flyback-Wandler können durch Überwachung der primärseitigen Signalform und Verwendung des Knickpunktes zur Ermittlung des Nulldurchgangs des Sekundärstroms auch primärseitig geregelt werden. Dies macht den Optokoppler überflüssig und verringert die Komponentenanzahl noch weiter.

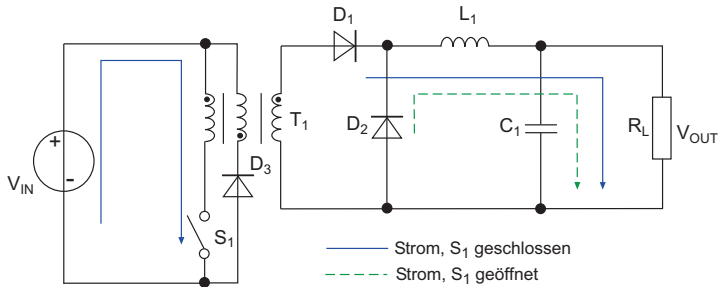
Der Nachteil besteht darin, dass der Transformatorkern sorgfältig ausgewählt werden muss. Der Luftspalt-Kern soll nicht gesättigt werden, obwohl ein positiver DC-Durchschnittsstrom durch den Transformator fließt; der Wirkungsgrad kann drastisch sinken, wenn die magnetische Hysterese zu groß wird. Auch Wirbelstromverluste in den Wicklungen können wegen den hohen Spitzenströmen zu einem Problem werden. Diese beiden Effekte beschränken den praktischen Taktfrequenzbereich dieser Topologie. Und schließlich führt die große induktive Spitze an der Primärwicklung, die sich ergibt, wenn S1 ausgeschaltet ist, zu einer starken Belastung des Schalt-FET.

1.2.2.2 Vorwärts-DC/DC-Wandler

Obwohl der Vorwärtswandler der Flyback-Topologie zu ähneln scheint, funktioniert er auf vollkommen andere Weise. Die Eingangsspannung wird als Funktion des Windungsverhältnisses des Transformators in eine geregelte Ausgangsspannung umgewandelt. Abb. 1.25 zeigt den vereinfachten Schaltkreis und die damit verbundenen Spannungs- und Stromwellenformen.

Wie in der Flyback-Topologie, wenn Schalter S_1 geschlossen ist, fließt der Strom I_{S1} durch die Transformatorprimärwicklung T_1 mit der Induktivität L_P und der Stromanstiegsgeschwindigkeit V_{IN}/L_P . Der steigende Primärstrom induziert den Sekundärstrom im Transformator T_1 infolge der magnetischen Kopplung zwischen Primär- und Sekundärwicklungen mit einem Spannungswert V_{IN}/N . Der Sekundärstrom fließt durch die Gleichrichterdiode D_1 und die Ausgangsinduktivität L_1 und wächst mit der Geschwindigkeit $V_{IN}/(L_1 N)$. Dieser Strom fließt auch in die Last R_L und den Ausgangskondensator C_1 . So steigt die Spannung am Kondensator C_1 , bis die obere Regelgrenze überschritten ist und das 'Stop'-Signal gegeben wird (die Rückmeldung erfolgt normalerweise durch einen Optokoppler). Der Controller der primären Seite bewirkt dann, dass sich S_1 öffnet, der Strom von der Spannungsquelle wird unterbrochen.

Die Rücksetzwicklung bei Diode D_3 verhindert den Abbau des Magnetfelds des Transformators und lässt den Strom stattdessen mit der gleichen Geschwindigkeit absinken, mit der er zuvor angestiegen ist, als S_1 noch geschlossen war. Wenn sich S_1 öffnet, setzt ein Polaritätswechsel in der Sekundärwicklung ein, der negative Strom verringert sich mit der Geschwindigkeit V_{OUT}/L_1 und fließt durch die Fangdiode D_2 und Induktivität L_1 und schließlich in die Last und C_1 . Die Spannung an C_1 verringert sich, bis die untere Regelgrenze erreicht ist. Ein 'Start'-Signal wird gesendet, S_1 schließt sich wieder, und ein neuer Zyklus beginnt.



$$V_{OUT} = V_{IN} \frac{1}{N} \delta \quad V_{IN} > V_{OUT} \text{ or } V_{IN} < V_{OUT}$$

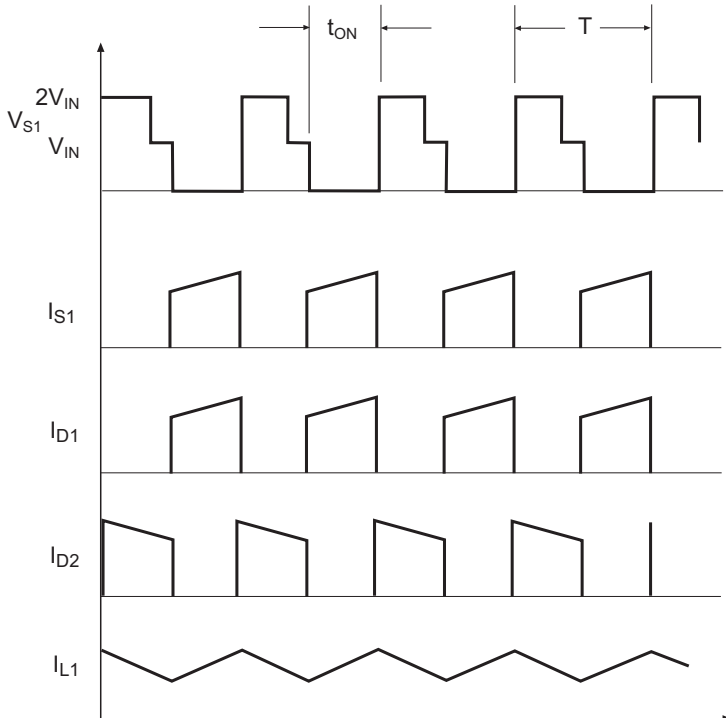


Abb. 1.25: Vereinfachtes Schema und Charakteristik des Vorwärtswandlers

Die entsprechenden Energiegleichungen sind:

Für den AN-Zustand: $Energy_{IN} = \left(\frac{V_{IN}}{N} - V_{OUT} \right) t_{ON}$, wobei $N = \text{Transformationsverhältnis}$

Für den AUS-Zustand: $Energy_{OUT} = V_{OUT} t_{OFF}$

Nach Umordnung: $\left(\frac{V_{IN}}{N} - V_{OUT} \right) t_{ON} = V_{OUT} (T - t_{ON})$

$$\frac{V_{OUT}}{V_{IN}} = \frac{\delta}{N}$$

Gleichung 1.14: Übertragungsfunktion des isolierten Vorwärtswandlers

Im Unterschied zum Flyback-Wandler überträgt ein Vorwärtswandler die Energie von der Primär- zur Sekundärwicklung fortlaufend über den Transformator, statt Energie-Pakete im Transformator kernspalt zu speichern; daher kann der Kern auf einen Luftspalt und die damit verbundenen Verluste und die abgestrahlte EMI verzichten. Der Kern kann auch eine höhere Induktivität aufweisen, da Hystereseverluste nicht so kritisch sind. Die geringeren Spitzenströme reduzieren Wicklungs- und Diodenverluste und führen zu einer niedrigeren Welligkeit von Eingangs- und Ausgangsstrom. Bei gleicher Ausgangsleistung ist der Vorwärtswandler deshalb effizienter.

Der Nachteil ist in den höheren Kosten der Komponenten und der Notwendigkeit einer gewissen Mindestlast zu sehen, die erforderlich ist, um zu verhindern, dass der Wandler bei einer entsprechend dramatischen Änderung in der Übertragungsfunktion in den diskontinuierlichen Betrieb übergeht.

1.2.2.2.3 Aktive Klemmvorwärtswandler

Eine Variation des Vorwärtswandlers besteht in der Verwendung einer aktiven Klemme (FET) zur Rücksetzung des Transformators anstelle einer einzelnen Wicklung. Der vereinfachte Schaltkreis ist unten erkennbar:

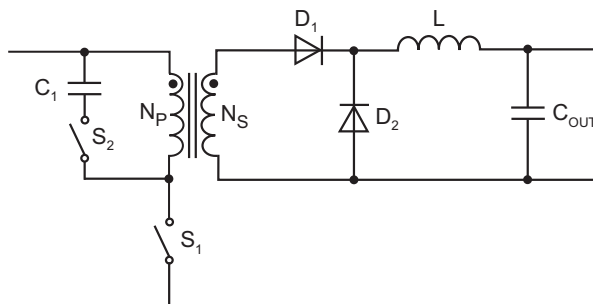


Abb. 1.26: Aktiver Klemmvorwärtswandler

S_2 wird mit einem phasenverschobenen PWM-Signal mit ausreichender Phasenverschiebung angesteuert, wodurch die beiden Transistoren nicht gleichzeitig eingeschaltet werden. Die Wellenformen ähneln denen des Vorwärtswandlers, mit Ausnahme dessen, dass die Spannung an S_1 eine Rechteckwelle darstellt. Die im Ausgang fließenden Ströme sind gleich. Der Grund dafür, dass die Ausgangswellenformen gleich sind, liegt darin, dass das Magnetfeld nicht abgebaut wird, wenn sich S_1 öffnet, sondern allmählich gedämpft wird, da der Strom in den Primärwicklungen immer noch durch C_1 und S_2 fließen kann. Deshalb ist die Übertragungsfunktion dieselbe.

Das Hinzufügen einer aktiven Klemme hat verschiedene Vorteile. Eine Rücksetzwicklung des Transformators ist nicht erforderlich, und die Spannung an S_1 erreicht ihren Spitzenwert bei V_{IN} und nicht bei $2 \times V_{IN}$ wie bei der Standard-Topologie. Der Gesamtwirkungsgrad ist höher, da Diodenverluste vermieden werden und nur der Entmagnetisierungsstrom durch S_2 fließt. Und was noch wichtiger ist: Die aktive Klemme lässt einen Betrieb über 50% des Impuls-Pausen-Verhältnisses mit höheren Windungsverhältnissen zu, ohne Verluste bei hohen Impulsspannungen an S_1 hinnehmen zu müssen.

Der Nachteil der aktiven Klemme liegt darin, dass ein zweites PWM-Signal erzeugt werden muss und S_2 einen hochspannungsseitigen Treiber benötigt. Es gibt jedoch viele Controller-ICs, die die notwendigen Zeitschaltungen und hochspannungsseitigen Treiber bereits integriert haben. Der Klemmkondensator C_1 hat einen hohen Wellenstrom und es muss sichergestellt sein, dass er nicht überhitzt. Der Strom im Klemmkondensator kann wie folgt angenähert dargestellt werden:

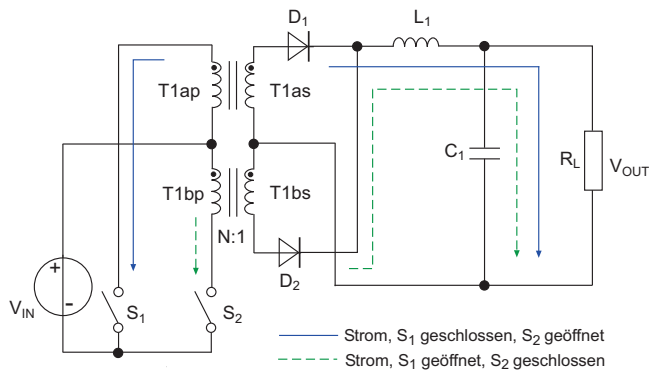
$$I_{C,Clamp(rms)} \approx \frac{V_{IN} \delta}{f_{SW} L_{MAG}} \sqrt{\frac{1 - \delta}{2}}$$

wobei L_{MAG} = Magnetisierungsinduktivität des Transformators.

Gleichung 1.15: Annäherung für Klemmkondensatorstrom

1.2.2.2.4 Push-Pull-Wandler

Der Push-Pull-Wandler wandelt eine Eingangsspannung in eine niedrigere geregelte Ausgangsspannung um und benötigt zum Betrieb einen Transformator mit geteilter Wicklung. Abb. 1.27 zeigt den vereinfachten Schaltkreis und die damit verbundenen Spannungs- und Stromwellenformen.



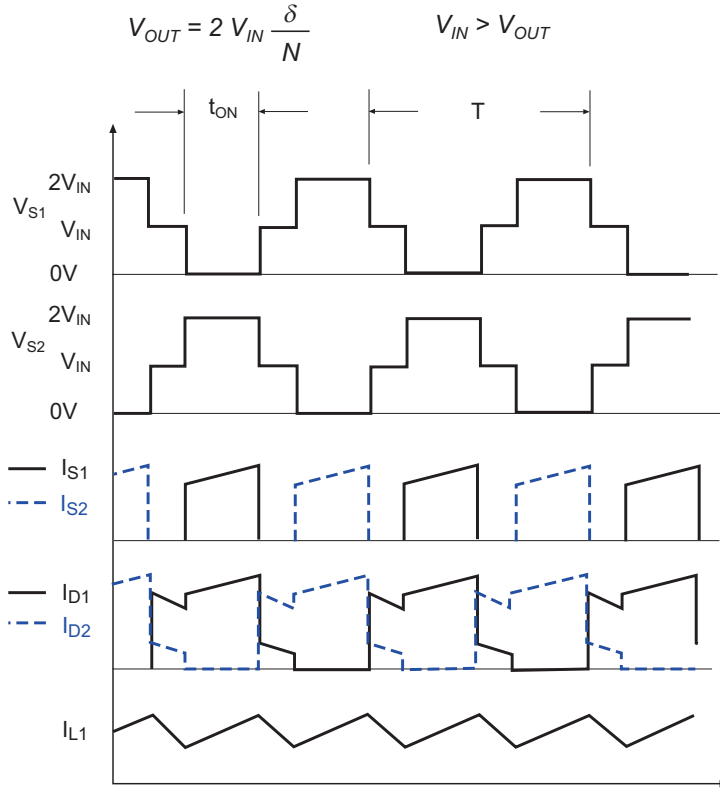


Abb. 1.27: Vereinfachtes Schema und Charakteristik des Push-Pull-Wandlers

Wenn der Schalter S_1 geschlossen ist, erhöht sich der Strom durch die Transformatorprimärwicklung mit der approximierten linearen Anstiegsgeschwindigkeit $V_{IN}/L_{T1,AP}$. Gleichzeitig wird eine Spannung V_{IN}/N in der Sekundärwicklung $T_{1,AS}$ infolge der Kopplung der Primär- und Sekundärwicklung des Transformators aufgebaut. Der Sekundärstrom, der durch die Gleichrichterdiode D_1 und die Induktivität L_1 fließt, steigt linear mit der Geschwindigkeit $(V_{IN}/N - V_{OUT})/L_1$ an. Dieser Strom fließt auch in die Last R_L und lädt den Ausgangskondensator C_1 auf. Wenn S_1 geöffnet ist, vollzieht sich ein Polaritätswechsel, wobei Diode D_1 die negative Spannung an der Sekundärwicklung $T_{1,AS}$ blockiert. Der Strom setzt jedoch seinen Fluss durch L_1 mittels Diode D_2 von der invertierten Sekundärwicklung $T_{1,BS}$ fort. Der Strom verringert sich nun linear proportional zu V_{OUT}/L_1 . S_2 wird dann geschlossen, und der Zyklus beginnt von vorne, jedoch über Sekundärwicklung $T_{1,BS}$ die den Strom versorgt, während S_2 geschlossen ist. Um die Übertragungsfunktion zu erhalten, werden folgende Energiegleichungen verwendet:

Für den AN-Zustand: $Energy_{IN} = \left(\frac{V_{IN}}{N} - V_{OUT} \right) t_{ON}$, wobei $N = \text{Transformationsverhältnis}$

Für den AUS-Zustand: $Energy_{OUT} = V_{OUT} t_{OFF}$, wobei $t_{OFF} = T/2 - t_{ON}$

Der Wert $T/2$ wird verwendet, da während der PWM-Zykluszeit zwei Schalter im Einsatz sind, weshalb die Energie, die während der EIN-Zeit jedes Transistors in der Zeit T erzeugt wird, halbiert wird.

$$\begin{aligned} \text{Nach Umordnung:} \quad & \left(\frac{V_{IN}}{N} - V_{OUT} \right) t_{ON} = V_{OUT} (T/2 - t_{ON}) \\ \text{oder} \quad & \frac{V_{OUT}}{V_{IN}} = \frac{2 \delta}{N} \end{aligned}$$

Gleichung 1.16: Übertragungsfunktion des Push-Pull-Wandlers

Da das Impuls-Pausen-Verhältnis sowohl für S_1 als auch für S_2 nahezu 50% beträgt, ist es unbedingt erforderlich sicherzustellen, dass die beiden Schalter nicht gleichzeitig eingeschaltet werden können, sonst würden sehr hohe Kurzschlussströme (Shoot-through) fließen. Deshalb ist die entsprechende Totzeit zwischen der Ausschaltung eines Schalters und der Schließung des anderen erforderlich.

Ein anderes Problem, das im Push-Pull-Wandler vorkommen kann, ist eine Verschiebung des magnetischen Flusses (Flux Walking). Da der Push-Pull-Wandler den gesamten Bereich der BH-Kennlinie des Transformators verwendet, könnte der kleinste Unterschied in den Arbeitscharakteristiken der Schalter (Übersteuerungsspannungen, Schaltzeiten usw.) zu einer Unausgeglichenheit im magnetischen Fluss führen. Der Versatz dieser Flussverschiebung ist leider kumulativ, da diese im Transformator am Ende jedes Schaltzyklus nicht vollständig auf null gesetzt werden kann. Der Offset vom Vorzyklus wird somit zum Ausgangspunkt des folgenden Zyklus. Das Kernmaterial des Transformators kann schließlich gesättigt werden und die Energieübertragung noch weiter aus dem Gleichgewicht bringen. Da ein gesättigter Kern kein induktives Verhalten mehr zeigt, können einer oder beide Schalter durch hohe Ströme in der Primärwicklung zerstört werden. Dieses Problem kann vermieden werden, indem der Strom von Zyklus zu Zyklus gemessen und beschränkt wird.

Andererseits kann die Push-Pull-Topologie durch einen Transformator gleicher Größe die doppelte Leistung übertragen, da der Push-Pull-Wandler im Gegensatz zum Vorwärtswandler, der nur den ersten Quadranten verwendet, beide Quadranten der BH-Kurve des Transformators nutzt. Das trägt zur Kosteneffizienz dieser Topologie bei, die für höhere Ausgangsleistungen oder zur Herstellung deutlich kleinerer DC/DC-Wandler geeignet ist.

Da das Impuls-Pausen-Verhältnis für den maximalen Wirkungsgrad typischerweise auf etwa 50% festgelegt wird, ist das Eingangs-Ausgangsspannungsübersetzungsverhältnis durch das Windungsverhältnis des Transformators bestimmt. Deshalb wird der geregelte Push-Pull-Wandler mit geregelter Eingangsspannung optimalerweise als Bus-Wandler verwendet.

1.2.2.2.5 Halb- und Vollbrückenwandler

Eine Topologie, die dem Push-Pull-Wandler ähnelt, weisen die Halb- und Vollbrückenwandler auf, die zwei oder mehrere Schalter verwenden, um den Strom durch die Transformatorprimärwicklung zu steuern, welche – im Gegensatz zum Push-Pull-Wandler – keinen mittigen Primärwicklungsanschluss mehr benötigt (aber immer noch den mittigen Sekundärwicklungsanschluss verwendet).

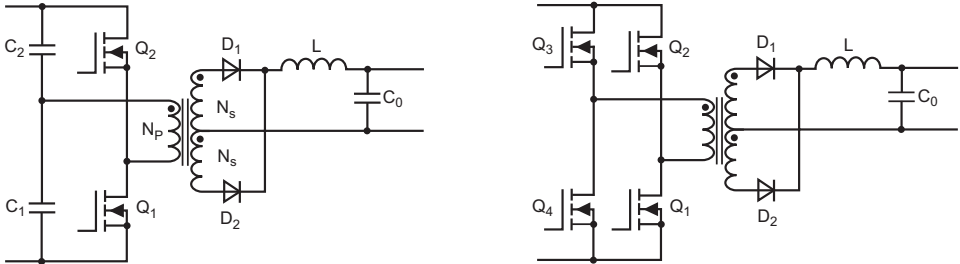


Abb. 1.28: Halb- und Vollbrückenwandler

Die Halbbrücke verwendet die beiden Kondensatoren C_1 und C_2 , um eine Spannungsteilung zu erzielen, sodass ein Ende der Primärwicklung bei $V_{IN}/2$ verbleibt. Die beiden Schalter S_1 und S_2 schließen dann abwechselnd das andere Ende der Wicklung an $V_{IN}+$ oder an die primärseitige Masse. Da die Spannung durch Primärwicklung $|V_{IN}/2|$ nicht überschreitet, wird das Übertragungsverhältnis im Vergleich zum Push-Pull-Wandler halbiert:

$$\frac{V_{OUT}}{V_{IN}} = \frac{\delta}{N}$$

Gleichung 1.17: Übertragungsfunktion des Halbbrückenwandlers

Die Vorteile der Halbbrücken- gegenüber der Push-Pull-Topologie bestehen darin, dass die Schalter V_{IN} statt $2 \times V_{IN}$ standhalten müssen und dass das Problem des Flux-Wal-king beseitigt ist, da die Primärwicklung eine einzelne Wicklung ist. Der Gesamtwirkungs-grad ist in den meisten Fällen höher, weshalb sich die Halbbrückentopologie für Hochleistungssysteme eignet, während die vereinfachte Transformator-konstruktion diese Topologie ideal für die Anwendung von Planartransformatoren macht. Der Nachteil ist der hohe, pulsierende Strom durch C_1 und C_2 , die deshalb sorgfältig dimensioniert werden müssen, damit sie nicht überhitzen. Das Impuls-Pausen-Verhältnis ist ebenfalls in den meisten Fällen auf 45% begrenzt, um Shoot-through zu vermeiden (S_1 und S_2 sind zu gleicher Zeit angeschaltet). Schließlich ist ein hochspannungsseitiger Treiber für S_2 notwendig, was die Komponentenkosten erhöht.

Die Nachteile der Halbbrücke können mit der Vollbrückentopologie, die vier Schalter, aktiviert in der Folge $S_3 + S_1$: EIN, $S_2 + S_4$: AUS und dann $S_2 + S_4$: EIN, $S_3 + S_1$: AUS verwendet, beseitigt werden, sodass die Primärwicklung in jedem Schaltzyklus immer die gesamte Eingangsspannung sieht.

Die Vollbrückentopologie besitzt alle Vorteile der Halbbrückentopologie, jedoch keinen ihrer Nachteile. Das Timing der Schalteransteuerung ist komplizierter und zwei pimärseitige Treiber sind notwendig, weshalb Vollbrückenkonstruktionen typischerweise für Hochleistungsanwendungen verwendet werden, bei denen zusätzliche Komponentenkosten weniger bedeutsam sind. Die Übertragungsfunktion der Vollbrücke ist somit identisch mit der des Push-Pull-Wandlers.

1.2.2.2.6 Bus- oder ratiometrischer Wandler

Der Bus-Wandler, auch ratiometrischer Wandler genannt, nimmt eine Sonderposition unter den isolierten DC/DC-Wandlern ein. Der Bedarf an solchen Wandlern stammt aus komplizierten Telekommunikations-Stromversorgungssystemen, die viele verschiedene Versorgungsspannungen enthalten. Statt eine separate Stromversorgung für jede Spannung zu bauen, wurde das Konzept einer Zwischenkreisarchitektur (IBA, engl. Intermediate Bus Architecture oder DBA, engl. Distributed Power Architecture) erfunden. Hierbei wird die Primärversorgung zuerst in eine isolierte DC-Zwischenversorgung umgewandelt, die dann zur Speisung anderer, nicht isolierter DC/DC-Wandler verwendet werden kann, welche nahe an der zu versorgenden Last sitzen können (POL, engl. Point of Load).

Der Bus-Wandler hat ein fixes Wandlungsverhältnis, normalerweise 4:1; daher der alternative Name ratiometrischer Wandler (aus dem Englischen, ratiometer = Verhältnismesser). Dies bedeutet, dass sich die Ausgangsspannung proportional zur Eingangsspannung ändert, was aber nicht ins Gewicht fällt, da die darauffolgenden POL-Abwärtswandler einen weiten Eingangsspannungsbereich besitzen. Sie sind stattdessen für den maximalen Wirkungsgrad optimiert und liefern sogar bei sehr hohen Lastströmen 97% oder mehr.

Bus-Wandler können mit Vorwärts- oder Push-Pull-Topologien unter Verwendung der Halbbrücken- oder der Vollbrückenschaltung hergestellt werden, jedoch mit fixen, auf den maximalen Wirkungsgrad abgestimmten Impuls-Pausen-Verhältnissen. Darüber hinaus wird häufig Synchrongleichrichtung verwendet, um die Ausgangsdioden zu ersetzen und weitere Verluste zu verringern.

In der Praxis werden häufig zwei Zwischenkreisspannungen verwendet. Der AC-Eingang des Speisernetzes wird zuerst in 48VDC umgewandelt, was durch Akkus abgesichert wird, um auch bei Netzausfall eine Stromversorgung zu gewährleisten. Die 48-V-Spannung wird dann im Verhältnis 4:1 ratiometrisch abwärts umgewandelt, um eine 12-V-Schiene für POL-Wandler sicherzustellen, die wiederum 5-V- und 3,3-V-Versorgungen auf der Leiterplattebene zur Verfügung stellt (Abb. 1.29).

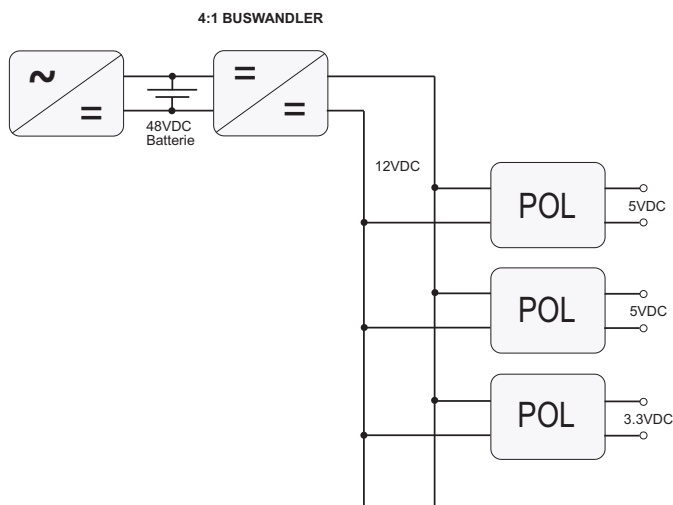


Abb. 1.29: Vereinfachte IBA-Regelung

1.2.2.2.7 Ungeregelter Push-Pull-Wandler

Die Push-Pull-Topologie findet auch in den isolierten unregulierten DC/DC-Wandlern eine breite Verwendung. Wenn die Eingangsspannung geregelt ist, ist die Push-Pull-Topologie ein kostengünstiges Verfahren der Erzeugung hoher, niedriger, invertierter oder bipolarer Spannungen, da allein das Transformatorwindungsverhältnis das Übersetzungsverhältnis bestimmt. Abb. 1.30 zeigt den Schaltkreis eines unregulierten Push-Pull-Wandlers, mit induktiver Rückkopplung eines Freischwingoszillators (Royer-Topologie).

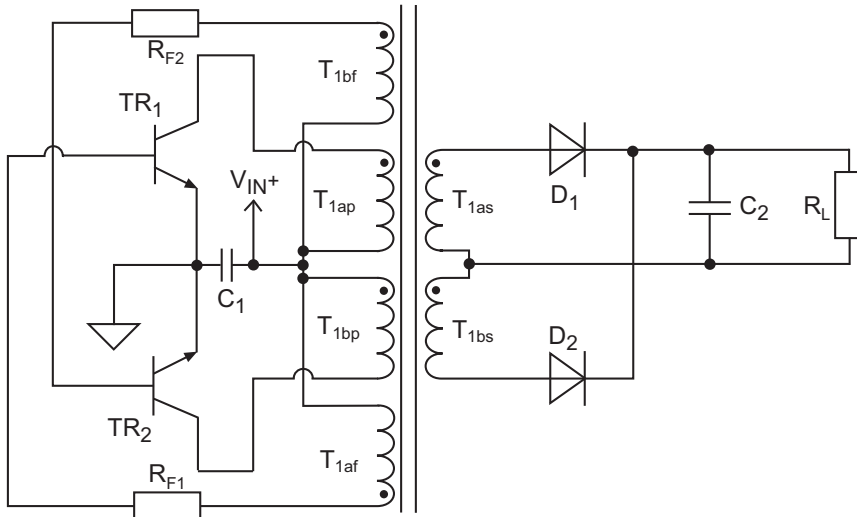


Abb. 1.30: Ungeregelter Push-Pull Wandler

Wie man dem Schema entnehmen kann, ist der Schaltkreis symmetrisch. Das Anlegen der Eingangsspannung legt die Basiselektroden beider Transistoren über die strombegrenzenden Widerstände R_{b1} und R_{b2} an V_{IN+} an, wobei sich der Transistor zuerst mit dem niedrigsten V_{BE} -Wert anschaltet. Obwohl T_{R1} und T_{R2} zum Beispiel Transistoren gleichen Typs sind, reagiert T_{R1} aufgrund technologischer Toleranzen etwas schneller. Der Strom fließt durch $T_{1,ap}$ und speist den Transformator, wobei er positive Ströme in $T_{1,as}$ und $T_{1,bf}$ und negative Ströme in $T_{1,bs}$ und $T_{1,af}$ erzeugt. Der in $T_{1,af}$ erzeugte negative Strom schaltet T_{R1} aus, wobei der Strom in $T_{1,ap}$ unterbrochen wird, während der positive Strom in $T_{1,bf}$ T_{R2} einschaltet. Wenn sich T_{R2} einschaltet, fließt der Strom nun durch $T_{1,bp}$ und speist den Transformator wieder, erzeugt jetzt aber positive Ströme in $T_{1,bs}$ und $T_{1,af}$ und negative Ströme in $T_{1,as}$ und $T_{1,bf}$. Der in $T_{1,b}$ erzeugte negative Strom schaltet T_{R2} aus, wobei der Strom in $T_{1,bp}$ unterbrochen wird, während der positive Strom in $T_{1,af}$ T_{R1} wieder einschaltet. Der Wandler ist somit ein Freilaufoszillator mit Transformatorenkopplung, der sich schnell auf ein 50% Impuls-Pausen-Verhältnis, der effizientesten Arbeitscharakteristik, einschwingt.

Das oben gezeigte Schema ist fast vollständig, es fehlen nur einige passive Komponenten, um einen vollständig funktionierenden DC/DC-Wandler zu ergeben. Somit ist dieser Typus eines Wandlers der kostengünstigste, da er nur mit ca. 10 Bauteilen realisiert werden kann. Die Bauform des Wandlers kann sehr klein sein. RECOM bietet den RNM-Wandler mit einer Gehäusegröße von nur 8,3 x 8,3 x 6,8mm an, der ungeachtet seiner geringen Größe immer noch 1W Ausgangsleistung und 2000VDC Isolation zwischen Eingang und Ausgang bietet.

Es gibt bei dieser Ausführung keinen Rückführkreis vom Ausgang zum Eingang. Daher schwingt der Wandler mit 50% Impuls-Pausen-Verhältnis, ungeachtet dessen, ob am Ausgang eine Last anliegt oder nicht, er ist unregelt. Im Lastbetrieb werden die Schaltspitzen, die am Ausgang wegen parasitärer Effekte eintreten, stark gedämpft und wirken nicht wesentlich auf die Ausgangsspannung ein. Im lastfreien Betrieb werden die Spitzen jedoch von den Ausgangsdioden gleichgerichtet, und wesentlich höhere Ausgangsspannungen, als die Berechnung des Transformator-Windungsverhältnis bestimmen würde, können auftreten. Eine typische Abweichungskurve der Ausgangsspannung vs. Last ist in Abb. 1.31 dargestellt:

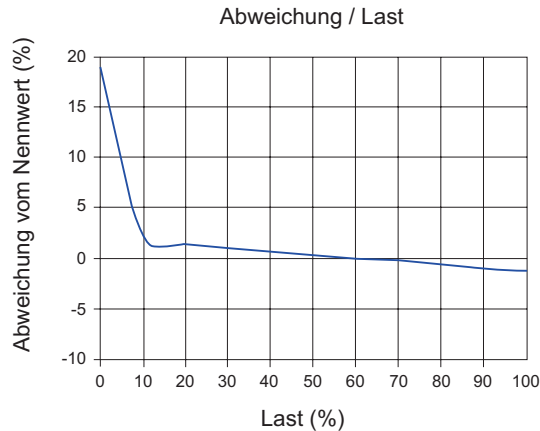


Abb. 1.31: Typische Ausgangsspannungsabweichungskurve eines unregelteten Wandlers

Wie der Ausgangsspannungsabweichungskurve zu entnehmen ist, sollen Lasten kleiner als 10% mit unregelteten Wandlern vermieden werden. Wenn Null-Last-Betrieb eine Konstruktionsanforderung ist, kann entweder ein Blindlastwiderstand am Ausgang gesetzt werden, um ein Ansteigen der Ausgangsspannung zu vermeiden, wenn sich die Last im Leerlauf oder Niedriglastbetrieb befindet, oder eine Zener-Diode verwendet werden, um die Ausgangsspannung innerhalb sicherer Grenzen zu halten.

Außer unter 10% Lastbereich ist die Ausgangsspannungsabweichung überraschend flach. Im Beispiel oben liegt die Lastregelung innerhalb von $\pm 2\%$ für Lasten zwischen 10% und 100%. Das ist eine sehr beachtliche Ziffer, die mit einigen geregelten Wandlern vergleichbar ist. Für viele kostensensitive Anwendungen, bei denen Eingangsspannung und Last relativ konstant sind, ist ein unregelteter Wandler eine sehr wirtschaftliche Lösung, die typischerweise 30% günstiger ist als eine vergleichbare geregelte Alternative.

Andere Nachteile des unregelteten Wandlers sind die komplexere Transformator-konstruktion (sechs Wicklungen anstelle von vier) und die Tatsache, dass sich der Freischwingoszillator nicht ohne weiteres abschalten lässt. Dies bedeutet, dass Standardschutzmaßnahmen, wie z. B. Überhitzungs-, Überlastungs- oder Kurzschluss-schutz, in dieser Topologie meistens fehlen.

Das in Abb. 1.30 gezeigte Schema gilt für einen einfachen Ausgang. Durch Umkehren von D2 und Hinzufügen eines zweiten Ausgangskondensators kann daraus jedoch leicht ein DC/DC-Wandler mit Bipolarausgang konstruiert werden. Solche Wandler dienen der Versorgung bipolarer Spannungsschienen, die von manchen Anlogschaltkreisen benötigt werden. Ein +5V auf $\pm 12V$ -Wandler könnte zum Beispiel für die Generierung der positiven und negativen Versorgungsspannungen von Operationsverstärkerschaltkreisen aus einer einfachen Standard-5V-Schiene dienen. Die Tatsache, dass der Ausgang unregelt ist und sich abhängig von der Last ändert, ist aufgrund des weiten Versorgungsspannungsbereichs vieler Operationsverstärker (z. B. arbeitet der klassische Operationsverstärker 741 mit einem Versorgungsbereich von $\pm 5V$ bis $\pm 18V$) unerheblich. Die galvanische Trennung des Wandlers stellt dabei sicher, dass das digitale Rauschen an der 5V-Schiene nicht auf die analogen Versorgungsspannungen übertragen wird.

Asymmetrische bipolare Spannungen können durch eine unterschiedliche Anzahl der Windungen beider Sekundärwicklungen erzeugt werden. Der Wandler RECOM RH-121509DH erzeugt zum Beispiel +15/-9 V aus 12 V-Nenningangsspannung mit 4 kVDC Isolationsspannung. Solche Wandler werden in IGBT-Anwendungen benötigt, da die Treiber-ICs typischerweise eine isolierte asymmetrische bipolare Versorgung verlangen; und da die IGBT-Treiber an der Hochspannungs-IGBT-Versorgung direkt angeschlossen sind, muss der DC/DC-Wandler einer dauerhaft anliegenden Hochspannung an der Isolierbarriere standhalten. Das nächste Blockschaltbild zeigt einen typischen Anwendungsfall für solche Wandler. Da die Spannungen an jedem IGBT verschieden sind, ist eine separate isolierte Versorgung für jeden einzelnen Treiber notwendig.

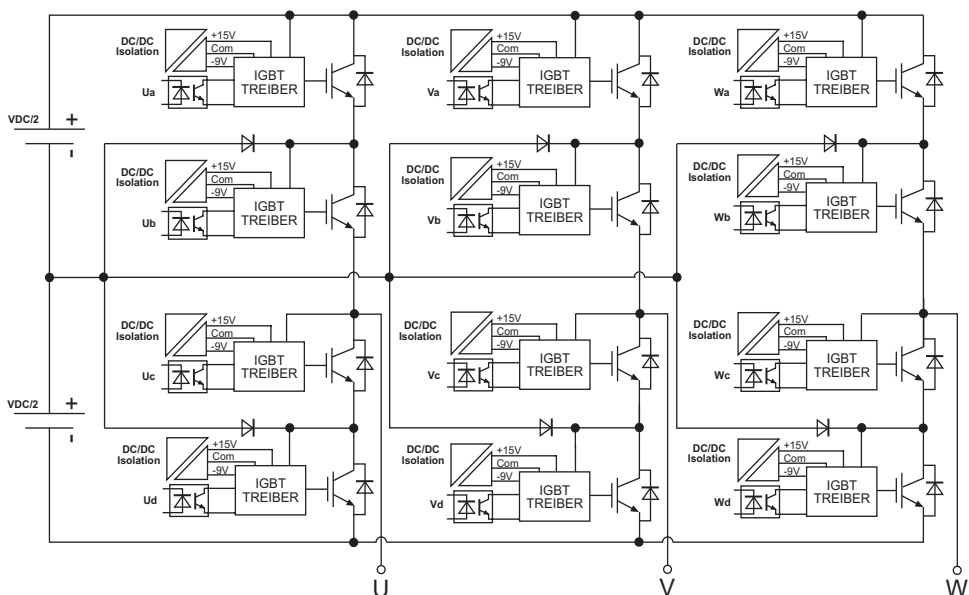


Abb. 1.32: Beispiel eines IGBT-3-Ebenen-Inverters, der 12 isolierte asymmetrische DC/DC-Wandler verwendet

1.2.3 Parasitäre Elemente und deren Effekte

Wie in diesem Kapitel oben erwähnt, setzen die vorhergehenden Beschreibungen von DC/DC-Wandlern ideale Komponenten voraus und ignorieren sämtliche parasitären Effekte. In der Realität haben Spulen auch kapazitive und resistive Anteile.

Die Wahl der im Wandler verwendeten Komponenten hat deshalb eine große Auswirkung auf dessen Arbeitscharakteristik. Kritische Komponenten wie z. B. Schalter und Gleichrichtungen, magnetische Komponenten sowie Filterkondensatoren beeinflussen sowohl die Schaltfrequenz als auch den Gesamtwirkungsgrad des Wandlers. In den vorausgehenden Kapiteln wurden Leistungsschalter, Gleichrichter-dioden, Transformatoren, Induktivitäten und Kapazitäten alle als ideale Komponenten betrachtet. Materielle Komponenten sind jedoch nicht ideal und haben parasitäre Eigenschaften, mit denen Entwickler von Wandlern vertraut sein müssen, um entsprechende Vorkehrungen treffen zu können.

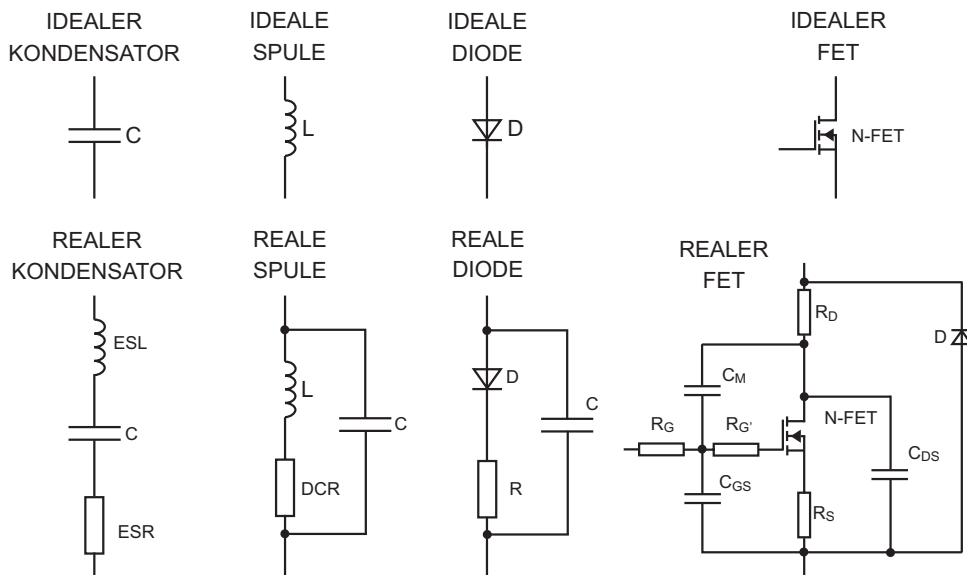


Abb. 1.33: Wandlerkomponenten mit typischen parasitären Elementen

Insbesondere Halbleiterschalter weisen viele nicht ideale Eigenschaften auf. FETs erfordern vom Treiberschaltkreis hohe Spitzenströme, insbesondere um die parasitäre Miller-Kapazität zwischen Gate and Drain aufzuladen und zu entladen. Dioden haben eine parallele äquivalente Kapazität, die ihre Schaltgeschwindigkeit sowie natürlich den internen Fluss-Spannungsabfall verzögert. Spulenverluste hängen sehr von der Wahl des Kernmaterials ab. Ihre DC-Widerstände (DCR) führen zu I^2R -Verlusten in den Wicklungen und haben interne Kupplungskapazitäten. Kondensatoren haben wesentliche parasitäre Effekte wie z. B. äquivalenten Serienwiderstand (equivalent series resistance, ESR) und äquivalente Serieninduktivität (equivalent series inductance, ESL). Alle diese Effekte sind frequenzabhängig, weshalb sich eine Spule bei hohen Frequenzen als Kondensator verhalten und sich ein Kondensator als Spule verhalten kann.

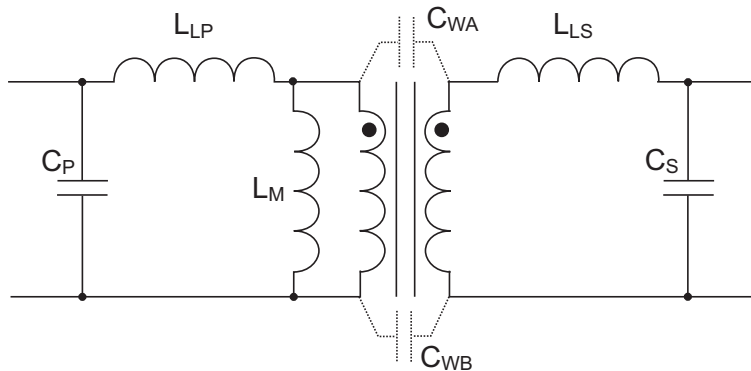


Abb. 1.34: Parasitäre Elemente des Transformators

Abb. 1.34 zeigt parasitäre Elemente, die bei einem Transformator auftreten können. C_{WA} und C_{WB} sind Koppelkapazitäten, C_s und C_p sind Primär- und Sekundärwicklungskapazitäten (üblicherweise unbedeutend, außer bei hochfrequenten Ausführungen), L_M ist die Magnetisierungsinduktivität des Kerns, und L_{LP} und L_{LS} sind Streuinduktivitäten. Diese parasitären Effekte haben großen Einfluss auf die Wandlerkonstruktion. Die Koppelkapazität verursacht Gleichtakt-EMC-Probleme, Sättigungseffekte infolge von L_M beschränken den Transformatorstrom und die Betriebstemperatur, und Streuinduktivitäten sind besonders problematisch, da sie den Wirkungsgrad reduzieren und EMI-Strahlung verursachen.

Streuinduktivitäten sind auch verantwortlich für die Spannungsspitzen, die jedes Mal auftreten, wenn sich der Strom in den Wicklungen rasch ändert. Solche Überspannungen belasten den Primärschalter und die Sekundärdioden. Daher müssen sie entweder so dimensioniert sein, dass sie dieser Spitzenspannung standhalten, oder mit parallelen Snubber-Schaltkreisen ausgestattet werden, um diese Energiespitzen abzufangen. Die Energie in den Spitzen und die Leistung, die der Snubber absorbieren muss, kann mit Hilfe folgenden Formeln berechnet werden:

$$E = \frac{1}{2} L_{LEAK} I_{LEAK}^2 \quad P = \frac{1}{2} L_{LEAK} I_{LEAK}^2 f$$

Gleichung 1.18: Energie- und Leistungsverluste in Schaltspitzen infolge von Streuinduktivität

Abb. 1.35 zeigt die Spannung an einem Schalt-FET in Flyback-Ausführung ohne Snubberschaltkreis zum Absorbieren der Schaltspitzen. Das Signal oben zeigt die Spannung am Schalt-FET. In diesem Beispiel wäre ein FET mit 600V Nennspannung erforderlich, auch wenn die Versorgungsspannung nur 160VDC beträgt.

Das andere Signal darunter zeigt den Strom durch die Ausgangsgleichrichterdiode. Es ist erkennbar, dass die Diode einer zusätzlichen Belastung ausgesetzt ist, da der Spitzenstrom infolge der parasitären Induktivität gegenüber der Spitze, hervorgerufen durch die Transformatoreninduktivität, um 50% höher ist. Die Diode wird sich aufgrund dieses parasitären Effektes also deutlich stärker erhitzen als im Idealfall.

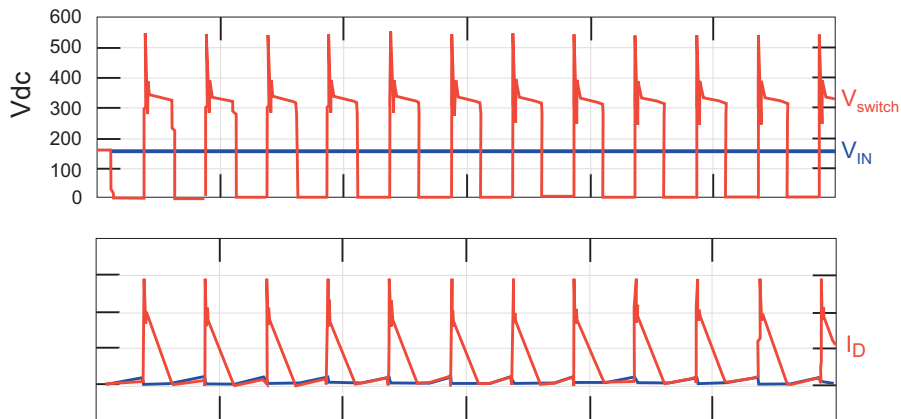


Abb. 1.35: Reale Schaltwellenformen in einem 160 VDC-zu-12 VDC-Flyback-Wandler. Oberer Leiter: Spannung am Schalter. Unterer Leiter: Gleichrichterdiodenstrom.

Durch den Einsatz von Snubber-Schaltkreisen wird ein Teil der Energie in den Spitzen absorbiert und dadurch die Überspannung an Schaltern und Dioden verringert. Ebenso werden die Selbsterwärmung und sowohl leitungsgebundene Störungen, als auch Strahlungsemissionen verringert. Der Snubber-Schaltkreis kann jedoch den Leistungsverlust, der durch die Spitzen verursacht wird, nicht beseitigen. Die Leistung, die sonst im Schalter oder der Gleichrichterdiode in Wärme umgesetzt würde, wird nun stattdessen von den Snubber-Schaltkreiswiderständen in Wärme umgesetzt. Da Widerstände jedoch passive Elemente sind und eine hohe Nennarbeitstemperatur haben, hat der Einsatz von Snubber-Schaltkreisen normalerweise ein positives Kosten-Nutzen-Verhältnis zur Folge. Abb. 1.36 zeigt die Verteilung der Snubber-Komponenten zum Absorbieren von Schaltspitzen im Flyback-Wandler.

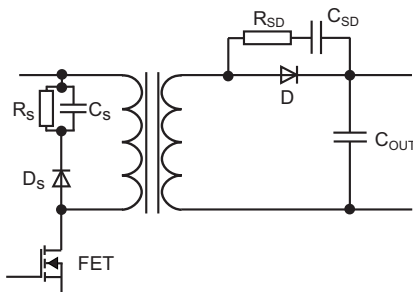


Abb. 1.36: Snubber-Komponenten im Flyback-Wandler

Außer den durch parasitäre Streuinduktivität verursachten Spitzen, zeigt jedes gekoppelte Reaktivsystem auch Resonanzfrequenzen. In den meisten transformatorbasierten Konstruktionen wird versucht, diese parasitären Elemente zu minimieren oder Taktfrequenzen auszuwählen, in denen Resonanz kein Problem darstellt.

Eine Quasiresonanz- oder Resonanz-Wandlerkonstruktion fördert die Resonanz jedoch bewusst durch Erhöhen der Wicklungsinduktivität oder den Einsatz zusätzlicher Induktivitäten. Lässt sich diese Resonanz steuern, kann eine sehr effiziente Wandlerkonstruktion realisiert werden.

1.2.3.1 Quasiresonanz-Wandler

Ein Quasiresonanz-(QR) Wandler kann mit jeder DC/DC-Topologie konstruiert werden. Da er aber üblicherweise mit einem Flyback-Schaltkreis verwendet wird, soll der Einfachheit halber nur der Flyback betrachtet werden.

Der Hauptunterschied liegt darin, dass das PWM-Timing beim QR-Wandler nicht allein von der Ausgangsspannung abhängt sondern ebenfalls vom Minimum des Schaltstromes. Der Flyback-Controller hat eine fixe PWM-Frequenz, die festlegt, wann der nächste Zyklus anfängt, der QR verwendet jedoch einen Freischwingoszillator.

Wie in der Standard-Flyback-Topologie schaltet der PWM-Controller in der QR-Topologie zum Speichern der Energie im Transformatorkern den Schalter auf EIN, und nachdem der Schalter auf AUS geht erfolgt die Energieübertragung zur Sekundärwicklung. In dem Moment, in dem der Strom in der Ausgangsgleichrichterdiode auf null gefallen ist, sind sowohl die Eingangs- als auch die Ausgangswicklung spannungsfrei. Jegliche Restenergie im Kern wird in die Primärwicklung zurückreflektiert und beginnt mit einer Frequenz, die abhängig ist von der Primärwicklungsinduktivität L_P und der Kapazität C_D (bestehend aus der Summe der Schalterkapazität, der Koppelkapazität zwischen den Wicklungen und weiteren parasitären Kapazitäten) zu schwingen.

$$f_{\text{RESONANCE}} = \frac{1}{2\pi \sqrt{L_P C_D}}$$

Gleichung 1.19: Resonanzfrequenz des Transformators im QR-Betrieb

Bei einer Induktivität der Primärwicklung von $500\mu\text{H}$ und einem C_D -Wert von 1nF beträgt die Resonanzfrequenz ca. 225kHz . Die Spannung am (offenen) Schalter ist die Versorgungsspannung der sich diese Resonanzschwingung überlagert. Wird nun der PWM-Zyklus genau in dem Moment rückgesetzt, in dem diese Spannung auf dem Minimum ist (engl.: valley switching), bedeutet das, dass die Wirkspannung am Schalter niedriger ist als die Versorgungsspannung. Dies führt dazu, dass der Schalter jetzt eine viel niedrigere Gesamtbelastung erfährt, aufgrund geringerer Einschaltspannung und niedrigerem Einschaltstrom, was einen messbaren Anstieg des Wirkungsgrades bewirkt.

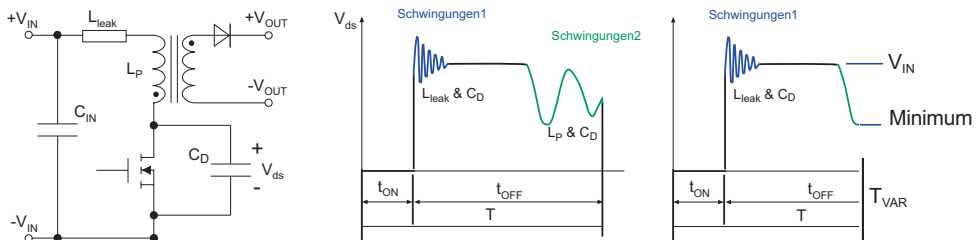


Abb. 1.37: Flyback-Topologie mit fixer PWM- und QR-Synchronisation

Ein anderer Vorteil des QR-Betriebs besteht darin, dass sich die PWM-Signal-Synchronisation mit jedem Zyklus je nach Genauigkeit des Minimum-Erkennungskreises leicht ändert. Dieser Jitter flacht das Spektrum der EMI ab und verringert die EMI-Spitzen. Eine Reduzierung um 10dB in den Pegeln der leitungsgebundenen Störung kann im Vergleich zu traditionellen Flyback-Schaltkreisen leicht erreicht werden. Ein Nachteil des QR-Betriebs liegt darin, dass die PWM-Frequenz lastabhängig ist und frequenzbeschränkende Maßnahmen oder eine Minimumabsicherung (sog. Valley-lockout circuits) notwendig sind, um den Wandler für den lastfreien Betrieb abzusichern.

1.2.3.2 Resonanzmoduswandler

Eine Weiterentwicklung des QR-Wandlers ist die Vollresonanzmoduswandlerkonstruktion. Der Resonanzmoduswandler (Resonant Mode, RM) kann mit Serienresonanz-, Parallelresonanz- oder Serienparallelresonanz-Topologien (auch bekannt als LCC-Topologie) realisiert werden. Da der Halbbrücken-LCC-Schaltkreis besondere Vorteile im Resonanzmodus bietet, wird der Einfachheit halber nur diese Topologie betrachtet.

Der Zweck des Resonanzmoduswandlers besteht darin, ausreichend zusätzliche Kapazität und Induktivität bereitzustellen, sodass der Resonator eine Null-Spannungsschaltung (ZVS = Zero Voltage Switching) ermöglicht. Die Vorteile der ZVS bestehen in außerordentlich niedrigen Verlusten.

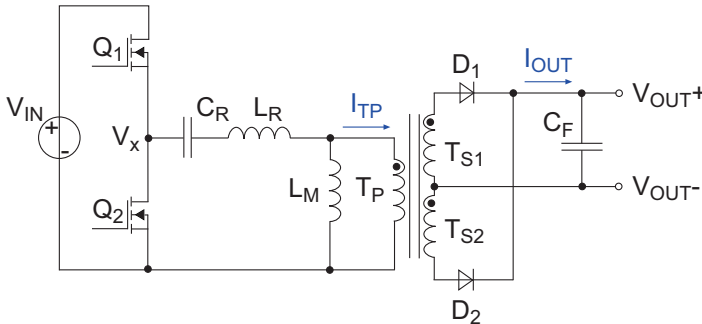


Abb. 1.38: Halbbrücken-LLC-Resonanzmodus

Diese Topologie hat zwei Resonanz-Frequenzen: Die erste ergibt sich aus dem aus C_R und L_R gebildete Serienresonanzkreis und die zweite aus dem Parallelresonanzkreis, gebildet von C_R und $L_M + L_R$. In der Regel werden L_M und L_R nebeneinanderliegend auf den Transformator gewickelt, um Streuinduktivitätseffekte zu verringern.

$$f_{\text{RESONANCE,SERIES}} = \frac{1}{2\pi\sqrt{L_R C_R}} \quad f_{\text{RESONANCE,PARALLEL}} = \frac{1}{2\pi\sqrt{(L_M + L_R) C_R}}$$

Gleichung 1.20: Doppelresonanzfrequenzen des LCC-Wandlers

Der Vorteil der Doppelresonanzen besteht darin, dass entsprechend der Last die eine oder die andere Priorität hat. Während der Serienresonanzkreis eine Frequenz aufweist, die mit fallender Last zunimmt, und der Parallelresonanzkreis eine Frequenz, die mit steigender Last zunimmt, hat ein gut konstruierter Serien-Parallelresonanzkreis eine stabile Frequenz über den gesamten Lastbereich. Die Schaltfrequenz und die Parameter L_R und C_R sind so gewählt, dass sich die Primärwicklung in Dauerresonanz befindet und eine fast perfekte sinusförmige Wellenform aufweist. Die zwei Halbbrückenschalter Q_1 und Q_2 werden in Gegenphase betrieben. Wenn die FETs aktiviert sind, ist die Spannung daran faktisch negativ. Die Gate-Drain-Spannung entspricht nur dem Spannungsabfall an der internen Diode, der Gate-Strom ist somit extrem niedrig. Wenn die Spannungen an den FETs positiv werden, sind diese schon angeschaltet und beginnen zu leiten, während die sinusförmige Spannung Null durchläuft.

In Kombination mit niedrigen Schaltverlusten und den aufgrund der sinusförmigen Erregungswellenform niedrigen Verlusten im Transformator sind Konversionswirkungsgrade über 95% erreichbar. Ein weiterer Vorteil besteht darin, dass EMI-Emissionen extrem niedrig sind, da der gesamte Leistungskreis sinusförmig ist.

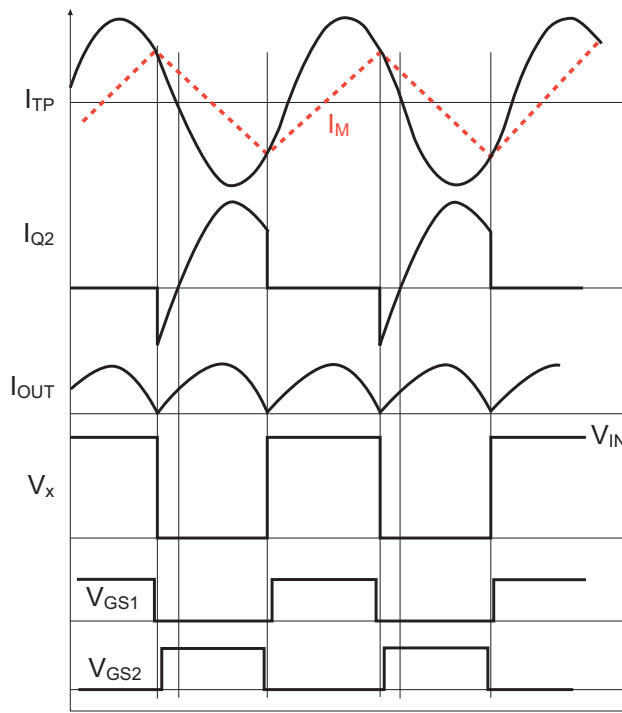


Abb. 1.39: Resonanzmodus-LCC-Charakteristik

Die Nachteile der LCC-Wandler-Topologie liegen darin, dass die erforderlichen Induktivitäten hoch sein können, um eine stabile Resonanzfrequenz mit gutem Q-Faktor zu erhalten (d. h. niedrige C_R). Der Wandler muss auch auf einen Betrieb unter dem maximal möglichen Verstärkungsfaktor abgestimmt sein, um einen Start ohne Probleme zu ermöglichen. Normalerweise stellt ein Arbeitsverstärkungsfaktor von 80 bis 90% des Maximums eine sichere Reserve dar.

Zusätzliche Maßnahmen zur Signalformung können für einen lastfreien Betrieb notwendig sein. Obwohl ein LLC-Lastbereich theoretisch Null-Last einschließt, können die Komponententoleranzen den Wandler ohne Last in der Praxis instabil machen. Schließlich erfordert die „nebeneinanderliegende“ Transformatorkonstruktion sorgfältige Ausführung, wenn Sicherheitsabstände (creepage- / clearance-distances) der Isolationsbarriere zu erfüllen sind.

1.2.4 Wirkungsgrad von DC/DC-Wandlern

Im Vergleich zu Linearreglern ist die Ermittlung des Wirkungsgrades von Sperrwandlern um einiges komplizierter. Der Linearspannungsregler hat leicht vorhersagbare DC-Verluste, wobei die größten Verluste im Vorwärts- bzw. Durchlasstransistor auftreten. Ein Sperrwandler hat jedoch nicht nur DC-, sondern auch AC-Verluste, die in den Schaltern und Komponenten für die Energiespeicherung auftreten. Zum Beispiel setzen sich die Gesamtverluste eines Schalters nicht nur aus den Verlusten in den EIN- und AUS-Zuständen, sondern auch aus den Verlusten im Übergang von EIN- in den AUS-Zustand und vice versa zusammen. Im Falle eines Transformators werden die Gesamtverluste aus der Summe von AC- (Kern), AC- (Wicklung) und DC- (Wicklung) Verlusten berechnet.

Die Verluste im Transformator Kern werden hauptsächlich durch die Wechselwirkung zwischen dem magnetischen Fluss und dem Kernmaterial (Hystereseverluste, Wirbelstromverluste) verursacht. Die Wicklungsverluste ergeben sich hauptsächlich aus dem Material der Transformatorwicklung (ohmsche Verluste, Skin-Effekte). Der Gesamteffekt ist eine Erhöhung der Temperatur im Transformator. Um den Wirkungsgrad des DC/DC-Wandlers zu berechnen, müssen die Verluste jedes Teils des Umwandlungszyklus durch Mitteln der Verluste über den ganzen PWM-Arbeitszyklus ermittelt werden.

Sperrwandler weisen hohe Wirkungsgrade auf, da der Schalter nur für kurze Zeit innerhalb des ganzen Schaltzyklus angeschaltet wird. Verluste in magnetischen, induktiven und kapazitiven Komponenten können durch sorgfältige Projektierung und durch geeignete Dimensionierung minimiert werden, sodass ein Umwandlungsfaktor von $> 96\%$ realisiert werden kann. Das bedeutet, dass nur 4% der aufgenommenen Leistung verloren gehen und in Wärme umgewandelt werden. Nicht isolierte Wandler sind in der Regel effektiver als ihre isolierten Pendanten, da weniger Komponenten an der Leistungsumsetzung beteiligt sind und keine Transformatorverluste auftreten. Ungeachtet eines höheren Integrationsgrades, können isolierte DC/DC-Wandler – je nach der Nennleistung – Wirkungsgrade von über 85% erreichen.

Eine der Hauptursachen für Verluste sind Ausgangsdioden. Wenn der Ausgangsstrom 1A und der Fluss-Spannungsabfall an der Diode 0,6 V beträgt, gehen allein in der Diode 600mW verloren. So verwenden DC/DC-Wandler für hohe Ausgangsströme häufig FETs mit synchroner Schaltung (Synchrongleichrichtung), um Gleichrichtungsverluste zu verringern.

Es mag überraschend erscheinen, dass Niedrigleistungswandler generell niedrigere Wirkungsgrade haben als Hochleistungswandler. Insbesondere in Anbetracht dessen, dass höhere I^2R -Verluste bei höheren Ausgangsströmen eintreten. Der interne Leistungsverbrauch der Schalt-Controller, Shuntregler und Optokoppler ("Systemverwaltungs"-Verbrauch) spielt jedoch eine wesentliche Rolle. Wenn der Gesamtsystemverwaltungsbedarf 1W beträgt, kann ein 10-W-Wandler keinen Wirkungsgrad von mehr als 90% haben, der maximal mögliche Wirkungsgrad eines 100-W-Wandlers würde jedoch 99% betragen. Die „Systemverwaltungs“-verluste erklären auch, warum alle DC/DC-Wandler in lastfreiem Betrieb einen Wirkungsgrad von 0% haben, während die Wandler immer noch Leistung verbrauchen, aber keine Ausgangsleistung erzeugen.

FETs verbrauchen mehr Leistung beim Schalten als im stabilen EIN- oder AUS-Zustand. Dies geschieht, weil die interne Gate-Kapazität auf- und entladen werden muss, um den Ausgang zu schalten. Gate-Impulsströme von 2A sind nicht untypisch. Auch in lastfreiem Betrieb schaltet der DC/DC-Wandler die FETs immer noch hunderttausend Mal pro Sekunde, weshalb ein DC/DC-Wandler auch ohne Last warm läuft.

1.2.5 PWM-Regelungsverfahren

Es gibt zwei Haupttypen der PWM-Regelung. Sie unterscheiden sich darin, wie die Rückführung in der Feedbackschleife ausgeführt wird oder was als Stellgröße dient. Eines der Regelungsverfahren ist die Spannungsregelung (Spannungsmodus), bei der das Impuls-Pausen-Verhältnis δ proportional zur Fehlerdifferenz zwischen der Ist- und der Vergleichsspannung ist. Bei der Stromregelung (Strommodus) ist das Impuls-Pausen-Verhältnis δ proportional zur Abweichung von einer Vergleichsspannung und einer mit einem Strom verbundenen Spannung, wobei der Strom in nicht isolierten Topologien der Strom durch den Leistungsschalter und in isolierten Wandlern der Strom in der Primärwicklung sein kann.

Der Konstantspannungsregler reagiert nur auf Änderungen der Ausgangsspannung und stellt das Impuls-Pausen-Verhältnis entsprechend ein. Da er den Laststrom oder die Eingangsspannung nicht direkt misst, muss er mit Sprüngen im Laststrom oder der Eingangsspannung auf die entsprechende Auswirkung in der Lastspannung warten. Diese Verzugszeit beeinflusst die Regelcharakteristik des Wandlers in der Form, dass für die Ausregelung immer ein oder mehrere Taktintervalle erforderlich sind. Um Instabilitäten oder Überspringen der Ausgangsspannung zu vermeiden, muss der Regelkreis kompensiert werden.

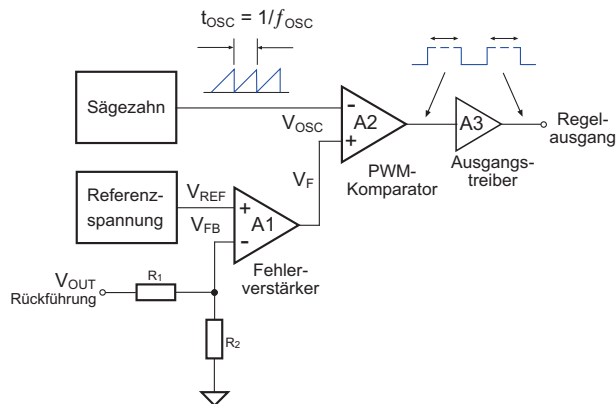


Abb. 1.40: Blockdiagramm eines Spannungsmodus-PWM-Kontrollers

Abb. 1.40 zeigt einen typischen Spannungsmodus-PWM-Controller. In diesem Schaltkreis ist A1 der Fehlerverstärker, A2 der PWM-Komparator, und A3 ist ein optionaler Ausgangstreiber, der als Schnittstelle zur Steuerung des Leistungsschalters verwendet wird. Der Sägezahngenerator erzeugt eine periodische Rampenspannung V_{OSC} , die linear von Null am Anfang des Taktes bis zu einer bestimmten Größe ansteigt, die dem maximalen Impuls-Pausen-Verhältnis am Ende entspricht. Der Abweichungsverstärker A1 misst die Differenz zwischen der höchstgenauen temperaturausgeglichenen Referenzspannung und einer abwärts geteilten Menge der Ausgangsspannung des DC/DC-Wandlers V_{FB} gleich $V_{OUT} R_2 / (R_1 + R_2)$.

Die Ausgangsspannung V_F des Abweichungsverstärkers A1 ist eine Größe proportional zur Differenz zwischen der Vergleichsspannung V_{REF} und der abwärts geteilten Ausgangsspannung V_{FB} . Wenn am Anfang jedes Taktes V_{FB} niedriger ist als V_{REF} , ist die Ausgangsspannung des Abweichungsverstärkers hoch. Wenn die Ausgangsspannung ansteigt, verringert sich V_F , bis sie die Anstiegsspannung V_{OSC} übersteigt, wonach A2 für den Rest des Zyklus unten bleibt.

Also existiert eine umgekehrte Abhängigkeit zwischen der Ausgangsspannung des DC/DC-Wandlers und dem Impuls-Pausen-Verhältnis. Negative Rückkopplung im Regelkreis ist ein stabiler Zustand.

Ein Spannungsmodus-PWM-Controller kann überschwingen, überkorrigieren und dann negativ überschwingen (unterschwingen), sodass die Ausgangsspannung ständig über und unter dem Sollpegel oszilliert. Deshalb wird die Rückführungsreaktion oft absichtlich verzögert, um dieses "Jagd"-Verhalten zu stoppen. Der Nachteil ist, dass der Wandler dann auf plötzliche Änderungen von Last oder Eingangsspannung langsamer reagiert.

Wenn zur Verringerung der Reaktionszeit für die Übergangscharakteristik (engl.: transient response) jedoch eine schnelle Reaktion der PWM-Regelung erforderlich ist, kann die alternative Stromregelung (Strommodusregelung) verwendet werden, um diesen Nachteil zu beseitigen.

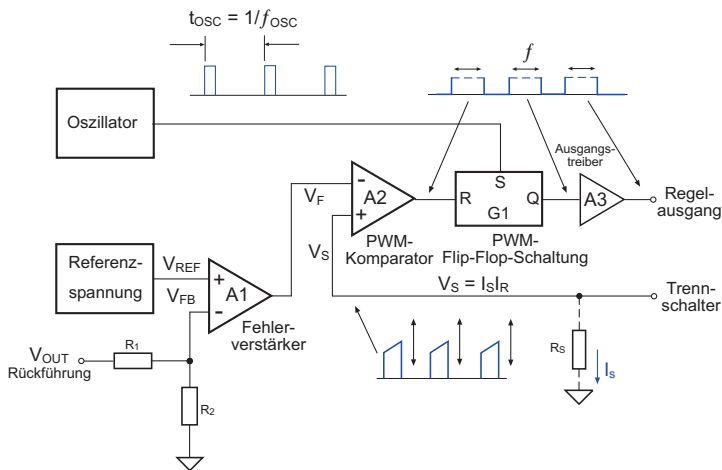


Abb. 1.41: Blockschaltbild eines Strommodus-PWM-Kontrollers

In einem stromgeregelten DC/DC-Wandler wird der Regelkreis auf zwei Rückkopplungsschleifen, den Innenregelkreis für Stromregelung und den Außenregelkreis für Spannungsregelung, aufgeteilt. Das Ergebnis ist, dass für jeden einzelnen Impuls nicht nur der Spannungswechsel am Ausgang, sondern auch der durch den Laststrom verursachte Wechsel ausgeglichen werden kann. Ein typischer stromgeregelter PWM-Controller ist in Abb. 1.41 gezeigt. Wie im vorherigen Schaltkreis, ist A1 der Fehlerverstärker, A2 der PWM-Komparator und A3 ein optionaler Ausgangstreiber, jedoch ist zusätzlich ein Flip-Flop G1 vorhanden. Ein Oszillator erzeugt Synchronisationsimpulse mit der Frequenz f , die normalerweise viel höher ist als die Frequenz f_{osc} . Am Beginn jedes Zyklus gibt dieser Impuls dem Flip-Flop ein Set-Signal. Wie im Spannungsmodus-PWM-Controller erzeugt der Fehlerverstärker A1 eine Ausgangsspannung, abhängig von der Differenz zwischen der Vergleichsspannung V_{REF} und der laut Verhältnis $R_2/(R_1 + R_2)$ geteilten Ausgangsspannung.

Am Beginn jedes Zyklus ist also das Flip-Flop gesetzt, und der PWM-Ausgang steigt an. Der Strom durch den Schalter ist AN , der Strom fließt also durch den Sensewiderstand oder Shunt R_S , wobei die Spannung $V_S = R_S I_S$ erzeugt wird.

Die Wandlerausgangsspannung beginnt anzusteigen, der Ausgangsstrom erhöht sich, bis die Sensespannung V_s dem Ausgang des Fehlerverstärkers A1 gleicht. Der PWM-Komparator schaltet ein, wobei er das Flip-Flop zurücksetzt und den PWM-Ausgang bis zum nächsten EIN-Impuls ausschaltet.

Die Strommodusregelung hat die gleiche inverse Abhängigkeit zwischen Ausgangsspannung und Impuls-Pausen-Verhältnis wie bei der Spannungsmodusregelung, ist also ein stabiles System. Sie besitzt jedoch den zusätzlichen Vorteil, dass der Außenregelkreis (Spannung) die Schwelle, bei der der Innenkreis den Primärspitzenstrom regelt, einstellt. Da der Eingangsstrom zum Ausgangsstrom proportional ist, bedeutet dies, wenn sich der Ausgangsstrom plötzlich ändert, ändert sich auch der Strom in der Primärwicklung, und die PWM beginnt innerhalb desselben Zyklus zu reagieren. So kann der spannungsbasierte Regelkreis zur Vermeidung des Einpendelns immer noch mit einer langsameren Reaktion realisiert werden, aber der Wandler reagiert immer noch fast sofort auf den Wechsel des Ausgangsstroms.

Ein Nachteil der Strommodusregelung ist der durch den zusätzlichen Sensewiderstand verursachte Verlust im Wirkungsgrad. Der Widerstand muss möglichst klein gehalten werden, um diesen Verlust zu minimieren, jedoch groß genug sein, um ausreichende Spannung zu entwickeln, sodass der PWM-Komparator reibungslos schalten kann. Der Strommodus-PWM-Komparator muss von höherer Qualität mit niedrigerer Eingangs-Offset-Drift und besserer thermischer Stabilität sein als ein Spannungsmodus-PWM-Komparator.

1.2.6 Regelung durch DC/DC-Wandler

1.2.6.1 Regelung von Mehrfachausgängen

Die meisten im Einsatz befindlichen DC/DC-Wandler sind vom unipolaren Typ, der nur eine einfache Ausgangsspannung versorgt. Wie zuvor gezeigt, ist die Regelung der Ausgangsspannung einfach und benötigt nur eine Rückkopplungsschaltung für den Fehlerverstärker.

In bipolaren DC/DC-Wandlern mit zwei symmetrischen Ausgangsspannungen von gegensinniger Polarität muss man einen Kompromiss eingehen, da nur eine Feedbackschleife realisiert werden kann. Wenn davon auszugehen ist, dass der positive Ausgang und negative Lasten ausgeglichen sind, kann das Regelproblem einfach durch Regelung der kombinierten Ausgangsspannung auf die Ebene des Unipolartransformators reduziert werden. Abb. 1.42 zeigt das Prinzip.

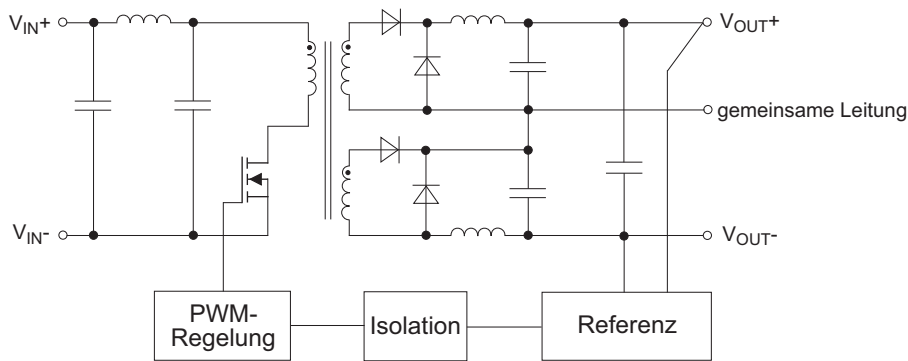


Abb. 1.42: Regelung eines bipolaren DC/DC-Wandlers

Ein Wandler mit $\pm 12\text{V}$ Ausgang regelt beispielsweise tatsächlich nur den kombinierten 24-V-Ausgang mit einer gemeinsamen, schwebenden Mittelanzapfung. Das bedeutet, dass sich, obwohl die Summe der beiden Ausgänge immer konstant bleibt, bei unsymmetrischer Last am DC/DC-Wandler aufgrund der unterschiedlichen Spannungsabfälle in jedem Zweig sich Ausgangsspannungen einstellen, die ihrerseits von den einzelnen Strombelastungen abhängen. Das kann zu unterschiedlichen $\pm V_{\text{OUT}}$ -Spannungen bezüglich der gemeinsamen Referenz führen. So könnte zum Beispiel ein am positiven Ausgang mit Volllast und 25% Last am negativen Ausgang belasteter $\pm 12\text{-V}$ -Ausgangswandler Ausgangsspannungen von $+13\text{V}$ und -11V bezüglich des gemeinsamen Common-Anschlusses haben, da nur die Summe beider Spannungen auf 24V ausgeregelt wird.

Der Anwender sollte deshalb prüfen, wie viel Asymmetrie oder Ungenauigkeit die zu versorgende Applikation vertragen kann. Die meisten bipolaren Anwendungen werden in Analogschaltkreisen verwendet, weshalb eine Schaltungsausführung mit guter PSSR (Power Supply Suppression Ratio) empfohlen wird. In manchen Anwendungen mit sehr asymmetrischen Lasten kann es notwendig sein, eine Dummy-Last oder Nachregelung zur Kompensation dieser Effekte einzusetzen.

Bei Zweifachausgangs-DC/DC-Wandlern, die mit zwei Ausgangsspannungen gleicher Polarität ausgestattet sind, funktioniert der einfache Trick der Regelung nach der Summe der Ausgänge nicht.

Eine Option besteht darin, nur einen Ausgang zu regeln (Hauptausgang) und den zweiten Ausgang (Hilfs- oder Auxiliaryausgang) ungeregelt auszuführen. In der Nähe der Nennbelastung kann das immer noch gut funktionieren, da beide Ausgänge bei wechselnder Eingangsspannung stabil bleiben. Abb. 1.43 zeigt das Schaltbild eines Zweifachausgangswandlers. Nur der Hauptausgang ist geregelt. Obwohl der Hilfsausgang ungeregelt ist, bleibt er noch immer proportional zur Hauptausgangsspannung, weil derselbe Primär-PWM-Controller gemeinsam genutzt wird.

Für manche Anwendungen, die eine niedrigere Spannungsgenauigkeit an der Hilfs- oder Auxiliaryspannungsversorgung als an der Hauptspannungsversorgung erfordern, stellt dies kein Problem dar. Ein Beispiel ist ein Hilfsstromversorgungsschaltkreis, der einen geregelten $+5\text{V}$ Hauptausgang für Logikbaugruppen und einen $+12\text{V}$ Hilfsausgang zur Versorgung eines Relais erfordert.

Gewisse Vorsicht ist allerdings bei der Entwicklung von Wandler-Kurzschlusschutz-Schaltkreisen geboten, um sicherzustellen, dass der Wandler sicher funktioniert, wenn der Hilfsausgang, nicht aber der Hauptausgang, kurzgeschlossen wird.

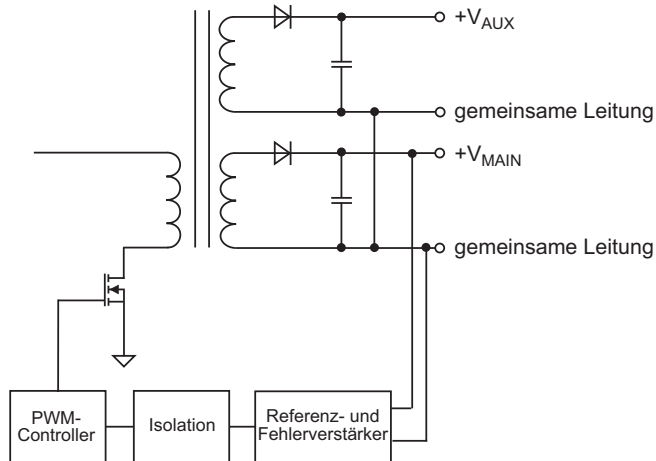


Abb. 1.43: Regelung mit zweifachen Ausgängen (Haupt- + Hilfsstromkreis)

Eine andere Option besteht darin, einen separaten Regelschaltkreis zum Nachregeln des Hilfsausgangs auf der Sekundärseite einzuschalten. Für Kleinleistungswandler, bei denen der Betriebswirkungsgrad nicht kritisch ist, besteht die einfachste Lösung darin, einen Linearspannungsregler zur Nachregelung des Hilfsausgangs einzusetzen und für einen gewissen Kurzschlusschutz zu sorgen (siehe Abb. 1.7). Das Transformatorübersetzungsverhältnis der Hilfswicklung muss sorgfältig gewählt werden, damit der Linearspannungsregler ausreichend „Spielraum“ hat, um korrekt über den vollen Lastbereich des Hauptausgangs zu regulieren. Für höhere Ausgangsströme oder Anwendungen, bei denen der Wirkungsgrad mehr Bedeutung hat als die Kosten, kann ein separater DC/DC-Wandler verwendet werden, um den Linearspannungsregler zu ersetzen.

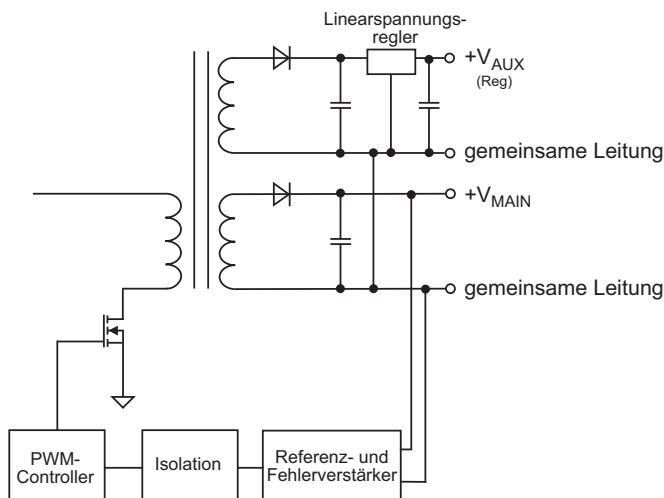


Abb. 1.44: Regelung mit zweifachen Ausgängen (Haupt- + nachgeregelter Hilfsstromkreis)

Die dritte Option besteht darin, die Ausgänge zu koppeln (engl.: stacked). Dies ist nützlich, wenn die Hilfsausgangsspannung in der Nähe der Hauptausgangsspannung liegt, z. B. $V_{AUX} = 5V$, $V_{MAIN} = 3,3V$. Abb. 1.45 zeigt einen Wandler mit gekoppelten Ausgängen. Der Vorteil dieser Anordnung besteht darin, dass der Strom, der in den Hilfsausgang fließt, vom Hauptausgang gemeinsam benutzt und so ebenfalls teilweise geregelt wird. Der Nachteil ist, dass die Wicklung und die Diode des Hauptkreises die Strombelastung beider Ausgänge erfahren.

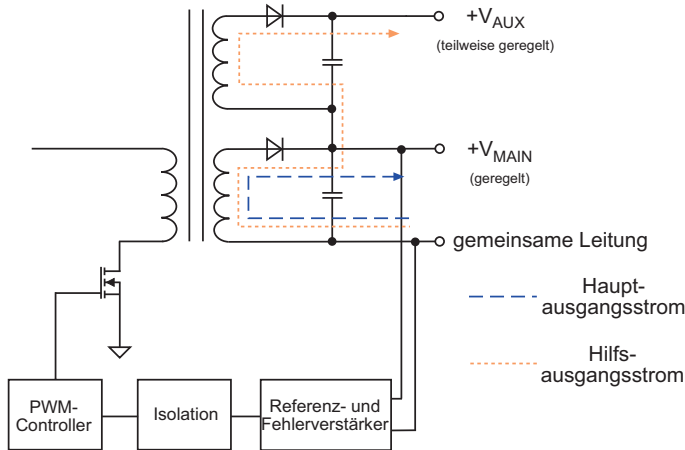


Abb. 1.45: Regelung mit zweifachen Ausgängen (gekoppelte Ausgänge)

1.2.6.2 Remote Current Sense Regulation

Gewöhnlich sorgt ein strom- oder spannungsbasierter Feedback-Kreis für eine ausgezeichnete Regelung der Ausgangsspannung. Bei hohen Ausgangsströmen kann sich die Performance der Regelung, aufgrund der Spannungsabfälle in Zuleitungen, Steckverbindern oder PCB-Leitungen, in der Praxis jedoch verschlechtern. Gerade bei Wandlern mit niedriger Ausgangsspannung kann sich dies dramatisch auswirken. Ein leitungsbedingter Spannungsabfall von z.B. 200mV kann bei einem Wandler mit 12V Ausgang durchaus noch zulässig sein, während dies bei einem 3,3V Ausgang weit außerhalb des Toleranzbereichs liegt.

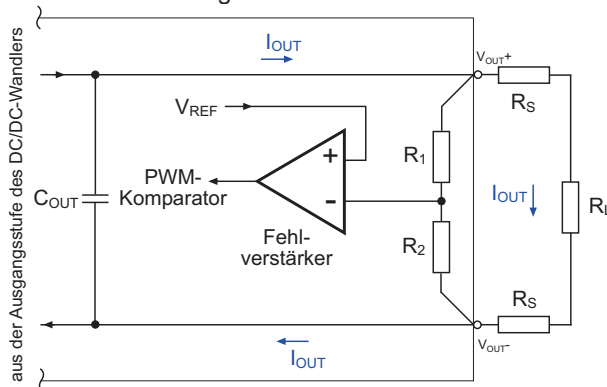


Abb. 1.46: Einfluss des parasitären Serienwiderstandes

Diese Spannungsabfälle können einen großen Einfluss auf die Lastregelung haben. Wie Abb. 1.46 zeigt, ist die Spannung an der Last um den Betrag $I_{OUT}^2 R_S$ gegenüber der Ausgangsspannung an den Anschlüssen des Wandlerausgangs verringert. Bei niedrigen Lasten ist der Wert R_S nicht so signifikant. Bei hohen Strömen steht der Spannungsabfall jedoch zum Serienwiderstand in quadratischem Verhältnis. So ist der Regelfehler quadratisch vom Ausgangsstrom abhängig. Deshalb reicht es nicht aus, nur die Ausgangsspannung anzupassen, um Verluste bei der Volllast auszugleichen, ohne bei niedrigen Lasten eine Überspannung an der Last zu riskieren.

Auch sind die Auswirkungen dieses Regelfehlers von mehreren anwendungsspezifischen Parametern abhängig, was eine Voraussage erschwert. Beispielsweise können sich die Kontaktwiderstände jeglicher Verbinder durch Oberflächenoxidierung, Verunreinigung, physischen Verschleiß oder thermischen Abbau, ändern. Folglich verändert sich mit der Zeit auch der Wert R_S . Darüber hinaus ist das Problem nicht bloß auf statische, ohmsche Verluste beschränkt. Besonders bei der hochdynamischen Strommodusregelung, kann eine mit einigem Abstand vom Wandler angeordnete dynamische Last, aufgrund reaktiver Komponenten der Zuleitung andere Lastregelungsfehler verursachen. Im ungünstigsten Fall könnten parasitäre Induktivitäten oder Kapazitäten Wandlerausgangsschwingungen und -instabilität hervorrufen, die dann zum Überspringen und zu Überspannungsschäden in der zu versorgenden Applikation führen können.

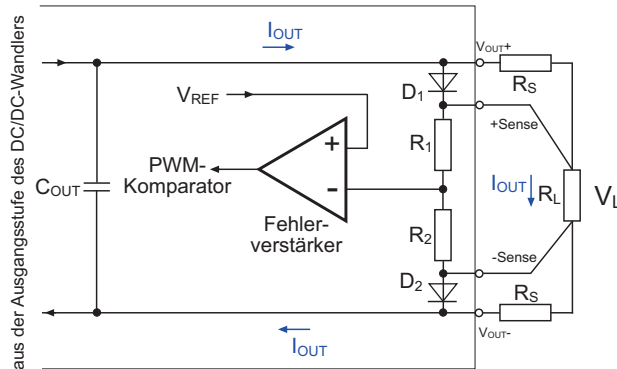


Abb. 1.47: Ableseausgänge zum Ausgleich von Ausgangsspannungsverlusten

Abb. 1.47 zeigt eine Lösung dieses Problems. Zwei zusätzliche Eingangspins werden zum Wandler hinzugefügt, um eine Kontrolle der Spannung am Punkt der Last zu ermöglichen. Diese Art des Anschlusses wird häufig als Kelvin-Brücke oder Sense-Eingang bezeichnet. Der Ausgang des Wandlers ist wie oben an der Last angeschlossen, zwei zusätzliche Anschlüsse wurden jedoch zwischen Feedbackkreis und Last gesetzt. Da der einzige Strom, der durch diese zusätzlichen Sense-Anschlüsse fließt, $V_L/(R_1 + R_2)$ beträgt und $R_1 + R_2$ normalerweise einige $k\Omega$ beträgt, ist der Sense-Strom sehr klein, und der Spannungsabfall entlang des Anschlusses ist ebenfalls entsprechend klein. So regelt der Wandler entsprechend der Istwertspannung, welche an der Last anliegt. Sämtliche lastabhängigen Effekte werden aufgrund des Serienwiderstands R_S ausgeglichen. Die niedrigen Sense-Ströme sind auch weniger empfänglich für dynamische parasitäre Effekte, weshalb die gesamte Lastregelung bei schnell veränderlichen Lasten stabiler ist. Die Dioden D_1 und D_2 ermöglichen es dem Wandler auch dann zu funktionieren, wenn die Sense-Pins nicht angeschlossen werden.

Der Aspekt der zulässigen Gesamtleistung darf jedoch nicht außer Acht gelassen werden. Dies trifft besonders auf Wandler mit niedrigen Ausgangsspannungen und hohen Ausgangsströmen zu. In solchen Fällen muss die Ist-Ausgangsspannung V_{OUT} für die Berechnung der Leistung verwendet werden und nicht die Nutzspannung V_L an der Last. Das folgende Verhältnis für V_L hält fest:

$$V_L = V_{OUT} - V_{S-} - V_{S+}$$

Gleichung 1.21: Verhältnis zwischen der Last- und Ausgangsspannung

Diese Gleichung hat weitreichende Konsequenzen. Man kann sofort erkennen, dass der Ausgleich der I^2R -Verluste seine Grenzen hat. Die Ausgangsspannung kann nicht beliebig vergrößert werden, da diese nicht den Ausgangs-Überspannungsschutz des Wandlers auslösen darf. Die Verlustleistung muss ebenfalls innerhalb der Grenzen des Wandlers bleiben, was die Ausgangsspannung bei Volllast ebenso einschränkt.

Wie in diesem Kapitel bereits besprochen, werden enorme Bemühungen von Herstellern von DC/DC-Wandlern unternommen, um den Wirkungsgrad so weit wie möglich in Richtung 100% zu erhöhen. Die Fähigkeit des DC/DC-Wandlers, die I^2R -Verluste auszugleichen, bedeutet jedoch, dass Systementwickler den Ausgangsanschlusswiderstand nicht mehr beachten müssen. Hier wird allerdings ein neuer Leistungsverlust im System geschaffen, der die mühsam erzielte hohe Effizienz durch die Hintertür „verdünnt“. Um diesen zusätzlichen Leistungsverlust, der durch den Spannungsabfall des Ausgangsanschlusses verursacht wird, zu berechnen, können wir Gleichung 1.22 nutzen:

$$P_{VD} = (R_{S+} + R_{S-}) I_{OUT}^2$$

Gleichung 1.22: Zusätzliche Leistungsverluste aufgrund der Sense-Rückführung

Wir können den RP60-4805S verwenden, um die Folgerungen dieser Gleichung zu veranschaulichen. V_{S+} und V_{S-} stellen die jeweilige Potentialdifferenz zwischen den Ausgangspins V_{OUT+} und V_{OUT-} (direkt an den Pins gemessen) und V_{L+} und V_{L-} (an den Lastklemmen gemessen) dar. Dieser Wandler liefert bis zu 12A bei 5V Ausgangsspannung (60W). Wenn die Last durch kupferne PCB-Leiter mit einer Länge von 10cm, einer Breite von 10mm und einer großzügigen Kupferstärke von 70 Mikron angeschlossen wird, würde der ohmsche Widerstand jedes Leiters $2,5m\Omega$ betragen, was einen Gesamtwiderstand des Ausgangsanschlusses von $5m\Omega$ ergibt. Der Leistungsverlust P_{VD} wäre dann:

$$P_{VD} = 0.005 \times 12 \times 12 = 0.72W$$

Dieser zusätzliche Leistungsverlust hat einen Effekt auf den Gesamtwirkungsgrad. Der RP60 hat einen Wirkungsgrad von 90%, was bedeutet, dass 6W bei Volllast intern gestreut werden. Der zusätzliche P_{VD} -Verlust von 0,72W stellt somit eine Erhöhung der Verluste des Gesamtsystems um 8,3% dar.

Einen Ausweg aus diesem Dilemma bietet das Point-of-load (POL)-Konzept. POL ist eine Strategie zum Minimieren der I^2R -Verluste. In diesem Fall wird der DC/DC-Wandler so nah wie möglich in die Nähe der Last positioniert, um Zuleitungen kurz und Verluste klein zu halten. Es werden mehrere Wandler verwendet, einer pro Last, statt eine zentralisierte Versorgung, die alle Lasten speist.

1.2.7 Beschränkungen im Hinblick auf den Eingangsspannungsbereich

Der Eingangsspannungsbereich (Input Voltage Range) eines DC/DC-Wandlers ist durch die Schaltkreistopologie sowie die verwendeten Komponenten bestimmt. Die Eingangsspannung und das Impuls-Pausen-Verhältnis des Wandlers sind umgekehrt proportional zueinander, weshalb eine Vergrößerung der Eingangsspannung eine Reduzierung des Impuls-Pausen-Verhältnisses hervorruft. Das minimale Impuls-Pausen-Verhältnis hängt wiederum vom maximalen Spitzenstrom des Leistungsschalters und dessen maximaler Nennspannung ab. Wenn das Impuls-Pausen-Verhältnis klein ist, ist der Spitzenstrom im Vergleich zum Durchschnittseingangsstrom hoch, und in den meisten Topologien tritt die höchste Schaltspannung ein, wenn der Strom dann schnell unterbrochen wird. Theoretisch kann der DC/DC-Wandler abwärts bis zum Impuls-Pausen-Verhältnis von 0% arbeiten; in der Praxis aber ist ein Wert von 5% bis 10% die Untergrenze aufgrund der Limitierungen durch die Anstiegsgeschwindigkeit der Ausgangsspannung (slew rate), Stabilität der Feedback-Kompensation und um alle negativen parasitären Effekte zu vermeiden. Dies begrenzt die maximale Eingangsspannung, des Wandlers.

Auch am oberen Ende des Impuls-Pausen-Verhältnissbereichs gibt es eine solche Beschränkung. Der Grenzwert des Impuls-Pausen-Verhältnisses ist durch die maximal zulässige Verlustleistung des Schalters und die Sättigung des Transformators oder andere Induktivitäten, sowie des Materials des Kerns beschränkt. Bei großen Werten des Impuls-Pausen-Verhältnisses ist der durchschnittliche Stromfluss am Maximum und die Verlustleistung im Schalter ist hoch. Um Sättigungseffekte zu vermeiden, brauchen Leistungsinduktivitäten und -transformatoren ebenfalls eine gewisse Abfallzeit, um das Magnetfeld im Kern wiederherzustellen. Theoretisch könnten alle DC/DC-Wandler bis zum 100%igen Impuls-Pausen-Verhältnis arbeiten, tatsächlich empfehlenswert ist jedoch eine Grenze von 85% bis 90%, die die Untergrenze der Eingangsspannung des Wandlers beschränkt.

Die Beschränkung des Impuls-Pausen-Verhältnissbereichs begrenzt somit den Eingangsspannungsbereich. Der DC/DC-Wandler-Eingangsbereich wird normalerweise durch das Verhältnis von Maximum-zu-Minimum der zulässigen Eingangsspannung beschrieben. Daher können isolierte Vorwärtswandler in der Regel innerhalb eines Eingangsspannungsbereichs von 2:1 oder 4:1 arbeiten.

Die Nenneingangsspannung orientiert sich an den Standardspannungen von Blei-Säure-Batterien von 12V, 24V und 48V, die in der Telekommunikationsindustrie, dem ersten großen Einsatzgebiet für DC/DC-Wandler, verwendet wurden.

2:1 Eingangsbereich

Nennspannung	Eingangsspannungsbereich
12V	9 - 18VDC
24V	18 - 36VDC
48V	36 - 72VDC

4:1 Eingangsbereich

Nennspannung	Eingangsspannungsbereich
24V	9 - 36VDC
48V	18 - 72VDC
11V	40 - 160VDC

Natürlich hängt die Wahl der Nennspannung vom Transformationsverhältnis und/oder der Wahl der Komponenten ab. Die oben genannten Bereiche haben sich in der heutigen Praxis jedoch etabliert. Einige Militärstandard-Wandler verwenden eine logischere Definition der 28V-Eingangsspannung, um den für den Einsatz mit Militär-Blei-Säure-Batterien mit zwei zusätzlichen Zellen für die 28V- statt 24V-Versorgung entwickelten 18V- bis 75V-Wandler zu beschreiben.

1.2.8 Synchrongleichrichtung

Wie oben erwähnt, liegt einer der Hauptgründe für den nach oben limitierten Wirkungsgrad eines Wandlers in der Verlustleistung, die in den Ausgangsdioden auftritt. Eine Leistungsdiode hat normalerweise einen Vorwärtsspannungsabfall von 500mV, was einem Leistungsverlust von 0,5W bei 1A entspricht. Spezielle Dioden – low-forward-voltage drop-Schottky-Dioden - können manchmal als Alternative für Kleinleistungswandler verwendet werden, sind aber relativ teurere Komponenten, wenn sie so dimensioniert sind, dass sie für hohe Ströme geeignet sind. Dennoch beträgt der Spannungsabfall an ihnen noch immer ca. 200mV, weshalb der Leistungsverlust immer noch signifikant sein kann. Ein Quantensprung in der Erhöhung des Wirkungsgrades war die Entwicklung der synchronen Gleichrichtung.

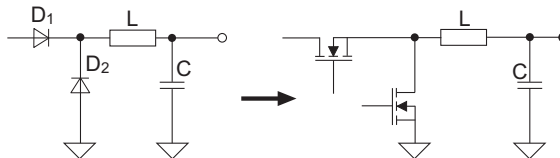


Abb. 1.48: Vergleich von Dioden- und Synchrongleichrichtung

Das Bild links zeigt ein typisches Schema mit Diodengleichrichtung. Darin fungiert Schaltkreis D_1 als Gleichrichter und D_2 als Freilaufdiode. Beide Dioden werden abwechselnd mit annäherungsweise gleichem Strom I_L geladen. Die Verluste infolge des Vorwärtsspannungsabfalls V_F in den Dioden sind gleich: $P_{VD} = V_F I_L$. Bei einer typischen Durchlassspannung V_F von 0,5V kann eine relative Verlustleistung von 0,5W pro A angenommen werden. Ein 3,3V/10A-Ausgangswandler hätte also einen Umsetzungsverlust von 15%, ohne andere Transformationsverluste in Betracht zu ziehen. Die in der Diode in Wärme umgewandelte Verlustleistung wäre 5W, sodass die Diode wahrscheinlich mit einem Kühlkörper montiert werden müsste, um einen halbwegs sinnvoll nutzbaren Arbeitstemperaturbereich zu erreichen.

Glücklicherweise können FETs als Stromgleichrichter verwendet werden, indem sie während des Vorwärtsteils des Zyklus eingeschaltet und während des Rückwärtsteils des Zyklus ausgeschaltet werden. Die Vorteile als schnellwirkende Schalter mit sehr niedrigem EIN-Widerstand $R_{DS,ON}$ macht sie als Gleichrichter ideal. Der Nachteil liegt darin, dass sie aktiv angesteuert werden müssen, sodass zusätzliche Synchronisierungs- und Treiberschaltungen erforderlich sind. Diese Schaltkreise müssen interne Spannungen auswerten, um die beiden FETs synchron mit der Ausgangswellenform an- und auszuschalten, daher der Name dieser Topologie. Im Vergleich dazu sind Dioden passive Bauteile, die für korrekten Betrieb keine weiteren Schaltungen benötigen. Jedoch macht der sehr niedrige $R_{DS,ON}$ von ca. 10mΩ den Nachteil komplizierter Schaltungen für Wandler mit hohem Ausgangsstrom mehr als wett.

In einigen Konstruktionen wird eine zusätzliche Sekundärwicklung verwendet, um ein sauberes Erregungssignal zu erzeugen.

Dioden haben normalerweise eine höhere maximal zulässige Durchbruchspannung als FETs. Weshalb man beim Redesign von existierenden Diodengleichrichtungsschaltungen zu Synchrongleichrichtung sicherstellen muss, dass die auftretenden Übergangsspannungen die V_{DS} -Grenzwerte der FETs nicht überschreiten.

1.2.9 Planartransformatoren

Planarmagnetika sind etwa seit den 80er Jahren bekannt. Da ihre Herstellung in der Vergangenheit jedoch sehr kostspielig war, gab es am Markt keine große Nachfrage nach diesem Typ, der nur für spezielle Anwendungsfälle zum Einsatz kam. Da mit der Zeit Mehrlagenleiterplatten kostengünstiger wurden und mittlerweile aufgrund ausgefeilterer Herstellmethoden weiter verbreitet sind, rückte diese Technologie wieder ins Licht der Öffentlichkeit.

Der Planartransformator nutzt mehrfache Kupferlagen als Wicklung für Transformatoren oder Induktivitäten. Um die erforderliche Anzahl von Windungen zu realisieren, werden zur Herstellung der Verbindung zwischen den Lagen verborgene Kontaktlöcher verwendet.

Es gibt eine praktische, sowie eine Kostengrenze für die Anzahl verwendbarer Lagen, weshalb Planarmagnetika hochfrequente PWM-Modulatoren und ebenso hochfrequente Treiberschaltkreise nutzen, um die Anzahl der Windungen gering zu halten und insgesamt weniger Lagen notwendig zu machen. Das durch die zum Einsatz kommenden höheren Frequenzen verursachte Hauptproblem besteht in dem sog. Skin-Effekt, der durch eine flache Konstruktion der Wicklung ausgeglichen werden kann. Bei Erhöhung der Frequenz bewegen sich die Ladungsträger immer mehr zur Oberfläche des Leiters, sodass dessen wirksamer Querschnitt reduziert wird, was die I^2R -Verluste erhöht. Dieser Effekt ist als Skin-Effekt bekannt.

Der Vorteil dieser Technologie besteht hauptsächlich in der außerordentlich flachen Konstruktion des Transformators, der aber dennoch in der Lage ist, große Energiemengen zu übertragen. Der DC/DC-Wandler kann also ein sehr niedriges Profil haben. Die übrigen Vorteile von Planarkonstruktionen, wie z. B. verbesserte Wärmeabfuhr in Wicklungen, reproduktionsfähige und mechanisch stabile Wicklungsstruktur, hochintegrierte Ausführung, niedrige Streuinduktivität und hohe Leistungsdichte, machen den Planartransformator zu einer notwendigen Komponente in der Technologie von Hochleistungs-DC/DC-Wandlern. Außerdem bedeuten Fortschritte in der Mehrlagenleiterplattenfertigung, dass die Epoxidharzisolierung zwischen den Wicklungen hohen Isolationsspannungen zwischen Primär- und Sekundärkreisen standhalten kann, was den technischen Anforderungen an eine Hauptisolierung von 2250VDC entspricht.

Für den Benutzer hat diese Technologie offenbar keine Nachteile, für den Hersteller sind Konstruktion und Fertigung aber ganz und gar nicht einfach zu bewerkstelligen.

Die mehrlagige Konstruktion besitzt eine hohe Koppelkapazität, die die hochfrequente PWM-Regelung komplizierter macht, die Kopplung zwischen Planartransformator und die damit verbundenen traditionellen Schaltkreise müssen sorgfältig gesteuert werden, um Abschlussverluste zu vermeiden, die unmittelbare Nähe von Kernzwischenraum und Lagenwicklung kann wesentliche Wirbelstromverluste hervorrufen, und infolge der unterschiedlichen Windungsanzahlverhältnisse muss ein spezieller Schaltkreis für jede Eingang-Ausgang-Kombination entwickelt und geprüft werden.

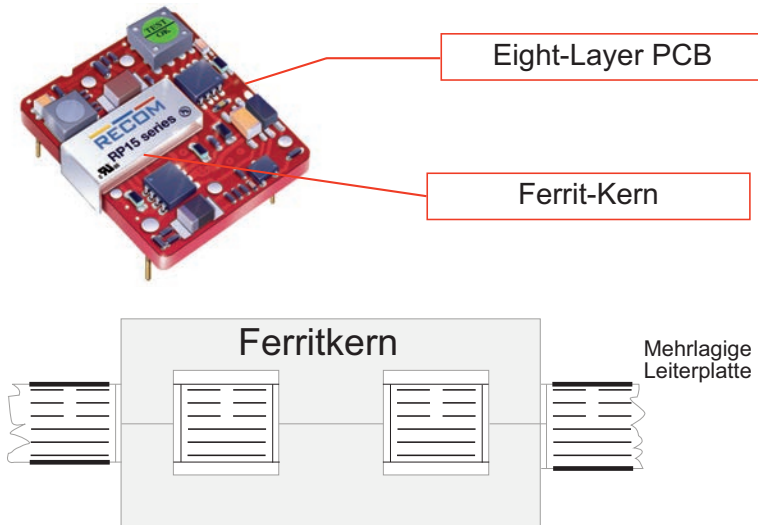


Abb. 1.49: Konstruktion eines transformatorbasierten Planarwandlers

Die Wicklungen in einem Planartransformator entsprechen den PCB-Leitern auf einer mehrlagigen Leiterplatte. Primär- und Sekundärwicklungen folgen typischerweise aufeinander, um die Streuinduktivität zu verringern. Eine Schwierigkeit bei Nutzung der PCB-Leiter ist der elektrische Anschluss am Ende der dem Kern am nächsten liegenden Wicklung. Eine verbreitete Lösung besteht darin, ein versenktes Kontaktloch – wie in Abb. 1.50 unten gezeigt – zu verwenden. Das Diagramm zeigt eine Wicklung mit sechs Windungen, die durch einen Oberleiter mit drei Windungen und einen Unterleiter mit drei Windungen gebildet wird, die durch ein Kontaktloch verbunden sind. Die zwei Enden der Wicklung können auf diese Weise bequem außerhalb des Kerns platziert werden.

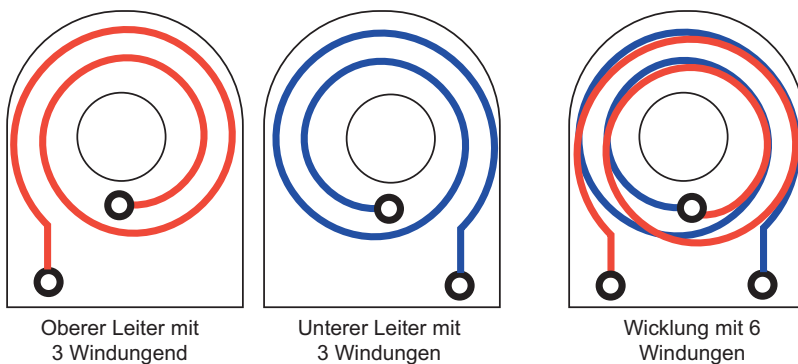


Abb. 1.50: Aufeinanderfolgende mehrlagige Planarleiter

1.2.10 Gehäusetypen von DC/DC-Wandlern

Der folgende Abschnitt beinhaltet eine Übersicht über einige standardisierte Gehäusetypen von DC/DC-Wandlern. Weitere Informationen zur Pin-Anordnung, zu Pin-Durchmessern und -Toleranzen bieten die Datenblätter des Herstellers.

Ein Kennwert, der normalerweise für den Gehäusotyp verwendet wird, ist die Art und Weise, wie die Anschlusspins angeordnet sind, somit besitzt ein SIL- (Single In Line) oder SIP- (Single In-line Pins) Gehäuse eine einzige Reihe von Pins und ein DIL- (Dual In Line) oder DIP- (Dual In-line Pins) Gehäuse zwei Reihen Kontakte. Das SIP-Gehäuse wird allgemein für isolierte Kleinleistungswandler ($< 2W$) und nicht isolierte DC/DC-Wandlern verwendet, während das DIP-Gehäuse häufiger für Kleinleistungswandler $< 10W$ verwendet wird. Der Hauptgrund für den Unterschied von Gehäusetypen ist das Gewicht des Wandlers. Ein Wandler für besonders kleine Leistungen $< 2W$ wiegt nur ca. 3g, weshalb durch Schwingungen verursachte Pinbrüche selten sind. Kleinleistungswandler $> 2W$ sind etwas schwerer. Sie wiegen ca. 30g und benötigen zwei Reihen von Pins, um die mechanische Festigkeit zu gewährleisten.

Hochleistungswandler ($> 10W$) können signifikante Ausgangsströme mit niedrigen Ausgangsspannungen haben und werden normalerweise mit Metallgehäusen ausgestattet, um eine zusätzliche Wärmeabfuhr zu ermöglichen. Daher benötigen sie dickere Anschlusspins, um höheren Strom leiten und zusätzliches Gewicht tragen zu können. Die Kennwerte, die üblicherweise für höhere Leistungswandler verwendet werden, sind die Gehäusegrößen in Zoll. Verbreitete Größen sind z. B. $1'' \times 1''$, $2'' \times 1''$, $1,6'' \times 2''$ und $2'' \times 2''$.

In den USA sowie in der Telekommunikationsbranche ist eine Gehäusegröße verbreitet, die auf dem "Ziegel"format basiert. Ein „Ziegel“ (engl. brick) entspricht den Maßen $2,4'' \times 6'' \times 0,5''$ bzw. $61mm \times 116,5mm \times 12mm$ und wird u. a. als Bezugsmaß für DC/DC-Wandler verwendet. Verschiedene Modelle im Ziegelformat werden in binären Brüchen von Ziegeln, wie z. B. Halb- (half-brick), Viertel- (quarter-brick), Achtel- (eighth-brick) oder Sechszehntel“ziegel“ (sixteenth-brick), angegeben.

Um eine Austauschbarkeit von Wandlern zwischen Herstellern zu ermöglichen, wurden verschiedene Versuche unternommen, die Wandlergrößen und Pin-outs zu standardisieren, insbesondere durch die DOSA (Distributed-power Open Standards Alliance) und die Konkurrenz POLA (Point Of Load Alliance), die sich beide auf digitale Point-Of-Load- und ratiometrische Wandler konzentrieren. Doch die meisten Hersteller bieten einfach ein Sortiment von Standard-Pinout-Optionen für jede der Wandlerserien an, um die Produkte ihrer Konkurrenz abzudecken und einen 1:1-Ersatz für Sublieferanten zu ermöglichen.

Da jedoch kein vereinbarter Standard für die Nummerierung der Pins existiert (manche beginnen in der oberen linken Ecke von oben gesehen, manche wählen den PCB-Standpunkt, manche schließen optionale Pins in die Nummerierungsfolge ein, die anderen machen es nicht), ist oft ein sorgfältiger Vergleich zwischen Datenblättern erforderlich, um Pin-Kompatibilität zu gewährleisten.

Im Folgenden stellen wir eine kleine Auswahl von Gehäusetypen aus dem Portfolio von RECOM vor, um einen Eindruck der Vielfalt verfügbarer Ausführungsoptionen zu geben.



SIP3



SIP3 (B case)



**gehäuselose SMD
Ausführung**



SIP4 (Micro)



SIP4



**gehäuselose
Ausführung mit Pins**



SIP7



SIP8



1" x 1" Metallgehäuse



DIP6



SMD



2" x 1" Metallgehäuse



DIP24



DIP24 SMD



DIP24 Metallgehäuse

2. Feedbackschleifen

2.1 Einleitung

Einige der wichtigsten Entwicklungskriterien in der Konstruktion von DC/DC-Wandlern sind Berechnungen und Methodiken, die zur Auslegung der Rückkopplungsschleifen der Regelung eingesetzt werden. Werden Schleifenparameter nicht korrekt berechnet, kann der Wandler Instabilität und Regelungsfehler aufweisen.

Die Funktion einer Rückkopplungsschleife in einem DC/DC-Wandler besteht darin, den Ausgang bei einer fixen Größe, die nur vom Referenzwert abhängt, zu halten, d. h. sie ist unabhängig von Last-, Eingangsspannungs- oder Umweltabweichungen. Das klingt einfach und ist unter statischen oder sich langsam ändernden Bedingungen tatsächlich auch relativ einfach zu erfüllen. Um jedoch dynamische oder stufenweise Änderungen zu bewältigen, wird die Rückkopplungsschleifenkonstruktion sehr kompliziert. Einer der wichtigsten Kompromisse, der gemacht werden muss, ist die Bilanz zwischen einem ausbalancierten Ausgang während des stationären Betriebs (niedrige periodische Ausregelungsabweichungen, kleine Totzone und hohe Regelgenauigkeit) und adäquater Reaktion auf dynamische Betriebsbedingungen (schnelle Reaktion auf Änderungen, kurze Ausregelzeit und niedriges Überspringen). Außerdem muss der Regelkreis unter allen Betriebsbedingungen, einschließlich des Niedriglast- oder sogar des lastfreien Betriebs, stabil sein. Die Rückkopplungsschleifenkonstruktion ist deshalb einer der Hauptfaktoren, die die Gesamtarbeitscharakteristik des Wandlers bestimmen.

2.2 Konstruktion mit offener Schleife

Nicht alle DC/DC-Wandler verwenden eine Rückkopplung. Der ungeregelte Push-Pull-Wandler (Royer Topology) aus dem Beispiel in Abb. 1.30 hat keinen Rückführungsschaltkreis. Der freilaufende Oszillator läuft bei einer Frequenz, die nur durch physikalische Eigenschaften des Transformators und die Eingangsspannung bestimmt ist, gemäß folgendem Verhältnis:

$$V_{IN} = 4 N_P B A_E f$$

Gleichung 2.1: Transformatorgleichung

wobei N_P die Anzahl der Primärwindungen ist, B ist der magnetische Fluss und A_E der Transformatorquerschnitt. Die Formel kann wie folgt zur Frequenz f umgeformt werden:

$$f = \frac{V_{IN}}{4 N_P B A_E}$$

Gleichung 2.2: Umgeformte Transformatorgleichung

Der Faktor „4“ unterscheidet sich von der Standard-Transformatorgleichung, die „4,44“ verwendet, weil der Royer-Oszillator eine Rechteckwelle und kein sinusförmiges Signal erzeugt. Die Ausgangsspannung hängt vom Windungsverhältnis der Windungsanzahl auf der Primärwicklung N_P zur Windungsanzahl auf der Sekundärwicklung N_S ab:

$$\frac{N_P}{N_S} = \frac{V_{IN}}{V_{OUT}}$$

Gleichung 2.3: Transformationsverhältnis

Aus diesen Abhängigkeiten ist ersichtlich, dass weder die Ausgangsspannung noch die Arbeitsfrequenz fixiert ist und dass beide von der Eingangsspannung abhängen. Deshalb sollen unregelmäßige DC/DC-Wandler idealerweise nur mit geregelten Eingangsspannungen verwendet werden.

In der Praxis gibt es „verborgene“ Rückkopplungs-Mechanismen, die die Arbeitsleistung des Royer-Oszillators gegenüber der theoretisch vorausgesagten verbessern. Alle Primär-, Sekundär- und Rückkopplungswicklungen weisen infolge von Streuinduktivität und Kopplungskapazitäten eine Wechselwirkung miteinander auf. Wicklungen können auch am Kern angeordnet werden, um diese Wechselwirkungen zu erhöhen oder zu verringern oder sogar um eine Wicklung von der anderen abzuschirmen. Unregelmäßige Wandler können z. B. durch Wickeln der Sekundärwicklung zwischen die Primär- und die Rückkopplungswicklungen kurzschlussicher realisiert werden, sodass kurzgeschlossene Ausgangswindungen eine Art Schirm gegen die Kopplung der Primär- zur Sekundärwicklung bilden. Der Wandler setzt seine Schwingungen fort, wenn der Ausgang kurzgeschlossen wird, jedoch erfolgt dies unter einer stark reduzierten Leistung, welcher die Schalttransistoren standhalten können. Die Temperatur des unregelmäßigen Wandlers erhöht sich beim Vollkurzschluss zwar, er bleibt jedoch funktionstüchtig. Ist der Kurzschluss beseitigt, setzt der Wandler seinen Normalbetrieb mit voller Leistung fort.

2.3 Geschlossene Schleifen

Die Abhängigkeit der Ausgangs- von der Eingangsspannung kann durch die Verwendung einer Rückkopplungsschleife eliminiert werden. Wichtig ist, dass ein Rückführkreis an einem Fehlerverstärker eingespeist wird, der den Ist-Ausgang mit dem gewünschten Wert vergleicht und zum Abgleich entsprechend korrigiert. Da die Korrektur immer in Gegenrichtung zum Fehler erfolgt (ist der Ausgang zu hoch, verringert sie ihn; ist der Ausgang zu niedrig, vergrößert sie ihn), wird die Rückführung als „negativ“ bezeichnet. Wenn die Rückführungsschleife „positiv“ ist, werden sämtliche Fehler verstärkt, und der Ausgang oszilliert oder stellt sich schnell auf den kleinsten oder größten Pegel ein. Eine der größten Herausforderungen bei der Entwicklung von Regelkreisen besteht darin, sicherzustellen, dass im transienten Betrieb keine positiven Rückkopplungsbedingungen auftreten können.

Der Vorzug der Rückkopplung besteht darin, dass nicht nur der Eingangsspannungswechsel, sondern auch alle Änderungen in der Ausgangsspannung infolge eines Lastwechsels ausgeglichen werden. Eine einzige Rückkopplungsschleife korrigiert beide Situationen. Der andere Vorteil von geschlossenen Rückkopplungsschleifen besteht darin, dass Eingang und Ausgang nicht dieselben Einheiten haben müssen. Eine Rückführungsschleife kann verwendet werden, um aus einer variablen Eingangsspannungsversorgung einen konstanten Stromausgang zu erzielen. Der Fehlerverstärker korrigiert den Ausgang in Abhängigkeit des Stromes und nicht des Spannungspegels (damit arbeitet er als Transkonduktanzverstärker).

Zur Veranschaulichung eines Feedback-Designs verwenden wir einen einfachen, nicht isolierten Buck-Wandler. Ein typisches Schaltbild könnte wie folgt aussehen:

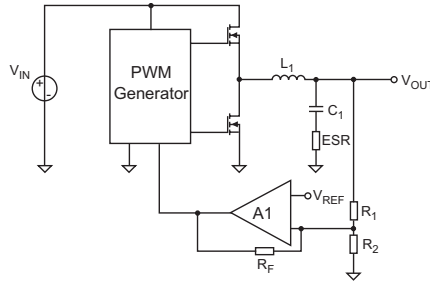


Abb. 2.1: Vereinfachtes Schaltbild des Buck-Wandlers

Was die Funktionsbaugruppen betrifft, kann Abb. 2.1 wie folgt reduziert werden zu:

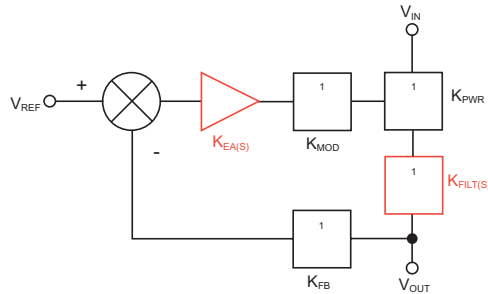


Abb. 2.2: Blockdiagramm der Feedbackschleife

Jede Funktionsbaugruppe besitzt einen Verstärkungsfaktor K . Die Leistungsschaltbauelemente (power switching elements, FETs) besitzen einen Verstärkungsfaktor von K_{PWR} , einen Ausgangsfilter gebildet aus L_1 und C_1 mit $K_{FILT(S)}$, und das Feedback-Element (Widerstandsteiler, gebildet aus R_1 und R_2) hat einen Verstärkungsfaktor K_{FB} . Das Gesamttrückmeldungssignal wird im Summierungspunkt mit der Vergleichsspannung V_{REF} verglichen und der Fehler, durch den die Abweichung ermittelnden Fehlerverstärker A_1 mit dem Verstärkungsfaktor $K_{EA(S)}$ verstärkt, um den PWM-Modulator, der den Verstärkungsfaktor K_{MOD} besitzt, zu regeln. Manche dieser Verstärkerbaugruppen haben eine hohe Verstärkung, und manche schwächen das Signal ab, der Gesamtverstärkungsfaktor des offenen Kreises (Summe aller Verstärkungen) ist jedoch positiv und beträgt üblicherweise etwa 1.000.

$$G_{OL} = K_{PWR} + K_{FILT(S)} + K_{FB} + K_{EA(S)} + K_{MOD}$$

Gleichung 2.4: Verstärkungsfaktor der offenen Schleife

Der einfache Schaltkreis (Abb. 2.1) besitzt durch den LC-Ausgangsfilter verursachte Resonanzen (Pole) bei folgender Frequenz:

$$f_{ZO} = \frac{1}{2\pi (ESR) C_1}$$

Gleichung 2.5: Winkelfrequenz des LC-Filters

$$f_{PO} = \frac{1}{2\pi \sqrt{L_1 C_1}}$$

Gleichung 2.6: ESR-C-Winkelfrequenz

Bei Frequenzen über f_{PO} verringert sich der Verstärkungsfaktor mit einer Geschwindigkeit von -40dB/Dekade aufgrund der LC-Charakteristik zweiter Ordnung des Ausgangsfilters. Der Punkt, in dem er die 0dB-Linie erreicht (Verstärkungsfaktor = 1 = 0dB), ist die Transitfrequenz (engl. crossover frequency) f_C . Bei der Frequenz f_{ZO} ändert der Effekt des RC-Filters erster Ordnung infolge der ESR des Filterkondensators die Steilheit der Verstärkungskurve zu -20 dB/Dekade. Das Diagramm t „Verstärkungsfaktor vs. Frequenz“ zeigt, dass sich Steilheit und Phase mit der Frequenz ändern:

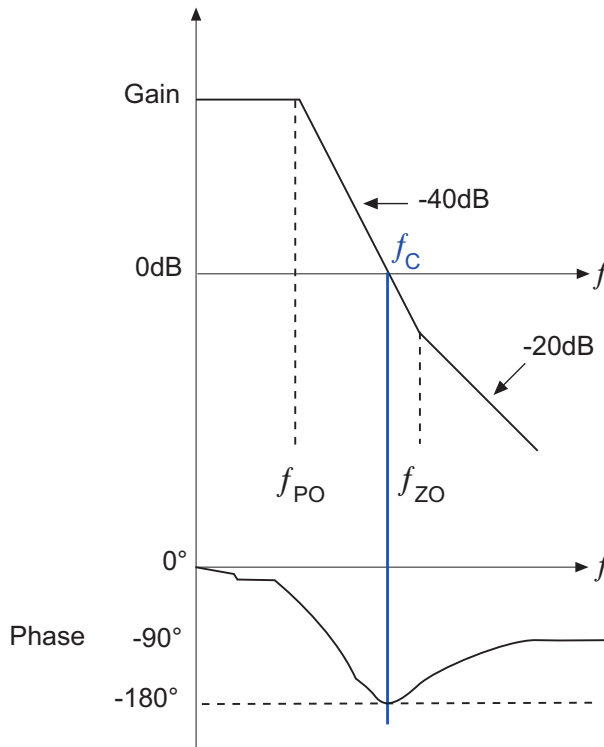


Abb. 2.3: Normalisiertes Amplituden- und Phasendiagramm der Abb. 2.1

Das Phasendiagramm ist die Phasenveränderung zusätzlich zu den durch den invertierten Eingang des Fehlerverstärkers A1 verursachten 180°.

Wie sich anhand des Phasendiagramms erkennen lässt, ist der Schaltkreis bei der Überschneidungsfrequenz instabil, da die Phasendrehung insgesamt -180° oder -360° beträgt. Dies bewirkt, dass der Wandler in den positiven Rückkopplungsbereich wechselt und der Ausgang zu schwingen beginnt.

Durch Erhöhung des Verstärkungsfaktors in der Fehlerverstärkerstufe kann die Frequenz, bei der der Gesamtverstärkungsfaktor gleich Null ist, in einen sichereren Bereich verschoben werden. Die Phasenreserve (die Differenz zwischen der vollen Phase und -180° bei der System- f_C) und die Amplitudenreserve (Verstärkung des Systems bei -180° Phase) bestimmen, wie stabil die Feedbackschleife ist (Abb. 2.4).

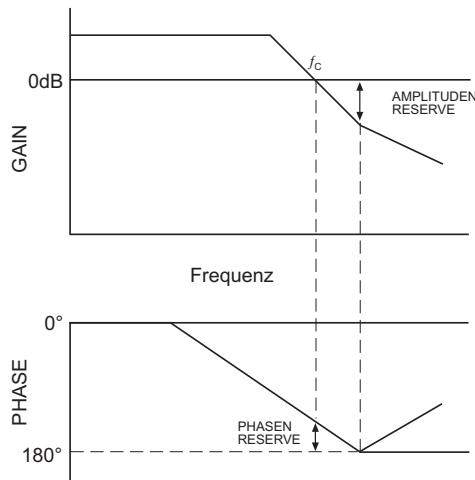


Abb. 2.4: Amplituden- und Phasenreserven

2.4 Kompensation von Rückkopplungsschleifen

Je weiter die ausgewählte Systemtransitfrequenz (engl.: system crossover frequency) von der Transitfrequenz der Leistungsstufe (engl: powerstage crossover frequency) entfernt ist, desto stabiler ist der Ausgang (er weist bessere Amplituden- und Phasenreserven auf), desto langsamer ist jedoch das Übergangsverhalten. Eine Phasenreserve von ca. 45° gewährleistet eine gute Reaktion mit geringem Überspringen, aber ohne Oszillation.

Nun gibt es 2 Möglichkeiten um einen sicheren Betrieb zu gewährleisten: Man kann einerseits den Verstärkungsfaktor des Fehlerverstärkers für alle relevanten Frequenzen ausreichend hoch wählen – die System-Transitfrequenz liegt dann in einem sicheren Bereich. Eine andere Möglichkeit ist eine frequenzabhängige Phasenverschiebung des Fehlerverstärkers anzustreben, was mittels Kompensation im Rückkopplungszweig des Fehlerverstärkers, welcher meist ein Operationsverstärker ist, erreicht wird.

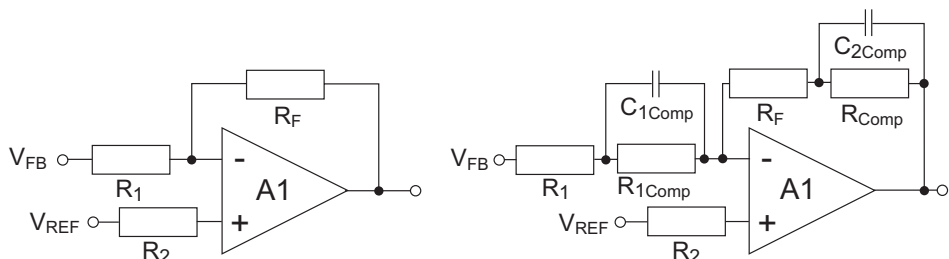


Abb. 2.5: unkompensierter (links) und kompensierter Fehlerverstärker (rechts)

Die Kompensationskomponenten können so gewählt werden, dass sich die Phase umkehrt und bei der kritischen Transitfrequenz zur Phasenreserve addiert, um so die Stabilität zu erhöhen. Dabei muss der Ausgangsfilter nicht so stark gedämpft sein, wodurch die Reaktion des DC/DC-Wandlers auf die Transiente beschleunigt wird, ohne übermäßiges Überspringen oder Oszillation zu riskieren.

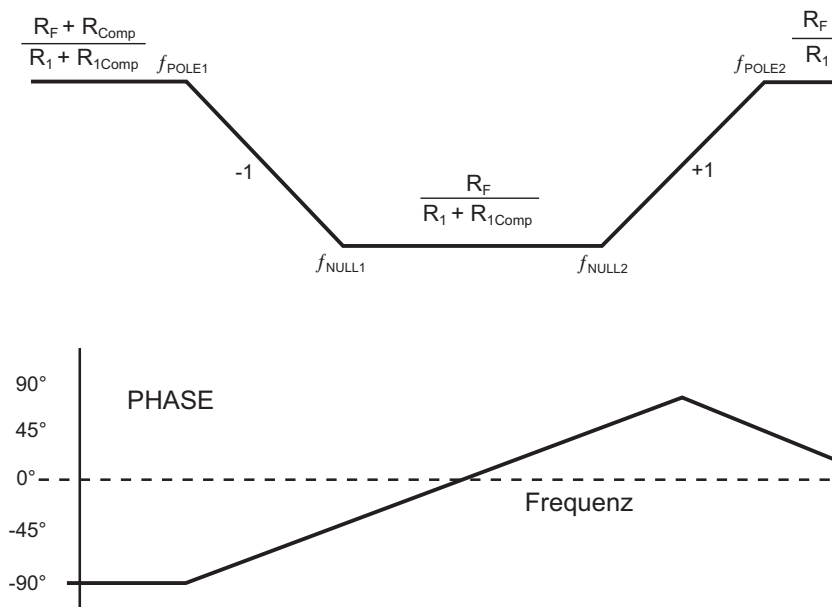


Abb. 2.6: Amplituden- und Phasenverhältnisse des kompensierten Fehlerverstärkerschaltkreises aus Abb. 2.5

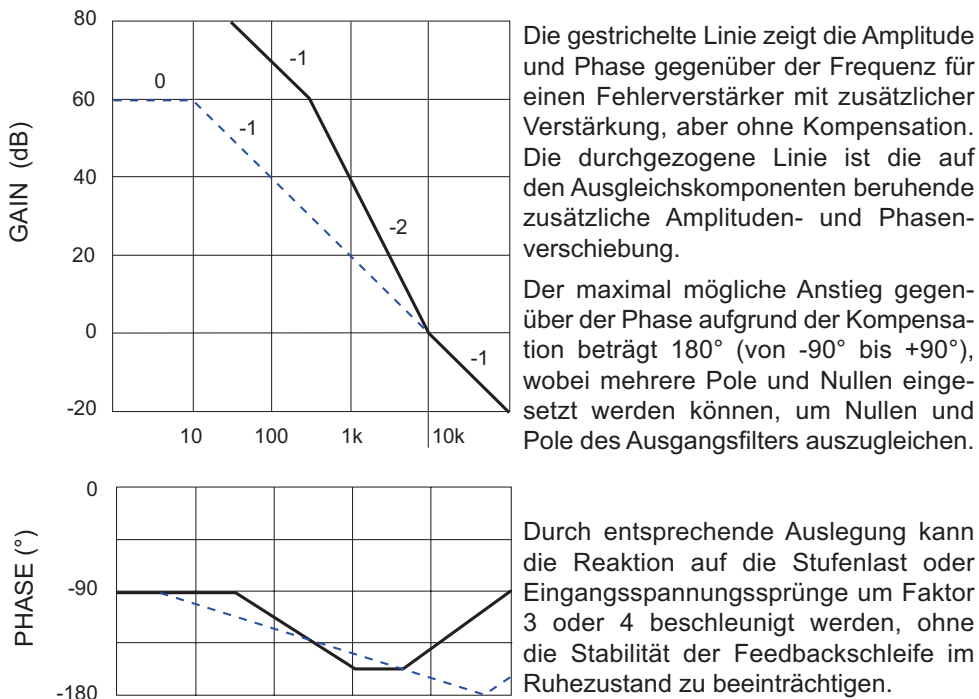


Abb. 2.7: Kompensierte (durchgezogene Linie) im Vergleich zur einpoligen (gestrichelte Linie) Feedbackschleifencharakteristik für den Schaltkreis aus Abb. 2.5

2.4.1 Instabilität der rechten Halbebene

In Topologien, die die Ausgangsinduktivität mit einem Dauerstrom durch eine Diode antreiben, wie z. B. Boost-, Buck-Boost-, Flyback- und Eintaktflusswandler, bewirkt die Leitungszeit der Diode einen Verzug in der Feedbackschleife. Nimmt die Last plötzlich zu, muss das Impuls-Pausen-Verhältnis vorübergehend vergrößert werden, um mehr Energie in die Spule zu übertragen. Ein hohes Impuls-Pausen-Verhältnis gibt der Diode jedoch weniger Zeit (t_{OFF}) zum Leiten, sodass sich der Diodendurchschnittsstrom während t_{OFF} faktisch verringert (Abb. 2.8, links). Da der Ausgangsstrom durch die Diode versorgt wird, verringert sich auch der Ausgangsstrom. Dieser Zustand bleibt erhalten, bis der durchschnittliche Spulenstrom langsam ansteigt und der Diodenstrom seinen richtigen Wert erreicht hat.

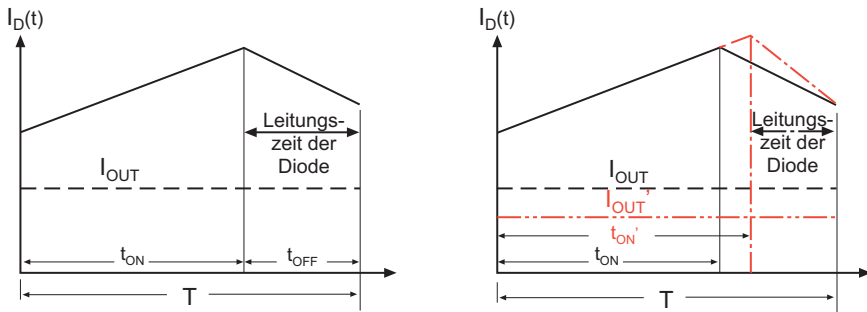


Abb. 2.8: RHP-Phänomen

Dieses Phänomen, bei dem sich der Diodenstrom zuerst verringern muss, bevor er zunehmen kann, ist als Instabilität der rechten Halbebene (RHP= Right Half Plane) bekannt, da der Ausgangsstrom mit dem Impuls-Pausen-Verhältnis vorübergehend 180° außer Phase ist. Zum Beispiel tritt eine vorübergehende zusätzliche Null in einem einfachen Boost-Wandler (Abb. 1.13) wie folgt ein:

$$f_{RHP,ZERO} = \frac{R_L}{2\pi L_1} (1 - \delta)^2$$

Gleichung 2.7: Nullberechnung für die rechte Halbebene (RHP)

Es ist fast unmöglich, RHP-Instabilität auszugleichen, insbesondere weil sich die Null mit dem Laststrom ändert. Die Lösung besteht in einer Feedbackschleife mit einer Transitfrequenz, die deutlich unter der niedrigsten Frequenz liegt, bei der RHP-Nullen entstehen (dies hat den Nachteil, dass die Reaktionszeit des DC/DC-Wandlers auf Stufenlastwechsel verringert wird), oder darin, im Buck-Boost-Wandler diskontinuierlichen Betrieb (DCM = discontinuous mode) zu verwenden, um das Gesamtproblem zu beseitigen.

2.5 Anstiegsausgleich oder slope compensation

Ein weiterer potentieller Grund für die Instabilität der Feedbackschleife ist die subharmonische oder Bifurkationsinstabilität. Die Grundursache ist der PWM-Komparator, der die Pegel des Feedbacksignals mit der Timing-Sägezahnspannung vergleicht (siehe Blockschaltbild 1.40).

Das Problem kann eintreten, weil die Energie in den Induktivitäten nicht mit jedem Schaltzyklus vollständig entladen wird, sodass der Strom zur falschen Zeit in den Feedbackkreis zurückfließt, oder einfach infolge von Schaltstörungen an den Eingängen der Komparatoren. Der Effekt ist der gleiche: Der PWM-Modulator erzeugt Bifurkations- oder Doppelschwebungen.

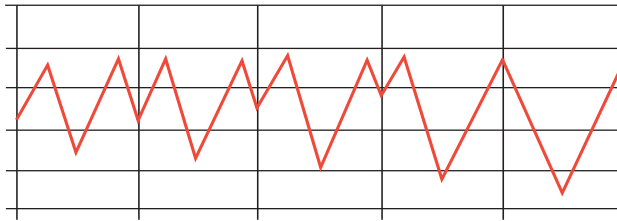


Abb. 2.9: Subharmonische Instabilitätswellenform

Die Lösung des Problems der subharmonischen Instabilität heißt Anstiegsausgleich. (engl.: slope compensation). Eine künstliche Rampenwellenform (üblicherweise aus dem Anstieg des Stroms in der Induktivität oder manchmal direkt aus der Spannung des Timing-Kondensators abgeleitet) wird der Rückkopplungsspannung hinzugefügt, um eine Fehltriggerung oder Retriggerung des PWM-Komparators zu vermeiden.

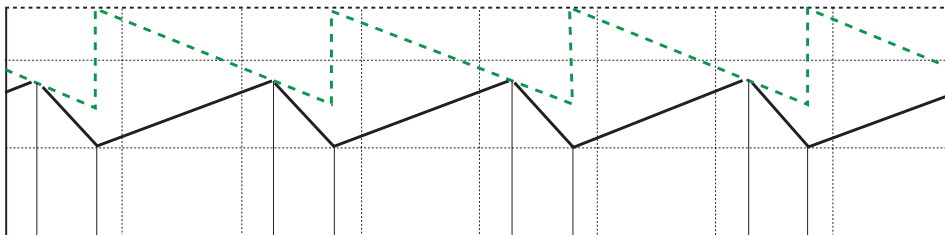


Abb. 2.10: Anstiegsausgleich (gestrichelte Linie) und Feedback-Signal (durchgezogene Linie)

2.6 Analyse der Schleifenstabilität in analogen und digitalen Feedbacksystemen

2.6.1 Experimentelle Ermittlung der analogen Schleifenstabilität

Die Stabilität von Rückkopplungsschleifen lässt sich unter Verwendung eines Bode-Plotters experimentell bestimmen. Um ein Störsignal in den Regelkreis zu senden, kann ein Sinusgenerator mit einem Übertrager verwendet werden. Die Frequenz der Sinuswelle wird dabei so lange erhöht, bis die Störung am Ausgang genauso hoch ist wie das Störsignal. Die Verstärkung beträgt dann 0dB, und somit muss die Störfrequenz – die Transitfrequenz der Rückkopplungsschleife – f_c sein. Die Phasenverschiebung zwischen dem Stör- und dem Ausgangssignal ist die Phasenreserve. Durch weitere Erhöhung der Frequenz, bis die Phasenverschiebung -180° beträgt, lässt sich die Amplitudenreserve ermitteln.

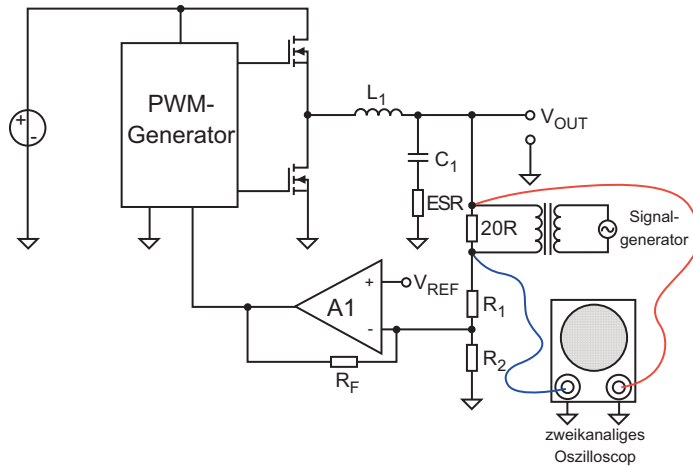


Abb. 2.11: Konfiguration für experimentelle Ermittlung des Bodediagramms

2.6.2 Ermittlung der analogen Schleifenstabilität unter Verwendung der Laplace-Transformation

Eine Alternative zur Versuchsmethode ist die mathematische Ermittlung der Pole und Nullen. Dazu muss die Übertragungsfunktion des Wandlers bekannt sein.

Für einen einfachen Buck-Wandler (Abb. 2.1) beträgt die Übertragungsfunktion:

$$G_S = \frac{1 + R_{ESR} C_1 s}{1 + (L_1/R_{LOAD} + R_{ESR} C_1) s + L_1 C_1 s^2}$$

Gleichung 2.8: Übertragungsfunktion für Abb. 2.1

Der Buchstabe „s“ weist darauf hin, dass die Variable frequenzabhängig ist. Die Übertragungsfunktion kann durch Verwendung der Laplace-Transformation (LT) gelöst werden, zuerst muss jedoch die Fourier-Transformation behandelt werden.

Die Fourier-Transformation (FT) ist eine spezielle Form der LT. Fourier legte fest, dass jedes periodische Signal die Summe sinusförmiger Signale unterschiedliche Frequenz, Phase und Amplitude ist. Die Transformation verschiebt sich von der Zeitebene zum Frequenzbereich (und umgekehrt). Das Ergebnis einer Fourier-Transformation an einem periodischen Signal ist dessen Spektrum. Abb. 2.12 ist die grafische Darstellung der Fourier-Reihe einer Rechteckwelle:

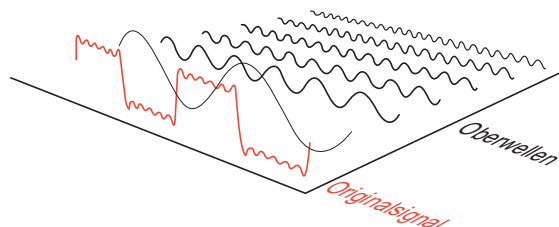


Abb. 2.12: Grafische Darstellung der Fourier-Reihe einer Rechteckwelle

Die Fourier-Transformation ist eine integrierte Funktion von der negativen bis zur positiven Unendlichkeit, die wie folgt beschrieben werden kann:

$$F(\omega) = \int_{-\infty}^{\infty} f(t) e^{j\omega t} dt$$

Gleichung 2.9: Fourier-Transformation

Abgebildet auf die S-Domäne, ist die Variable von FT $s = j \omega$. Die Ergebnisse sind nur imaginäre Variablen.

Die Laplace-Transformation ist eine erweiterte Menge von FT. Die Variable von LT befindet sich in der komplexen Ebene. Die Integration fängt bei Null anstatt $-\infty$ an. Sie kann ebenfalls zur Analyse stufenförmiger oder halbumendlicher (semi-infiniter) Signale wie Pulse oder exponential abklingende Reihen verwendet werden. Die Laplace-Transformation kann wie folgt beschrieben werden:

$$F(s) = \int_0^{\infty} f(t) e^{-st} dt$$

Abb. 2.13: Pol-Null-Diagramm in der S-Domäne mit typischen Wellenformen

Abgebildet auf die S-Domäne, ist die Variable von LT $s = \sigma + j\omega$.

Durch Verwendung der LT lässt sich die Feedbackschleife mathematisch simulieren und in der S-Ebene ein Pol-Null-Diagramm erzeugen. Die vertikale Achse ist mathematisch imaginär, die horizontale real. Je höher man sich entlang der imaginären Achse nach oben oder unten bewegt, desto schneller vollziehen sich die Schwingungen. Je weiter man sich entlang der negativen realen Achse bewegt, desto schneller ist das Abklingen, und je weiter entlang der positiven realen Achse, desto schneller die Zunahme.

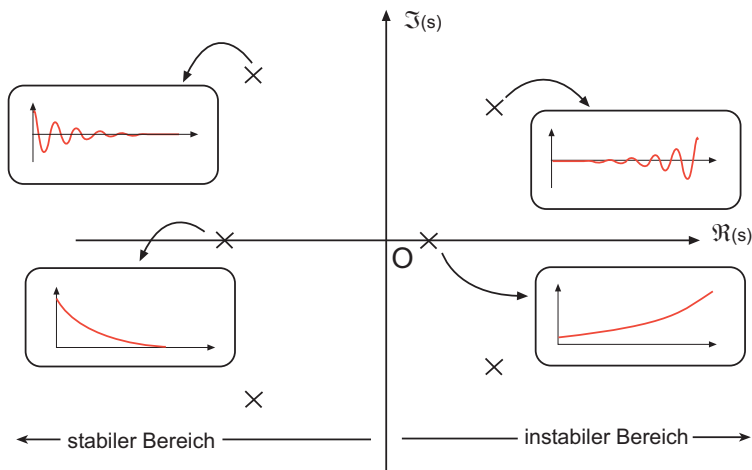


Abb. 2.13: Pol-Null-Diagramm in der S-Domäne mit typischen Wellenformen

Nullen liegen immer auf der realen Achse. Komplex konjugierte Polpaare in der linken Hälfte der s-Ebene werden kombiniert, um eine Reaktion zu erzeugen, die eine abklingende sinusförmige Schwingung der Form ist.

Ein Polpaar, das auf der imaginären Achse $\pm j\omega$ liegt (das heißt, ohne reale Komponente), erzeugt Schwingungen mit konstanter Amplitude.

Der Abstand eines Pols vom Referenzpunkt 0 zeigt an, wie gedämpft die Reaktion ist: je näher zum Referenzpunkt, desto langsamer die Geschwindigkeit der Dämpfung. Wenn ein Pol am Referenzpunkt liegt, bedeutet dies, dass das System als DC funktioniert.

Wenn der Pol auf der rechten Ebene (right hand plane) liegt, ist das System instabil (daher kommt der Begriff RHP-Instabilität, wie in Abschnitt 2.4.1 beschrieben).

2.6.3 Ermittlung der digitalen Schleifenstabilität mit Hilfe bilinearer Transformation

Wird ein digitaler Signalprozessor zur Kompensation der Feedbackschleife verwendet, kann die Stabilität der digitalen Schleife aus der Laplace-Transformation unter Verwendung einer weiteren Transformation zur Korrektur der Abtastfrequenz ermittelt werden.

Da das Eingangssignal in einem digitalen System nicht mehr zeitlich kontinuierlich ist, müssen die Werte der S-Ebene mit Hilfe der bilinearen Transformation (Tustin-Verfahren) in die zeitdiskreten Werte der z-Ebene umgewandelt werden.

Das Ergebnis dieser Funktion ist, dass der stabile Bereich der z-Ebene ein Kreis mit dem Radius = 1 (Einheitskreis) wird.

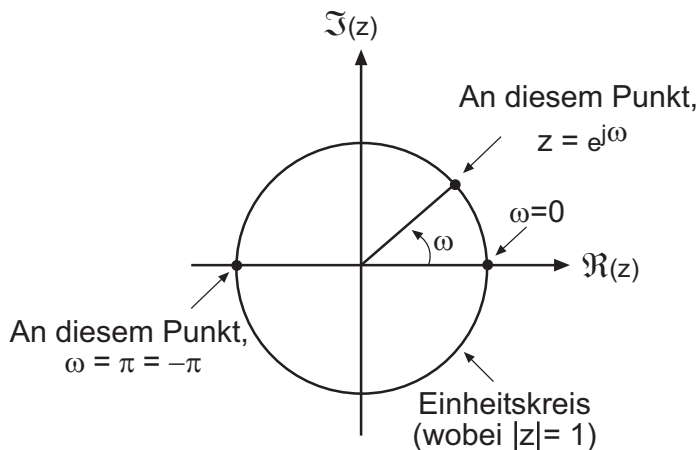


Abb. 2.14: Einheitskreis in der z-Ebene

Der äußerste rechte Rand des Kreises ($\omega = 0$) stellt DC dar. Der äußerste linke Rand des Kreises stellt die Faltungsfrequenz (engl.: aliasing frequency) dar. Alle Pole, die außerhalb des Kreises liegen, sind instabil. Die Pole der Feedbackschleife können jetzt in der z-Ebene punktwise aufgezeichnet werden, die Positionen der Pole repräsentieren die Reaktion normalisiert auf die Samplingrate, und nicht auf die kontinuierlichen Zeit wie in der s-Ebene.

Digitale Kompensation setzt erstens voraus, dass die DSP-Samplingfrequenz wesentlich größer ist als die Transitfrequenz des Systems, sodass sämtliche Simulationen akkurat sind. Es gibt zwei verbreitete Herangehensweisen zur Ermittlung der Ausgleichswerte: Umgestaltung (re-design) und direkte Gestaltung (direct-design). Bei der digitalen Umgestaltung wird ein lineares Modell eines Schaltwandlers eingerichtet, wobei der Ausgleich der Feedbackschleife auf konventionelle Weise in der S-Domäne erfolgt. Die Ergebnisse der analogen Kompensation werden in der Z-Domäne abgebildet, um die Darstellung der digitalen Kompensation zu vervollständigen. Im Falle der direkten Gestaltung wird das diskrete Modell eines Schaltwandlers vollständig simuliert und die Kompensation direkt in der Z-Domäne berechnet. Dies erfordert exakte Modelle der Analogteile unter Verwendung von Programmen wie Spice™ oder Matlab™.

Das Ergebnis beider Methoden ist dasselbe: eine in der Informationstabelle gespeicherte Matrix von Werten. DSP oder μC nehmen dann das digital umgesetzte Eingangssignal, fügen es in die Rechenmatrix ein und geben den resultierenden Wert entweder als analoges Signal oder, häufiger, als direkten PWM-Signal aus. In letzterem Fall werden der Komparator und PWM-Schaltkreis ebenfalls in digitaler Form umgesetzt. Dies beseitigt analoge Regelkreisfehler, die aus dem Anstiegsausgleich und der RHP-Instabilität entstehen. Wenn eine andere Reaktion der Kompensation der Rückkopplung erforderlich ist, um eine andersartige Funktionsweise hervorzubringen, kann der digitale Controller reibungslos zwischen den Informationstabellen umschalten, ohne einen der Ausgangswerte zurückzusetzen – eine Eigenschaft, mit der analoge Controller nicht mithalten können. Auf diese Weise müssen bei der Wahl der Ausgleichs-Charakteristik weniger Kompromisse gemacht werden.

Gerade diese Kompromisslosigkeit und Fähigkeit, schnell zwischen schnellwirkender Übergangsfunktion und stabilem Ausgang umzuschalten, machen die digitale Feedback-Schleife so attraktiv. Da der Preis der Mikrocontroller laufend fällt, verwenden immer mehr DC/DC-Wandler volldigitale oder hybride Feedback-Schleifen-Regler.

2.6.4 Digitale Feedback-Schleife

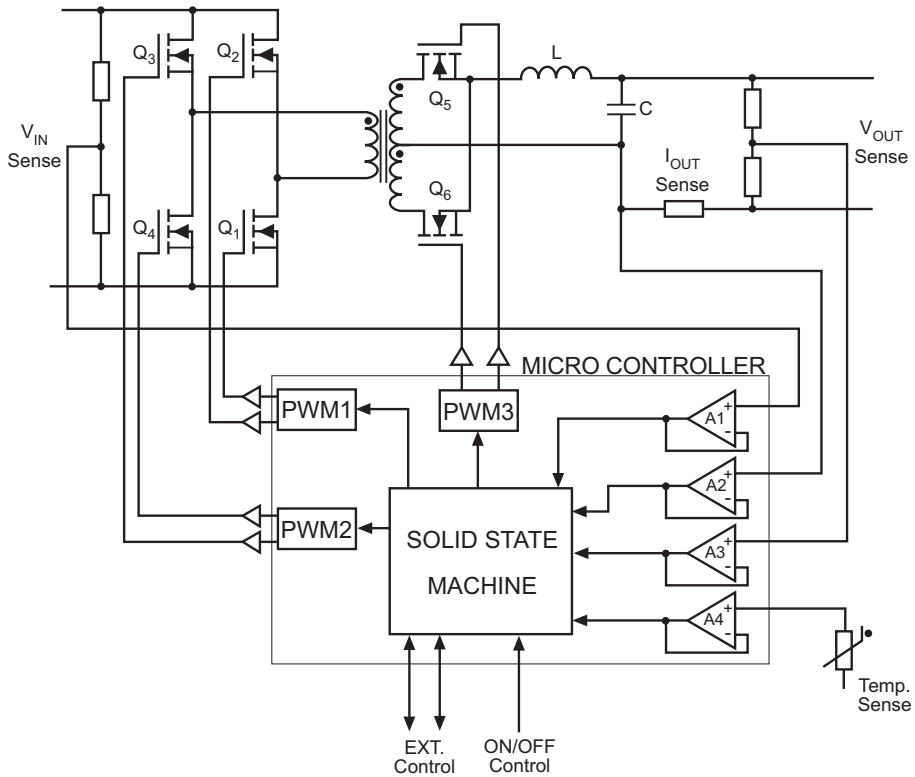


Abb. 2.15: Mikroreglergestützter DC/DC-Wandler

Der Schaltkreis oben (Abb. 2.15) zeigt einen vereinfachten mikrocontrollergestützten DC/DC-Wandler. Die komplette Steuerung, d. h. sowohl die der Vollweggleichrichterbrücken-Leistungsstufe als auch die Synchrongleichrichtung am Ausgang, erfolgt digital.

Der Mikrocontroller besitzt on-board integrierte Elemente des Operationsverstärkers. Das bedeutet, dass Messeingänge direkt am Mikrocontroller angeschlossen werden können. Da der Mikrocontroller Informationen bezüglich der Eingangsspannung, der Ausgangsspannung und des Ausgangsstroms erhält, ist es nicht erforderlich, Kurzschlüsse oder den Überlastbetrieb extern zu überwachen. Die Überwachung der Eingangsspannung ermöglicht sowohl ein kontrolliertes Anlaufen des Wandlers als auch einen programmierten Schutz vor Unterspannung mit adaptiver Hysterese. Der vierte Operationsverstärkereingang wird verwendet, um Überhitzung entweder innerhalb des DC/DC-Wandlers oder an der Last zu überwachen. Die Folgen einer Überhitzung sind laut vorgegebener Arbeitscharakteristiken in der Anwendung programmierbar, zum Beispiel Abschaltung und Neustart nach Abkühlung oder Leistungsbegrenzung zur Verringerung der Abwärme.

Der Außendatenanschluss ermöglicht es, den Betriebsmodus operativ (on-the-fly) zu aktualisieren oder verschiedene vorprogrammierte Performance-Varianten auszuwählen. Darüber hinaus ermöglicht der bidirektionale Bus Fehlermeldungen und Zustandsaktualisierungen.

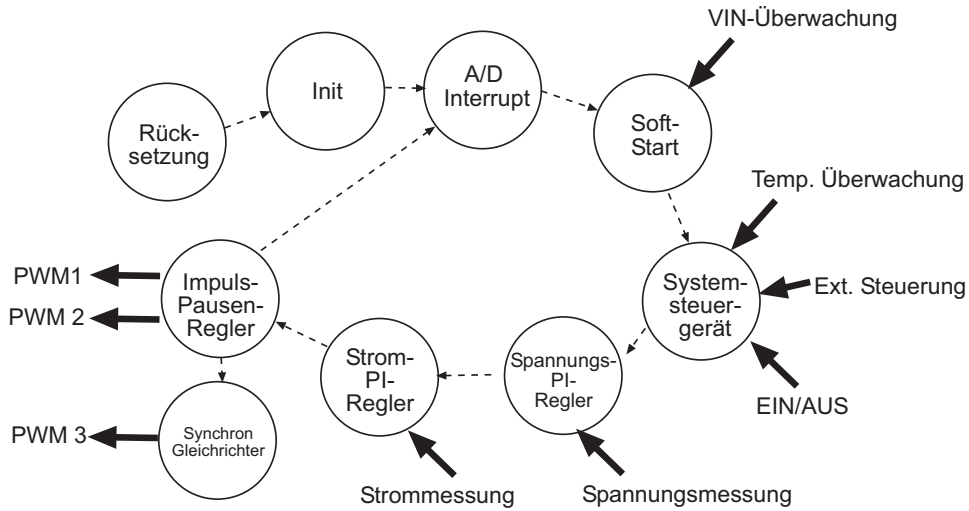


Abb. 2.16: Software-Flussdiagramm

Abbildung 2.16 stellt den internen Ablauf schematisch dar. Um die korrekte Ansprechzeit in Echtzeit zu berechnen, werden in verschiedenen Regler-Unterprogrammen Matrix-Nachschlagetabellen verwendet.

$$\frac{V_{OUT}}{V_{OUT}^*} = \frac{[(K_P R_A) + (K_i/s R_A)]}{s^2 LC + (sC R_A) + (K_P R_A) + (K_i/s) R_A}$$

wobei, V_{OUT} = Innenschleife

V_{OUT}^* = Außenschleife

P = Proportionalverstärkung Stromkompensator

und K_i und K_P können aus folgender Matrix abgeleitet werden:

$$\begin{bmatrix} \omega_1^2 & \omega_1 & 1 \\ \omega_2^2 & \omega_2 & 1 \\ \omega_3^2 & \omega_3 & 1 \end{bmatrix} \times \begin{bmatrix} C R_A \\ K_P R_A \\ K_i R_A \end{bmatrix} = \begin{bmatrix} \omega_1^3 LC \\ \omega_1^3 LC \\ \omega_1^3 LC \end{bmatrix}$$

Gleichung 2.11: Charakteristische Gleichung für Strombetriebsteuerung (CMC)

Laut Betriebszustand kann das Systemreglerprogramm unterschiedliche Matrixtabellen ein- oder ausschalten. Der Vorteil eines digitalen Reglers besteht auch in einer sehr verringerten Stückliste (BOM) sowie in der intelligenten Regelung von Ausgangsspannung und -strom.

3. Datenblatt-Parameter richtig verstehen

Jeder seriöse Hersteller liefert mit seinem Produkt ein technisches Datenblatt, das wenigstens wesentliche Performanceparameter, Abmessungen und die Pinbelegung detailliert aufzeigt. Jedoch ist für den Vergleich zweier verschiedener DC/DC-Wandler anhand der Datenblatt-Informationen häufig mehr Interpretation gefragt, als ein einfacher Vergleich von Zahlenwerten.

Das Problem besteht darin, dass viele Spezifikationen miteinander zusammenhängen, sodass eine Festlegung bestimmter Kenndaten erforderlich ist. Das heißt, dass spezifische Werte wie die Umgebungstemperatur oder Eingangsspannung während der Messung der betrachteten Performance-Werte konstant sein müssen. Zum Beispiel wird eine Lastkennlinie mit Nenneingangsspannung und einer Umgebungstemperatur von 25°C erstellt und ist über einen spezifizierten Lastbereich gültig. Aber unter Herstellern gibt es keine vereinbarten Normen über die Festlegung der Parameter, sodass einige von ihnen z.B. Regelparameter für den ganzen Lastbereich von 0% bis 100%, andere für den Lastbereich von 10% bis 100% und noch andere für den Lastbereich von 20% bis 80% der Last bestimmen. Dies bedeutet, dass eine Lastregelungs-Spezifikation von $\pm 5\%$ für den Lastbereich von 10% bis 100% nicht unbedingt schlechter sein muss als die des Konkurrenzwandlers mit einer Lastregelungs-Spezifikation von $\pm 3\%$ für den Lastbereich von 20% bis 100%. Ebenso ist ein Wandler mit Zuverlässigkeitsanforderungen von 1 Million Stunden nach MIL-HDBK-217E nicht obligatorisch zuverlässiger als ein Wandler mit nur 800 Tausend Stunden nach MIL-HDBK-217F oder ein anderer Wandler mit „nur“ 400 Tausend Stunden nach Bellcore/Telcordia.

Ist der Hersteller etwas „kreativ“ in der Anwendung dieser fehlenden Standardisierung, so kann er sein Produkt bestmöglich darzustellen. Ein klassisches Beispiel ist die Spezifikation der Ausgangsrestwelligkeit (output ripple and noise specification), die normalerweise in Millivolt Spitze zu Spitze (mVp-p) angegeben wird. Ist ein Wandler mit 50 mVp-p besser als einer mit 100mVp-p? Wohl nicht, wenn das Kleingedruckte auf der Rückseite des Datenblatts festlegt, dass die Messung des ersten Wandlers – um den Ausgang zusätzlich zu filtern – mit einem 47- μ F-Elektrolytkondensator parallel mit einem 0,1- μ F-MLCC an den Ausgangspins ausgeführt wurde, die Spezifikation des zweiten Wandlers aber ohne jegliche externe Beschaltung durchgeführt wurde. Um eine zuverlässige, reproduzierbare Messung zu erhalten, können in einigen Fällen zusätzliche Filterkomponenten erforderlich sein. Der Kunde sollte aber wissen, dass die Art und Weise, wie die Messung durchgeführt wird, den Messwert beeinflusst und der Vergleich zwischen zwei Wandler-Spezifikationen nur dann vollzogen werden kann, wenn beide bekannt sind. In vielen Fällen muss der Kunde die kritische Spezifikation, um die es geht, unter Verwendung des tatsächlichen oder erwarteten Betriebsmodus der Anwendung selbst messen. Zum Beispiel beinhalten Datenblätter normalerweise keine „Wirkungsgrad vs. Arbeitstemperatur“-Kurven (RECOM kann solche ausführlichen Informationen auf Anfrage liefern).

3.1 Messverfahren – DC-Charakteristik

Wie schon erwähnt, wird das elektrische Verhalten eines DC/DC-Wandlers durch viele unterschiedliche Kenndaten bestimmt, die im Datenblatt spezifiziert sind. Um einen Wandler schnell und wirksam zu charakterisieren und die Gültigkeit des Datenblattes zu prüfen, ist es häufig sinnvoll, eine Messmatrix zu verwenden, anhand derer verschiedene Last- und Eingangsspannungs-Kombinationen verglichen werden können.

Test	V_{IN}	I_{OUT}	V_{OUT}
1	$V_{IN,NOM}$	$I_{OUT,NOM}$	V_{O1}
2	$V_{IN,NOM}$	$I_{OUT,MIN}$	V_{O2}
3	$V_{IN,NOM}$	$I_{OUT,MAX}$	V_{O3}
4	$V_{IN,MIN}$	$I_{OUT,NOM}$	V_{O4}
5	$V_{IN,MIN}$	$I_{OUT,MIN}$	V_{O5}
6	$V_{IN,MIN}$	$I_{OUT,MAX}$	V_{O6}
7	$V_{IN,MAX}$	$I_{OUT,NOM}$	V_{O7}
8	$V_{IN,MAX}$	$I_{OUT,MIN}$	V_{O8}
9	$V_{IN,MAX}$	$I_{OUT,MAX}$	V_{O9}

$V_{IN,NOM}$	Nenningangsspannung
$V_{IN,MIN}$	Minimale Eingangsspannung
$V_{IN,MAX}$	Maximale Eingangsspannung
$I_{OUT,NOM}$	Nennausgangsstrom
$I_{OUT,MIN}$	Minimaler Ausgangsstrom
$I_{OUT,MAX}$	Maximaler Ausgangsstrom

wobei „ $I_{OUT,MIN}$ “, $\geq 0\%$ sein kann

Table 3.1: Messmatrix

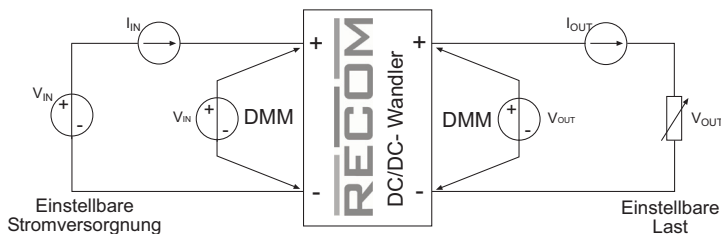


Abb. 3.1: Messanordnung

Um gute und zuverlässige Messwerte ermitteln zu können, sollte der Benutzer einige grundlegende Sicherheitsvorkehrungen zur Durchführung der Messungen treffen. Beim Vorbereiten des Prüfaufbaus vergewissern Sie sich, dass die Kontakte zum DC/DC-Wandler entsprechend niederohmig sind. Häufig haben Messklemmen variable Kontaktwiderstände. Daher ist die beste Prüfanordnung ein "Kelvin"-Kontakt – wie in Abb. 3.1 oben gezeigt, in der die Strom- und Spannungskreise separat an die Pins angeschlossen sind. Häufig ist es beim Einsatz der Multimeter, um zwei oder mehr Messgeräte anzuschließen, verlockend, die 4mm langen Steckverbinder in Messbuchsen aufeinanderzustapeln, was aber zu wesentlichen Messfehlern führen kann. Jedes Messgerät soll – wie oben gezeigt – separat an den Wandlerpins angeschlossen werden.

Um den DC/DC-Wandler zu belasten, können Hochlastwiderstände oder Hochlastregelwiderstände verwendet werden. Eine elegantere Lösung besteht aber darin, elektronische Lasten zu verwenden. Um den Strom korrekt zu regulieren, bedürfen manche Elektroniklasten jedoch einer Mindesteingangsspannung, sodass für Wandler-Ausgangsspannungen unter 4V Hochlastwiderstände häufig die einzige Alternative sind. Eine Laborstromversorgung liefert eine gut einstellbare Stromversorgung. Vergewissern Sie sich jedoch, dass sie die notwendige Spannung und Strom liefern kann, um alle Eingangsprüfungsanforderungen abzudecken. Um $V_{IN,MAX}$ zu liefern, könnte es notwendig sein, mehrere Stromversorgungen zusammenzuschalten. Die Strombegrenzung soll so bestimmt werden, dass ausreichend Spannung verfügbar ist, um den DC/DC-Wandler zu versorgen, auch bei der niedrigsten Eingangsspannung. Schließlich prüfen Sie die Polarität vor dem Einschalten, denn die meisten DC/DC-Wandler besitzen keinen Verpolungsschutz!

3.2 Messverfahren – AC-Charakteristik

Einfach ein Oszilloskop zu nehmen, einen Standardprüfkopf an den Wandler anzuschließen und die Ergebnisse vom Bildschirm abzulesen, ist meist keine sehr zuverlässige Methode. Insbesondere wenn Störungsmechanismen und ihre Wechselbeziehungen noch unbekannt sind. Differentialbetrieb- (differential mode - DM) und Gleichtakt- (common mode - CM) Effekte können die Ablesewerte verzerren. Kapitel 5 beschreibt die DM- und CM-Störungen detaillierter, vorerst aber ist es ausreichend zu wissen, dass ein einfacher Oszilloskop-Tastkopf DM-Störungen zum größten Teil ignoriert, da er symmetrisch ist und an beiden Anschlüssen gleichzeitig angelegt wird, sodass die DM-Komponente der AC-Messung am Oszillogramm fehlt.

Eine weitere Quelle von AC-Messfehlern ist die nutzbare Bandbreite des Oszilloskops. Oszilloskope weisen heute eine Eingangsbandbreite von 400MHz oder mehr auf. Eine nähere Untersuchung des Datenblattes zeigt jedoch, dass die Messung der Ausgangswelligkeit normalerweise innerhalb der Bandbreitengrenze von 20MHz ausgeführt wird – einerseits, weil der CM-Anteil über 20MHz nicht sehr wesentlich ist, da es mit einem kleinen Kondensator leicht herausgefiltert werden kann, und andererseits soll die Messung nicht vom Typus oder Hersteller des Oszilloskops abhängig sein. Ein Oszilloskop, das ohne die 20MHz-Bandbreitenbegrenzungs-Option verwendet wird, wird deshalb immer höhere Ablesewerte liefern.

Praktischer Hinweis

Auch kann der Prüfkopf selbst eine Fehlerquelle darstellen. Es sollte darauf geachtet werden, die Kabel zum Prüfkopf möglichst kurz zu halten. Im Idealfall berührt die Prüfkopfspitze den +Pin, und der Masse-Pin berührt den Ring. Die mitgelieferte Masseklemme darf auf keinen Fall verwendet werden, da die durch den Masseleiter gebildete Schleife eine Antenne darstellt, die erhebliche Störsignale einfangen kann.

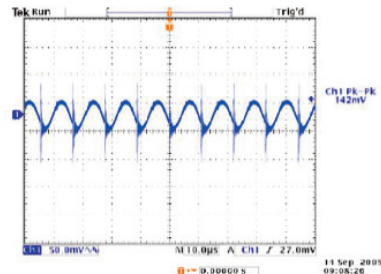
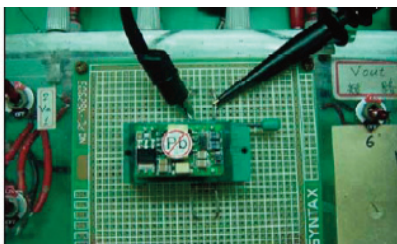


Abb. 3.2a: Unkorrekte Messung von AC-Signalen

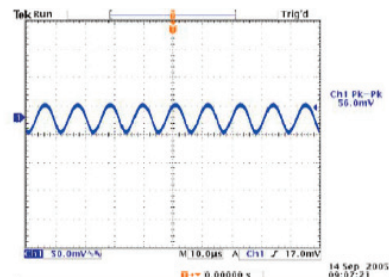
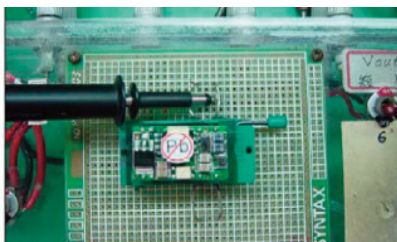


Abb. 3.2b: Korrekte Messung von AC-Signalen

Wenn kein Prüfkopf mit kurzen Kontakt-Strompfaden verwendet werden kann, ist der in Abb. 3.3 gezeigte Vorschlag von Nutzen. Die Impedanzanpassung mittels RC-Komponenten vermeidet zu hohe Reflexionssignale, die den Ablesewert verfälschen könnten.

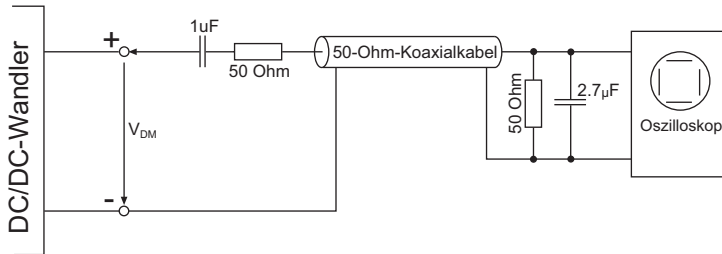


Abb. 3.3: Alternatives AC-Messverfahren

Praktischer Hinweis

Beachten Sie, dass die gemessene Wellenform durch den durch zwei 50-Ohm-Widerstände gebildeten Spannungsteiler halbiert wird, sodass die Werte des Oszillogramms verdoppelt werden sollten. Selbst bei passenden Komponenten sollte das Koaxialkabel möglichst kurz gehalten werden.

3.2.1 Messung des minimalen und maximalen Impuls-Pausen-Verhältnisses

In einigen Anwendungen wäre es nützlich, mehr über die interne Taktung des DC/DC-Wandlers zu wissen, da das Impuls-Pausen-Verhältnis-Signal häufig nicht direkt von außerhalb des Moduls zugänglich ist. Mit etwas Erfahrung kann jedoch die Interpretation des Eingangs- oder Ausgangsrauschens diese Informationen liefern.

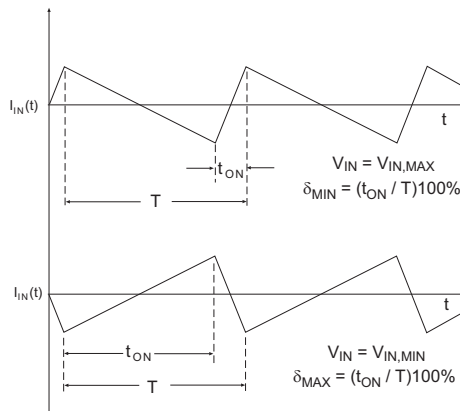


Abb. 3.4: Messung des Impuls-Pausen-Verhältnisses aus der Ausgangswellenform

Das minimale Impuls-Pausen-Verhältnis δ_{MIN} wird durch die Parameter $V_{\text{IN}} = V_{\text{IN,MAX}}$ und $I_{\text{OUT}} = I_{\text{OUT,MIN}}$, das maximale Impuls-Pausen-Verhältnis δ_{MAX} bei $V_{\text{IN}} = V_{\text{IN,MIN}}$ und $I_{\text{OUT}} = I_{\text{OUT,MAX}}$ bestimmt. Die Zeit T ist konstant, weil sie durch die Taktfrequenz des DC/DC-Wandlers festgelegt ist. Abb. 3.4 zeigt, wie das Impuls-Pausen-Verhältnis aus der Wellenform des Eingangsstroms ermittelt werden kann.

3.2.2 Ausgangsspannungsgenauigkeit

Die Charakteristik der Ausgangsspannungsgenauigkeit, auch Sollwertgenauigkeit genannt, beschreibt die spezifizierte Toleranz der Ausgangsspannung. Normalerweise wird der Parameter in Prozent der Nennausgangsspannung angegeben, normalerweise bei Raumtemperatur, Volllast und Nenneingangsspannung.

$$ACC_{V,OUT} = \frac{V_{OUT} - V_{OUT,NOM}}{V_{OUT,NOM}} 100\%$$

Gleichung 3.1: Ausgangsspannungsgenauigkeit

Abweichungen vom Sollwert der Ausgangsspannung ergeben sich aus Komponententoleranzen, insbesondere im Widerstandsteiler, der die Ausgangsspannung auf die Referenzspannung des PWM-Komparators absenkt (siehe Abb. 1.46). Für Ausgangsspannungen über 1,5VDC wird üblicherweise eine Bandgap-Spannungsreferenz von 1,22V verwendet (eine Bandgap-Spannungsreferenz verwendet zwei PN-Übergänge, die so angeordnet sind, dass der Temperaturkoeffizient eines Übergangs denjenigen des anderen ausgleicht und somit eine sehr stabile Referenzspannung schafft). Bei einer Ausgangsspannung von 5V hat der Widerstandsteiler ein Verhältnis von 3:1, sodass die Ausgangsspannungsgenauigkeit $\pm 3\%$ beträgt, wenn Widerstandstoleranz 1% zugrunde gelegt wird. Anstelle des berechneten Widerstandswerts kann außerdem der nächstliegende Standardwert-Widerstand verwendet werden, was einen weiteren Fehler mit sich zieht.

Einige geregelte Wandler haben einen Trim-Bereich, mit dem die Ausgangsspannung innerhalb eines bestimmten Bereichs angepasst werden kann, normalerweise $\pm 10\%$. In diesem Fall wird die Spezifikation mit offenem Abgleichpin (unbenutzt) verwendet.

Ungeregelte Wandler haben eine Ausgangsspannung, die lastabhängig ist. Wenn die Nennausgangsspannung bei 100% Last definiert würde, wäre die Ausgangsspannung für alle Lasten unter 100% höher als die Nennspannung, was den Nutzlastbereich des Wandlers reduzieren könnte. Deshalb wird die Ausgangsspannung normalerweise zwischen ca. 60% bis 80% der Last festgelegt (siehe Abb. 1.31). Im Volllastbetrieb liegt die Ausgangsspannung somit immer etwas unter V_{NOM} .

3.2.3 Temperaturkoeffizient der Ausgangsspannung

Obwohl die innere Bandgap-Spannungsreferenz eine sehr stabile Spannung über den Betriebstemperaturbereich aufweist, kann es zu Abweichungen kommen. Der Gesamttemperaturkoeffizient (TC) der Ausgangsspannung wird als relative Abweichung der Ausgangsspannung bei Ober- und Untergrenze des Betriebstemperaturbereichs im Vergleich zur Nennausgangsspannung bei Raumtemperatur bestimmt. Normalerweise werden die Daten in %/C oder ppm/K angegeben (ppm = parts per million = Teile je Million). Üblicherweise sind Temperaturkoeffizienten bei niedrigen Temperaturen positiv und bei hohen Temperaturen negativ.

Um den TC zu bestimmen, ist eine Thermokammer, die die notwendigen Umgebungstemperaturen erzeugen kann, erforderlich. Bei Raumtemperatur T_{RT} wird die Nennspannung $V_{OUT}(T_{RT})$ unter Nennlast gemessen, nachdem der DC/DC-Wandler, um eine Temperaturstabilisierung zu erzielen, für eine Wartezeit von 20 Minuten warmgelaufen ist.

Ein entsprechendes Verfahren findet bei den oberen und unteren Grenzwerten des Betriebstemperaturbereichs Verwendung. Die Berechnung von TC erfolgt laut unten stehender Gleichung.

$$TC(T) [\%/^{\circ}\text{C}] = \frac{\Delta V_{OUT}(T)}{V_{NOM} \Delta T} = \frac{V_{OUT}(T) - V_{OUT}(T_{RT})}{V_{OUT}(T_{RT}) |T_{RT} - T|} 100\%,$$

oder

$$TC(T) [\text{ppm}/^{\circ}\text{K}] = \frac{\Delta V_{OUT}(T)}{V_{NOM} \Delta T} = \frac{V_{OUT}(T) - V_{OUT}(T_{RT})}{V_{OUT}(T_{RT}) |T_{RT} - T|} 10^6$$

Gleichung 3.2: Berechnung des Temperaturkoeffizienten

Ein typischer TC Wert beträgt $\pm 0,02\%/^{\circ}\text{C}$. Dies bedeutet: Wenn die Ausgangsspannung einen Nennwert von 25°C hat, reduziert sich die Spannung um 1 % bei $+75^{\circ}\text{C}$ und erhöht sich um 1 % bei -25°C .

3.2.4 Lastregelung

Die Lastregelung (engl.: load regulation) wird als die maximale Abweichung der gemessenen Ausgangsspannung über den zulässigen Lastbereich von der Mindestlast (minimum load - ML) bis zur Volllast (full load - FL), angegeben in Prozent, definiert. Die Eingangsspannung wird normalerweise am Nennwert V_{NOM} konstant gehalten. Es ist zu beachten, dass für viele Wandler eine minimale Last erforderlich ist, um eine korrekte Regulierung zu gewährleisten, und Überlastschutz kann vorhanden sein, weshalb es nicht zulässig ist, die Lastregelungskurve über ML oder FL hinaus zu extrapolieren.

Normalerweise ändert sich die Ausgangsspannung mit dem Laststrom linear, weshalb nur zwei Messpunkte innerhalb des angegebenen Lastbereichs zur Berechnung der Lastregelung notwendig sind. Wenn die Spezifikationen mit einem Nennstrom von beispielsweise 80 % des Maximums erstellt werden, so kann die Lastregelung auf drei verschiedene Arten gemessen werden: die Ausgangsspannungen unter ML und unter Halblast (Halblast = $(ML + FL)/2$), die Ausgangsspannungen unter FL und unter Halblast oder Ausgangsspannungen unter ML und FL, wobei alle in etwa dasselbe Ergebnis ergeben sollten. Diese unterschiedlichen Berechnungen sind aufgrund der in Abb. 3.1 gezeigten Prüfanordnung und des in Tabelle 3.1 dargestellten Messverfahrens möglich. Die folgende Gleichung zeigt das Ergebnis, das auf Messungen, die mit ML und FL durchgeführt wurden, basiert:

$$REG_{LOAD} = \frac{V_{OUT,ML} - V_{OUT,FL}}{V_{OUT,FL}} 100\% = \frac{V_{03} - V_{02}}{V_{02}} 100\%$$

Gleichung 3.3: Berechnung der Lastregelung

Wenn das Datenblatt die Ausgangsspannungsmessgenauigkeit (OVA)-Kurve bei 50% der Last bestimmt und eine Lastregelung von $\pm 1\%$ festgelegt wird, beträgt die relative Änderung der Ausgangsspannung unter Volllast -1%, und unter der minimalen Last beträgt sie +1%. Somit kann die gemessene Ausgangsspannung um bis zu 1% oberhalb oder unterhalb der OVA-Kurve liegen. Das kann verwirrend sein, wenn die Ausgangsspannungsmessgenauigkeit unter Volllast bestimmt wird, da dann das Lastregelungsergebnis nur negativ sein kann, gemeinhin ist immer noch ein $\pm$ -Prozentsatz angegeben.

Dies bedeutet, dass wenn die Lastregelung bei $\pm 1\%$ (und OVA = 100% Last) bestimmt wird, kann die gemessene Spannung nur gleich sein wie die OVA-Kurve oder bis zu 1% darunter liegen. Das ist tatsächlich doppelt so genau wie die erste Definition. Ungeregelte Wandler verwenden die Abweichung vs. Last, gemessen mit Nenn- V_{IN} , um zu bezeichnen, wie sich die Ausgangsspannung bei den jeweiligen Lastsituationen ändert, da sie weder konstante Last- noch Netzregelung (line regulation = Ausregelung von Änderungen der Eingangsspannung) haben.

3.2.5 Kreuzregelung

Dieser Parameter trifft nur für Wandler mit bipolaren oder Mehrfachausgängen zu. Ein Ausgang wird unter Volllast betrieben, und der andere hat eine niedrigere Last, normalerweise 25% (Kleinlast). Dann werden die Lasten so geschaltet, dass der erste Ausgang 25% und der zweite 100% Last hat. Die Kreuzregelung in Prozent wird aus derjenigen der beiden Gleichungen abgeleitet, die den höchsten Wert ergibt:

$$REG_{CROSS,1} = \frac{V_{OUT1,LL} - V_{OUT2,FL}}{V_{OUT2,FL}} 100\%$$

$$REG_{CROSS,2} = \frac{V_{OUT2,LL} - V_{OUT1,FL}}{V_{OUT1,FL}} 100\%, \text{ wobei } FL = \text{Volllast, } LL = \text{Kleinlast}$$

Gleichung 3.4: Berechnung der Kreuzregelung

3.2.6 Netzregelung (line regulation)

Die Netzregelung, besser bezeichnet mit dem englischen Begriff „line regulation“, da der eingedeutschte Begriff im DC-Bereich verwirrend und die wörtliche Übersetzung „Eingangsspannungsänderungsausregelungskoeffizient“ wenig anwendungsfreundlich erscheint, ist die Abweichung der Ausgangsspannung infolge der Änderungen der Eingangsspannung innerhalb deren minimalen (VL) und maximalen (VH) Grenzen. Die Last bleibt konstant, normalerweise bei maximalem Strom. Die line regulation wird als die prozentuale Abweichung der Ausgangsspannung bezüglich des Ausgangsspannungsnennwertes bestimmt. Wie bei der Lastregelung ändert sich die Ausgangsspannung mit der Eingangsspannung linear, deshalb können die Nenneingangsspannung (VN) und die Differenz zwischen VL und VH verwendet werden, um diesen Parameter zu bestimmen. Die folgende Gleichung basiert auf den Messungen der Abweichung der Ausgangsspannung über den vollen Eingangsspannungsbereich VL zu VH.

$$REG_{LINE} = \frac{V_{OUT,VH} - V_{OUT,VL}}{V_{OUT,VN}} 100\% = \frac{V_{09} - V_{06}}{V_{03}} 100\%$$

Gleichung 3.5: Berechnung der line regulation

Die Nenneingangsspannung wird normalerweise ungefähr in der Mitte des Eingangsspannungsbereichs definiert. Wenn die line regulation im Datenblatt also z. B. als $\pm 1\%$ festgelegt wird, bedeutet dies, dass eine Änderung der Eingangsspannung von VN zu VH ein Ansteigen der Ausgangsspannung von +1% und eine Änderung der Eingangsspannung von VN zu VL ein Absinken der Ausgangsspannung von -1% bewirkt.

Ungeregelte Regler können, wie im Namen bereits klar festgelegt, keine Regelung der Ausgangsspannung bei Änderung der Eingangsspannung sicherstellen. Für eine festgelegte Last nimmt die Ausgangsspannung mit der zunehmenden Eingangsspannung zu und mit der abnehmenden Eingangsspannung ab. Das Verhältnis ist jedoch normalerweise nicht 1:1. Eine Änderung der Eingangsspannung von 1% ruft nicht unbedingt eine entsprechende Änderung von 1% im Ausgang hervor. Die Wirkung der Eingangsspannung auf die Ausgangsspannung wird bei Volllast im Format „x %/1% von V_{IN} “ bestimmt. Wenn die line regulation eines unregulierten Wandlers z. B. als 1,2%/1 % von V_{IN} angegeben wird, nimmt die Ausgangsspannung um 1,2% für jede Vergrößerung der Eingangsspannung um 1% zu.

3.2.7 Ausgangsspannungsgenauigkeit im ungünstigsten Fall

Der ungünstigste Fall ist die Kombination der Grenzwerte der Ausgangsspannungsgenauigkeit, der Lastregelung über den verwendeten Lastbereich, der line regulation im verwendeten Eingangsspannungsbereich und der Temperaturkoeffizient. Da sich die Fehler aufaddieren, wirkt sich die Reihenfolge, in der die Berechnung durchgeführt wird, aus. Normalerweise kann jeder Fehler jedoch separat bearbeitet werden, um eine erste Annäherung der Ausgangsspannungsgrenzen zu ermitteln:

$$V_{OUT,MIN} = V_{OUT,NOM} [1 - ACC_{V,OUT} - REG_{LOAD} - REG_{LINE} - TC |T_{RT} - T_{MAX}|]$$

$$V_{OUT,MAX} = V_{OUT,NOM} [1 + ACC_{V,OUT} + REG_{LOAD} + REG_{LINE} + TC |T_{RT} - T_{MIN}|]$$

Gleichung 3.6: Ausgangsspannung im ungünstigsten Fall

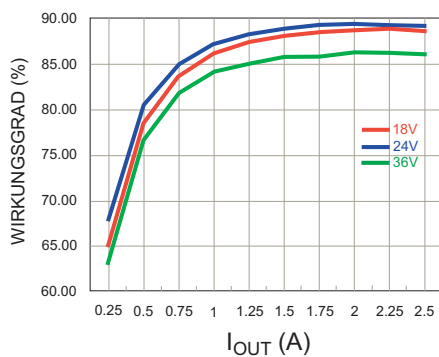
Wenn zum Beispiel die Nennausgangsspannung 5V, die Ausgangsspannungsgenauigkeit = $\pm 2\%$, die Lastregelung $\pm 0,5\%$, die line regulation $\pm 0,3\%$ und der Temperaturkoeffizient über dem Arbeitstemperaturbereich $+1,2\%$ / $-1,3\%$ beträgt, dann:

$$V_{OUT,MIN} = 5 \times (1 - 0.02 - 0.005 - 0.003 - 0.013) = 4.795V$$

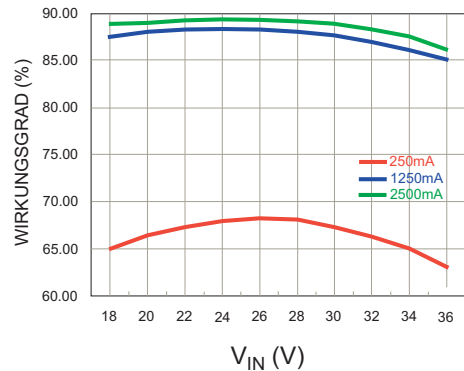
$$V_{OUT,MAX} = 5 \times (1 + 0.02 + 0.005 + 0.003 + 0.012) = 5.200V$$

3.2.8 Berechnung des Wirkungsgrades

Der Wirkungsgrad wird durch das Verhältnis der Ausgangs- zur Eingangsleistung angegeben. Im Leerlauf (Last = 0) ist der Wirkungsgrad immer gleich Null. Typisch ist die Angabe in Prozent, er kann aber auch als normalisierte Zahl (≤ 1) angegeben werden. Normalerweise werden die Daten unter mehreren Bedingungen, wie Nenneingangsspannung und Volllast, bereitgestellt. Zur Veranschaulichung der Komplexität dieses Parameters zeigt Abb. 3.5 die Wirkungsgradkurven desselben DC/DC-Wandlers unter verschiedenen Messbedingungen.



Wirkungsgrad vs. Last für verschiedene Eingangsspannungen



Wirkungsgrad vs. Eingangsspannung für unterschiedliche Lasten

Abb. 3.5: Typische Wirkungsgradkurven des Wandlers (RP30-2405S)

3.2.9 Eingangsspannungsbereich

Der nutzbare Eingangsspannungsbereich des DC/DC-Wandlers wird durch die Eingangsspannungsgrenzen V_L und V_H bestimmt, zwischen denen der Wandler korrekt mit garantierter geregelter Ausgangsspannung arbeitet. Die meisten Wandler arbeiten auch außerhalb dieses Bereichs, können dann aber nicht alle Datenblatt-Spezifikationen erfüllen. Hochleistungswandler können eine Unterspannung (UVL)-Sperrschaltung, eine sogenannte undervoltage-lockout-Schaltung, haben, die den Wandler ausschaltet, wenn sich der Eingangsstrom infolge einer zu niedrigen Eingangsspannung zu sehr erhöht. Das soll den Wandler vor zu hohen primärseitigen Strömen schützen. Manchmal sind „absolute Maximalwerte“, sog. „absolute maximum ratings“, im Datenblatt angegeben, um die maximale Eingangsspannung anzugeben, die der Wandler zulässt, bevor für die Komponenten im Wandler zerstörende Spannungsgrenzwerte überschritten werden.

Praktischer Hinweis

Der im Datenblatt angegebene Eingangsspannungsbereich gilt für Gleichspannungen. Häufig können aufgrund von Einschwingvorgängen höhere Eingangsspannungen auftreten, ohne dass Schäden entstehen. So besitzt z. B. der Schaltregler R-785.0-0.5 mit einem Ausgang von 5V einen Eingangsspannungsbereich von 6,5V bis 32V und eine absolute maximale Eingangsspannung von 34V, hält aber kurzzeitig einer Überspannung von 50V/100 ms, sowie einem schnellen Übergangsbetrieb von 1000 V/50µs, stand.

3.2.10 Eingangsstrom

Der Eingangsstrom besteht aus zwei Komponenten, einer DC-Komponente (Durchschnittswert des Eingangsstroms) und einer AC-Komponente, dem sog. back ripple current. Die Messung des back ripple currents wird detailliert in Kapitel 5.2.1 behandelt.

Die DC-Komponente des Eingangsstroms besteht ihrerseits aus zwei Komponenten, der Eingangsruhestrom (eng.: bias current) und dem Eingangsstrom infolge der Last. Um die Eingangsruhestromaufnahme zu ermitteln, kann die Last einfach getrennt werden. Dieser Eingangsruhestrom wird normalerweise auch als lastfreier Ruhestrom (I_Q) (engl.: quiescent current) bezeichnet. Dieser Strom entsteht, weil auch im Leerlauf der Taktgenerator im Wandler immer noch läuft und Leistung infolge der verschiedenen Schaltverlust- und Blindleistungen anfällt und interne Spannungsregler und Spannungsreferenzen arbeiten müssen, auch wenn kein Laststrom fließt. Der Eingangsruhestrom hängt von der Eingangsspannung und von der Umgebungstemperatur ab, weshalb I_Q normalerweise bei $V_{IN,NOM}$ und 25°C Raumtemperatur gemessen wird. Wandler, die eine Ein-/Aus-Funktion- oder einen Stand-by-Betrieb haben, können den Eingangsruhestrom noch weiter verringern, indem die inneren Oszillatoren und Regler sowie die Ausgangsleistungsstufe gesperrt werden. Deshalb ist I_{OFF} immer kleiner als I_Q .

Praktischer Hinweis

Der lastabhängige Teil des Eingangsstroms ist nicht immer so einfach zu interpretieren. Vor allem hängt er von der Eingangsspannung ab, deshalb gilt das Verhältnis „minimale Eingangsspannung = maximaler Eingangsstrom“. Zudem ist das Verhältnis „Wirkungsgrad gegen Last“ nichtlinear (siehe Abb. 3.5) und erschwert so die Untersuchung. Der Wirkungsgrad ist daher eine komplexe Funktion von Ausgangsstrom und Eingangsspannung. Wenn ein Entwickler den maximalen Eingangsstrom berechnen will, sollte er die mögliche Mindesteingangsspannung und die maximale Last, die in dieser Situation vorkommen könnten, sowie den Wirkungsgrad des Wandlers unter diesen Bedingungen kennen (zum Beispiel, indem er die Wirkungsgradwerte aus einem Diagramm wie dem in Abb. 3.5 gezeigten abliest). Es ist häufig nicht korrekt, die Berechnung unter Annahme des im Datenblatt angegebenen Zahlenwertes des Wirkungsgrades bei Volllast durchzuführen, insbesondere nicht wenn die Lasten von der Volllast abweichen.

3.2.11 Kurzschluss und Überlastung

Der Ausgangskurzschlussstrom (shortcircuit current - S/C) ist der Ausgangsstrom, der fließt, wenn die Ausgangsanschlusspins kurzgeschlossen sind. Normalerweise wird ein Kurzschluss als eine Verbindung definiert, die einen Widerstand $<1\Omega$ oder einen ausreichend niedrigen Nebenwiderstand aufweist, sodass die resultierende Ausgangsspannung unter 100mV liegt. Bei einem Wandler mit einfachem Ausgang wird der Kurzschlusstest zwischen V_{OUT+} und V_{OUT-} durchgeführt. Bei einem Wandler mit bipolarem Ausgang kann der Kurzschlusstest zwischen V_{OUT+} und V_{OUT-} , V_{OUT+} und der gemeinsamen Leitung oder V_{OUT-} und der gemeinsamen Leitung durchgeführt werden.

Unregelte Niederleistungs-DC/DC-Wandler sind häufig nicht kurzschlussfest, d.h. nicht gegenüber kurzgeschlossenem Ausgang geschützt. Es ist in der Branche üblich, die Kurzschlussfestigkeit für 1 Sekunde zu definieren. Dies ist üblicherweise die Zeit, die angenommen wird, bis sich im Inneren die Komponenten überhitzen und durchbrennen. Vor der Durchführung eines Kurzschlusstests sollte sich der Entwickler also zuerst vergewissern, ob der Wandler kurzschlussfest ist und wenn ja, welche Art des Schutzes verwendet wird: Leistungsbegrenzung mit Übertemperaturabschaltung, Stromfoldback-Schutz oder Hiccup-Schutz.

Überlastungsschutz ist nicht dasselbe wie Kurzschlusschutz. Wenn der Ausgangsstrom das Limit (normalerweise von 110% bis 150% des Nennausgangsstroms) überschreitet, dann lässt der strombegrenzte DC/DC-Wandler zu, dass sich die Ausgangsspannung verringert, um den Strom stabil an diesem Grenzwert zu halten. Bei einer weiteren Erhöhung der Last sinkt die Ausgangsspannung proportional. Der Wandler befindet sich in einem Modus der die Ausgangsleistung auf einen konstanten Wert beschränkt, anstelle einer konstanten Ausgangsspannung. Wenn die Überlast beseitigt ist, kehrt der Wandler zum Normalbetrieb zurück. Wenn die Überlastsituation jedoch andauert, führt die erhöhte innere Verlustleistung zu einer Überhitzung des Wandlers und hat entweder dessen Ausfall oder thermische Abschaltung zur Folge.

Wenn der Ausgang kurzgeschlossen wird, bleibt der Ausgangsstrom durch den definierten Grenzwert immer noch beschränkt, die Ausgangsspannung jedoch ist sehr niedrig und beträgt bei einem idealen Kurzschluss theoretisch Null, praktisch aber einige Millivolt. Die Ausgangsleistung ist dann auch nahe Null, und der Wandler kann auf unbestimmte Zeit arbeiten, vorausgesetzt, die inneren Komponenten sind für diese Überlastsituation ausgelegt. Somit kann ein Wandler im Überlastbetrieb ausfallen, einen unbestimmten Kurzschluss jedoch unbeschädigt überstehen. Eine Abwandlung des Strombegrenzungsschutzes stellt der Stromfoldback-Schutz dar (Abb. 3.6).

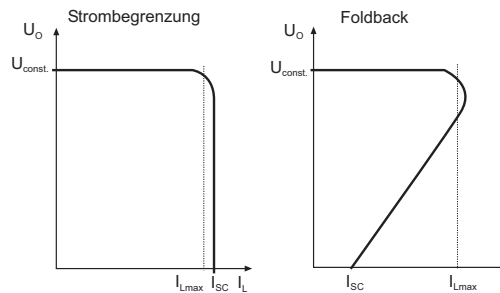


Abb. 3.6: Ausgangsstrombegrenzung und Foldback-Charakteristik

Wenn der Ausgangsstrom den definierten Grenzwert überschreitet, wird der Grenzwert auf einen sehr viel niedrigeren Wert zurückgesetzt. Der DC/DC-Wandler befindet sich im Leistungsbegrenzungsmodus, aber bei einer niedrigeren Leistung als im Normalbetrieb. In der Regel muss dieser Modus zurückgesetzt werden, indem der Wandler von der Versorgung getrennt wird. Während Strombegrenzung oder Stromfoldback sehr wirksame Kurzschlusschutzverfahren für niedrige bis mittlere Leistungs-DC/DC-Wandler sind, können sie für Hochleistungswandler ungeeignet sein.

Ein 1W Wandler mit einer Strombegrenzung von 150%, muss mit einer zusätzlichen Verlustleistung von 500mW während Überlastungs- oder Kurzschlussituationen zurechtkommen, aber ein 50W Wandler muss 25W zusätzliche Verlustleistung umsetzen. Dies würde die inneren Komponenten schnell überhitzen, aber deren Überdimensionierung, um diesen hohen Kurzschlussstrom zu überstehen, wäre unwirtschaftlich.

Die Lösung des Problems besteht darin, den Hiccup-Schutz zu verwenden. Wenn der Ausgangsstrom den definierten Grenzwert überschreitet, schaltet der Wandler sofort ab. Nach einer kurzen Verzögerung versucht der Wandler, wieder zu starten. Wenn der Ausgangsstrom die Grenze immer noch überschreitet, schaltet der Wandler wieder ab, und der Zyklus wiederholt sich.

Der Vorteil des Hiccup-Schutzes besteht darin, dass der Wandler am folgenden Hiccup wieder selbsttätig startet, wenn die Funktionsunfähigkeit beseitigt ist. Ein weiterer Vorteil dieser Schutzart besteht darin, dass, obwohl der kurze Ausgangsstromimpuls kurzzeitig eine hohe innere Verlustleistung hervorruft, die lange Wartepause zwischen den „Hiccups“ dazu führt, dass die inneren Komponenten wieder abkühlen können und der Wandler auch bei einem vollständigen Kurzschluss nicht heiß läuft.

Die Nachteile des Hiccup-Schutzes bestehen darin, dass hochkapazitive Lasten den Hiccup-Mechanismus auslösen können und der Wandler im hochkapazitiven Lastenbetrieb nicht startet. Ein weiterer Nachteil ergibt sich, wenn der DC/DC-Wandler zur Spannungsversorgung in langen Kabelnetzen verwendet wird. Ein Kurzschluss an irgendeiner Stelle löst den Hiccup-Mechanismus aus und die Hiccup-Stromspitzen können es sehr erschweren, die genaue Lage des Kurzschlusses zu bestimmen.

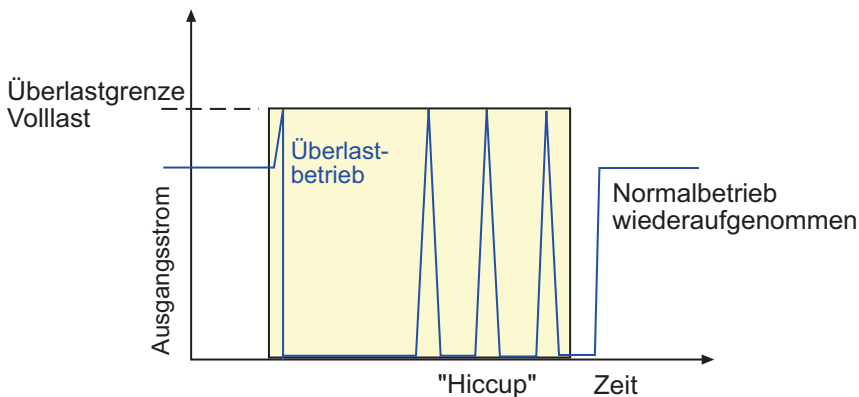


Abb. 3.7: Hiccup-Charakteristik

Praktischer Hinweis

Die einfachste Art, die Kurzschlussfunktion mit Hiccup-geschütztem DC/DC-Wandler zu testen, besteht darin, einfach zuzuhören. Ein Wandler mit Hiccup-Schutz erzeugt bei jedem Versuch zu starten, ein hörbares „Tick“. Um den Ausgang zu überwachen, kann alternativ ein Oszilloskop mit Strommessshunt verwendet werden.

Um die Strombegrenzungs- oder Strom-Foldback-Charakteristiken zu messen, können die in Abb. 3.8 gezeigten Prüfanordnungen verwendet werden. In der oberen Prüfanordnung wird ein Digitalmultimeter (DMM) im DC-Strommessmodus verwendet, und der innere Shuntwiderstand wird als Kurzschlusselement eingesetzt. Dies muss überwacht werden, um zu prüfen, dass die Spannung an den DC/DC-Wandler-Ausgangsklemmen 100 mV nicht überschreitet. Für größere Kurzschlussströme, die den DMM-Messbereich überschreiten oder einen größeren als 0,1V Spannungsabfall hervorrufen würden, sollte ein Außenstrommessshunt verwendet werden. Der Shuntwiderstandswert wird so gewählt, dass $R_S < 0.1V/I_{SHUNT}$ und $P_V > 0.1V I_{SHUNT}$.

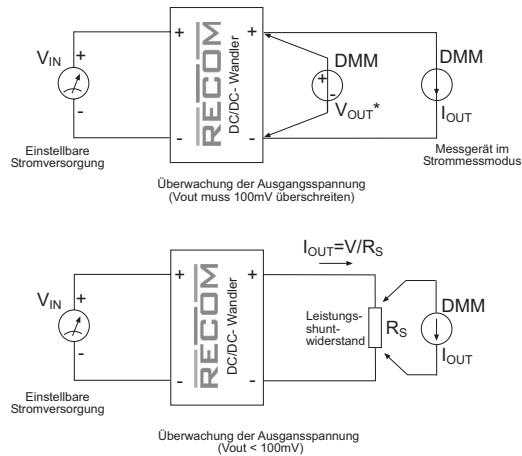


Abb. 3.8: Messung der Kurzschluss-Charakteristik

3.2.12 ON/OFF-Steuerung (Remote ON/OFF-Control)

Bei vielen Systemen ist ein Ein- und Ausschalten des DC/DC-Wandlers durch ein externes Signal wünschenswert. Dies kann aus Effizienzgründen zur Reduzierung des Energieverbrauchs, zur Steuerung einer An- und Abschaltsequenz oder aus Sicherheitsgründen erforderlich sein. Deshalb haben viele DC/DC-Wandler einen Steuereingang (ON/OFF-Pin), mit dem der Wandler in den Ruhezustand geschaltet werden kann. Der ON/OFF-On ist ein mit geringer Leistung ansteuerbarer Eingang, da jeder offene Kollektorausgang oder NPN-Transistor verwendet werden kann, um den Wandler zu steuern, weil er nur wenige Milliampere des Treiberstroms benötigt, um sogar einen Hochleistungswandler ein- oder auszuschalten.

Praktischer Hinweis

Die Art der Steuerlogik sollte beachtet werden. Positive Logik bedeutet, dass der Wandler bei hohem Pegel oder logischem ,1'-Signal AN und bei niedrigem Pegel oder logischem ,0'-Signal AUS ist. Der Steuereingang wird intern hochgezogen, sodass der Wandler AN ist, wenn er unbeschaltet bleibt. Diese Variante ist allgemein weiter verbreitet, weil der Wandler aktiv ist, wenn der Regelpin nicht benötigt wird.

Negative Logik bedeutet, dass dieser Wandler bei hohem Pegel oder logischem ,1'-Signal AUS und bei niedrigem Pegel oder logischem ,0'-Signal AN ist. Der Steuereingang wird intern hochgezogen, sodass der Wandler AUS ist, wenn er unbeschaltet ist. Dieser Regelungstypus ist in einer sicherheitskritischen Anwendung nützlich, bei der der Wandler nur eingeschaltet werden kann, wenn zuvor bestimmte Außenbedingungen erfüllt sind.

Für isolierte Wandler sollte das Datenblatt auch festlegen, welchen anderen Pin diese ON/OFF-Funktion als Referenz hat. In den meisten Fällen ist das Referenzpotential die Masse des Primärkreises, aber einige Wandler haben die ON/OFF-Funktion an der Ausgangsseite und die Sekundär- V_{OUT} als Referenz. Wenn das Einschaltsignal an der Primärseite entsteht, muss um den Ausgang zu schalten ein Trennungselement wie z.B. ein Optokoppler verwendet werden.

Um unerwünschtes wiederholtes Schalten bei langsam ansteigendem Steuersignal zu vermeiden, sollte das Ansprechverhalten des ON-/OFF-Eingangs eine Hysterese aufweisen. Beispielsweise, wenn eine externe RC-Verzögerungsschaltung am ON/OFF-Pin verwendet wird, um den Wandler, erst zu starten wenn sich die Eingangsspannung stabilisiert hat (Abb. 3.9). Das Datenblatt legt die Ansprechspannung V_{REMOTE} als Schwellenwert fest. Typische Werte sind logische ‚0‘, wenn $0 < V_{\text{REMOTE}} < 1,2\text{V}$, und logische ‚1‘, wenn $3,5 < V_{\text{REMOTE}} < 12\text{V}$. Dies bedeutet, dass sich der Wandler bei Erhöhung der Spannung V_{REMOTE} einschaltet, wenn die Spannung 1,2V überschreitet, und sich bei Absenkung der Spannung V_{REMOTE} ausschaltet, wenn die Spannung unter 3,5V abfällt. Die Diode im Zeitverzögerungsschaltkreis stellt sicher, dass sich der für das Timing verantwortliche Kondensator schnell entlädt, wenn die Eingangsspannung ausgeschaltet wird, sodass die Zeitverzögerung konstant bleibt, wenn die Versorgungsspannung wieder angelegt wird.

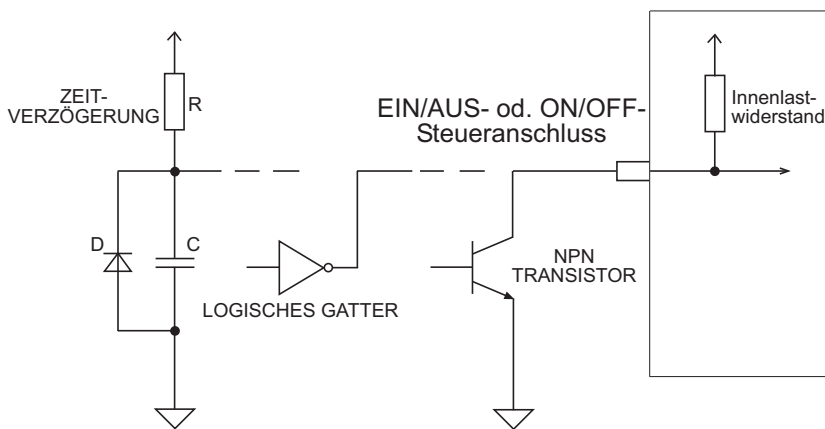


Abb. 3.9: Verschiedene Methoden zum Antrieb des Ein-/Aus-Regelpins

3.2.13 Isolationsspannung

In isolierten DC/DC-Wandlern sind die Primär- und Sekundärseite durch Transformator und Optokoppler getrennt, und zwar so, dass es keinen Gleichstromweg zwischen den beiden Schaltkreisen gibt, was als galvanische Trennung bezeichnet wird. Die Isolationsspannung beschreibt, bis zu welcher Potentialdifferenz diese Isolationsbarriere wirkt. Eine Prüfhochspannung ist entweder als DC-Spannung oder als Effektivwert der Wechselspannung bestimmt, wobei kein wesentlicher Strom fließen sollte, wenn diese Prüfhochspannung zwischen Primär- und Sekundärseiten angelegt wird.

Praktischer Hinweis

Bei diesen Prüfungen muss ein Hochspannungstester (HiPot) mit einem Strombegrenzungskreis eingesetzt werden, da gefährliche Hochspannungen verwendet werden. Führen Sie keine Hochspannungsprüfung auf dem elektrostatisch abgeschirmten Arbeitstisch durch, da die Oberfläche so bearbeitet wurde, dass sie elektrisch leitfähig ist. Vergewissern Sie sich, dass der Hochspannungstester eine Not-Aus-Taste besitzt, und stellen Sie sicher, dass der Erdanschluss am Tester ordnungsgemäß ausgeführt ist.

Das Prüfobjekt (DUT) muss von jedem Teil, das der Bediener unbeabsichtigt berühren könnte, entsprechend isoliert werden, und der Tester sollte über automatische Entladekreise zum Prüfspannungsentladen nach Abschluss der Prüfungen verfügen. Befolgen Sie sämtliche Herstelleranweisungen genau

Während der Wandler gleichstrommäßig galvanisch getrennt ist, fließt ein Leckstrom, wenn eine AC-Isolationsprüfung durchgeführt wird. Der AC-Strom fließt durch die kapazitive Kopplung zwischen den Wicklungen im Transformator und durch die EMC-Entstörkondensatoren, die an den Isolationsbarrieren platziert werden. Deshalb sollte für die AC-Hochspannungsprüfung nicht nur die Effektivspannung, sondern auch der zulässige Leckstrom festgelegt werden. Typische Grenzwerte sind 1mA oder 3mA, da jegliche Leckströme, die höher sind als diese, den Wandler auf Dauer zerstören könnten.

Infolge des AC-Leckstroms beansprucht die AC-Hochspannung die Isolationsbarriere stärker als eine äquivalente Ruhe-DC-Spannung. Die Einwirkung nimmt mit der Frequenz und Spannung zu, weil:

$$I_{LEAK} = \frac{dV(t)}{dt} C_{LEAK}$$

Gleichung 3.7: AC-Leckstrom

Praktischer Hinweis

Aus diesem Grund sollte der Wandler mit einem Nennwert von 1kVDC/1 Sekunde nur mit 700 VAC/1 Sekunde geprüft werden (wenn ein Signal von 50Hz verwendet wird). Dies klingt logisch, wenn man bedenkt, dass die Spitzenspannung mit einer Sinusspannung von 700 VACRMS 980 V beträgt. Mit Erhöhung der Frequenz nimmt dieser Leckstrom ebenfalls zu. Ein Prüfsignal von 100Hz erzeugt einen doppelt so hohen Leckstrom wie ein Prüfsignal von 50Hz. Ein Wandler, der also einen Nennwert von 1kVDC/1 Sekunde besitzt, sollte unter Verwendung eines Testsignals mit 100Hz nur mit 360VAC/1 Sekunde geprüft werden. Glücklicherweise ist eine Hochspannungsprüffrequenz von 50Hz Industriennorm, und obwohl die meisten Hersteller die verwendete Testfrequenz nicht erwähnen, kann man davon ausgehen, dass beim Vergleich der Isolationswerte im Datenblatt 50Hz verwendet wurden. RECOM stellt auf der Website ein nützliches Instrument zum Isolationsspannungsvergleich dar.

Auch bei länger als 1 Sekunde dauernden Prüfungen ist die Äquivalenz zwischen der DC- und der AC-Hochspannungsprüfung nicht so einfach. Eine 60 Sekunden dauernde Prüfung bringt aufgrund der Teilentladung (engl.: partial discharge - PD) genannten Phänomens, sehr viel mehr Belastung der Isolationsbarriere mit sich. Die Teilentladung beschreibt den kurzzeitigen Durchschlag infolge der hohen angelegten Spannung zwischen Koppelkreisen innerhalb des Isoliermittels infolge von Hohlstellen oder Brüchen.

Betrachten wir das nun anhand eines konventionellen Kupferlackdrahtes, wie er in Trafos üblicherweise verwendet wird. (Abb. 3.10). Normalerweise wird der Isolierlack in mehreren Stufen aufgetragen, weshalb es Unterbrechungen zwischen den Schichten und Hohlstellen innerhalb der Isolierung kommen kann.

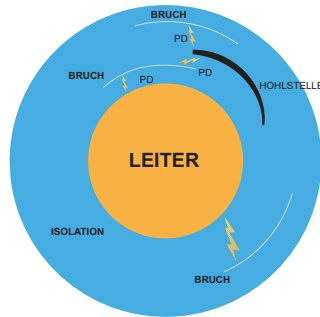


Abb. 3.10: Querschnitt eines Kupferlackdrahtes mit Teilentladungen (PD)-pfaden in der Isolation

Innerhalb der Isolationsschicht fließt nun kurzzeitig ein Strom, während Teilentladungen auftreten. Der Leiter ist immer noch isoliert, allerdings mit verringerter Dicke der Isolationsschicht. Die Hochspannung kann von einer Schwachstelle zu der folgenden überspringen, bis sie letztlich einen vollen Ein- oder Ausgangsisolationsdurchbruch hervorruft.

Das Stichwort hier heißt „letztlich“, denn es braucht Zeit, bis sich PD-Störungen vereinen, um einen Vollausschlag zu verursachen. Je länger die Hochspannung angelegt ist, desto wahrscheinlicher ist es, dass der Ausfall eintreten kann. Somit ist eine Hochspannungsprüfung mit einer Dauer von 1 Minute viel belastender als die typische 1-sekündige Werksprüfung. Ein Wandler mit einem Nennwert von 1000VDC/1 Sekunde sollte nur bei 500VAC/1 Min geprüft werden, um die Wahrscheinlichkeit solcher sich ansammelnder PD-Störungen zu verringern.

Praktischer Hinweis

Infolge des Risikos, dass während der Hochspannungsprüfung am Wandler ein permanenter Schaden verursacht wird, gibt es zwei wichtige praktische Punkte, die bei der Durchführung der Prüfungen einer besonderen Sorgfalt bedürfen. Erstens ist es wichtig, nicht zuzulassen, dass sich innerhalb des Wandlers Spannungsgradienten entwickeln, da diese die Nenndurchschlagspannungen der inneren Komponente schnell überschreiten könnten. Deshalb sollten vor Durchführung einer Hochspannungsprüfung alle Eingänge und alle Ausgänge kurzgeschlossen werden. Da Hochspannungsprüfungen die Wandlerisolation beanspruchen und einen kumulativen Schaden an der Isolation hervorrufen, ist es zweitens empfehlenswert, die Hochspannung für jede weitere Prüfung um 20% zu verringern.

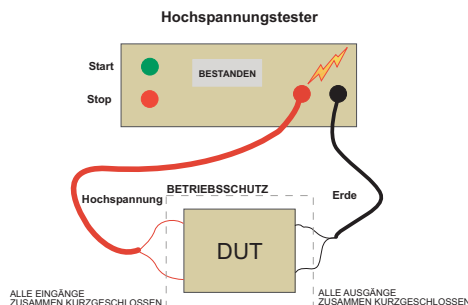


Abb. 3.11: Hochspannungsprüfung (Durchschlagsfestigkeit)

Der Vorteil einer Hochspannungsprüfung besteht darin, dass alle potentiellen Kurzschlussstromwege geprüft sind, wenn eine Hochspannung an den Ein- und Ausgang der Isolationsbarriere angelegt wird, sodass die Gesamtdurchschlagsfestigkeit des Wandlers durch ein positives Ergebnis garantiert werden kann. Der Nachteil der Prüfungen besteht darin, dass ein fehlerhaftes Ergebnis zerstörend ist – der Wandler muss dann getauscht werden.

Es gibt eine Alternative zur Hochspannungsprüfung – den PD-Tester. Die Prüfeinrichtung überwacht, die durch die Teilisolationsdurchschläge hervorgerufenen Spannungsspitzen und zeigt sie auf einem oszilloskopähnlichen Bildschirm als „scheinbare Ladung“ oder äquivalente Ladungsträgerinjektion, welche eine entsprechende Testspannung hervorrufen würde. Scheinbare Ladung wird in Pikocoulomb gemessen, weshalb es sich hier um ein sehr schonendes Prüfverfahren handelt. Der Vorteil der PD-Prüfung besteht darin, dass die Häufigkeit der PD-Ereignisse überwacht werden kann, da die Prüfspannung erhöht ist und ein potentieller Isolationsdurchbruch vorausgesagt werden kann, bevor er auftritt; somit kann die Prüfung angehalten werden, bevor der Wandler beschädigt wird.

Die Ergebnisse der PD-Prüfung sollten sorgfältig ausgewertet werden, da es vorkommen kann, dass zunächst viele falsche Werte abgelesen werden, bevor ein gültiges Ergebnis ermittelt werden kann. Deshalb ist eine "Einschwing"zeit notwendig, in welcher sich die Ladungen ausgleichen können; die Ablesungen sollten nur während der letzten 10 Sekunden der Prüfung vorgenommen werden (siehe Abb. 3.12). Für die Betriebsprüfung kann eine berichtigte Version des Prüfprotokolls mit einer Prüfdauer von nur 1 Sekunde und einer Spannung von $1,875 \times V_{\text{RATED(RMS)}}$ verwendet werden.

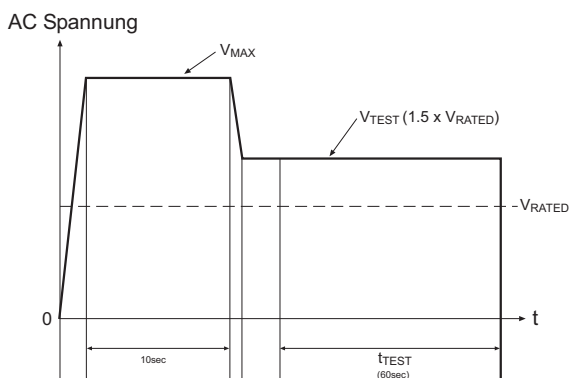


Abb. 3.12: Teilentladungstest (PD)

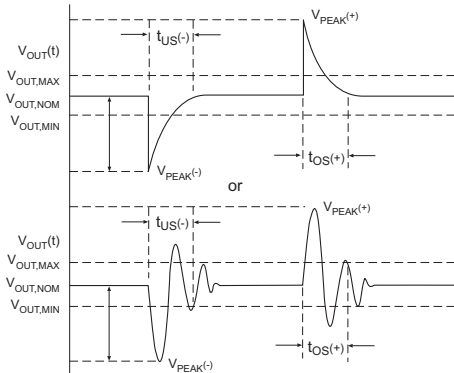
3.2.14 Isolationswiderstand und -kapazität

Ein- und Ausgangswiderstand und -kapazität müssen mit einem AC-Signal gemessen werden. Der Isolationswiderstand wird in der Regel mit 500VDC mit Hilfe eines Widerstandsmessers gemessen und beträgt üblicherweise $10 \text{ G}\Omega$ oder mehr. Die Isolationskapazität muss bei einer hohen Frequenz von 1MHz gemessen werden, um den Einfluss von onboard-Filterkondensatoren auszuschließen. Die Isolationskapazität kann je nach Aufbau des Wandlers zwischen 5pF und 1500pF variieren. Wie bei allen Messungen kleiner Ströme, können Luftfeuchtigkeit und Temperatur die Ergebnisse stark beeinflussen.

3.2.15 Dynamische Belastungskennlinie

Die Dynamische Belastungskennlinie (engl.: Dynamic Load Response - DLR) bestimmt die Reaktion des DC/DC-Wandlers auf eine sprunghafte Laständerung. Sie kann auf zwei verschiedene Weisen definiert werden: durch die Zeit, die die Ausgangsspannung benötigt, um innerhalb des vorgegebenen Toleranzbereichs zum Nennwert der Ausgangsspannung zurückzukehren, oder als maximale Abweichung der Ausgangsspannung von der Nennausgangsspannung. Um eine vollständige Definition der DLR zu erhalten, müssen beide Werte bekannt sein. In den meisten Datenblättern wird jedoch nur die Einschwingzeit angegeben. Darüber hinaus verwenden einige Hersteller Lastwerte zwischen 25% und 100%, andere Werte zwischen 50% und 100% Last, und wieder andere Hersteller nennen nur "25% stufenweise Änderung", ohne die Grenzen anzugeben. Daher ist ein direkter Vergleich zwischen Datenblättern unterschiedlicher Hersteller nicht möglich. Die einzige Alternative um gesicherte Werte zu erhalten, besteht darin den Wandler selbst zu prüfen.

Bei allen Wandlern tritt bei plötzlicher Lastabnahme ein Überspringen und bei plötzlicher Lastzunahme ein Unterschwingen auf. Die Einschwingzeit (die längste Zeit von $T_{\text{OVERSHOOT}}$ oder $T_{\text{UNDERSHOOT}}$) hängt hauptsächlich von der Kompensationsschaltung im PWM-Controller ab. Die Kompensation muss einen Kompromiss finden, zwischen einer schnellen Reaktion auf eine sprunghafte Änderung und einer möglichen Überreaktion in Form eines schwingenden Ausgangs. Das ideale Ansprechverhalten ist "aperiodisch", d. h. der Wert der Ausgangsspannung wird lediglich einmal in eine Richtung abgelenkt.



Das obere Signal ist aperiodisch, das untere Signal zeigt einen schwingenden Ausgang infolge einer schlecht gedämpften Kompensationsschaltung.

Abb. 3.13: Mögliche Reaktionen des geregelten Wandlers auf eine sprunghafte Laständerung

Viele Elektroniklasten verfügen über eine eingebaute Sprungfunktion, um zwischen zwei voreingestellten Lasten automatisch umzuschalten. Falls Sie jedoch keinen Zugriff auf eine Elektroniklast haben, zeigt Abb. 3.14, wie ein Testaufbau für dynamische Belastung aus zwei Lastwiderständen und einem durch einen Rechteckgenerator getriebenen FET einfach aufgebaut werden kann.

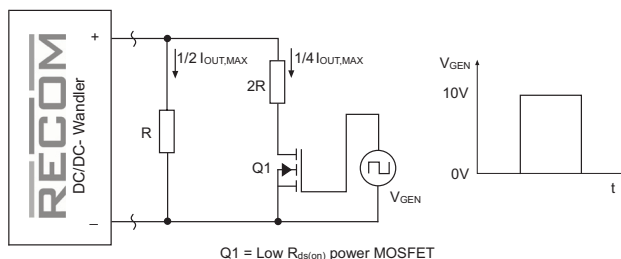


Abb. 3.14: Testaufbau zur Prüfung des dynamischen Belastungsverhaltens

Es gibt einige Anwendungen, bei denen sowohl eine sehr konstante Ausgangsspannung als auch eine hohe Ansprechgeschwindigkeit bei einem Lastsprung ohne jegliches Schwingen der Ausgangsspannung erforderlich sind. Zum Beispiel ergeben sich bei vielen Digitalschaltungen schnelle Laständerungen, die Ausgangsspannung muss jedoch dennoch exakt geregelt sein. Wenn die Laständerung vorausgesagt oder ermittelt werden kann, ist es möglich, die Kompensationsschaltung während der Laständerung von "langsam" auf "schnell" umzuschalten.

Mit analogen Reglern lässt sich dies nicht so leicht realisieren, mit einem digitalen Controller kann DLR jedoch programmierbar gemacht werden. Diese Fähigkeit ist eine der größten Vorteile der digitalen Regelung gegenüber der analogen.

So wie sich die Ausgangsspannung mit plötzlicher Laständerung ändert, ändert sich die Ausgangsspannung auch, wenn sich die Eingangsspannung plötzlich ändert. Da sprunghafte Änderungen der Eingangsspannung nur in wenigen Anwendungen vorkommen, wird dieser Parameter in den Datenblättern selten definiert. Gegebenenfalls lässt sich ein Testaufbau verhältnismäßig leicht realisieren, wenn die Laborstromversorgung über einen externen Steuereingang, der an einen Rechteckgenerator angeschlossen werden kann, verfügt.

3.2.16 Ausgangsrestwelligkeit

Alle DC/DC-Wandler haben einen gewissen Anteil an Ausgangsrestwelligkeit (engl.: Output Ripple/Noise). Der Hauptanteil an Restwelligkeit entsteht durch die Lade- und Entladeströme in der Ausgangsfilterkapazität; und je nach Topologie entspricht ihr Frequenzwert üblicherweise entweder der Arbeitsfrequenz oder der doppelten Arbeitsfrequenz.

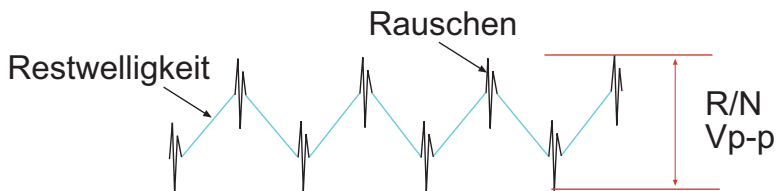


Abb. 3.15: Ausgangsrestwelligkeit

Bei jeder Zustandsänderung der Schaltelemente rufen parasitäre Effekte Spannungsspitzen (Rauschen) auf dem Hauptanteil der Restwelligkeit hervor und überlagern diese. Dies tritt an jeder Spitze und in jedem Tal der Restwelligkeit auf. Die Frequenzen der Schalttransienten liegen normalerweise in MHz-Bereichen und besitzen Größenordnungen, die über der Frequenz der Ausgangsrestwelligkeit liegen. Die Gesamtwellenform ergibt den Wert der Ausgangsrestwelligkeit (R/N); dieser wird üblicherweise in Millivolt Spitze-Spitze (mVp-p) gemessen wird.

ANMERKUNG: Das korrekte Messverfahren ist in Kapitel 4.3 beschrieben.

Außerdem überlagert eine viel langsamere Schwingung, die durch die Ausgangsspannungsregelung hervorgerufen wird, diese Ausgangsrestwelligkeit. Bei konstanter Last und Eingangsspannung mäandert die Ausgangsspannung mit einer Frequenz im einstelligen Hz-Bereich innerhalb des Toleranzbereichs nach oben und nach unten.

Dieser "Scheren"-Effekt wird durch die Hysterese im Regelkreis verursacht und wird in der Regel in dem Wert der Ausgangsrestwelligkeit in den Datenblättern ignoriert, da er Teil der Spezifikation der Messgenauigkeit der Ausgangsspannung ist und deshalb nicht in die Spezifikation der Ausgangsrestwelligkeit einfließt.

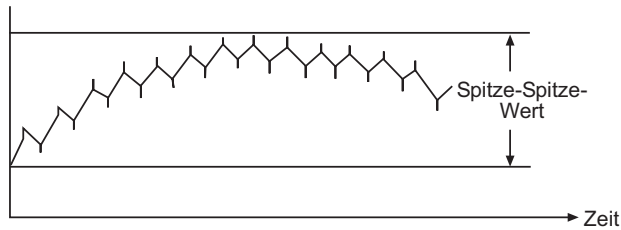


Abb. 3.16: "Scheren" der Ausgangsspannung

Es scheint einfach, die Ausgangsrestwelligkeit durch Hinzufügen von Ausgangskondensatoren zu verringern zu wollen. Jedoch ist es fast unmöglich, diese vollständig herauszufiltern – obwohl man den Spitze-Spitze-Wert auf diese Weise ein wenig verringern kann. Tatsächlich ist in Wandlern, die Takt-für-Takt-Regelung verwenden, ein gewisses Maß an Ausgangswelligkeit für die entsprechende Regelung notwendig.

Praktischer Hinweis

Ein effektiveres Verfahren zur Verminderung der Ausgangswelligkeit besteht darin, den Ausgang mit einem linearen Spannungsregler nachzuregulieren. Die Störspannungsunterdrückung (PSRR-Wert) eines linearen Spannungsreglers ist sehr hoch (bis zu 70dB), was ihn zu einem sehr effektiven Welligkeitsfilter macht.

3.3 Zum Verständnis thermischer Parameter

3.3.1 Introduction

Da das thermische Design eine wichtige Rolle bei der Optimierung der Systemperformance spielt, ist für die Wahl des am besten geeigneten DC/DC-Wandlers eine korrekte Einschätzung des Wärmezustands des Systems unerlässlich. Um eine umfassende Aussage über das thermische Verhalten des Wandlers zu machen, sollte das technische Datenblatt nicht nur die Grenzwerte der zulässigen Umgebungsbetriebstemperatur, sondern auch das Temperatur-Derating, die interne Verlustleistung, die maximale Gehäusetemperatur und die thermische Impedanz bestimmen. Die Verlustleistung kann anhand des Wirkungsgrades berechnet werden. Fehlt jedoch die thermische Impedanz im Datenblatt oder muss diese unter realen Einsatzbedingungen der Anwendung genau bekannt sein, dann muss sie in Wärmekammertests ermittelt werden.

Selbst unter den geregelten Umgebungsbedingungen einer Wärmekammer erfordert die Ermittlung des Temperaturverhaltens modularer DC/DC-Wandler sehr sorgfältige Messverfahren. Da sogar sehr geringe Luftströme die Messergebnisse wesentlich verzerren, sollte das Prüfobjekt (Device Under Test, DUT) – um Luftzug vom

Kammer-Luftzirkulationsgebläse zu vermeiden – innerhalb der Kammer in einen Karton platziert werden. Die Umgebungstemperatur innerhalb des Kartons sollte mit einem geeichten Sensor gemessen werden, der so positioniert wird, dass die vom Wandler erzeugte Wärme die Ablesung nicht direkt beeinflusst. Die Temperatur des Wandlergehäuses sollte an der heißesten Stelle ($T_{C,MAX}$) wie vom Hersteller definiert oder anhand von Infrarotkamerabildern ermittelt, gemessen werden. Bei sehr kleinen Wandlern kann auch die Befestigung des Temperatursensors selbst die Ergebnisse durch Ableitung zusätzlicher Wärme vom Wandler entlang der Sensorleitungen beeinflussen, weshalb ein Temperatursensor mit möglichst kleiner Kontaktfläche eingesetzt werden sollte.

Da die Eigenerwärmung eines Kleinleistungswandlers eine unwesentliche Wärmequelle darstellt, kann es bei diesen besonders schwierig sein, zuverlässige Messergebnisse zur thermischen Impedanz zu erhalten. Der zulässige Betriebstemperaturbereich wird dann in den meisten Fällen durch die Temperaturgrenzwerte interner Komponenten bestimmt. Dies kann durch Befestigung der Temperaturmessgeber an den kritischen Komponenten gemessen werden, um den Temperaturanstieg gegenüber der Umgebung zu messen und dann die Sicherheitsgrenzen durch Extrapolieren der Ablesewerte – in 10°C -Umgebungstemperaturschritten – zu berechnen. Für gekapselte Wandler müssen die Temperaturmessgeber vor dem Vergießen angebracht werden, um genaue Ablesewerte zu ermitteln.

Bei Hochleistungswandlern kann die thermische Impedanz durch Messung des Temperaturanstiegs unter freier Konvektion (stehende Luft = keine Luftströmung) und Berechnung der internen Verlustleistung bestimmt werden. Um die Wärmeübertragungszahl für unterschiedliche Luftstromgeschwindigkeiten bei erzwungener Konvektion zu berechnen, kann auch die thermische Impedanz verwendet werden.

Schließlich beeinträchtigen auch niedrige Temperaturen das Betriebsverhalten eines Wandlers. Die untere Betriebstemperaturgrenze hängt von einem von drei Faktoren – je nachdem, welcher zuerst auftritt – ab: dem minimalen Temperaturnennwert der verwendeten Bauteile; einer verringerten Verstärkung oder Offset der Vorspannung des PWM-Schaltkreises, die das Starten des Wandlers verhindern; oder einem mechanischen Fehler, der durch fehlangepasste Wärmeausdehnung verursacht wird.

3.3.2 Thermische Impedanz

Um die Analyse des Wärmezustands des DC/DC Wandler zu beginnen, muss zunächst die thermische Impedanz betrachtet werden.

Thermische Impedanz oder Wärmewiderstand ist ein Maß der Wirksamkeit des Wärmeaustauschs zwischen einer internen Wärmequelle, wie dem Transformatorkern oder Halbleiterübergang, und der Umgebung. Betrachten wir zum Beispiel den Schalt-FET. Die Wärmequelle ist der Halbleiterübergang TJ. Die Wärme wird zum Transistorgehäuse (TB), dann durch das Vergussmaterial zum Wandlergehäuse (TC) und anschließend vom Gehäuse zur Umgebung (TAMB) abgeleitet. Jede dieser Stufen hat eine thermische Impedanz θ , gemessen in $^{\circ}\text{C}/\text{W}$, und einen Wärmewiderstand RTH, gemessen in $^{\circ}\text{K}/\text{W}$. Diese beiden Werte sind in der Praxis vollständig austauschbar.



Abb. 3.17: Thermische Impedanzkette

Aus den oben genannten thermischen Impedanzen kann der Benutzer nur die letzte Impedanz in der Kette, θ_{CA} , beeinflussen, da die anderen beiden Impedanzen konstruktionsbedingt als Teil des Wärmedesigns des Wandlers festgelegt sind.

Der Temperaturanstieg des Wandlers kann unter Verwendung folgender Gleichung berechnet werden:

$$\Delta T_{RISE} = P_{DISS} \theta_{CA}, \text{ where } P_{DISS} = \frac{P_{OUT}}{\eta} - P_{OUT}$$

Gleichung 3.8: Berechnung des Temperaturanstiegs

Beispiel: Der Wandler RECOM RP15-4805SA hat eine Ausgangsleistung von 15W, einen Wirkungsgrad von 88% und eine thermische Impedanz zwischen Gehäuse und Umgebung von 18,2°C/W. Die maximal zulässige Gehäusetemperatur beträgt 105°C. Die Verlustleistung = 15/0,88 - 15 = 2,04W und der zugehörige Temperaturanstieg des Gehäuses über der Umgebungstemperatur = 37°C. Somit beträgt die maximal zulässige Umgebungstemperatur 105°C - 37°C = 68°C.

Wenn die thermische Impedanz nicht bekannt ist, kann sie durch Messung ermittelt werden. Für einen Näherungswert ist eine Wärmekammer nicht notwendig. Ein geeigneter Testaufbau ist in Abb. 3.18 dargestellt. Wie bei allen thermischen Messungen, sollte vor Durchführung jeglicher Ablesung ausreichend lange abgewartet werden, damit ein Temperatúrausgleich stattfinden kann.

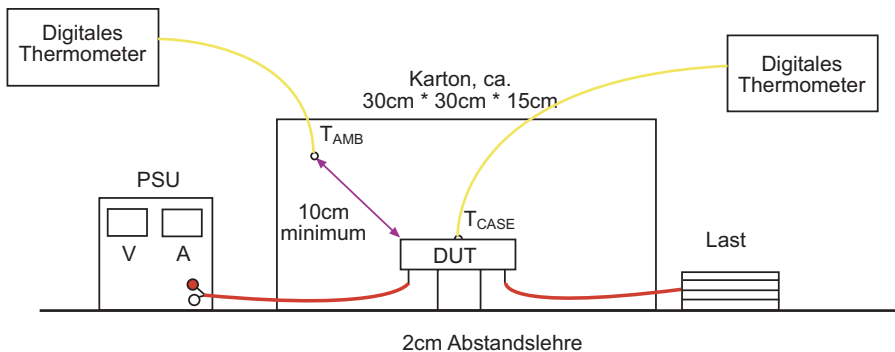


Abb. 3.18: Testaufbau zur Messung der thermischen Impedanz zwischen Gehäuse und Umgebung

Die thermische Impedanz kann aus der Umordnung der Gleichung 3.8 abgeleitet werden:

$$\theta_{CA} [^{\circ}\text{C/W}] = \frac{\Delta T_{RISE}}{P_{DISS}}$$

Da die Verlustleistung (Differenz zwischen der Eingangs- und Ausgangsleistung) bekannt ist, lässt sich die thermische Impedanz aus der Messung des Anstiegs der Gehäusetemperatur des Wandlers gegenüber der Umgebung bestimmen.

3.3.3 Temperatur-Derating

Alle DC/DC-Wandler geben Leistung intern in Form von Wärme ab und werden daher wärmer als ihre Umgebung. Solange diese zusätzliche Wärme in die Umgebung abgeleitet werden kann, kann der Wandler bei voller Leistung arbeiten. Mit Anstieg der Umgebungstemperatur wird es jedoch immer schwieriger, diese überschüssige Wärme abzugeben. Bei einer bestimmten Umgebungstemperatur erreicht der Wandler seine maximale Temperaturgrenze, und jeder weitere Anstieg der Umgebungstemperatur muss durch Verminderung der innerhalb des Wandlers als Verlustleistung umgesetzten Energiemenge durch eine Lastverminderung ausgeglichen werden. Diesen Vorgang nennt man Temperatur-Derating.

Abb. 3.19 zeigt das Beispiel einer Derating-Kurve, wiederum für den im vorherigen Beispiel verwendeten Wandler RECOM RP15-4805SA. Der Wandler kann bei voller Leistung über einen Umgebungstemperaturbereich von -40°C bis $+68^{\circ}\text{C}$ eingesetzt werden. Wenn die Anwendung die Spezifikation besitzt, dass die maximale Umgebungstemperatur höchstens 85°C betragen darf, muss die Last des Wandlers auf 55% verringert werden, um einen Betrieb im Bereich von -40°C bis zu $+85^{\circ}\text{C}$ zu ermöglichen.

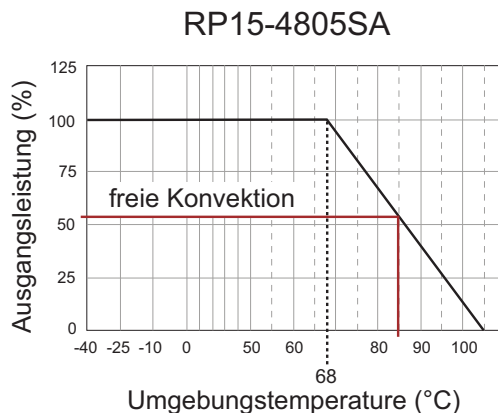


Abb. 3.19: Beispiel einer Derating-Kurve (DC/DC-Wandler der RP15-Serie)

Es gibt einen praktischen Grenzwert für die mögliche Derating-Zahl. Die Derating-Kurve setzt voraus, dass der Wirkungsgrad stabil bleibt, während sich die Last verringert, was jedoch bei kleinen Lasten nicht der Fall ist. Tatsächlich ist ein Derating, das wesentlich unterhalb einer Last von 40% liegt, nicht zielführend, da eine Verminderung der Verlustleistung infolge einer Lastsenkung, durch Erhöhung der Verlustleistung infolge eines niedrigeren Wirkungsgrades, negiert wird. Abb. 3.20 zeigt, wie die Kurve der Verlustleistung immer flacher wird, bei kleiner Last sogar anfangen kann, wieder anzusteigen.

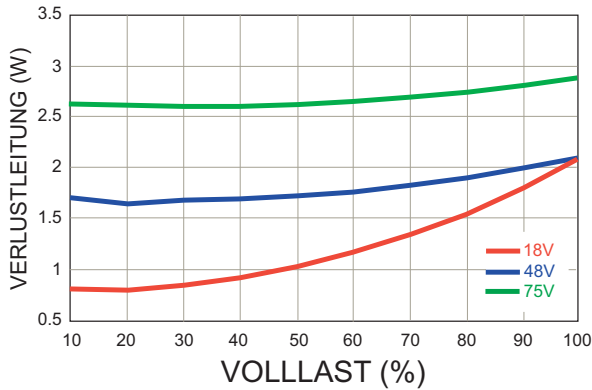


Abb. 3.20: Interne Wärmeverluste mit Last und V_{IN}

Durch Montieren eines Kühlkörpers an einem Wandler lässt sich die Wärmeabfuhr an die Umgebung verbessern (Verringerung von θ_{CA}) und somit der maximale Betriebstemperaturbereich erweitern. Obwohl dieses Element als Kühlkörper bezeichnet wird, absorbiert es keine Wärme. Die Wärme, die an den Kühlkörper abgegeben wird, muss letztlich doch noch an die Umgebung abgeleitet werden. Die Funktion des Kühlkörpers besteht deshalb darin, die wirksame Oberfläche des Wandlers zu vergrößern. Der in diesem Beispiel verwendete RP15-4805SA ist ein sehr kompakter Wandler mit einer Gehäusegröße von 1"× 1". Somit ist der Kühlkörper, der in den Wandler eingebaut werden kann, ebenfalls sehr kompakt und vergrößert die Fläche nicht wesentlich. Generell bringen Aufsteckkühlkörper eine Erhöhung des maximalen Arbeitstemperaturbereichs von lediglich 5 bis 10°C, es sei denn, dass sie selbst mittels anderer Verfahren, wie z. B. erzwungener Konvektion, gekühlt werden.

RP15-4805SA mit Kühlkörper

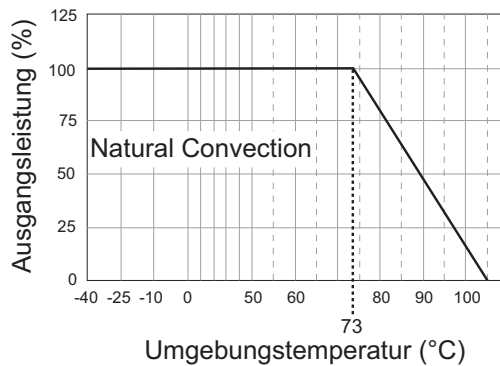


Abb. 3.21: Derating-Kurve mit kleinem Aufsteckkühlkörper

Ferner können Kühlkörper nur bei Wandlern mit Metallgehäuse oder auf einer sog. Baseplate aufgebauten Wandlern eingesetzt werden. Das Montieren eines Kühlkörpers auf einem Wandlergehäuse aus Kunststoff ist nicht zielführend, da der Kühlkörper eine schlechte thermische Anbindung an das Kunststoffgehäuse aufweist und die Konvektion blockiert. Zusammenfassend kann gesagt werden, Derating ermöglicht den maximalen Betriebstemperaturbereich um zusätzliche 10 bis 15°C zu erweitern.

Darin besteht jedoch auch schon die Grenze bei vielen Anwendungen. Wärmeabfuhr mit Kühlkörpern kann hilfreich sein. Ist der Kühlkörper jedoch ähnlich dimensioniert wie der Wandler, kann nur eine zusätzliche Erweiterung des maximalen Betriebstemperaturbereiches um 5 bis 15°C erwartet werden.

Um den maximalen Temperaturbereich wesentlich zu erweitern, ist eine Zwangskühlung erforderlich.

3.3.4 Zwangskühlung

Durch Hinzufügen von thermischer Advektion zur freien Konvektion verringert erzwungene Konvektionskühlung (Lüfterkühlung) die thermische Impedanz θ_{CA} . Advektion hängt von der Masse der Luft, die am Wandler pro Sekunde vorbeifließt, sowie von der Wirbelströmung des Luftstroms ab. Wenn die freie Konvektion eine normalisierte thermische Impedanz von 1,00 besitzt, verringert sich die thermische Impedanz mit Erhöhung des laminaren Luftstroms (angegeben in linearen Fuß pro Minute, LFM) wie folgt:

10 LFM (freie Konvektion)	1,00
100 LFM	0,67
200 LFM	0,45
300 LFM	0,33
400 LFM	0,25
500 LFM	0,20

Tabelle 3.2: Normalisierte Verminderung der thermischen Impedanz mit zunehmendem Laminarluftstrom

Betrachten wir noch einmal den Wandler, der in unserem Beispiel zur Ermittlung der Derating-Kurve in Abb. 3.19 verwendet wurde. Die Temperaturanstiegsgleichung ist immer noch gültig: $\Delta T_{RISE} = P_{DISS} \theta_{CA}$, und die interne thermische Verlustleistung ist immer noch dieselbe, nämlich 2 W. Mit freier Konvektion betrug der θ_{CA} -Wert 18,2°C/W, was einen Temperaturanstieg von 37°C ergibt und zu einer maximal zulässigen Betriebstemperatur bei Volllast von 68°C führt. Mit 100LFM erzwungener Konvektion würde der θ_{CA} -Wert mit 0,67 multipliziert bzw. 12,2°C/W betragen, wobei sich ein reduzierter Temperaturanstieg von 24,4°C ergibt, was zu einer maximal zulässigen Betriebstemperatur bei Volllast von 81°C führt.

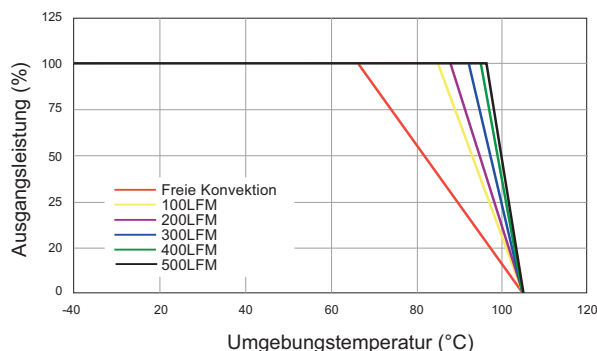


Abb. 3.22: Auswirkung des zunehmenden Luftstroms auf die Derating-Kurve

3.3.5 Leitungs- und Strahlungskühlung

Es gibt noch andere Transportmechanismen für thermischen Übergang als Konvektion oder Advektion. Wärme kann von einem Wandler durch Leitung und Abstrahlung abgeleitet werden.

Wärmeleitung ist die Wärmeübertragung von einem Objekt auf ein anderes über ein Temperaturgefälle durch Direktkontakt. Der Transportmechanismus besteht in Phononen oder in Energieübertragung von einem Molekül zu einem anderen durch Kristallgitterschwingungen. Die Übertragungsgeschwindigkeit, oder thermische Strömung, hängt vom Temperaturgefälle und der Wärmeleitfähigkeit des Materials (gemessen in $\text{Wm}^{-1}\text{C}^{-1}$) ab.

Wandler, die mit einer Baseplate ausgestattet sind, nutzen Wärmeleitung als primäres Kühlverfahren. Alle Wandler können jedoch vom Wärmeaustausch durch Leitung über Anschlusspins in PCB-Leiterbahnen profitieren. Die thermische Impedanzkette bei einem über Baseplate gekühlten Wandler ist in Abbildung 3.23 dargestellt:

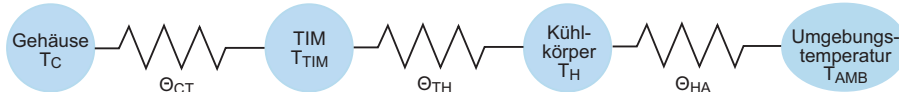


Abb. 3.23: Thermische Impedanz eines über Baseplate gekühlten Wandlers

Da der primäre Wärmestrom beim Direktkontakt entsteht, ist es entscheidend, dass die Kontaktfläche zwischen der Baseplate und dem Kühlkörper möglichst groß ist. Selbst sehr kleine Fehler führen zu Luftspalten, über die kein Wärmeaustausch durch Leitung stattfinden kann. Wenn die Baseplate und der Kühlkörper nicht bis auf 5 mil (0,125 mm) vollkommen plan sind, ist die Wärmeleitung stark beeinträchtigt. Um einen guten physikalischen und thermischen Kontakt zu gewährleisten, ist daher der Einsatz eines thermischen Schnittstellenmaterials (TIM), z. B. Wärmeleitpaste oder Gap-Pad, üblich.

Der Strahlungswärmeaustausch ist die Wärmeübertragung eines Körpers mittels Infrarotstrahlung. Bei der fühlbaren Sonnenwärme handelt es sich einzig und allein um Strahlungswärme, da das Vakuum des Weltraums jeglichen Wärmeaustausch durch Leitung und Konvektion blockiert. Ein Wandler kann auch im Vakuum Energie durch Strahlungswärme abgeben. Die niedrige Gehäusetemperatur des Wandlers von ca. 300 bis 400K führt jedoch dazu, dass diese Art des Wärmetransports im Vergleich zum Wärmeaustausch durch Leitung oder Konvektion sehr gering ist. Bei großen Höhen über Meeresniveau wird Kühlung sowohl mittels Leitung als auch mittels Abstrahlung immer wichtiger, da der Hauptnachteil der Wärmekonvektion darin besteht, dass sie von der Luftstrommasse abhängt und somit mit fallendem Luftdrucks sinkt.

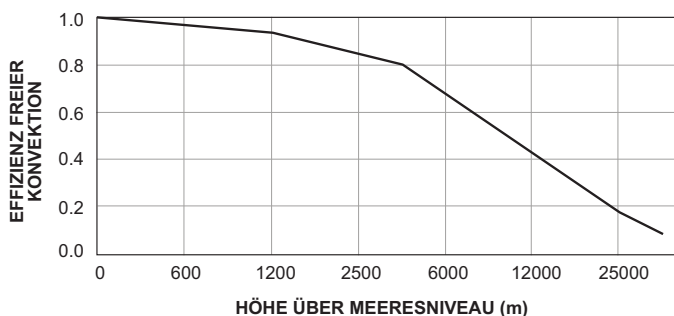


Abb. 3.24: Effizienzminderung der Konvektionskühlung mit Höhe über Meeresniveau

4. DC/DC-Wandler-Schutzmaßnahmen

4.1 Einleitung

Wie im Vorwort erwähnt, besteht eine der Funktionen des DC/DC-Wandlers darin, die Anwendung zu schützen. Im einfachsten Fall erfolgt dieser Schutz durch eine Anpassung der Last an die primäre Stromquelle und einer Stabilisierung der Ausgangsspannung gegen Eingangsüber- und Unterspannungen. Der DC/DC-Wandler ist aber auch ein wesentliches Element, das den Kurzschlusschutz des Systems gewährleistet. Die Begrenzung einer Ausgangsüberlastung und Kurzschlusschutz beispielsweise schützen nicht nur den Wandler vor Schäden durch mögliche Fehlerzustände in der Last, sondern können auch die Last durch die Ausgangsleistungsbegrenzung während eines Funktionsausfalls vor weiteren Schäden schützen. In einer Anwendung mit mehreren identischen Schaltkreisen oder Kanälen, wobei jeder separat durch individuelle DC/DC-Wandler versorgt wird, beeinflusst ein Fehler in einem Ausgangskanal nicht die anderen Ausgänge, was das System gegenüber Einzelfehlern tolerant macht. Andere Schutzfunktionen des Wandlers, wie die Übertemperaturabschaltung, werden vor allem entwickelt, um den Wandler vor dauerhaften durch interne Komponentenüberhitzung hervorgerufenen Schäden zu schützen. Ein Sekundäreffekt besteht jedoch in einem Abschalten, wenn die Umgebungstemperatur sehr stark ansteigt, es werden somit auch die Komponenten in der versorgten Anwendung vor Ausfall aufgrund Überhitzung geschützt.

Die Isolation zwischen Ein- und Ausgang trennt Erdschleifen auf, entfernt die Störquelle und erhöht die Systemzuverlässigkeit durch den Schutz der Anwendung vor Schäden aufgrund von Transienten. Beseitigung von Rückkopplungsverzerrungen der Stromversorgung ist ein wichtiger Aspekt des DC/DC-Wandlerschutzes. Betrachten wir zum Beispiel einen Hochleistungs-DC Motor Speed Controller. Der Drehzahlregelkreis benötigt eine stabile, rauschfreie Versorgung, um die Motordrehzahl stufenlos zu regulieren. Aber die vom Motor gezogenen hohen DC-Ströme können wesentliche Einschwingvorgänge auslösen, die in den Drehzahlregelkreis einwirken und Synchronisationsstörung oder Instabilität hervorrufen könnten. Ein isolierter DC/DC-Wandler kann nicht nur eine stabile, rauscharme Versorgung für den Drehzahlregelkreis bereitstellen, sondern auch einen Schutz des Motors bieten, da Störungen der Feedbackschleife und dadurch ausgelöste unerwünschte und sprunghafte Steuersignale die zugehörige Antriebskette zerstören könnten.

DC/DC-Wandler sind jedoch auch aus elektronischen Bauteilen aufgebaut, die wie jeder elektronische Schaltkreis, der außerhalb seiner Spannungs-, Strom- und Temperaturgrenzen betrieben wird, für Störungen anfällig sind. In diesem Kapitel werden Schutzmaßnahmen untersucht, die notwendig sein können, um den Wandler selbst vor Schaden zu bewahren.

4.2 Verpolungsschutz

DC/DC-Wandler sind nicht gegen falsche Polung geschützt. Ein Vertauschen der V_{IN+} und V_{IN-} -Pole wird mit hoher Wahrscheinlichkeit unmittelbar zum Ausfall führen, weshalb darauf geachtet werden soll, dass jegliche Eingangspins oder Akkuanschlüsse richtig gepolt angeschlossen werden. Wenn die Primärspannungsversorgung ein Transformator ist, könnte eine Störung der Gleichrichterdiode einen negativ werdenden Ausgang hervorrufen, was dann auch den Ausfall der DC/DC-Wandler zur Folge haben wird.

Der Hauptgrund, warum DC/DC-Wandler ausfallen, wenn sie verpolt sind, ist die Gehäusediode im FET. Diese Substratdiode leitet, wenn sie verpolt angeschlossen ist, und erlaubt einen sehr hohen Stromfluss I_R , was zur Zerstörung von Komponenten auf der Primärseite führen kann. Um diese potentielle Gefahr zu vermeiden, sind mehrere Varianten möglich.

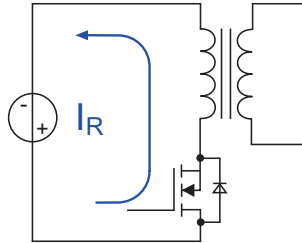


Abb. 4.1: Elektrischer Strom bei Verpolung

4.2.1 Verpolungsschutz durch Seriendiode

Der einfachste Weg, einen DC/DC-Wandler vor Schäden durch Verpolung zu schützen, besteht darin, eine Seriendiode hinzuzufügen. Abb. 4.2 zeigt den Schaltkreis. Wenn die Versorgungsspannung verpolt wird, sperrt die Diode D_1 den negativen elektrischen Strom, sodass kein unzulässiger Fehlerstrom durch den Eingangsstromkreis des DC/DC-Wandlers fließen kann. Es ist offensichtlich, dass durch Ersetzen der Diode durch einen Brückengleichrichter, der Wandler unabhängig von der Eingangsspannungspolarität funktioniert.

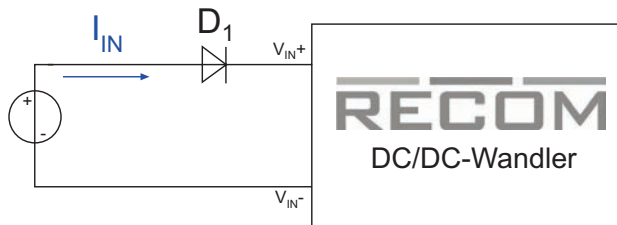


Abb. 4.2: Verpolungsschutz durch Seriendiode

Besonders bei niedrigen Eingangsspannungen hat ein Schutz durch Seriendioden den Nachteil, dass die Vorwärtsspannung an der Diode abfällt. Je nach Auswahl der Diode ist ein Durchlassspannungsabfall von 0,2 V bis 0,7V mit einem zugehörigen Leistungsverlust von $= V_F \times I_{IN}$ zu erwarten, was sowohl den Wandlungswirkungsgrad als auch den nutzbaren Eingangsspannungsbereich verringert. Wenn der Eingangsstrom 1A beträgt, entsteht an einer Standardleistungsdiode mit $V_F = 0,5V$ eine Verlustleistung von 0,5W, was ungefähr einem Viertel der Verlustleistung eines typischen 15-W-Wandlers entspricht und somit den Gesamtwirkungsgrad um 20% verringert.

In einigen Anwendungen ist der Spannungsabfall an der Diode ein Vorteil. In Rallyeautos wird zur Erhöhung der Helligkeit der Scheinwerfer häufig eine 16-V-Batterie verwendet. Die Lichtmaschine wird so modifiziert, dass sie 11 bis 20V außerhalb des Bereichs eines standardmäßigen 9- bis 18-V-DC/DC-Wandlers liefert. Durch Einsatz von drei in Reihe geschalteten Dioden kann der nutzbare Eingangsspannungsbereich herabgesetzt werden, um dem Standard-18-V-Eingangsspannungsbereich eines DC/DC-Wandlers zu entsprechen.

4.2.2 Verpolungsschutz durch Überbrückungs- oder Shuntdiode

Eine Alternative zur Seriendiode bietet der Verpolungsschutz durch Überbrückungs- oder Shuntdioden. Ein Durchlassspannungsabfall an der Diode wird vermieden, jedoch muss die Primärspannungsversorgung entweder über einen Überlastungsschutz oder eine Schmelzsicherung verfügen (Abb. 4.3). Obwohl diese Anordnung auf den ersten Blick besser erscheinen könnte als eine Lösung mit Seriendioden, hat sie in der Praxis mehrere Nachteile. Ein Nachteil besteht darin, dass die Spannung am Wandler bei Verpolung zwar auf $-0,7\text{V}$ beschränkt ist, diese niedrige Negativspannung jedoch bereits ausreichen kann, um einige Wandler zu zerstören. Zweitens ist die Auswahl der Schmelzsicherung keine triviale Aufgabe (siehe Abschnitt 4.3), und deren Wirkung auf die Performance wird häufig unterschätzt. Tatsächlich handelt es sich bei einer Schmelzsicherung um einen Widerstand, der speziell entwickelt wurde, um bei einem bestimmten Strom durchzubrennen. Wie bei allen Widerständen gibt es daran einen Spannungsabfall, der stromabhängig ist. Eine Schmelzsicherung kann einen Spannungsabfall haben, ähnlich oder sogar höher als der Durchlassspannungsabfall einer Diode (siehe folgende Abschnitt).

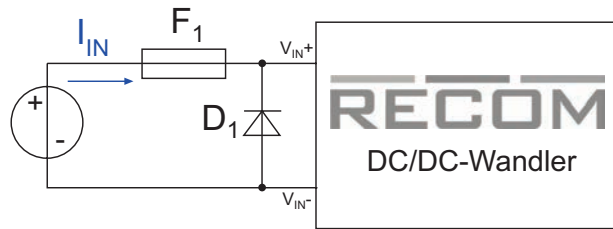


Abb. 4.3: Verpolungsschutz durch Überbrückungs- oder Shuntdioden

4.2.3 Verpolungsschutz mit P-FET

Die dritte Variante des Verpolungsschutzes besteht darin, einen P-FET einzusetzen. Der FET ist die teuerste Lösung; im Vergleich aber zum Wert eines Wandlers ist dies immer noch preiswert. Der FET muss ein P-Kanal MOSFET mit einer internen Gehäusediode sein, ansonsten würde diese Lösung nicht funktionieren. Die maximal zulässige Gate-Source-Spannung V_{GS} sollte größer als die maximale Versorgungsspannung oder verpolte Versorgungsspannung sein. Der FET sollte auch über einen extrem niedrigen ON-Widerstand $R_{DS,ON}$ verfügen, ca. $50\text{m}\Omega$ ist ein annehmbarer Kompromiss zwischen Komponentenpreis und Wirksamkeit. Mit korrekt angeschlossener Versorgung wird der FET voll geöffnet vorgespannt, und sogar mit einem Eingangsstrom über einem Ampere zeigt er einen Spannungsabfall von nur einigen Millivolt.

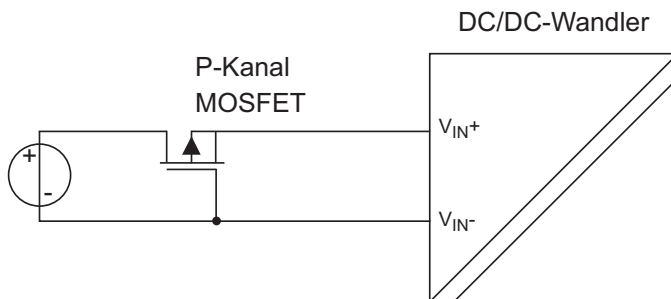


Abb. 4.4: Verpolungsschutz mit P-FET

Verpolungs- schutz	Versorgungs- spannung*	Wandler- eingang- spannung	Wandlerein- gangsstrom	V _{OUT} (V) I _{OUT} (mA)	Leistungs- aufnahme	Ausgangs- leistung	Wirungs- grad
Kein Schutz	9.0V	9.0V	1561mA	11.98V 1000mA	14.05W	11.98W	85.3%
1: Serien- diode (1N5400)	9.7V	8.5V	1660mA	11.98V 1000mA	16.10W	11.98W	74.4%
2: Überbrü- ckungsdiode + 3A Schmelz- sicherung	9.1V	8.5V	1667mA	11.98V 1000mA	15.17W	11.98W	78.9%
3: P-FET (IRF5305)	9.0V	8.9V	1572mA	11.98V 1000mA	14.15W	11.98W	84.7%

* 9 V oder die minimale Eingangsspannung für einen stabilen geregelten Ausgang, je nachdem, welche die höhere ist.

Tabelle 4.1: Messwerte unter Verwendung des Wandlers Recom RP12-1212SA für verschiedene Verpolungsschutzverfahren

Um die Unterschiede zwischen drei verschiedenen Verpolungsschutzverfahren zu untersuchen, wurden unter Verwendung eines 12-W-Wandlers mit Volllast für den ungünstigsten Fall von 9V Eingang Messungen bei einem Nenneingangsstrom von 1,5A durchgeführt. Wie aus Tabelle 4.1 ersichtlich, ist der Wirkungsgrad der P-FET-Lösung dem der Variante ohne Verpolungsschutz sehr ähnlich.

4.3 Eingangssicherung

Ob als Überstromschutz (ausfallsicher) ohne Überbrückungsdiode, oder als Verpolschutz mit Überbrückungsdiode verwendet, muss eine Eingangssicherung so gewählt werden, dass sie beim Eingangsstrom im ungünstigsten Fall während des Normalbetriebs nicht durchbrennt. Da der Schmelzdraht mit zunehmendem Alter brüchig wird, sollte der Sicherungsnennwert im Hinblick auf eine lange Lebensdauer mindestens 1,6 Mal so hoch sein wie der höchste Eingangsstrom. Die Einschaltstromspitze während des Startens des Wandlers ist wesentlich höher, als der danach dauerhaft fließende Eingangsstrom, weshalb die Sicherung träge sein muss, um ein unerwünschtes Durchbrennen beim Einschalten zu vermeiden. Die Kombination von hohem Durchlassstrom und langsamer Ansprechzeit bedeutet auch, dass die Diode so dimensioniert sein muss, dass sie dem Strom während eines Ausfalls durch Verpolung standhält, und das Netzteil muss ebenso im Stande sein, so viel Strom bereitzustellen, dass die Sicherung schnell durchbrennen kann.

Eine Schmelzsicherung ist nur einmalig verwendbar. Wenn das Netzteil aus Versehen verpolt angeschlossen wird, muss die Schmelzsicherung ersetzt werden, bevor der Wandler wieder einsetzbar ist. Das könnte ein Vorteil sein, wenn die Anwendung dauerhaft vom Stromnetz getrennt bleiben muss, bis die Ursache des Ausfalls durch ein Reparaturteam beseitigt ist; für viele andere Anwendungen wäre es jedoch von Vorteil, die Anwendung ausfallunempfindlich (Auto-Recovery) zu machen. Eine Alternative zu einer konventionelle Schmelzsicherung besteht darin, einen selbstrückstellenden Schutz, wie eine Polymer-PTC-Schmelzsicherung (PPTC), zu verwenden.

Dabei handelt es sich um ein Bauteil, das einem (PTC) Widerstand mit positivem Temperaturkoeffizienten, dessen Widerstand sich mit Temperaturanstieg erhöht, ähnlich ist. Unter Ausfallbedingungen erhitzt sich eine PPTC-Schmelzsicherung bis deren innerer Aufbau schmilzt, wobei sie zu einem sehr hohen Widerstand wird, der den Wandler – mit Ausnahme eines minimalen Haltestroms – effektiv trennt. Wenn der Überstrom durch diese Sicherung nun wegfällt, kühlt diese ab und wird wieder niederohmig – sie stellt sich somit automatisch zurück.

4.4 Ausgangs-Überspannungsschutz

Der Überspannungsschutz (OVP) kann an der Ausgangs- oder Eingangsseite eines DC/DC-Wandlers wirken. An der Ausgangsseite besteht die Funktion des OVP darin, die zu versorgende Anwendung vor einem Regelungsfehler zu schützen. Viele Wandler verwenden Zener-Dioden als Spannungsbegrenzer, auch "clamping" genannt, um zu gewährleisten, dass die Ausgangsspannung eine bestimmte Grenze nicht überschreiten kann. Die Schwierigkeit besteht im Einstellen des korrekten Clamp-Spannungspegels. Eine Zener-Diode beginnt bereits etwas Ableitstrom (leakage current) zu ziehen, auch wenn der wesentlich höhere Auslösewert noch nicht erreicht ist, setzt man aber die Zener-Dioden-Spannung zu hoch fest, ist kein sinnvoller Überspannungsschutz mehr gegeben. Ein annehmbarer Kompromiss ist in der Regel eine Clamp-Spannung, die 10 % über der Nennausgangsspannung liegt. Natürlich wird die Zener-Diode bald ausfallen, wenn ein Gesamtregelungsfehler vorliegt, da sie nur eine begrenzte Menge an Leistung aufnehmen kann. Jedoch ist solch ein clamping immer noch nützlich, um kurzzeitige Leistungsspitzen, die unter bestimmten Betriebsbedingungen auftreten könnten, abzufangen.

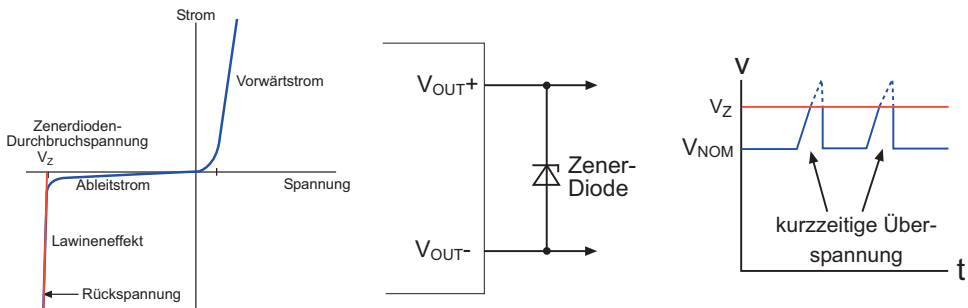


Abb. 4.5: Funktion der Zener-Diode als Spannungsklemmung

4.5 Eingangsüberspannungsschutz

Die Funktion des an der Eingangsseite des Wandlers wirkenden OVP besteht darin, den Wandler gegen Eingangsüberspannungs-Transienten oder Spitzen abzusichern, sodass er den EMV-Richtlinien und anderen Sicherheitsnormen und Performance Standards entspricht.

Da heute eine steigende Anzahl von Geräten und elektronischen Systemen im Einsatz ist, nimmt die Eintrittshäufigkeit von Spannungsspitzen in Stromversorgungen zu. Es existieren viele Normen und Richtlinien, die sowohl die Menge der verursachten Störungen, die ein Gerät erzeugen kann, als auch die Anzahl von Störungen, denen ein Gerät standhalten muss (Störfestigkeit), bestimmen.

Die Störfestigkeitsprüfungen decken Spannungsschöße, schnelle transiente elektrische Störgrößen/Bursts und ESD (elektrostatische Entladung) ab und sind mittlerweile so komplex, dass kaum ein Teil des elektronischen Geräts der Prüfung ohne umfangreiche Eingangs-OVP-Schaltungen standhalten kann.

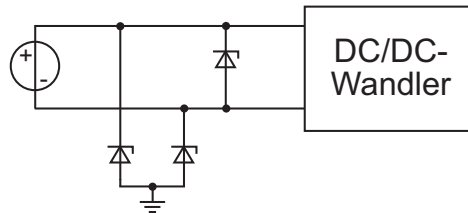


Abb. 4.6: ESD-Schutz

Da alle DC/DC-Wandler eine primäre Stromquelle brauchen, kann angenommen werden, dass die AC/DC-Spannungsversorgung der meisten Anwendungen mit Eingangsfiltern und Schutz vor leitungsgebundenen Überspannungen (welche im ungünstigsten Fall durch Blitzschlag ausgelöst werden) ausgestattet sind. Ein typischer Überspannungsschutz verwendet eine Kombination von Bauteilen wie Gasentladungsröhren, Metalloxid-Varistoren und Funkenstrecken, um entweder die Energie der Spitze zur Erde abzuleiten oder die Energie für längere Zeit zu abzubauen, um das Entstehen zu hoher Spannungsspitzen zu vermeiden. Die Energie, die ein Blitzeinschlag enthält, ist so hoch, dass Begrenzerdioden eine merkliche Verschlechterung mit jedem Impuls erleiden, weshalb sie austauschbar sein müssen.

Daher brauchen DC/DC-Wandler in der Regel vor durch Blitzschlag ausgelösten Hochspannungsschößen eingangsseitig nicht geschützt zu werden. Eine mögliche Ausnahme stellen netzferne stromversorgte Systeme, wie Photovoltaikanlagen dar, die schon einen Überspannungsschutz gegen Blitzschlag erfordern. Für die meisten Anwendungen ist es nicht nötig, den Schutz vor Blitzschlag an der Ausgangsseite bereitzustellen. Ausnahmefälle sind hier Stromversorgungssysteme in Industrieanlagen mit besonders hohem Sicherheitsanspruch, wie Raffinerien, oder Außenbeleuchtungssysteme mit langen, offen liegenden Kabeltrassen. Da DC/DC-Wandler jedoch direkt an die Primärstromquelle angeschlossen werden, sind sie der vollen Energie jeglicher Überspannung, die tatsächlich auftreten kann, ausgesetzt. Daher müssen häufig mehrere Arten von Überspannungsschutz vorgesehen werden.

In den nachfolgenden Abschnitten befassen wir uns mit den Grundlagen des OVP-Schutzes. Die Schutzmaßnahmen müssen stets im Zusammenhang mit der Quellenimpedanz gesehen werden. Je niedriger die Quellenimpedanz ist, desto mehr Energie enthalten die Spitzen der Überspannung und umso aufwendiger und teurer ist es, den Wandler davor zu schützen. Es gibt zwei Hauptschutzmethoden: Crowbar und Spannungsklemmung.

4.5.1 SCR-Crowbar-Schutz

Da es Ausnahmen von der Regel gibt, dass DC/DC-Wandler keinen Schutz vor Blitzschlag benötigen, betrachten wir eine Methode des DC-Überspannungsschutzes, die verwendet werden kann, um die Eingangs- oder Ausgangsseite eines DC/DC-Wandlers vor sehr energetischen Spitzen zu schützen, bevor wir die Besprechung allgemeinerer Spannungsklemmschaltungen fortsetzen, die zum Schutz der Eingangsseite vor

regelmäßig auftretenden Spitzen, sowie Schäden durch Transienten-Überspannungen, verwendet werden können

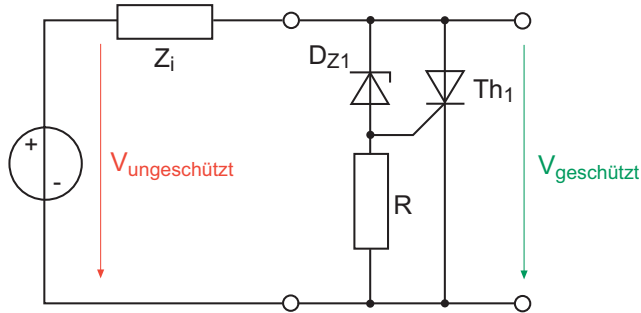


Abb. 4.7: Thyristor-Crowbar-Schaltkreis

Ein Crowbar reagiert auf eine Überspannung, indem er Stromkreise, an welchen die Überspannung auftritt, kurzschließt. Das meist verbreitete Schutzverfahren ist der Siliciumgleichrichter (SCR)-Crowbar. Ein SCR ist im Prinzip ein Thyristor, der gezündet wird, wenn eine voreingestellte Spannung überschritten wird, und der dann im Durchlasszustand bleibt, bis der durch ihn fließende Strom unter einen Haltestromgrenzwert sinkt. Das Schaltbild ist in Abb. 4.7 gezeigt. Die Zener-Diode D_{Z1} stellt die Triggerspannung ein. Z_i stellt die Impedanz der Zuleitung dar.

Der Vorteil des SCR-Crowbars an der Ausgangsseite eines mit Hiccup kurzschluss-geschützten DC/DC-Wandlers besteht darin, dass der Strom durch den Hiccup-Schaltkreis automatisch unterbrochen und der Thyristor zurückgesetzt wird, wenn der Ausgang kurzgeschlossen wird. Der Nachteil eines „Crowbars“ an der Eingangsseite besteht darin, dass der Thyristor sowohl den Kurzschlussstrom der Primärstromquelle als auch den Überspannungskurzschlussstrom absorbieren soll. Deshalb setzt er sich nach dessen Auslösung nicht automatisch zurück und muss beim Einsatz mit einer Eingangssicherung oder PPTC versehen werden, um die Versorgung zu unterbrechen und sowohl den Thyristor als auch die Primär-Stromquelle vor Dauerüberlastung zu schützen. Der eingangsseitige Thyristor-Schaltkreis ist derselbe wie in Abb. 4.7, außer dass Z_i durch eine Schmelzsicherung ersetzt werden kann.

4.5.2 Clamping-Elemente

Clamping-Schutzelemente sind Bauteile, deren Widerstand sich nicht linear mit der angelegten Spannung ändert – ab einem bestimmten Übergangspunkt erhöht sich der Strom durch sie exponentiell. Im Unterschied zu Thyristoren benötigen Clamps keine Rücksetzung, das heißt sie kehren in ihren ursprünglichen Zustand zurück, ohne dass die Versorgung unterbrochen werden muss.

4.5.2.1 Varistor

Ein Varistor ist ein nichtlinearer Widerstand (VDR), dessen Widerstand sich je nach angelegter Spannung ändert. Es gibt verschiedene Varistortypen, einschließlich Selen- und Siliciumkarbidtypen. Bei den am häufigsten verwendeten, handelt es sich jedoch um Metalloxid-Varistoren (MOV). Ein MOV besteht aus vielen mikroskopischen Schichten von ZnO , die zusammengepresst und dann gesintert wurden.

An den Korngrenzen werden die einem Halbleiterübergang ähnlichen Übergangseffekte hervorgerufen, sodass die interne Struktur eines VDR mit Hunderten von in Matrizenform in Seriell- und Parallelkreisen gegeneinander geschalteten Dioden verglichen werden kann. Ist die angelegte Spannung kleiner als die Zenerspannung der Dioden, fließt ein sehr geringer elektrischer Strom, wenn aber die Zenerspannung überschritten wird, tritt ein massiver Stromanstieg auf. Infolge der Kombination von sehr vielen pn-Übergängen in Reihe kann die Zenerspannung sehr hoch – bis auf mehrere Hundert Volt erhöht werden. Da sich die Dioden in gegeneinander geschalteten Paaren befinden, ist der Effekt symmetrisch, und ein MOV schützt sowohl vor positiven als auch negativen Überspannungen.

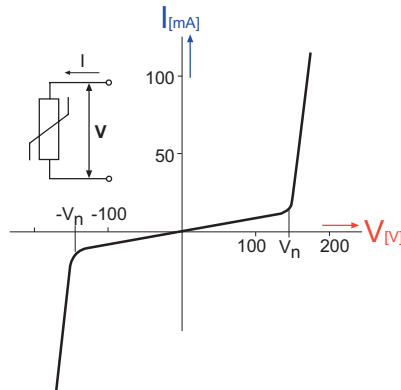


Abb. 4.8: Strom-Spannungs-Kennlinie des Varistors

Die in Abb. 4.8 gezeigte Strom-Spannungs-Kennlinie folgt dem Potenzgesetz, wie die Gleichung 4.1 zeigt:

$$I = k V^\alpha$$

Gleichung 4.1: VDR-Charakteristik

wobei k eine bauteilspezifische Konstante ist und α die Krümmung nach dem Knickpunkt darstellt. Die typischen Werte für verschiedene Schutzkomponenten:

- $\alpha = 35$ Überspannungs-Suppressordioden
- $\alpha = 25$ MOVs
- $\alpha = 8$ Selenzellen
- $\alpha = 4$ Siliciumkarbid-VDRs

MOVs haben eine kurze Ansprechzeit, weshalb sie auch sowohl transiente als auch andauernde Spitzen unterdrücken können, sind aber nicht schnell genug, um die ESD-Überspannungen in Submikrosekunden-Zeiträumen zu unterdrücken. Darüber hinaus können sie durch repetitive Überspannungsimpulse beschädigt werden, da jegliche Inhomogenität in der internen Kornstruktur lokale Überhitzungseffekte verursacht, die zur allmählichen Verschlechterung der Performance führen (erhöhter Ableitstrom). Mehrschichtige MOVs (MLVs) sind ein Versuch, diese Verschlechterung hinauszuzögern, sodass das Bauteil einer größeren Anzahl interner Störungen standhalten könnte, ohne vollständig auszufallen. Wird die interne Verlustleistung jedoch viel zu hoch, schmelzen alle MOVs und fallen komplett aus. Deshalb sollte ein MOV immer mit einer Eingangssicherung verwendet werden. Das Energie-Rating (in Joule) ist ein Hinweis auf die erwartete Lebensdauer eines MOV bezüglich repetitiver Spitzen und ein wichtiger Auswahlfaktor.

4.5.2.2 Suppressordiode

Im Unterschied zum VDR wird der durch eine Suppressordiode gebotene Schutz durch einen Einfach-pn-Übergang gewährleistet, verfügt jedoch über einen viel größeren Querschnitt für den Stromweg. Suppressordioden bezeichnet man auch als TVS (Transient Voltage Suppressors), Silizium-Lawinendioden (SAD)-Suppressors oder anderer herstellerspezifischer Bezeichnungen. Die unipolare Strom-Spannungs-Kennlinie ist dieselbe wie die einer Zener-Diode (siehe Abb. 4.9), Suppressordioden sind jedoch so beschaffen, dass sie ein viel höheres Verhältnis von Spitzen-zu-Durchschnittsleistung erlauben.

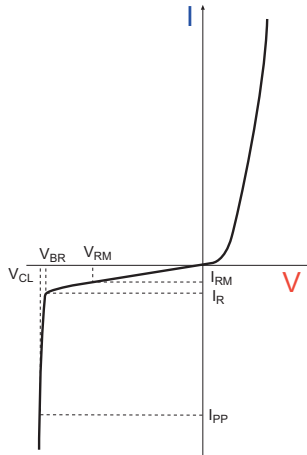


Abb. 4.9: Strom-Spannungs-Kennlinie einer unipolaren Suppressordiode

Wie in Abb. 4.9 gezeigt, verhält sich eine Suppressordiode im ersten Quadrant (oben rechts) wie eine normale Diode in Vorwärtsrichtung und im dritten Quadrant (unten links) wie eine Zener-Diode in Sperrrichtung. Der dritte Quadrant wird durch drei Paare von Werten bestimmt; die Nennspannung V_{RM} (Stand-off-Sperrspannung) beim Sperrstrom I_{RM} , der die zusätzliche Belastung der Versorgung infolge des Ableitstroms bezeichnet, die Zenerspannung V_{BR} beim Sperrstrom I_R , am Knipunkt der Kennlinie sowie die Clamping-Spannung V_{CL} , spezifiziert beim maximal zulässigen Strom I_{PP} . Die Suppressordiode sollte so gewählt werden, dass die normalen Arbeitsspannungen dem Wert V_{BR} nahe kommen, diesen jedoch nicht überschreiten. Darüber hinaus kann ein Strombegrenzungswiderstand erforderlich sein, um zu gewährleisten, dass I_{PP} nicht überschritten wird.

Da eine Suppressordiode eine unipolare Diode ist, kann sie nur auf positive Überspannungen reagieren. Deshalb enthalten die meisten TVS-Gehäuse zwei gegeneinander geschaltete Suppressordioden, um als bipolare TVS-Diode sowohl positive als auch negative Spitzen zu beschränken. Der Vorteil von Suppressordioden gegenüber MOVs besteht darin, dass sie nicht durch repetitive Spitzen beeinträchtigt werden und niedrigere Zenerspannungen mit genaueren VBR-Werten haben; folglich können sie sowohl Niederspannungs- als auch Signalleitungen schützen.



Abb. 4.10: Bipolares TVS-Symbol

4.5.3 OVP unter Verwendung mehrerer Bauteile

Die Sperrcharakteristik der einzelnen Komponenten deckt häufig nicht alle Anforderungen des Überspannungsschutzes ab. Um die gewünschten Gesamteigenschaften zu erreichen, kann es deshalb erforderlich sein, verschiedene Bauteile parallel zu schalten. Wie in vorhergehenden Abschnitten gezeigt, sind Varistoren oder Suppressordioden für OVP in vielen DC/DC-Anwendungen geeignet, jedoch sind in manchen Fällen Kombinationen beider notwendig, um den Eingang eines DC/DC-Wandlers angemessen zu schützen. MOVs weisen eine hohe Stromdurchlässigkeit, aber auch hohe Clamping-Spannungen auf. Andererseits haben TVS-Dioden sehr kurze Schaltzeiten (im Nanosekundenbereich), und der VBR kann herabgesetzt werden, jedoch ist auch die Nennleistung eingeschränkt. Die allgemeine Regel lautet: Je schneller ein Schutzelement anspricht, desto geringer die Leistung, die es bewältigen kann. Für eine volle OVP heißt dies, dass der Schutzmechanismus so in Reihe geschaltet werden muss, dass das Element, das den höchsten Strom aufnehmen kann, auch das erste in der Reihe sein muss. Abb. 4.11 zeigt eine typische Anordnung:

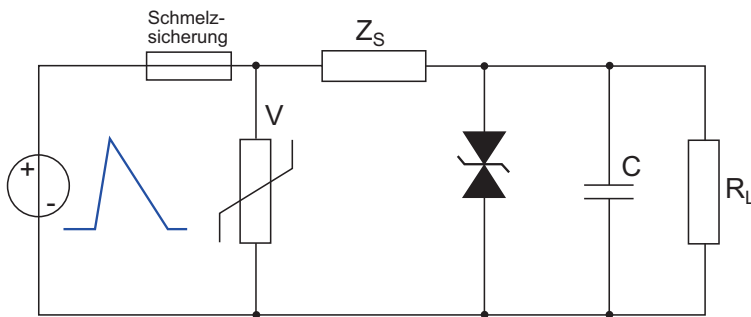


Abb. 4.11: Aus mehreren Schutzelementen gebildete OVP

Abb. 4.11 zeigt ein aus mehreren Stufen gebildetes OVP-Netzwerk. Wenn sich der MOV überhitzt und ausfällt, schützt die Schmelzsicherung vor Kurzschluss, ansonsten absorbiert der MOV am Eingang die meiste Energie der Überspannung. Während der Zeit, die der MOV für eine Reaktion benötigt, wird die Eingangsspannung durch das TVS-Element mit Strombegrenzung, die die Serienimpedanz Z_s gewährleistet, beschränkt. Schließlich hilft der Eingangskondensator, jegliche restliche pulsformige Energie zu absorbieren.

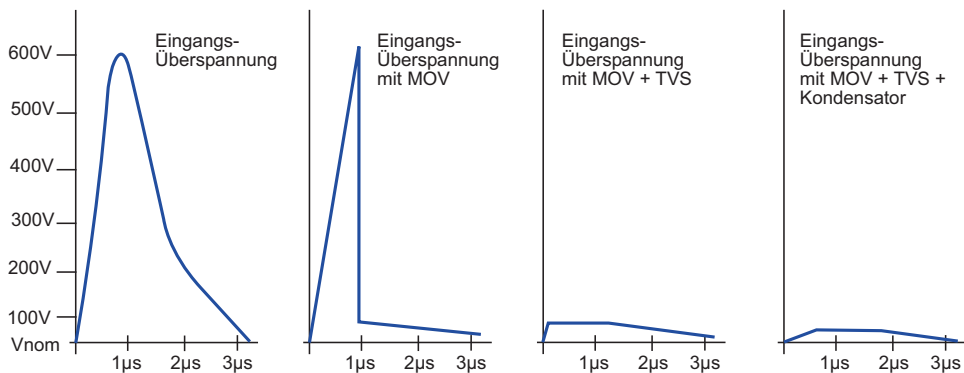


Abb. 4.12: Mit OVP-Schaltkreis aus Abb. 3.11 verbundene Wellenformen

Wenn die Eingangsspitzen besonders energiereich sind, können TVS-Bauteile nebengeschaltet werden, um den Strom auf mehrere Schutzelemente aufzuteilen. Es wird nicht empfohlen, mehrere MOVs parallel zu schalten, da dies die Chance eines Ausfalls erhöhen würde. Gegebenenfalls ist es besser, ein einzelnes Element mit einem höheren Joule-Wert auszuwählen.

Z_S kann ein Widerstand sein, was kostengünstig ist. Es sollte jedoch beachtet werden, dass er in Serie mit dem Eingang geschaltet ist, sodass er auch für den regulär auftretenden DC-Dauerstrom ausgelegt sein muss, wodurch der Gesamtwirkungsgrad reduziert wird. Eine bessere, aber auch teurere Wahl ist eine Drossel mit einem Serienwiderstand im Bereich von 100m Ω .

4.5.4 OVP-Standards

Die Datenblatt-Performance der OVP-Bauteile gibt nur theoretische Werte wider und können auch nur theoretisch sein, da der praktische Effekt einer Schutzbeschaltung u. a. auch von der Stabilität der Komponenten des DC/DC-Wandlers sowie vom gesamten realen Aufbau abhängt. Sogar kleine PCB-Streuinduktivitäten und Impedanzen können das Ergebnis erheblich beeinflussen. Daher ist eine praktische Prüfung erforderlich, um das schaltungsinterne Verhalten der OVP zu überprüfen und dessen tatsächliche Performance zu verifizieren.

Da es wenig praktikabel ist, im Test auf eventuell zufällig auftretende Überspannungstransienten und -spitzen zu warten, wurden sowohl auf nationaler als auch auf internationaler Ebene verschiedene Prüfnormen festgelegt. Die internationale Norm IEC 61000-4-5 definiert beispielsweise eine „Überspannung“ als von einem Hochspannungsimpulsgeber mit 2 Ω Quellenimpedanz (von Eingang zu Eingang) oder 12 Ω Quellenimpedanz (von Eingang zu Erde) bereitgestellte Transiente mit einer Anstiegszeit von 1,2 μ s, die wieder zu 50% ihres Spitzenwertes innerhalb von 50 μ s abfällt. Die Spitzenspannung dieses 1,2/50 μ s-Impulses kann, je nach der Installationsklasse des Produktes, zwischen 0,5kV und 4kV gewählt werden. Obwohl es möglich ist, einen eigenen Überspannungsprüfer (surge tester) einzusetzen (die Norm beinhaltet entsprechende Anweisungen), empfiehlt es sich, ein kalibriertes Test-Setup mit den Normvorgaben entsprechendem Verhalten zu erwerben.

Level	Leerlaufspannung (kV)
1	± 0.5
2	± 1
3	± 2
4	± 4
x	Sonderfall

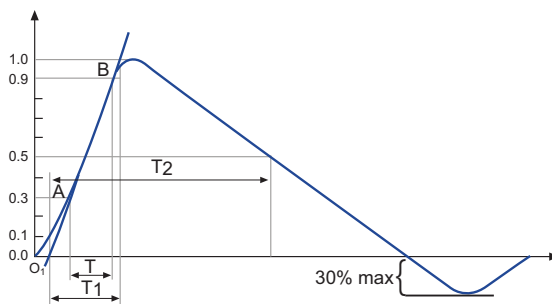


Tabelle 4.2: IEC 61000-4-5 Messpegel

Die Norm definiert auch die Wirkungen, die nach solch einem Überspannungsfall auftreten können:

Klasse	Ergebnis
A	Normale Funktion
B	Vorübergehender Funktionsausfall, Wiederherstellung normaler Funktion erfolgt automatisch
C	Vorübergehender Funktionsausfall, erfordert manuelle Rücksetzung zur Wiederherstellung normaler Funktion
D	Permanenter Verlust der Funktion oder Performance

Tabelle 4.3: IEC 61000-4-5 Leistungsgrad

Es gibt vergleichbare internationale Normen, die Störfestigkeit gegen schnelle transiente elektrische Störgrößen/Bursts (z. B.: IEC 61000-4-4: Wellenform 5/50ns, wiederholt bei 5kHz für 15ms oder bei 100kHz für 0,75ms) und elektrostatische Entladungs-Spannungspegel (ESD) definieren.

4.5.5 OVP durch Abschaltung

Die Auswahl des Prüfprotokolls hängt stark von der Anwendung des Endverbrauchers ab, darüber hinaus gibt es weitere OVP-Prüfnormen, die anwendungsspezifisch sind. Die Bahnnorm EN50155 fordert beispielsweise eine Störfestigkeit gegen Stoßspannungen von 140% der Nenneingangsspannung für 1 Sekunde. Solch eine lang andauernde Stoßspannung kann ohne Abbau überschüssiger Leistungsmengen nicht auf einfache Weise beschränkt werden. Eine Lösung besteht darin, den Eingang für die Dauer der Überspannung abzuschalten, um den DC/DC-Wandler zu schützen. Für diese Aufgabe sind kundenspezifische Controller-ICs verfügbar, die eine Überspannungs-Detektorschaltung sowie einen FET-Gatetreiber enthalten, die die Versorgungsspannung in $<1 \mu\text{s}$ (Abb. 4.13) ausschalten können.

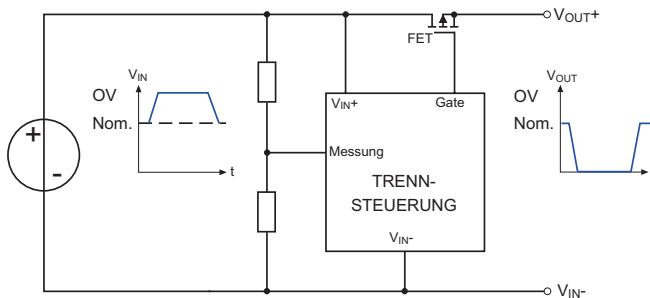


Abb. 4.13: OVP-Trennschutz

Das OVP-Trennverfahren ist nicht nur für den Langzeit-Überspannungsschutz sinnvoll, es ist auch eine der wenigen zuverlässigen Schutzschaltungen bei sehr niedrigen Eingangsspannungen. Eine Eingangsspannung des DC/DC-Wandlers z. B. von 1,2V kann nicht einfach unter Verwendung konventioneller OVP-Bauteile geschützt werden, da der Temperaturkoeffizient eine wesentliche Fehlerquelle darstellt: Entweder fangen die Dioden bei der Nennspannung an zu leiten oder die Clamping-Spannung wird höher als benötigt.

Natürlich besteht der Nachteil des Trennschutz-OVPs darin, dass dem DC/DC-Wandler während der Dauer der Überspannungssituation die Eingangsleistung entzogen wird. Für kurze Trennzeiten kann die Eingangsspannung zum DC/DC durch Hinzufügen eines entsprechend großen Kondensators am Eingang erhalten werden, für lange Trennzeiten kann jedoch eine Notstromumschaltung oder ein Supercap-Speichersystem erforderlich sein. Im folgenden Abschnitt wird diese Lösung untersucht.

4.6 Spannungseinbruch und -unterbrechung

In Stromverteilungssystemen können plötzliche Lasterhöhungen wesentliche Spannungsabfälle verursachen. Im Idealfall sollten diese kurzfristigen Rückgänge die nachfolgenden Stromversorgungskomponenten nicht beeinflussen. Die typische Lösung zum Schutz eines DC/DC-Wandlers vor Eingangsspannungseinbruch und -unterbrechung besteht darin, ausreichend Energie im Kondensator zu speichern, um den Wandler während der Zeit des Spannungseinbruches betriebsbereit zu halten. Abb. 4.14 zeigt einen einfachen Schaltkreis.

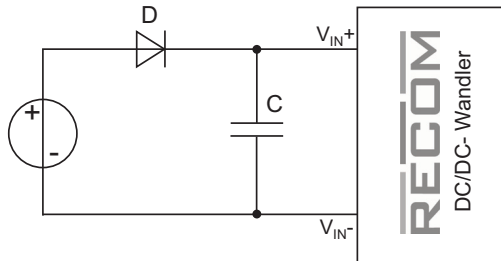


Abb. 4.14: Überbrückung von Eingangsspannungseinbrüchen und -unterbrechungen

Der Schaltkreis besteht aus einer Entkoppeldiode D und einem oder mehreren Kondensatoren C. Der Kondensator C wird im Normalbetrieb auf Arbeitsspannung $V_{IN} - V_{D_{iode}}$ geladen. Bei einem Eingangsspannungseinbruch oder einer -unterbrechung sperrt die Diode und verhindert das Entladen des Kondensators über die Versorgung, sodass die gesamte im Kondensator C gespeicherte Energie nun für den DC/DC-Wandler verfügbar ist. Die Spannung am Kondensator beginnt nun an zu sinken, da er sich in den DC/DC-Wandler entlädt. Dieses Verhältnis lässt sich jedoch nur auf komplizierte Weise berechnen, da der DC/DC-Wandler eine konstante Leistung aufnimmt und der Eingangsstrom somit umgekehrt proportional zur Eingangsspannung ist.

Die in einem Kondensator gespeicherte Energie E_C ist gleich der Kapazität C, multipliziert mit der zum Quadrat genommenen Spannung V_C , wobei V_C gleich der Eingangsspannung V_{IN} minus des Spannungsabfalls an der Diode D ist.

$$E_C = \frac{1}{2} C V_C^2$$

Wenn bei $t_0 = 0$ die Eingangsspannung unterbrochen wird, beginnt die Spannung am Kondensator laut der Kondensatorentladungsgleichung exponentiell abzufallen:

$$V_C(t) = V_C(t = 0) e^{-t/RC}$$

Der aufgeladene Kondensator kann bis zu der Zeit t_1 entladen werden. Zeit t_1 ist die Zeit, zu der die Kondensatorspannung V_C der minimalen Eingangsspannung $V_{IN,MIN}$ des DC/DC-Wandlers gleicht. Die im Kondensator verbleibende Energie ist dann:

$$E_C(t_1) = \frac{1}{2} C V_{IN,MIN}^2$$

Die Energie, die zur Sicherstellung der Eingangsspannung im Laufe des Zeitintervalls $t_0 - t_1$ erforderlich ist, beträgt somit:

$$E_{BACK} = E_C(t_0) - E_C(t_1) = 0.5 C (V_{IN}^2 - V_{IN,MIN}^2)$$

Diese Energie muss die im Laufe der Backupzeit t_{BACK} notwendige Eingangsspannung bereitstellen. Die erforderliche Eingangsleistung kann aus der Ausgangsleistung und dem Wirkungsgrad berechnet werden und ergibt:

$$t_{BACK} = \frac{E_{BACK} \eta}{P_{OUT}} = \frac{C (V_{IN}^2 - V_{IN,MIN}^2) \eta}{2 P_{OUT}}$$

Gleichung 4.2: Berechnung der Backupzeit

Nach Umformung dieser Gleichung ergibt sich die erforderliche Stützkapazität:

$$C_{BACK} = \frac{2 t_{BACK} P_{OUT}}{(V_{IN}^2 - V_{IN,MIN}^2) \eta}$$

Gleichung 4.3: Berechnung des Stützkondensators

Diese Gleichungen besagen, dass je größer der Stützkondensator ist, desto länger ist die Backupzeit. Da große Kondensatoren jedoch viel Platz einnehmen, gibt es häufig physikalische Beschränkungen im Hinblick auf die Größe von C . Die Gleichungen besagen jedoch auch, dass die gespeicherte Energie proportional zu V_{C2} ist. Je breiter der Eingangsspannungsbereich des DC/DC-Wandlers also ist, desto besser. Der DC/DC-Wandler muss so gewählt sein, dass die Nenn- V_{IN} nahe der maximalen Eingangsspannung des Wandlers liegt, um die maximale Backupzeit zu erzielen. Auch ein hoher Wirkungsgrad oder Last-Derating sind hilfreich.

Die einfache Schaltung in Abb. 4.14 besitzt zwei Nachteile. Der Spannungsabfall an der Diode D ist ein zusätzlicher Verlust, der den Wirkungsgrad während des Normalbetriebs verringert, und der hohe Einschaltspitzenstrom zum Aufladen des großen Stützkondensators kann ein Problem für die Primärspannungsversorgung darstellen. Diese beiden Probleme können durch die in Abb. 4.17 gezeigte Variante der Trennsteuerung, die bei Unterspannung anstelle von Überspannung ausschaltet und über eine Soft-Start-Funktion zum Verringern des Einschaltspitzenstroms verfügt, gelöst werden.

4.7 Einschaltspitzenstrom-Begrenzung

Häufig wird dem Problem des Einschaltspitzenstroms (engl.: inrush current) zu wenig Aufmerksamkeit gewidmet. Um leitungsgebundene Störungen zu verringern, verfügen alle DC/DC-Wandler über eine interne Filterschaltung. Bei diesem Filter handelt es sich zumindest um einen einfachen Eingangskondensator, häufiger jedoch um einen RC- oder LC-Tiefpassfilter oder π -Filter.

Gute Filterkondensatoren verfügen über einen sehr niedrigen äquivalenten Serienwiderstand (ESR), was bedeutet, dass sie fast einen Kurzschluss an den Eingangs- klemmen darstellen, wenn Spannung an den ungeladenen Kondensator angelegt wird. Ein MLCC-Kondensator kann über einen ESR unter 100mΩ verfügen. Der Einschalt- spitzenstrom I_{IR} ist eine Erscheinung, die nur beim Starten auftritt. Die Spitzenströme können jedoch um Größenordnungen höher sein als der Eingangsstrom im regulären Dauerbetrieb. Da die Eingangskondensatoren fast einen Vollkurzschluss darstellen, ist der Strom nur durch die Impedanz der Anschlussleitungen (Z_L) und den Innenwiderstand der Stromversorgung (Z_{IS}) beschränkt.

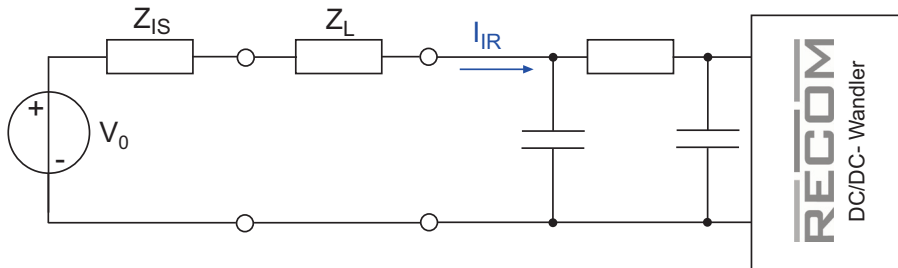


Abb. 4.15: Einschaltspitzenstrom-Modell

Neben dem Einschaltspitzenstrom infolge der Eingangsfilterkondensatoren versucht auch der DC/DC-Wandler zu starten. Der Transformator wird mit Strom versorgt, und die Kondensatoren in der Last werden aufgeladen. Alle diese Energieströme überlagern sich, sodass üblicherweise mehrere Einschaltstromspitzen und -abfälle auftreten, bevor der Eingangsstrom stabilisiert ist. Abb. 4.16 zeigt das Beispiel eines 2-W-Wandlers mit einem normalen Eingangsstrom von 80mA, aber einem Einschaltspitzenstrom von fast 8A. Obwohl solcher Einschaltspitzenstrom für einen Kleinleistungswandler bedrohlich hoch erscheint, dauert er nur 10µs an.

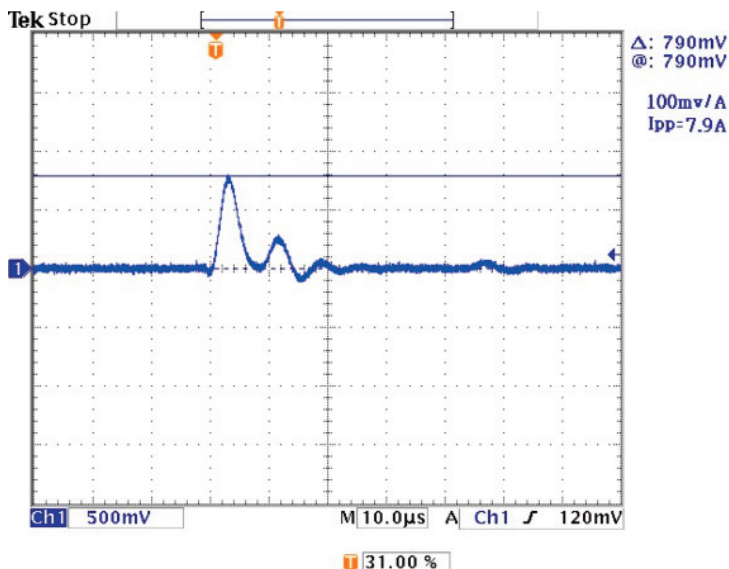


Abb. 4.16: Beispiel für Einschaltspitzenstrom

In Point-of-Load-Architekturen werden viele DC/DC-Wandler parallel mit Zwischenkreis-Busversorgung angeschlossen. Daher gibt es viele parallel geschaltete niedrig-ESR-Eingangsfilterkondensatoren, die zu extrem hohen Einschaltspitzenströmen führen können, falls keine entsprechenden Vorkehrungen getroffen werden.

Um den Einschaltspitzenstrom unter Kontrolle zu halten, können Wandler in komplexen Systemen in Reihe geschaltet werden. Andernfalls kann eine Softstart-Schaltung wie in Abb. 4.17 gezeigt zur Reduzierung des Einschaltspitzenstroms verwendet werden:

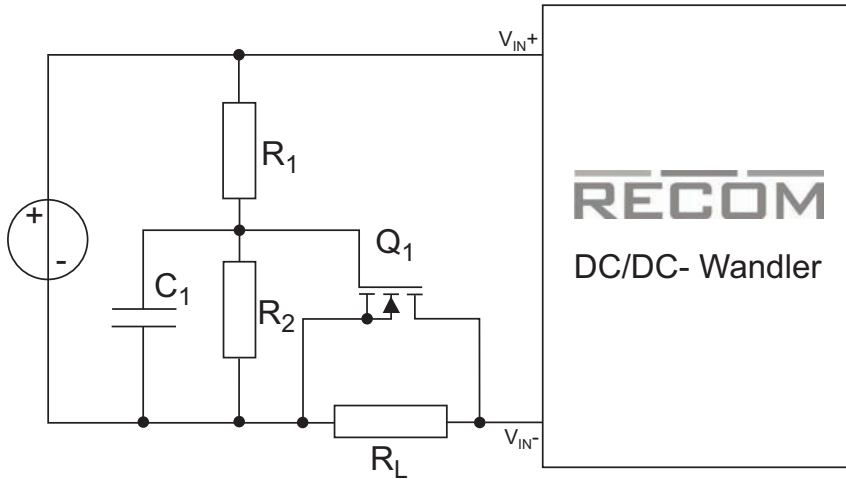


Abb. 4.17: Einschaltspitzenstrom-Begrenzer (Soft-Start)

Der Einschaltspitzenstrom-Begrenzer funktioniert durch Kurzschließen eines Strombegrenzungswiderstandes R_L erst, nachdem sich der Eingangsstrom stabilisiert hat. Der Feldeffekttransistor Q_1 ist ein N-Kanal-MOSFET, der durch das aus R_1 , R_2 und C_1 gebildete RC-Netzwerk gesteuert wird. Beim Einschalten der Versorgung ist C_1 noch ungeladen und die Gate-Spannung niedrig gehalten, sodass Q_1 AUS ist. Die Eingangskapazität des DC/DC-Wandlers wird durch R_L langsam geladen und der Einschaltspitzenstrom wird so gering gehalten. Inzwischen lädt sich C_1 durch R_1 , bis am Gate von Q_1 die Spannung $V_{IN} \times R_2 / (R_1 + R_2)$ erreicht wird. Diese Spannung muss ausreichend gewählt werden, um den FET durchzuschalten, der dann den Serienwiderstand R_L kurzschließt. Bei kleinen Werten von C_1 kann die Gate-Kapazität C_G ein wesentlicher Faktor sein und kann unter Verwendung der Ladezeitkonstante $\tau = (R_1 \parallel R_2) (C_1 \parallel C_G)$ berechnet werden. Der FET muss so gewählt werden, dass er den Eingangsstrom im ungünstigsten Fall kontinuierlich leiten kann (worst case: maximale Ausgangslast mit einer minimalen Eingangsspannung). R_1 und R_2 sollen so dimensioniert sein, dass die Gate-Spannung höher ist als der spezifizierte minimale Wert von V_{GS} bei minimaler Eingangsspannung.

Die Auswahl von R_L hängt vom Anwender ab. In der Regel sind jedoch einige Ohm ausreichend, um den Einschaltspitzenstrom auf ein annehmbares Niveau zu verringern, ohne dem DC/DC-Wandler so viel Strom zu entziehen, dass er nicht ordnungsgemäß starten kann. Die Norm ETSI ETS 300 132-2 definiert den maximal zulässigen Einschaltspitzenstrom als 48 x den Nenneingangsstrom.

Bei dem in Abb. 4.16 gezeigten Beispiel würde ein Serienwiderstand von 6Ω ausreichen, um den Wandler ETSI-konform zu machen. Da es sich um einen Kleinleistungswandler handelt, beträgt der Verlust durch den Widerstand während des Normalbetriebs in der Praxis nur 40mW , und die Einschaltsspitzenstrom-Begrenzerschaltung in Abb. 4.17 würde überflüssig.

In einigen Anwendungen kann ein NTC als Einschaltsspitzenstrom-Begrenzer verwendet werden. Der NTC verfügt ursprünglich über einen hohen Widerstand, der den Einschaltspitzenstrom beschränkt. Da er sich bei Stromdurchfluss erwärmt, verringert er mit der Zeit seinen Widerstand, was einen Anstieg des DC/DC-Wandlereingangsstroms zur Folge hat. Der Hauptnachteil besteht darin, dass der NTC kontinuierlich Leistung umsetzen muss, damit er warm genug bleibt, um seinen niederohmigen Zustand aufrechtzuerhalten.

4.8 Lastbegrenzung

Eine weitere Möglichkeit, den Einschaltspitzenstrom zu verringern, besteht darin, die Last am Wandler während des Startens zu verringern. Dies senkt den lastabhängigen Anteil des Einschaltspitzenstroms und behält nur den infolge der Eingangsfilterkapazität vorhandenen Anteil. Es gibt zwei Hauptvarianten von Lastsenkung: Ausgangs-Soft-Start und Ausgangs-Lastzuschaltung.

Der Ausgangs-Soft-Start arbeitet mit Wandlern mit einer Ausgangsspannungs-abgleichfunktion und eignet sich hauptsächlich nur für ohmsche Lasten. Der Ausgangsstrom ist zur Ausgangsspannung proportional. Wurde nun die Ausgangsspannung ursprünglich auf „low“ gesetzt, bleibt der Ausgangsstrom auch niedrig, weshalb auch der Einschaltspitzenstrom niedriger ausfällt. Hat sich der Anlauf stabilisiert wird nun die Ausgangsspannung bis zur gewünschten Arbeitsspannung linear angehoben.

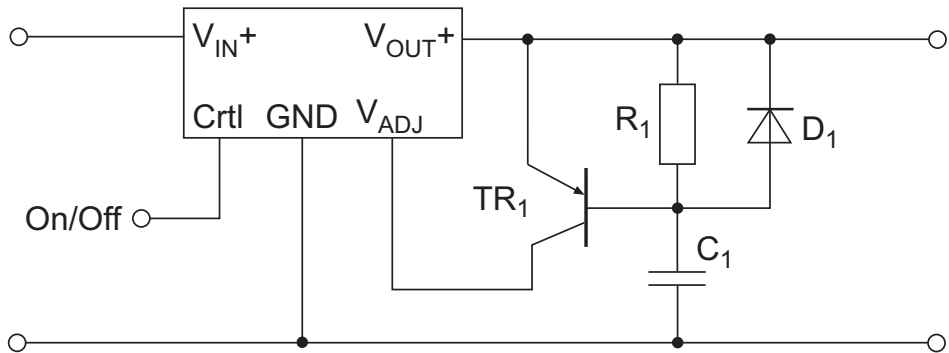


Abb. 4.18: Ausgangs-Soft-Start

Abb. 4.18 zeigt ein Beispiel dieses Verfahrens unter Verwendung eines RECOM Buck-Schaltreglers der R-6112x-Serie. Wird die Versorgungsspannung am Wandler angelegt, ist der Kondensator C_1 anfangs noch ungeladen. Der PNP-Transistor wird voll eingeschaltet, und der V_{ADJ} -Pin wird auf die V_{OUT+} -Spannung gezogen. Dies bewirkt, dass die Ausgangsspannung auf die minimale Ausgangsspannung eingestellt wird.

Mit dem Wandler R-6112x mit einem 12V @ 1A Nennausgang wird beispielsweise die Ausgangsspannung bei 275mA auf 3,3V bzw. über ein Viertel der Volllast heruntergeregelt. Da C_1 durch R_1 geladen wird, sinkt der Strom durch T_{R1} , und die Ausgangsspannung erhöht sich linear und erreicht schließlich die volle Nennausgangsspannung. Diode D_1 gewährleistet, dass C_1 schnell entladen wird, wenn der Wandler nach Abschaltung für den nächsten Ausgangs-Soft-Start bereit ist.

Die zweite Methode einer Einschaltspitzenstrombegrenzung ist die Lastzuschaltung. Dieses Verfahren arbeitet mit beliebigen Wandlern oder Lasttypen. Die Ausgangslast wird erst angelegt, nachdem sich die Ausgangsspannung stabilisiert hat. Deshalb hat der Einschaltspitzenstrom eine doppelte Spitze; einmal mit Schalter EIN und einmal mit Last EIN. Das verteilt den gesamten Einschaltspitzenstrom über eine längere Zeit und verringert den maximalen Spitzenstrom. Der Ausgangslastschalter ist eine Variante der in Abb. 4.19 gezeigten Einschaltspitzenstrom-Begrenzerschaltung. Der Feldeffekttransistor Q_1 ist ein N-Kanal-MOSFET, der durch das mit R_1 , R_2 , R_3 und C_1 gebildete RC-Netzwerk gesteuert wird. Beim Einschalten ist C_1 noch ungeladen und die Gate-Spannung des MOSFET ist niedrig, sodass Q_1 AUS ist. Dann wird C_1 über R_1 aufgeladen, bis das Gate des Q_1 die Spannung $V_{RL} \times R_2 / (R_1 + R_2)$ erreicht hat. Diese Spannung muss ausreichend hoch sein, um den FET sicher einzuschalten und die Last an den Wandlerausgang niederohmig anzuschalten. R_3 ist ein hoher Widerstand, der mit Kondensator C_1 den sich ergebenden Ausgangsspannungseinbruch infolge des plötzlichen Zuschaltens der Last abfiltert. R_1 und R_2 sind so zu dimensionieren, dass die Gate-Spannung höher ist als die spezifizierten minimalen Werte für V_{GS} bei Nennausgangsspannung. Diode D_1 gewährleistet, dass sich C_1 schnell entlädt, wenn der Wandler ausgeschaltet ist, und für den nächsten Einschaltzyklus bereit ist.

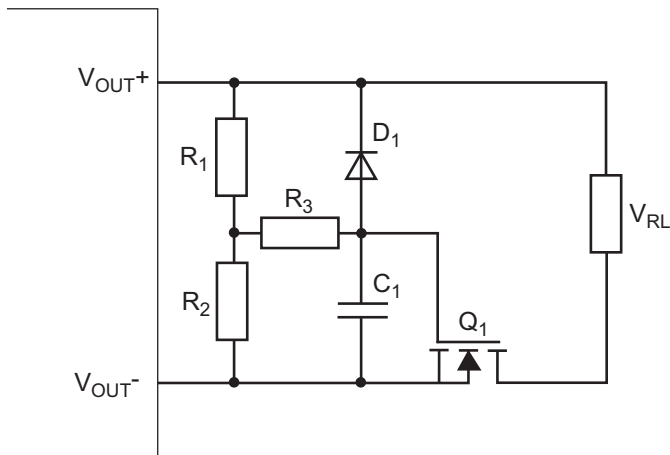


Abb. 4.19: Zuschaltung des Lastausgangs

4.9 Unterspannungsabschaltung

Wenn die Eingangsspannung viel zu niedrig ist, kann der Eingangsstrom die Grenzwerte der im DC/DC-Wandler verbauten Bauteile überschreiten. Deshalb verfügen manche Wandler über eine interne Steuerschaltung, die den Wandler im Fall einer zu niedrigen Eingangsspannung sperrt. Dieser Schaltkreis heißt Unterspannungsabschaltung (engl.: Under Voltage Lockout - UVL).

Der praktische Nutzen des UVL-Schaltkreises sollte nicht unterschätzt werden. Nehmen wir zum Beispiel eine Anwendung, die 12W Ausgangsleistung mit einer 12V Versorgung benötigt. Bei der Nenneingangsspannung würde eine Versorgung mit 1A ausreichen, sodass vielleicht eine Primärstromquelle von 1,5A spezifiziert werden könnte. Der Wandler verfügt in der Regel über einen Eingangsspannungsbereich von 9 bis 18V, sodass er beim Einschalten zu arbeiten anfängt, sobald sich die Eingangsspannung über 9V linear erhöht, und 1,3A verbraucht. Ohne UVL-Schaltkreis kann der Wandler jedoch schon bei 7V einen Startversuch unternehmen, obwohl dies außerhalb der Spezifikation liegt und verbraucht 1,7A. Dies liegt über dem Wert, der durch die Strombegrenzung der Versorgung begrenzt würde. Das Netzteil und der Wandler können einige Zyklen lang wechselwirken, bevor der Wandler schließlich ordnungsgemäß startet. Inzwischen erhält die Last mehrere unregelmäßige Spannungsschöße, die mitunter die Anwendung zerstören könnten.

Somit schützt eine UVL-Funktion nicht nur den DC/DC-Wandler, sondern auch Last und Primärstromquelle vor unzulässigen Betriebszuständen. Falls kein DC/DC-Wandler mit eingebauter UVL-Funktion verfügbar ist, kann die in Abb. 4.20 gezeigte externe Schaltung verwendet werden, um den Wandler zu sperren, bis sich die Eingangsspannung stabilisiert hat. Ein Op-Verstärker mit eingebauter Referenzspannung wie der LM10 ist eine geeignete Wahl.

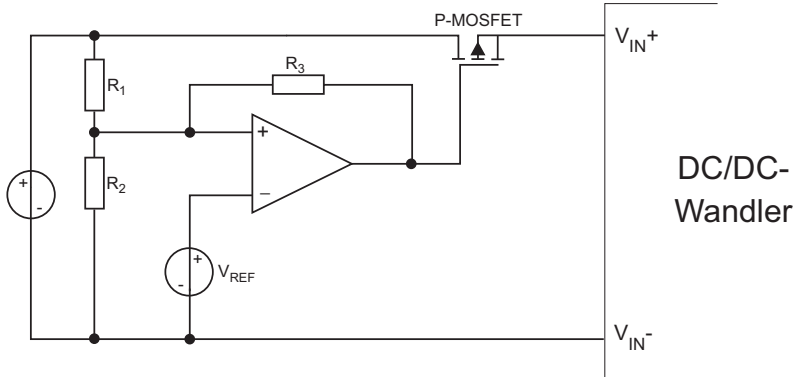


Abb. 4.20: Beispiel eines UVL-Schaltkreises

$$\text{Schalter EIN: } V_{UVL} = V_{REF} \left(\frac{R_1 (R_2 + R_3) + R_2 R_3}{R_2 R_3} \right)$$

$$\text{Schalter AUS: } V_{UVL} = V_{REF} \left(\frac{R_1 (R_2 + R_3) + R_2 R_3}{R_2 (R_2 + R_3)} \right)$$

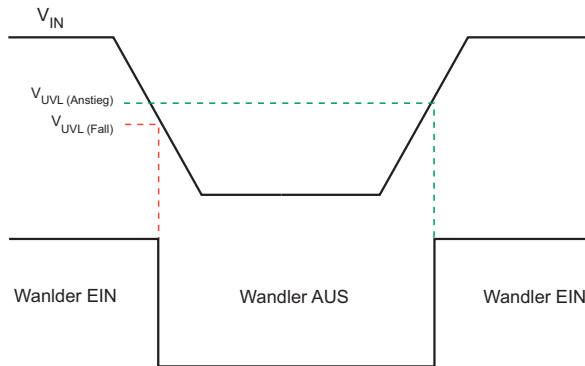


Abb. 4.21: Unterspannungsabschaltfunktion

5. Befilterung von Wandlereingang und -ausgang

5.1 Einteilung

Alle DC/DC-Wandler haben eine Restwelligkeit in der Ausgangsspannung infolge des Auf- und Entladens des Ausgangskondensators mit jedem Taktimpuls des internen Oszillators. Diese Ausgangswelligkeit hat je nach Topologie eine Frequenz, die entweder dieselbe oder doppelte Wandlerschaltfrequenz ist und in der Regel im 100-bis 200-kHz-Bereich liegt. Der Welligkeit überlagert sind Schaltspannungsspitzen mit einer viel höheren Frequenz, üblicherweise im MHz-Bereich.

Auch der Eingangsstrom besitzt eine Welligkeit mit zwei Komponenten: eine DC-Komponente, die lastabhängig ist, und eine AC-Komponente, die als „Rückwelligkeitsstrom“ (eng.: back ripple current) oder reflektierter Strom, der durch den vom internen Oszillator hervorgerufenen pulsierenden Strom hervorgerufen wird, bezeichnet wird. Kleinere Hochfrequenzstörungen, die den Schaltspannungsspitzen entsprechen, überlagern diese kombinierte Wellenform. Im Allgemeinen verursacht der DC-Strom solange keine Probleme, wie die Primärspannungsversorgung entsprechend dimensioniert ist. Jedoch können die AC-Stromimpulse in anderen Teilen des Systems infolge von Streuinduktivitäten und -kapazitäten in Leiterbahnen, Anschlussleitungen und Steckverbindungen Störungen hervorrufen. Darüber hinaus führt der Eingangsstrom zu einem Spannungsabfall an den Zuleitungen infolge des resultierenden Widerstands. Bei pulsierendem Strom pulsiert auch der Spannungsabfall, wobei die Zuleitungen wie eine strahlende Antenne wirken.

Deshalb müssen sowohl die Eingangs- als auch Ausgangswelligkeit z. B. durch Verwendung von externen Filtern verringert werden, jedoch aus verschiedenen Gründen: Der Ausgangsfilter soll die Ausgangsspannung glätten, während der Eingangsfilter weitere Störungen verringern soll. Dimensionierung und Bauteil-Auswahl dieser Filter sind keineswegs trivial, weil sowohl die Eingangs- als auch Ausgangs-Störungen unterschiedliche Frequenzen enthalten und sie sowohl asymmetrische (differential mode - DM) als auch symmetrische (common mode - CM) Störungen enthalten.

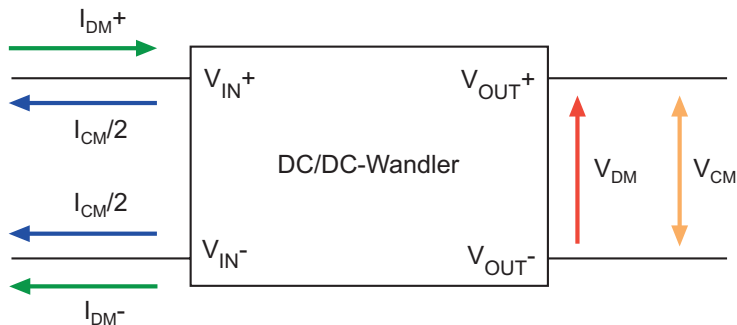


Abb. 5.1: Schaltbild von durch DC/DC-Wandler erzeugten Interferenzen

Die Applikationsingenieure von Recom werden immer wieder gefragt, warum die notwendigen Ein- und Ausgangsfilter nicht einfach in den Wandlern miteingebaut werden. Die Antwort ist, dass einfache Filter in allen Konstruktionen integriert sind, um für die meisten Applikationen Wandler mit einem annehmbaren Grad an Eingangs- und Ausgangswelligkeit zu liefern. Man könnte für einen höheren Preis bessere integrierte Filterung anbieten, was aber wiederum andere Anwender benachteiligen würde, die keine bessere Performance benötigen als die vom Standardprodukt gebotene.

Darüber hinaus besitzen viele Produktfamilien sehr kleine Bauformen die schlichtweg nicht genügend Raum im Wandler bieten, um größere Induktoren und Kondensatoren zu integrieren. Viele Anwender, die nur über wenig Platz auf der Leiterplatte der Applikationen verfügen, schätzen preiswerte und kleine DC/DC-Module und akzeptieren somit lieber den Kompromiss, dass Restwelligkeit und Rauschen etwas höher sind. Anwender, die besonders niedrigere Restwelligkeitswerte benötigen, sehen zusätzliche externe Filter vor, die den speziellen Anforderungen angepasst sind. Dies führt in der Regel zu besseren Ergebnissen als integrierte Standardfilter die ein sehr breites Spektrum abdecken sollen. Häufig ist es auch eine Frage der Kosten, da gerade bei kostensensitive Applikationen der hohe Bfilterungsgrad nicht erforderlich ist. Hier werden meist kostengünstigere DC/DC-Wandler bevorzugt.

5.2 Rückwelligkeitsstrom (back ripple current)

Der Eingangs-Rückwelligkeitsstrom wird in Milliampere Spitze-Spitze (mAp-p), in der Regel bei der Nenneingangsspannung und Volllast, definiert. Aber bevor er gefiltert werden kann, muss dieser zuerst messtechnisch ermittelt werden.

5.2.1 Messung des Rückwelligkeitsstroms

Eine Messung des Eingangsstroms mit einem DMM-Amperemeter (digitaler Multimeter) liefert einen Effektivwert, ohne eine Aussage über die Signalforn des pulsierenden Rückwelligkeitsstroms zu ermöglichen. Eine Messung des Eingangsstroms mit einer Oszilloskop-Stromzange erzielt häufig nicht die besten Ergebnisse. Dies beruht auf der hohen DC-Komponente des Eingangsstroms, die das Kernmaterial der Stromzange sättigt, sodass die AC-Komponente der Restwelligkeit nicht mehr ordentlich aufgelöst werden kann.

Die Lösung besteht darin, einen Präzisionsstrommesswiderstand, einen Shunt, einzusetzen, und um den Strom durch Messung der Spannung am Shunt zu ermitteln; hierbei ist jedoch Vorsicht geboten. Manche niederohmige Leistungswiderstände weisen eine Wickelstruktur auf, mit entsprechend hoher Längsinduktivität und einer damit einhergehenden Beeinflussung des Messergebnisses verursacht durch den Messaufbau (getreu dem Messtechniker-Grundsatz: Wer misst misst Mist). Für Messungen des AC-Anteils des Rückwelligkeitsstroms müssen deshalb Shuntwiderstände mit extrem niedrigen Längsinduktivitäten ($< 0,1\mu\text{H}$) eingesetzt werden. Metallstreifenwiderstände können solche Werte liefern.

Das Messverfahren selbst ist jedoch ebenso von allergrößter Bedeutung, da es leicht zu erheblichen Fehler kommen kann. Erstens sollte der Shuntwiderstand einen kleinen Widerstand besitzen, damit er auf die Eingangsspannung am Wandler keinen allzu großen Einfluss hat. Beim Einsatz eines $0,1\Omega$ -Shunts ermöglicht eine typische 5mV/Teilung Oszilloskop-Einstellung nur die Auflösung von Strömen von 50mA . Zweitens müssen die Verbindungen zum Tastkopf so kurz wie möglich gehalten werden, damit sie über Zuleitungen (z.B. über den Ground-Clip) keine Störstrahlung aufnehmen können. Abb. 5.2 zeigt, wie man den Shuntwiderstand korrekt mit dem Tastkopf berührt, und Abb. 5.3 zeigt durch korrekte und nicht korrekte Messverfahren erzielte unterschiedliche Ablesewerte:

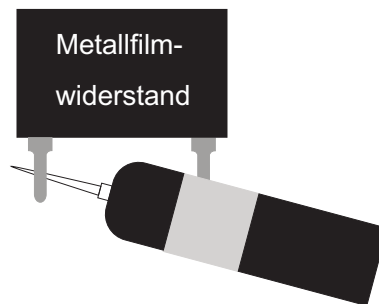


Abb. 5.2: Korrektes Meßverfahren

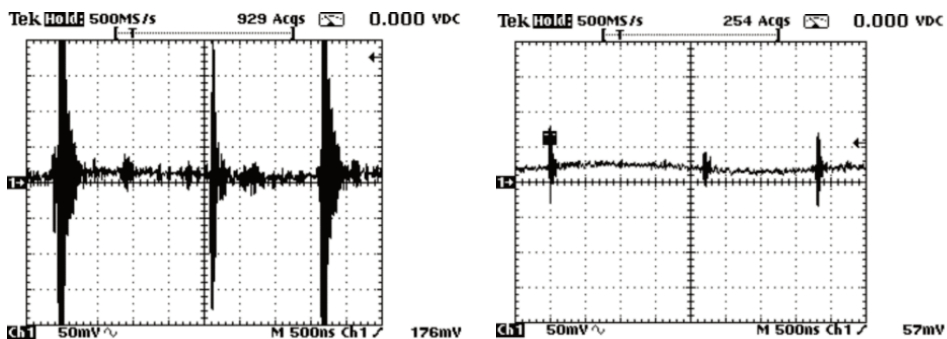


Abb. 5.3: Nicht korrekte (links) und korrekte (rechts) Ergebnisse für denselben Rückwelligkeitsstrom

5.2.2 Gegenmaßnahmen zur Reduktion des Rückwelligkeitsstroms

Die einfachste Weise, den Rückwelligkeitsstrom zu verringern, besteht darin, einen Elektrolyt- oder Tantalkondensator mit einem niedrigen ESR an den Eingangspins des DC/DC-Wandlers hinzuzufügen. Der Kondensator liefert die Energie für den pulsierenden Welligkeitsstrom mit viel niedrigerer Impedanz, als die primäre Stromquelle das durch die Leitungsimpedanz könnte. Somit versorgt die primäre Stromquelle die DC-Komponente des Eingangsstroms und der Kondensator puffert einen großen Teil der AC-Komponente des Eingangsstroms, sodass der AC-Strom, der von der primären Stromquelle fließt, nachfolgend wesentlich verringert wird. Abb. 5.4 veranschaulicht die Konzeption:

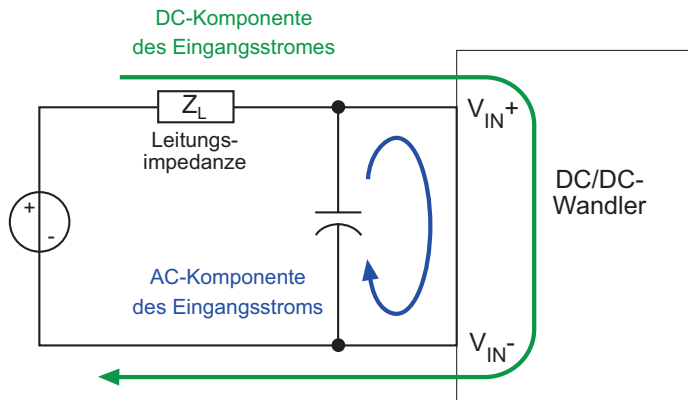


Abb. 5.4: Verminderung des Rückwelligkeitsstroms mit Eingangskondensator

Die folgenden Oszilloskop-Signale zeigen die Auswirkung eines Eingangskondensator auf den Rückwelligkeitsstrom eines DC/DC-Wandlers. Die Kurven wurden mit einem 1- Ω -Shuntwiderstand erstellt, um ein besser aufgelöstes Signalbild zu erhalten:

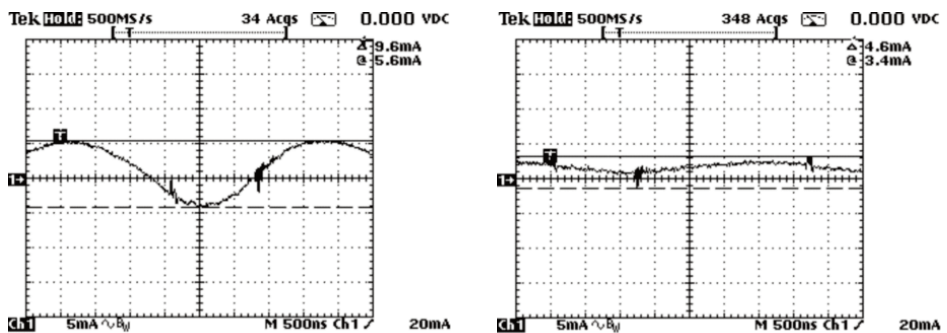


Abb. 5.5: Wirkung eines 47- μ F-Kondensators an den Eingangsklemmen des DC/DC-Wandlers

Der Rückwelligkeitsstrom wurde durch das Hinzufügen eines 47- μ F-Kondensators mit einem ESR-Wert von 400m Ω @ 100kHz mehr als halbiert. Wird ein teurerer Kondensator mit einem ESR von nur 35m Ω verwendet, ist es schwierig, die Restwelligkeit am Oszilloskop zu messen, da nur die Störspitzen infolge der Schaltstörungen als leicht messbare Störung übrig bleiben.

Praktischer Hinweis

Eine Alternative zur Nutzung von sehr teuren Ultra-low-ESR Kondensatoren besteht darin, zwei parallel geschaltete Standardkondensatoren zu verwenden. Ein einziger hochwertiger 47- μF -Kondensator könnte beispielsweise durch zwei 22- μF -Standardkondensatoren je mit 230m Ω ersetzt werden, was einen äquivalenten 44- μF -Kondensator mit einem ESR von 115m Ω ergibt (Abb. 5.6).

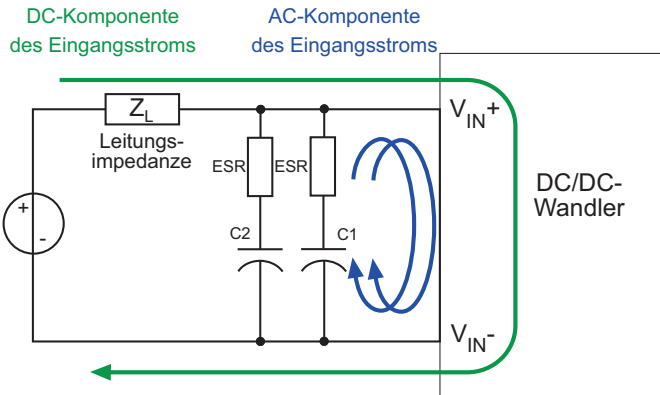


Abb. 5.6: Verminderung des Rückwelligkeitsstroms mit zwei parallel geschalteten Eingangskondensatoren

Die Wirkung auf den Rückwelligkeitsstrom ist bei Verwendung von zwei preiswerten Kondensatoren nicht wesentlich schlechter als mit einem teurem Ultra-low-ESR Kondensator (Abb. 5.7).

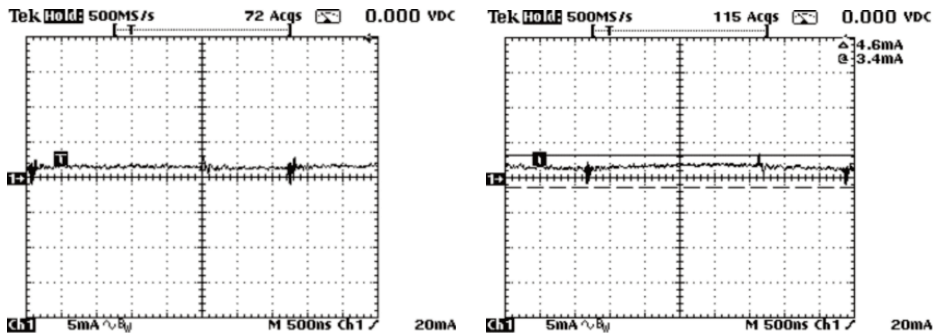


Abb. 5.7: Welligkeitsstromvergleich zwischen einem Ultra-low-ESR-47- μF Kondensator und zwei parallel geschalteten Standard-22- μF -Kondensatoren

Die verbleibenden hochfrequenten Stromspitzen werden durch Schaltstörungen des Wandlers verursacht. Dieses Rauschen erscheint gleichzeitig sowohl an den $V_{\text{IN}+}$ - als auch den $V_{\text{IN-}}$ -Anschlussleitungen des Wandlers und kann somit mit einem Kondensator am Eingang nicht ausgefiltert werden. Diese Art von CM-Störungen kann nur durch eine Eingangsgeleitetaktrossel beseitigt werden (siehe unten).

Bei niedrigen Eingangsspannungen kann ein Mehrschichtkeramikkondensator (MLCC) anstelle eines Elektrolytkondensators verwendet werden.

Ein hochwertiger MLCC hat einen ESR-Wert von ca. $3\text{m}\Omega@100\text{ kHz}$ und stellt somit einen auszeichneten Rückwelligkeitsstrom-Kondensator dar. Es muss darauf geachtet werden, dass die Eingangsspannung die maximal zulässige Spannung des MLCC nicht überschreiten kann, da es sonst zu einem internen Lichtbogenüberschlag kommen kann, was zu einem Totalausfall des MLCCs führen würde. Daher sollten MLCCs nur an geregelten primären Stromquellen oder überspannungsgeschützten Versorgungen eingesetzt werden.

5.2.3 Auswahl des Eingangskondensators

Im vorhergehenden Beispiel wurde ein $47\text{-}\mu\text{F}$ -Kondensator verwendet, um den Rückwelligkeitsstrom zu verringern. Woher stammte jedoch der Wert $47\mu\text{F}$? Offensichtlich ist, dass je größer die Kapazität, desto mehr Energie kann geliefert werden, um den Wandler zu speisen. Größere Kondensatoren haben auch infolge der größeren internen Fläche zwischen den Elektroden-schichten niedrigere ESR-Werte. Aber Elektrolytkondensatoren mit großer Kapazität benötigen mehr Platz und sind teuer. Der Auswahlprozess hängt deshalb von den Kostenaspekten genauso wie von der zu erwartenden Wirkung ab. Typische Eingangskondensatorwerte können zwischen $22\mu\text{F}$ und $22\mu\text{F}$ variieren, weshalb es sich bei $47\mu\text{F}$ um einen allgemein praktizierten Kompromiss und guten Startwert für erste Tests handelt.

Wichtiger als der Kapazitätsnennwert ist die Reaktion des Kondensators auf den AC-Strom. Der AC-Strom, der durch den Kondensator fließt, generiert Wärme. Überschreitet die Temperatur des Kondensators die Grenzwerte laut Spezifikation, verringert sich die Lebenszeit des Kondensators immens. Im Extremfall kann der Elektrolyt im Kondensator die Siedetemperatur erreichen, was einen unmittelbaren Ausfall des Kondensators zur Folge hat.

Praktischer Hinweis

Es ist sehr schwierig, den AC-Welligkeitsstrom im Kondensator zu messen, da das Hinzufügen eines Mess-Shuntwiderstands in Reihe einen wesentlichen Einfluss auf das Ergebnis hat und die realen Bedingungen (ohne Shunt) grob verfälscht. Besser ist es den gesamten Rückwelligkeitsstrom ohne externe Kondensatoren und dann erneut mit den eingesetzten Kondensatoren zu messen. Die Differenz beider Messergebnisse ergibt dann den Welligkeitsstrom, der in den Kondensatoren fließt.

Wenn der ESR des Kondensators und die Taktfrequenz des Wandlers f bekannt sind, kann alternativ die restliche Restwelligkeit der Eingangsspannung infolge der Leitungs-impedanz Z_L gemessen und der Welligkeitsstrom berechnet werden, durch folgende Gleichung:

$$I_{\text{RIPPLE}} = \frac{V_{\text{RIPPLE}}}{\sqrt{\text{ESR}^2 + \left(\frac{1}{2\pi f C}\right)^2}}$$

Gleichung 5.1: Berechnung des Kondensatorwelligkeitsstroms

Im Kondensatordatenblatt sind Spezifikationen für den maximal empfohlenen Welligkeitsstrom enthalten. Der begrenzende Faktor ist der Temperaturanstieg durch die im Kondensator entstehende Verlustleistung.

Die infolge des Welligkeitsstroms im Kondensator entstehende Verlustleistung ist:

$$P_{C,DISS} = I_{RIPPLE}^2 ESR$$

... und der daraus resultierende Temperaturanstieg:

$$T_{RISE} = P_{C,DISS} kA$$

Gleichung 5.2: Berechnung des Kondensatortemperaturanstiegs infolge des Welligkeitsstroms

wobei kA die Wärmeleitung des Kondensators ist, d. h. Wärmeimpedanz k mal die Fläche des Kondensators A . Wärmeleitung wird in $^{\circ}C/W$ gemessen.

Praktischer Hinweis

Die Messung des Rückwelligkeitsstroms ist schwierig. Daher ist es manchmal einfacher, die Kondensatortemperatur zu messen und den Welligkeitsstrom aus dem Temperaturanstieg abzuleiten.

5.2.4 Eingangsstrom von parallel geschalteten DC/DC-Wandlern

Es gibt mehrere Applikationen, bei denen es erforderlich ist, mehrere DC/DC-Wandler an einer einfachen Primärspannungsversorgung parallel zu betreiben. Die am meisten verbreiteten sind Point-of-Load (POL)- und redundante (N+1) Stromversorgungssysteme. Jeder DC/DC-Wandler erzeugt dabei seinen eigenen Rückwelligkeitsstrom, die sich alle zu einer Gesamtstromaufnahme addieren und gegenseitig überlagern.

Betrachten wir zwei identische DC/DC-Wandler mit einer Nenn-Schaltfrequenz von 100kHz. Infolge der Fertigungstoleranzen könnte der eine über 100kHz und der andere über 120kHz verfügen. Eine FFT-Analyse würde drei Frequenzlinien zeigen: 100kHz, 120kHz und die Differenz von 20kHz. Eine solche Niederfrequenz-Überlagerung oder Schwebungsfrequenz ist extrem schwierig auszufiltern.

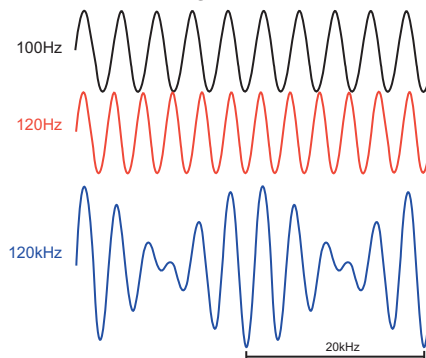


Abb. 5.8: Schwebungsfrequenzstörungen

Die Schwebungsfrequenzstörungen kann man vermeiden, indem jeder Eingang des DC/DC-Wandlers (Abb. 5.9) individuell befiltert wird. Der LC-Filter sperrt die Schwebungsfrequenzstörungen zwischen den einzelnen Wandlern.

Die Induktivitäten müssen für hohen DC-Strom ausgelegt sein, weshalb typische Werte für L sehr niedrig sind: $22\mu\text{H}$ bis $220\mu\text{H}$. Außerdem muss auch ein Kondensator an der Versorgung platziert werden. Der Filterungseffekt von LC-Tiefpassfiltern ist bidirektional. Daher ist der durch C_{MAIN} -L-C gebildete resultierende π -Filter hilfreich, um Interferenzen noch stärker zu vermindern.

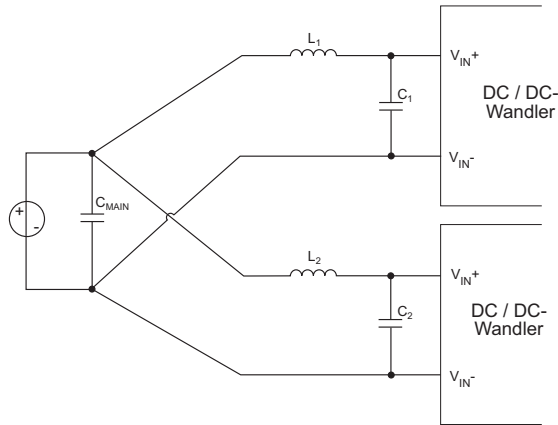


Abb. 5.9: Filterung von Schwebungsfrequenzstörungen

**Praktischer
Hinweis**

Es ist wichtig, dass die Eingangskondensatoren C_1 und C_2 so nah wie möglich an den Eingangs-Pins des Wandlers platziert werden. Selbst sehr kurze Zuleitungslängen der PCB-Leiterbahnen zwischen den Kondensatoren und dem Wandlern verringern die Wirksamkeit eines Filters. Der gemeinsame V_{IN} -Anschluss sollte massiv sein und eine möglichst niedrige Impedanz aufweisen. Alle Verbindungen sollten sich an den Anschlüssen der Versorgung (als Sternpunkte) treffen, um weitere Überlagerungseffekte zu vermeiden.

5.3 Ausgangsbefilterung

Wie in Abschnitt 2 erwähnt, haben alle DC/DC-Wandler-Ausgangsspannungen einen gewissen Anteil an Ausgangsrestwelligkeit.

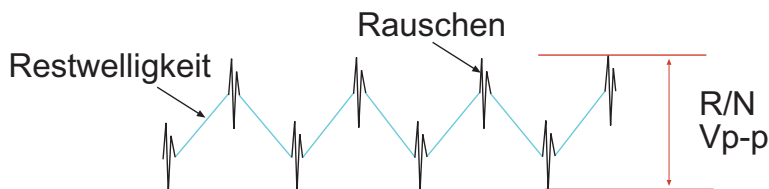


Abb. 5.10: Ausgangswelligkeit

Die Filterung der Ausgangswelligkeit erfordert zwei verschiedene Verfahren, da Restwelligkeit (Ripple) eine asymmetrische (differentiale) Störung ist, während das Rauschen (Noise) eine symmetrische (Gleichtakt-) Störung darstellt.

5.3.1 Differenzmodus-Ausgangsbefilterung

Die einfachste Methode zur Verringerung der Ausgangsrestwelligkeit besteht darin, einen zusätzlichen Kondensator am Ausgang hinzuzufügen (Abb. 5.11). Der externe Kondensator C_{EXT} ist mit dem internen Kondensator C_{OUT} parallel geschaltet.

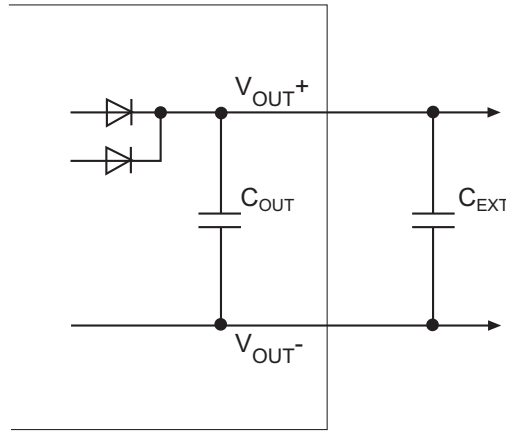


Abb. 5.11: Filterung der Ausgangsrestwelligkeit unter Verwendung eines externen Kondensators

Die Wirksamkeit dieses Verfahrens zur Verringerung der Ausgangsrestwelligkeit in mV, $V_{RIPPLE,p-p}$, hängt von der Gesamtkapazität, dem Ausgangsstrom und der Taktfrequenz des Wandlers laut Gleichung 5.3 ab.

$$V_{RIPPLE,p-p} = \frac{I_{OUT} \cdot 1000}{2 \cdot f_{OPER} \cdot (C_{OUT} + C_{EXT})}$$

Gleichung 5.3: Berechnung der Ausgangsrestwelligkeit

Wie aus dieser Gleichung ersichtlich ist, ist das Hinzufügen des externen Kondensators zur Verringerung der Welligkeit nur bis zu einer gewissen Grenze sinnvoll. Wenn beispielsweise ein vollweggleichgerichteter Wandler eine Ausgangskapazität von $22\mu F$ mit einem Strom von 1A und einer Taktfrequenz von 100kHz besitzt, würde die Ausgangswelligkeit ohne jegliche externe Beschaltung 226mVp-p betragen. Das Hinzufügen eines externen $22\mu F$ -Kondensators halbiert die Restwelligkeit auf 112mVp-p. Wenn die erforderliche Ausgangsrestwelligkeit wiederum die Hälfte davon beträgt, nämlich 56 mVp-p, ist $90\mu F$ Gesamtkapazität, d. h. ein externer Kondensator von $68\mu F$ notwendig. Eine weitere Verminderung der Restwelligkeit auf 20mVp-p würde fast $2500\mu F$ externe Kapazität erfordern. Solch eine hohe Ausgangskapazität könnte bei dem DC/DC-Wandler jedoch zu Anlaufproblemen führen, die Reaktion der Spannungsanstiegsgeschwindigkeit auf beliebig schnelle Laständerungen beeinträchtigen und das Wiederanlaufen des Wandlers nach einem Ausgangskurzschluss verzögern.

Eine praktikablere Lösung um niedrigere Ausgangsrestwelligkeit zu erzielen besteht im Hinzufügen einer Induktivität, sodass diese gemeinsam mit dem externen Kondensator einen Tiefpassfilter bildet:

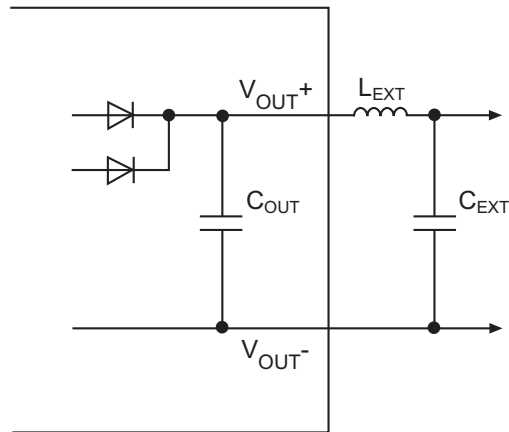


Abb. 5.12: Ausgangsrestwelligkeitsbefilterung unter Verwendung eines externen LC-Filters

Durch das Hinzufügen einer Induktivität ergibt sich folgende Berechnung der Ausgangsrestwelligkeit:

$$V_{\text{RIPPLE,p-p}} = \frac{I_{\text{OUT}} 1000}{2 f_{\text{OPER}} (C_{\text{OUT}} + \sqrt{L_{\text{EXT}} C_{\text{EXT}}})}$$

Gleichung 5.4: Berechnung der Ausgangsrestwelligkeit mit LC-Filter

Unter Verwendung des vorherigen Beispiels, wenn L_{EXT} z. B. $100\mu\text{H}$ beträgt, kann eine Ausgangsrestwelligkeit von 20mVp-p mit einem Ausgangskondensator C_{EXT} von nur $645\mu\text{F}$ erreicht werden. Das ist eine deutliche Verbesserung gegenüber $2500\mu\text{F}$ ohne Induktivität. Es muss darauf geachtet werden, dass die Induktivität für den Ausgangsstrom ausgelegt ist.

Wenn der interne Aufbau und die Bauteilwerte im Inneren des DC/DC-Wandlers nicht bekannt sind, lautet die effektive Faustregel, die Eckfrequenz des LC-Filters bei $1/10$ der Taktfrequenz des Wandlers festzulegen. Dies ergibt meist eine ausreichende Reduzierung der Ausgangsrestwelligkeit und hält die Kosten für die Filterkomponenten gering:

$$f_c = f_{\text{OPER}}/10 = \frac{1}{2\pi} \sqrt{\frac{1}{LC}}$$

Gleichung 5.5: Faustregelberechnung des Ausgangs-LC-Filters

Die Grenzfrequenz f_c ist der Punkt im Amplitudendiagramm des Filters, ab welchem das Störsignal bereits um -3dB unterdrückt, d. h. bereits um 30% gedämpft, ist. Da der LC-Filter ein Tiefpassfilter zweiter Ordnung ist, die Dämpfungskurve also mit -40dB pro Dekade abfällt, werden Frequenzen, die zehnmal höher sind als die Grenzfrequenz, um Faktor 100 gedämpft.

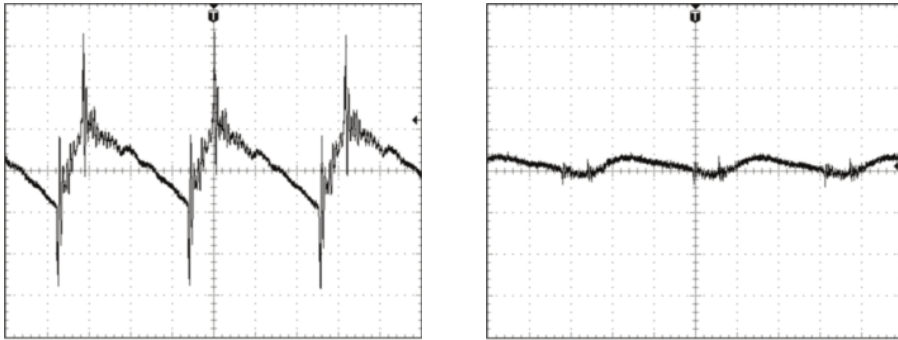


Abb. 5.13: Vorher-/Nachher-Ergebnisse eines effektiven Ausgangsrestwelligkeits-Filters

5.3.2 Gleichtakt-Ausgangsbefilterung

Wie oben erwähnt, bestehen die Ausgangsstörungen sowohl aus asymmetrischen als auch symmetrischen Anteilen. Die Restwelligkeit stellt im Wesentlichen eine Differentialstörung dar und das Rauschen eine Gleichtaktstörung. Da ein symmetrisches Rauschsignal an allen Ausgängen gleichzeitig vorhanden ist, kann es von keiner Ausgangskapazität „bemerkt“ werden, und das Hinzufügen einer Ausgangs-LC-Filterung führt nicht zu einer Verringerung dieser Interferenz. Gleichtaktstörungen stellen kein Problem dar, wenn die Last vollkommen symmetrisch, linear und isoliert wäre. Jegliche Nichtlinearitäten in der Lastkennlinie oder Stromwege zurück zur Erde „richten“ jedoch die Gleichtaktstörungen „gleich“ und erzeugen daraus resultierende Differentialstörungen. Daher ist es notwendig ebenso diese Gleichtaktstörungen zu beachten. Es gibt zwei Möglichkeiten, um Gleichtaktstörungen zu reduzieren: „Kurzschließen“ des Rauschens durch einen niederohmigen Pfad oder die Verwendung von Gleichtaktdrosseln.

Gleichtaktausgangsrauschen wird in den meisten Fällen durch Schaltspannungsspitzen an der Eingangsseite, die über die Koppelkapazität des Transformators an den Ausgang übertragen werden (Abb. 5.14), hervorgerufen. Um diese Störungen zu verringern, muss ein Strompfad zurück zur Eingangsseite geschaffen werden. Da der Ausgang galvanisch getrennt ist, muss ein Rückstrompfad über externe Kondensatoren ausgewählt werden, um für diese hochfrequenten Störungen eine niedrige Impedanz zu bieten.



5.3.3 Gleichtaktdrosseln

Bei einigen Anwendungen ist es nicht zulässig, Kondensatoren zur Unterdrückung der Gleichtaktstörungen an der Isolationsbarriere zu platzieren. Medizinische Geräte besitzen beispielsweise strenge Vorgaben was Leckströme angeht, die durch Vorhandensein eines niederohmigen Pfades an der Isolationsbarriere für Hochfrequenzen auftreten können. Bei solchen Anwendungen muss stattdessen eine Gleichtaktdrossel eingesetzt werden. Die Besonderheit der Gleichtaktdrossel besteht darin, dass sie zwei Wicklungen besitzt, die in entgegengesetzter Richtung auf denselben Kern gewickelt sind (Abb. 5.15).



Obwohl die Gleichtaktströme I_S in die gleiche Richtung fließen, generieren sie infolge der gegenläufigen Wicklungen einen Gesamtmagnetfluss im Kern. Deshalb dämpft die Impedanz des Kerns diese Gleichtaktströme wirksam. Die gegenläufig fließenden Vorwärts- und Sperrströme I_N erzeugen kein Gesamtmagnetfeld und werden daher nicht gedämpft. Dies zeigt den Vorteil, dass der Kern auch bei hohen Differenzmodusströmen nicht gesättigt wird. Weshalb ein Kernmaterial mit hoher Permeabilität verwendet werden kann, um das CM-Rauschen auszufiltern, ohne Gefahr zu laufen eine Überhitzung infolge des durch die Spule fließenden DM-Stroms zu bewirken.

Abb. 5.16 zeigt eine Ausgangsgleichtaktdrossel, die gemeinsam mit einem DC/DC-Wandler verwendet wird. Eine Wicklung ist im Pfad der V_{OUT+} -Leitung und die andere in der V_{OUT-} -Rückleitung. Die Impedanz der Gleichtaktdrossel wird so gewählt, dass ihr Maximum dicht an der Frequenz mit dem energetisch stärksten Anteil der Gleichtaktstörungen liegt (üblicherweise im Bereich von 10 bis 100MHz). Gleichtaktdrosseln dämpfen das durch die hohe Permeabilität des Kernmaterials erzeugte CM-Rauschen jedoch ohnehin über einen sehr großen Frequenzbereich.

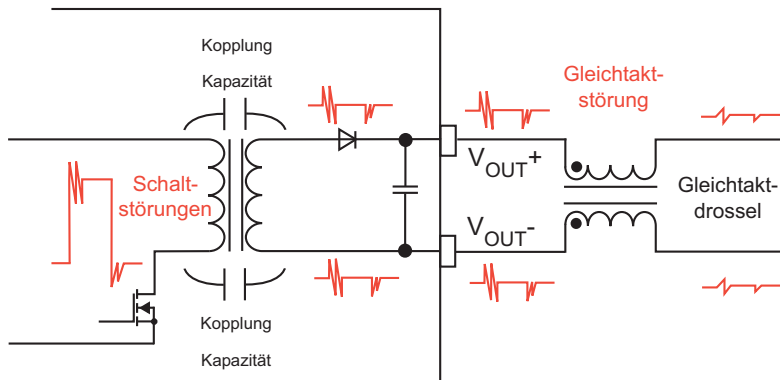


Abb. 5.16: Gleichtaktdrossel als DC/DC-Ausgangsfiler

Das Prinzip der Gleichtaktunterdrückung unter Verwendung einer Gleichtaktdrossel kann auch auf bipolare Ausgangswandler erweitert werden. CM-Rauschen erscheint an allen drei Ausgangsanschlüssen gleichzeitig und ist daher mit Standard-CM-Drosseln mit nur zwei Wicklungen besonders schwer auszufiltern. Die Lösung besteht darin, eine Gleichtaktdrossel mit drei Wicklungen einzusetzen. Ein positiver Nebeneffekt einer dreifachen Gleichtaktdrossel liegt wiederum darin, dass sie durch das Hinzufügen von zwei zusätzlichen Kondensatoren auch zum Ausfiltern von DM-Rauschen verwendet werden kann.

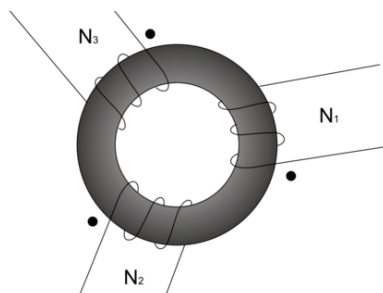


Abb. 5.17: Gleichtaktdrossel mit drei Wicklung

Diese drei Wicklungen werden separat auf den Kern gewickelt und, um etwas Streuinduktivität L_i zwischen den Wicklungen zu erreichen, getrennt gehalten. Bei der Auswahl des Kernmaterials ist es wichtig, eine hohe Permeabilität zu wählen, sodass die Anzahl der Windungen und somit der Wicklungswiderstand gering bleibt. Zur Berechnung der Induktivität werden die folgenden Verhältnisse verwendet:

$$\begin{aligned} N &= N_1 = N_2 = N_3 && \text{Anzahl der Windungen} \\ LC &= L_1 = L_2 = L_3 && \text{Induktivität} \\ LC &= N^2 A_L && \text{Wicklungsinduktivität} \end{aligned}$$

Gleichung 5.6: Berechnung der Wicklungsinduktivität

Der Induktivitätsfaktor AL ist die Induktivität je Windung [nH/N^2] und hängt von der Konfiguration der Spule und vom Kernmaterial ab. Die Streuinduktivität zwischen den Wicklungen L_S beträgt in der Regel ca. 3% der Wicklungsinduktivität L_C und kann genutzt werden, um hochfrequente Differenzmodusstörungen auszufiltern, wenn zwei zusätzliche Kondensatoren verwendet werden.

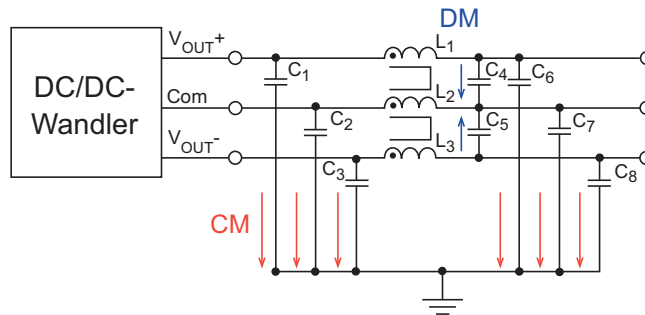


Abb. 5.18: Gleichtakt-Dreifachdrossel als kombinierter DC/DC-Ausgangsfiler

Die Kondensatoren $C_1 - C_3$ gewährleisten einen niederohmigen Pfad um das CM-Rauschen gegen Masse abzuleiten. Hochspannungs-Keramikplattenkondensatoren in der Größenordnung von 1 bis 10nF sind geeignet, obwohl auch MLCC-Kondensatoren verwendet werden können, wenn die Isolations-Prüfspannung niedrig ist. Je nach innerer Struktur des DC/DC-Wandlers können C_1 und C_3 vernachlässigt werden. Die Kondensatoren C_4 und C_5 bilden einen Differenzmodus-Tiefpassfilter in Kombination mit der Streuinduktivität zwischen den Wicklungen L_1/L_2 und L_2/L_3 . Die Kondensatoren C_4 und C_5 verfügen in der Regel über eine Kapazität in der Größenordnung von $>1\mu F$. MLCCs sind hierfür eine gute Wahl. Jedes CM-Rauschen das seinen Weg durch die Drossel über die Querkapazität zwischen den Wicklungen findet, kann durch den zweiten Satz von CM-Kondensatoren von C_6 bis C_8 zur Masse abgeleitet werden. Die Wicklungsinduktivität für die Drossel beträgt normalerweise einige Hundert Millihenry, somit beträgt die Streuinduktivität zur Berechnung des DM-Filters 5 bis 10 μH .

Die folgenden Berechnungen können verwendet werden:

$$\text{Differenzmodus: } C_{DM} = C_4 = C_5$$

$$f_{r,DM} = \frac{1}{2\pi\sqrt{L_S C_{DM}}} = \frac{1}{2\pi\sqrt{0.03 L_C C_{DM}}}$$

$$\text{Gleichtakt: } C_{CM} = C_1 = C_2 = C_3 = C_6 = C_7 = C_{8V}$$

$$f_{r,CM} = \frac{1}{2\pi \sqrt{L_C C_{CM}}}$$

Gleichung 5.7: Berechnung der Parameter einer Gleichtakt-Dreifachdrossel

5.4 Vollfilterung

Die Gleichtaktdrossel kann auch gegen CM-Interferenzen, die an der Primärseite entstehen, eingesetzt werden. Da die Differenzmodus-Störungen am Eingang in Bezug auf Gleichtaktstromstörungen sehr hoch (Einschaltspitzenstrom und reflektierter Welligkeitsstrom, engl.: inrush current bzw. back ripple current) sein können, könnte man meinen, sich keine Gedanken über den CM-Eingangsstrom machen zu müssen. Um jedoch die EMV-Konformität zu gewährleisten, ist gerade dies häufig erforderlich. Ein voll gefilterter DC/DC-Wandlerschaltkreis ist in Abb. 5.19 dargestellt:

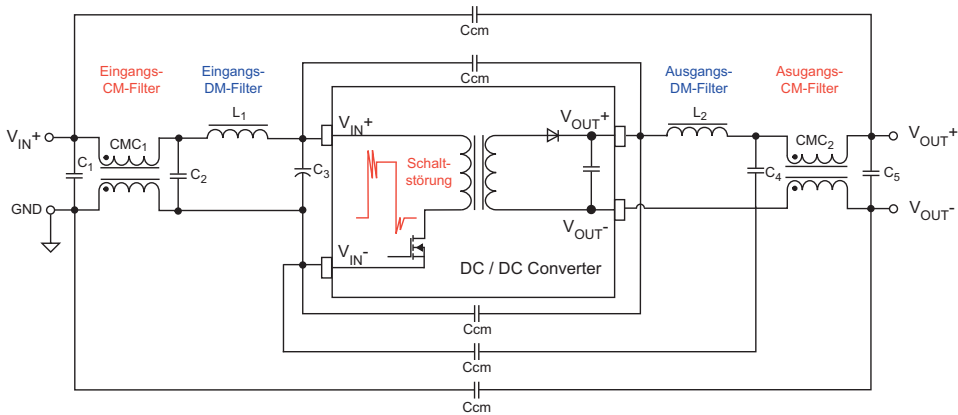


Abb. 5.19: Voll gefilterter DC/DC-Wandler

Es muss betont werden, dass in vielen Anwendungen nicht alle in Abb. 5.18 gezeigten Komponenten erforderlich sein müssen. Der Vollfilter sollte nur nach Bedarf bestückt werden, da jegliche zusätzlichen Komponenten den Gesamtwirkungsgrad verringern. Bei einigen Anwendungen sind nur der Eingangskondensator C_3 und ein oder mehrere CM-Kondensatoren CCM für die Gewährleistung der EMV-Konformität ausreichend.

Um die Anzahl der eingesetzten Bauteile gering zu halten, kann eine Gleichtaktdrossel durch Änderung der Anschlussanordnung als DM-Spule verwendet werden. Dies bedeutet, dass es möglich ist, $CMC_1 = L_1$ und $CMC_2 = L_2$ zu wählen. Dies ist besonders zu empfehlen, wenn SMD-Drosseln verwendet werden, da dann nur zwei Bestückungsrollen für alle vier Induktivitäten vonnöten sind.

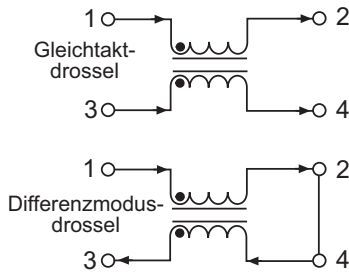


Abb. 5.20: Eine als DM-Induktivität verwendete Gleichtaktdrossel

5.4.1 Filter-PCB-Layout

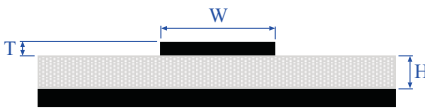
Das Layout der PCB-Leiterbahnen ist für die Eingangs- oder Ausgangsfilterleistung von entscheidender Bedeutung. Wie schon erwähnt, muss der Eingangskondensator möglichst nahe an den Eingangspins montiert werden. Da der ESR eines hochwertigen Kondensators in mΩ-Bereich liegt, muss die Impedanz jeglicher Verbindung zwischen dem Kondensator und den Wandlerpins ebenso in mΩ-Bereich liegen, um eine ausreichend gute Filterwirkung zu erreichen. Gleichung 5.8 zeigt die Berechnung des Leiterbahnwiderstands:

$$\text{Leiterzug-} \quad = \text{spezifischer Widerstand} \quad \frac{\text{Länge}}{\text{Stärke} \times \text{Breite}} \quad [1 + (\text{TempCo} \times (\text{Temp} - 25))]$$

Gleichung 5.8: Berechnung des Leiterbahnwiderstand

Eine typische PCB hat eine Kupferstärke von 35µm, sodass eine 1mm breite und 1cm lange Leitung bei 25°C einen DC-Widerstand von knapp 5mΩ besitzt, der bei +85°C auf 6mΩ ansteigt (spezifischer Widerstand von Kupfer = 1,7 x 10⁻⁶ Ω/cm und TempCo = +0,393%/°C).

Zusätzlich zum DC-Widerstand muss auch die AC-Leiterbahnimpedanz betrachtet werden. Eine PCB-Leiterbahn weist sowohl eine Induktivität als auch eine Eigenkapazität in Bezug auf andere Leiterbahnen und die Bauteile auf. Das kann zu unerwarteten Ergebnissen führen, da Störungen kapazitiv oder induktiv zwischen Leiterbahnen, Leiterbahn-Schichten und Bauteilen auf der PCB auftreten. Eine obere PCB-Leiterbahn hat beispielsweise beim Überqueren einer anderen Leiterbahn an der PCB-Unterseite oder innerhalb einer mehrschichtigen PCB einen charakteristischen Leitungswiderstand Z_0 und eine Kapazität C_0 laut Gleichung 5.9:



$$Z_0 = \frac{87}{\sqrt{\epsilon_r + 1.41}} \ln \left(\frac{5.98H}{0.8W + T} \right) \text{ ohms} \quad C_0 = \frac{0.67 (\epsilon_r + 1.41)}{\left(\frac{5.98H}{0.8W + T} \right)} \text{ [pF/inch]}$$

Für ein typisches PCB $\epsilon_r = 4$, $H = 30\text{mil}$ (0,76mm) und $T = 1,37\text{mil}$ (35µm)

Gleichung 5.9: Berechnung der Leiterbahnimpedanz und -kapazität

Daher ist es wichtig, dass die in den Filterschaltungen verwendeten PCB-Leiterbahnen nicht die anderen Signalleiterbahnen auf der anderen Seite oder in einem anderen Layer queren oder zu nahe an sie anschließen. Im Idealfall sollte ein zweiseitiges oder mehrschichtiges Layout verwendet werden, sodass ein Masselayer unterhalb der Filterkomponenten platziert werden kann. Ist die PCB nur einseitig, sollten die Anschlüsse möglichst kurz und breit gehalten werden.

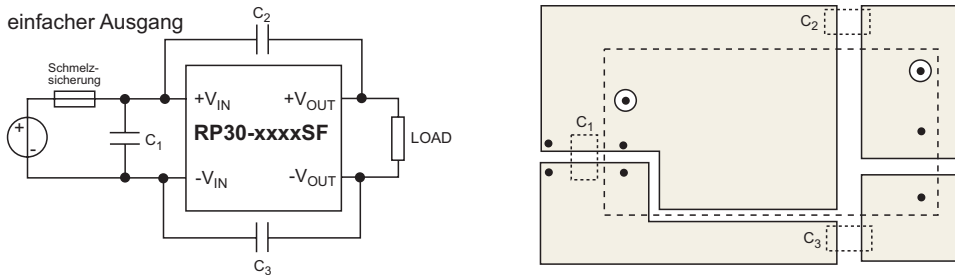


Abb. 5.21: Beispiel eines einfachen Klasse-A-Filters und PCB-Layouts (Serie RP30-SF)

Filterkomponenten müssen auch als reale und nicht nur als ideale Bauteile betrachtet werden. Dies bedeutet, dass bei hohen Frequenzen die Streuinduktivität eines Kondensators oder die Fremdkapazität einer Induktivität, die führende Rolle bei der Bestimmung des Filterverhaltens übernehmen kann. Mit anderen Worten Kondensatoren, die sich wie Induktivitäten zu verhalten beginnen, und umgekehrt. Widerstände können sich entweder wie Induktivitäten oder wie Kondensatoren verhalten. Durch geeignete Bauteilauswahl können diese Probleme reduziert oder vollständig vermieden werden. Das wichtigste Dimensionierungskriterium ist die Resonanzfrequenz, ab der sich das Verhalten ändert. Abb. 5.22 zeigt eine Kurve für Impedanz vs. Frequenz für eine kapazitive Komponente.

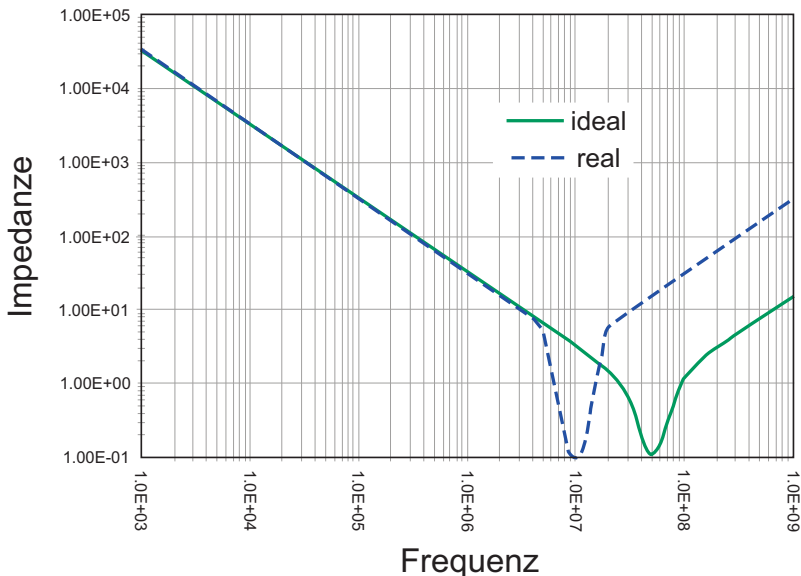


Abb. 5.22: Resonanzfrequenz eines Kondensators

Die durchgezogene Linie zeigt das Übertragungsverhalten des Bauteils an sich, für einen 4,7-nF-Kondensator mit einem parasitären ESR von $0,01\Omega$ und einer parasitären $ESL = 2,5nH$. Die gestrichelte Linie zeigt dasselbe Bauteil, wobei zusätzlich eine wenig optimale Anbindung des Bauteils simuliert wurde. Die Anbindung fügt einen zusätzlichen ESR von $50\text{ m}\Omega$ und eine ESL von $50nH$ hinzu. Wie im Diagramm erkennbar, verursacht diese schlechte Verbindung eine geringere Resonanzfrequenz. Das bedeutet, dass der Kondensator anfängt, sich wie eine Induktivität zu verhalten, und das bei einem Zehntel der berechneten Resonanzfrequenz.

Für den PCB-Entwickler hat das in Abb. 5.23 gezeigte Verhalten zur Folge, dass der ESL -Wert für den GND-Anschluss des Kondensators so niedrig wie möglich gehalten werden muss. Es genügt nicht, nur eine elektrische Verbindung zum Masselayer mit einem einzelnen Kontaktloch herzustellen. Es sind Mehrfach-Durchkontaktierungen notwendig, um sowohl die DC-Impedanz als auch die AC-Impedanz zu reduzieren.

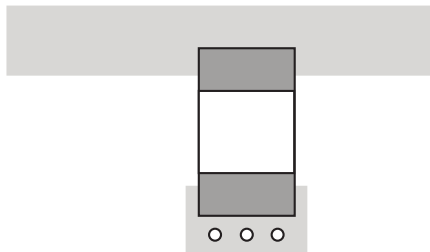


Abb. 5.23: GND-Anschluss mittels Mehrfachdurchkontaktierungen

Für Induktivitäten hat die Länge der Zuleitung geringere Bedeutung, da eine lange Leiterbahn lediglich die Gesamtinduktivität erhöht; es ist jedoch generell anzuraten, jede Störung möglichst nahe an der Quelle zu unterdrücken.

Bei allen Filterlayouts sollte auf die fließenden Ströme geachtet werden. Jeglicher in einer Schleife fließender Strom erzeugt ein elektromagnetisches Feld, das Störungen in anderen Teilen der Schaltung hervorrufen kann. Im Idealfall sollte eine sternförmige Masseanbindung vorgesehen werden, bei der alle Masseströme zu einem einzigen Punkt zurückfließen. Ist eine Schleife unvermeidbar, sollte der Schleifenbereich möglichst klein gehalten werden.

Mit einem guten PCB-Layout und den richtig gewählten Bauteilen können die Ergebnisse jedoch erstaunlich sein.

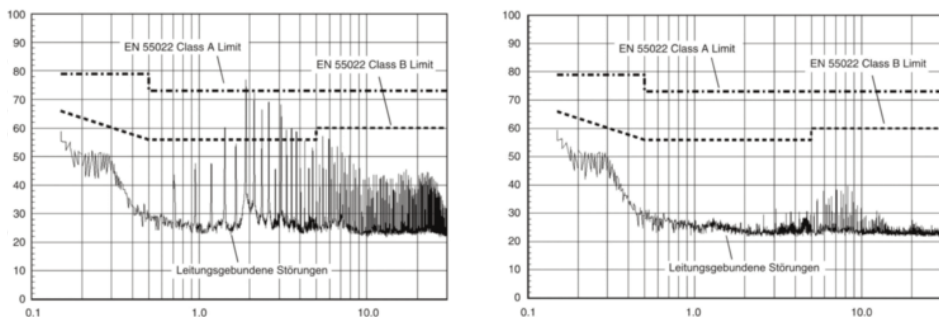


Abb. 5.24: Beispiel von Störungen vor und nach Filterung

6. Sicherheit

Die Hauptziele verschiedener Sicherheitsstandards und -vorschriften bestehen darin Sach- oder Personenschaden zu vermeiden, indem man Schutzgrade gegen folgende potentielle Gefahren definiert:

- Elektrischer Stromschlag
- Gefährliche Energie
- Brand und Rauchentwicklung
- Körperverletzung
- Gefährdung durch Strahlung und Chemikalien

Die Begriffe "Gefahr" und "Risiko" werden häufig austauschbar verwendet. Dies lässt sich differenzieren, indem man eine Gefahr als ein potentielles Risiko betrachtet. Eine Netzleitung kann beispielsweise eine gefährliche Spannung führen, die Leitung an sich ist jedoch immer noch sicher zu handhaben, da die Drähte isoliert sind. Wird die Isolierung allerdings beschädigt oder ist minderwertig ausgeführt, dann ist es gefährlich, die Leitung zu berühren.

Wie in der Einleitung zu diesem Buch erwähnt, besteht ein wichtiger Nutzen von DC/DC-Wandlern darin, die Sicherheit der Anwendungen, in denen sie eingesetzt werden, zu erhöhen. Wenn der DC/DC-Wandler ein sicherheitszertifiziertes Produkt ist, kann der Geräteentwickler den Wandler wie eine "Blackbox" behandeln und sich darauf verlassen, dass der Hersteller des DC/DC-Wandlers adäquate interne Sicherheitsmaßnahmen getroffen hat, um die Sicherheitsvorschriften zu erfüllen.

Dies bedeutet nicht, dass der Geräteentwickler deshalb nicht mehr für die Benutzersicherheit seiner Applikation verantwortlich ist, da er immer noch die verkehrsübliche Sorgfalt beim Erkennen potentieller Risiken aufbringen und die notwendigen Schritte vornehmen muss, um davor zu schützen. Diese Aufgabe ist jedoch wesentlich einfacher zu erfüllen, wenn der DC/DC-Wandler schon entsprechend zertifiziert ist. Fällt ein DC/DC-Wandler beispielsweise infolge eines internen Kurzschlussfehlers aus, darf er sich überhitzen, sollte aber nicht in Flammen aufgehen.

Die zum Aufbau des Wandlers eingesetzten Materialien dürfen daher nicht entflammbar und müssen selbstauslöschend sein. Versäumt es der Geräteentwickler jedoch, eine entsprechende Absicherung dieser Art des Schadenseintritts sicherzustellen (zum Beispiel durch Nichtbeschränkung des Eingangsstroms in den Wandler), könnte sich der DC/DC-Wandler immer noch so stark erhitzen, dass er eine Entzündung einer anderen Komponente oder eines anderen Materials hervorruft und Brandentwicklung auslösen kann. Der Entwickler ist daher auch dann für die Folgen eines Komponentenausfalls verantwortlich, wenn die Komponente selbst sicherheitszertifiziert ist.

Die Richtlinien der Sicherheitszertifizierungen tendieren dazu, diese Verantwortlichkeit des Geräteentwicklers hervorzuheben, indem sie Sicherheitstechnik auf Grundlage der Gefahrenanalyse (HBSE= Hazard-Based Safety Engineering) und Risikomanagement (RM) in den gesamten Prozess der Sicherheitszertifizierung mit einbeziehen.

Dies ist eine Hauptänderung in der Vorgehensweise im Vergleich zu konventionellen Sicherheitsstandards der Elektrotechnik wie 60950 oder ETS300, die sich einfach auf die Sicherheit des DC/DC-Wandlers konzentrieren, ohne ein Folgerisiko in der Endanwendung beim Ausfall des Wandlers zu berücksichtigen. Dies ist auch ein Grund, warum die meisten Hersteller von DC/DC-Wandlern darauf hinweisen, dass die Produkte generell nicht zum Einsatz in sicherheitskritischen Anwendungen geeignet sind.

Das HBSE-Verfahren besteht aus den folgenden vier Hauptschritten:

- 1) Identifizieren von Gefahrenquellen im Produkt (z. B. Energiequellen)
- 2) Einstufung der Risiken (z. B. Klasse 1: kein Verletzungsrisiko und kaum entzündbar; Klasse 2: geringes Verletzungsrisiko und entzündbar; oder Klasse 3: hohes Verletzungsrisiko und Entzündung wahrscheinlich)
- 3) Identifizieren von geeigneten Sicherheitsmaßnahmen (z. B. gefährliche Hochspannungen unzugänglich machen, Strombegrenzung)
- 4) Definieren von Sicherheitsmaßnahmen (z.B.sicherstellen,dass gefährliche Hochspannungen nur mit Hilfe von Werkzeugen zugänglich sind; maximale Ströme sind während des Normalbetriebs und unter Ausfallbedingungen abgesichert)

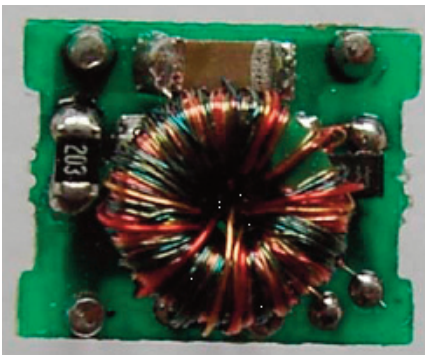
6.1 Stromschläge

Die meisten isolierten DC/DC-Wandler werden in Anwendungen mit einer netzbetriebenen AC/DC-Primärspannungsversorgung verwendet. Wenn diese primäre Stromquelle dadurch ausfällt, dass an ihrer Ausgangsklemme gefährliche Hochspannungen vorhanden sind, hat der DC/DC-Wandler die Funktion, den Benutzer vor Stromschlägen zu schützen. Mit anderen Worten, fällt die Primärwicklungsisolierung am AC/DC-Wandler aus, sollte die Sekundärwicklungsisolierung am DC/DC-Wandler den Benutzer vor Stromschlägen schützen. Dieses Konzept zweier unabhängiger Schutzarten ist die Grundlage vieler Sicherheitsvorschriften. Generell gilt, wenn der Schaltkreis unzugänglich ist (für den Zugriff sind Werkzeuge erforderlich), dann ist eine einzige Isolierungsschicht vertretbar, ansonsten sind mindestens zwei Schutzarten erforderlich.

6.1.1 Isolationsklasse

Drei Hauptisolationsklassen sind in den Sicherheitsvorschriften definiert.

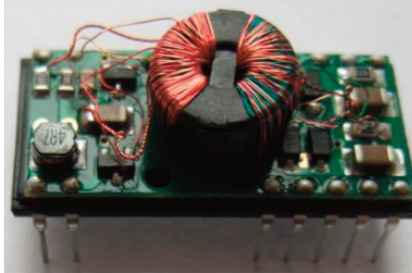
1: Funktionale Isolierung: Die Isolierung ist für die Funktion des Wandlers ausreichend, erfüllt die entsprechenden Anforderungen an Sicherheitsabstände und stellt während eines Ausfalls keine Brandgefahr dar; die Isolierung ist jedoch nicht ausreichend, um Schutz gegenüber Stromschlag zu gewährleisten.



Die meisten DC/DC-Wandler gehören dieser Klasse an, da sie durch nicht-gefährliche Hochspannungen versorgt werden. Ein funktional isolierter Wandler gewährleistet beschränkten Schutz vor Stromschlag im Falle eines Fehlers der primären Stromquelle. Dies ist jedoch kein zuverlässiger Schutz vor einer dauerhaft gefährlichen Eingangsspannung. Abb. 6.1 zeigt ein Beispiel eines funktional isolierten Wandlers. Die Eingangs- und Ausgangswicklungen sind übereinander gewickelt und nur durch die Lackisolierung des Kupferlackdrahtes isoliert. Trotz dieser einfachen Konstruktionsart können Isolierungsspannungen von bis zu 4 kVDC erreicht werden.

Abb. 6.1: Beispiel funktioneller Isolierung

2: Basis- oder zusätzliche Isolierung: erfüllt die Anforderungen funktioneller Isolierung, enthält aber zusätzliche Isolierung mit einer Stärke von mindestens 0,4mm zum Basisschutz vor Stromschlag und weist größere innere Sicherheitsabstände auf als die funktionale Isolierung.



Basisisolierte DC/DC-Wandler verfügen normalerweise über eine physikalische Isolationschicht, sodass die Isolierung nicht nur von der Lackisolation der Transformatorwicklungen abhängt. Das in Abb. 6.2 gezeigte Beispiel hat einen Ringkern in einem Kunststoffgehäuse, das mit einem Abstandshalter quer durch das Mittelloch ausgestattet ist, sodass die Ein- und Ausgangswicklungen physikalisch voneinander getrennt sind. Der Ferritkern wird als leitend angenommen, weshalb das Gehäuse auch jede Wicklung vom Kern isoliert und die Primär- und Sekundärwicklungen durch den Abstandshalter voneinander getrennt werden.

Abb. 6.2: Beispiel für Basisisolierung

3: Doppelte oder verstärkte Isolierung: Die Isolierung erfüllt die Anforderungen der Basisisolierung, enthält aber physikalische Mehrfach-Isolationsschichten, um verbesserten Schutz vor Stromschlag zu gewährleisten. Jede Schicht muss eine Stärke von mindestens 0,4mm aufweisen und mit größeren inneren Sicherheitsabständen versehen sein.

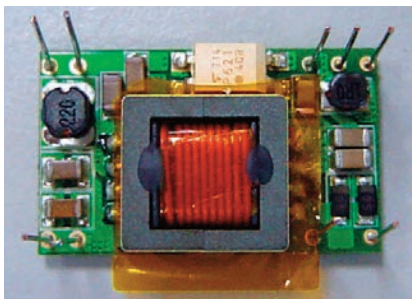


Abb. 6.3 zeigt ein Beispiel eines verstärkt isolierten Wandlers. Die Ausgangswicklungen verwenden dreifach isolierte Drähte, die Mylarfolie gewährleistet einen erhöhten Kriechweg zwischen dem Eingang und Ausgang und bietet so eine physikalische Isolationsschicht. Solche DC/DC-Wandler sind imstande, gefährlichen Langzeit-AC-Spannungen standzuhalten (Betriebsspannung = 250VAC) und gewährleisten bis zu 10kVDC Isolierung.

Abb. 6.3: Beispiel verstärkter Isolierung

6.1.2 Stromgrenzwerte für den menschlichen Körper

Die Definition einer durch Stromfluss verursachten Verletzung ist nicht einfach. Der elektrische Widerstand des menschlichen Körpers beträgt bei 110VDC ca. $2\text{k}\Omega$ und nimmt mit zunehmender Spannung ab. Dieser Wert ist jedoch von unterschiedlichen Faktoren abhängig. Der Hautwiderstand ist wesentlich höher, als der Widerstand von Körperorganen. Menschen mit besonders trockener Haut können einen Widerstand von bis zu $100\text{k}\Omega$ haben. Der Vollkontakt, der einem Kontakt einer typischen ganzen Handfläche von ca. 8cm^2 entspricht, hat einen niedrigeren elektrischen Widerstand als ein Teilkontakt z. B. durch eine Fingerspitze mit einer Fläche von nur $0,1\text{cm}^2$.

Wenn der elektrische Strom jedoch in einem kleinen Kontaktpunkt konzentriert wird, was eine lokale Verbrennung auslösen kann, oder die Haut besonders feucht ist, beträgt der Körper-Innenwiderstand weniger als 1kΩ. Schließlich sind AC-Ströme gefährlicher, da die Haut wie ein Isolator zwischen der Kontaktfläche an der Hautoberfläche und dem darunter liegenden Gewebe agiert. Daher ist der AC-Widerstand niedriger als der DC-Widerstand.

Um die Definition einer durch Stromfluss hervorgerufenen Verletzung zu klären, werden folgende Stromgrenzwerte definiert, um zu vermeiden dass ein Mensch einen Stromschlag erhält: 2mA bei DC oder einem Spitzenwert von 0,7mA AC oder 0,5mA_{RMS} bei 50Hz AC. Die Schwellenwerte des Stroms, der durch den menschlichen Körper fließt, sind in Tabelle 6.1 aufgeführt.

Wirkung des Stroms	Strom	Electrosicherheit (HBSE) Klasse
Minimale Reaktion	<0.5mA	ES1
Schreckreaktion, aber keine Verletzung	Up to 5mA	ES2
Muskelkontraktion, nicht in der Lage loszulassen	Up to 10mA	ES3
Herzdefibrillation, Organverletzung, Tod	>10mA	

Tabelle 6.1: Stromgrenzwerte durch den menschlichen Körper

Wenn die Ausgänge eines DC/DC-Wandlers aufgrund ihrer Konstruktion bis 60VDC oder 42,4VAC beschränkt sind, besitzen die Wandlerausgänge lediglich eine Schutzkleinspannung (SELV), und es sind keine weiteren Schutzmaßnahmen erforderlich, um zu verhindern, dass der Benutzer einen Stromschlag vom Ausgang erhält. Die Ausnahme dieser Regel bilden die Grenzen von Fernmeldenetzspannungen (TNV), die SELV-Spannungen überschreiten können, aber durch eine Dauer von max. 200ms beschränkt sind und deren Kontakte für den Benutzer unzugänglich sein müssen. TNV-Spannungen können Spitzenwerte von bis zu 120VDC oder 71VAC haben. Jegliche Ausgangsspannungen, die höher sind als die SELV oder TNV werden als gefährlich betrachtet und es müssen, um einen zufälligen Kontakt des Benutzers auszuschließen, entsprechende Sicherheitsvorkehrungen getroffen werden.

Spannungen in Stromkreisen unter SELV-Grenzen	TNV-1
Spannungen in Stromkreisen über den SELV-, aber unterhalb der TNV-Grenzen, keine Eingangsüberspannungen möglich	TNV-2
Spannungen in Stromkreisen über den SELV-, aber unterhalb der TNV-Grenzen, transiente Eingangsüberspannungen möglich (bis zu 1,5kVDC)	TNV-3

Tabelle 6.2: TNV-Definitionen

Netzspannungen werden als gefährlich eingestuft, weshalb die Isolationsfestigkeit jedes Wandlers mit einer AC/DC-Versorgung ausreichend sein muss, um den Benutzer in jedem Fall vor Stromschlag durch diese gefährliche Hochspannung während eines Ausfalls zu schützen. Die Spitzenspannung an einer Versorgung von 230VAC beträgt $230\sqrt{2} = 325V$. Daher ist jeder DC/DC-Wandler mit einem Isolationsnennwert von mindestens 500 VDC dazu geeignet. Vor zehn Jahren wurden viele Wandler tatsächlich so entwickelt, gebaut und geprüft, dass sie dieser Spitzenspannung widerstehen. Aufgrund von Standardisierungen in der Spezifikation beträgt der minimale Isolierungsnormwert heute jedoch 1000VDC, während medizinische Anwendungen mindestens 2000VDC fordern, und häufig mindestens 3000VDC gefordert werden.

Es gibt Anwendungen, bei denen am DC/DC-Wandler sehr hohe Spannungen vorhanden sein könnten, wie zum Beispiel in Röntgenapparaten, Laser-Stromversorgungen, Elektrovakuumgeräten mit Ionenpumpen und IGBT-Ansteuerungen. Andererseits kommen bei der überwiegenden Mehrzahl von DC/DC-Wandlern an der Isolierungsstrecke nie mehr als 48VDC vor.

Die folgenden Beschreibungen beziehen sich auf Industrie-, Telekommunikations- und Computer-Sicherheitsvorschriften. Standards für medizinische Sicherheit haben zusätzliche Anforderungen, mit denen wir uns am Ende dieses Kapitels gesondert befassen werden.

6.1.3 Schutz vor Stromschlägen

Die Sicherheitsvorschriften beinhalten drei primäre Arten von Schutzmaßnahmen gegenüber Stromschlag:

- Durchschlagsfestigkeit
- Luftstrecke (clearance separation)
- Kriechstrecke (creepage separation)

Der Sicherheitstest zur Durchschlagfestigkeit kann mit einer DC- oder AC-Prüfspannung (Spitzen-AC-Spannung ist gleich der Ruhe-DC-Spannung) durchgeführt werden. Die Isolierung muss dieser Spannung 60 Sekunden lang ohne Zerstörung widerstehen. Der Vorteil eines AC-Tests besteht darin, dass sowohl positive als auch negative Spannungsgradienten am Wandler anliegen. Der Nachteil besteht darin, dass der AC-Blindstrom irrtümlicherweise als Durchschlag gedeutet werden könnte, wenn ein EMV-Kondensator an der Isolierungsbarriere verbaut wurde. Im Zweifelsfall verwenden Sie deshalb lieber DC.

Isolationsklasse	Prüfspannung (DC)	Prüfspannung (AC)
funktionell	1000V/60 Sekunden	707VAC _{RMS}
Basisisolierung	1000V/60 Sekunden	707VAC _{RMS}
verstärkt	2000V/60 Sekunden	1414VAC _{RMS}

Tabelle 6.3: Sicherheitstest zur Durchschlagfestigkeit eines DC/DC-Wandlers (nicht medizinische Anwendungen)

Luftstrecke, oder engl. clearance, ist die kleinste Entfernung, die die Eingangsseite von der Ausgangsseite per „Luftlinie“ trennt. Manchmal wird er auch als Lichtbogenstrecke bezeichnet. Kriechweg ist die kleinste Entfernung, die die Eingangsseite von der Ausgangsseite entlang einer Oberfläche trennt. Er wird manchmal auch als Kriechstrecke bezeichnet. Abb. 6.4 stellt diese Entfernungen schematisch dar.

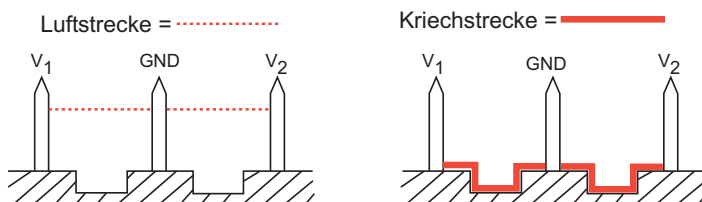


Abb. 6.4: Luftstrecke/Kriechstrecke

Die minimalen Luft- und Kriechstrecken werden in den Sicherheitsvorschriften entsprechend der Spannung am Wandler, der verwendeten Materialien und der Betriebsumgebung definiert.

Die Abstandstrecken werden durch die Normen „in Luft, unter 2000m“ definiert und hängen von der Eingangsspannung und der Isolationsklassifikation des Wandlers ab. Da ein vollgekapselter DC/DC-Wandler keine Luft enthält, sind die Abstände einfach die Strecken zwischen den Eingangs- und Ausgangspins des Wandlers. Bei einem Open-Frame-DC/DC-Wandler wird die innere Trennstrecke Transformatorwicklung nicht in die Abstandsberechnung mit einbezogen. Allerdings würde der Abstand von der Transformator-Primärwicklung bis zu den Sekundärpins oder von einer Primärwicklung bis zum angrenzenden Bauteil auf der Sekundärseite verwendet, wenn einer dieser beiden Abstände kürzer wäre, als der Sicherheitsabstand zwischen Eingang und Ausgang am PCB.

DC (AC) Spannung									
Isolationsklasse	12 (12)	36 (30)	75 (60)	150 (125)	300 (250)	450 (400)	600 (500)	800 (66)	VDC (VAC)
funktionell	0.4	0.5	0.7	1.0	1.6	2.4	3	4	mm
Basisisolierung	0.8	1	1.2	1.6	2.5	3.5	4.5	6	mm
verstärkt	1.6	2	2.4	3.2	5	7	9	13	mm

Tabelle 6.4: Minimale Luftstrecken für verschiedene Isolationsklassen

Die minimalen Kriechwegstrecken werden durch die Arbeitsspannung, die Oberflächenleitfähigkeit der verwendeten Materialien und den Verschmutzungsgrad definiert. Als Kriechstrecke auf der PCB wird der kürzeste Abstand zwischen den primär- und sekundärseitigen Leiterbahnen gemessen.

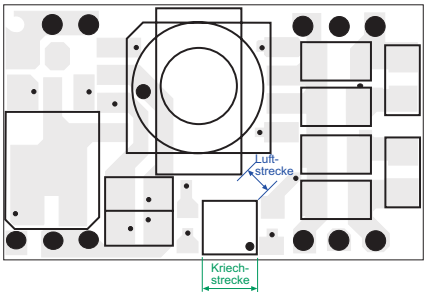


Abb. 6.5: PCB-Layout eines DC/DC-Wandlers mit minimalen Kriech- und Luftstrecken

Die Vergleichszahl der Kriechwegbildung (Comparative Tracking Index, CTI) fügt entsprechend der Oberflächenleitfähigkeit des Isolierstoffes, der den Eingang vom Ausgang trennt, zum Kriechweg einen Koeffizienten hinzu, in der Regel der PCB des Wandlers.

Isoliermaterial I:	$600 \leq \text{CTI}$
Isoliermaterial II:	$400 \leq \text{CTI} < 600$
Isoliermaterial IIIa:	$175 \leq \text{CTI} < 400$
Isoliermaterial IIIb:	$100 \leq \text{CTI} < 175$

Tabelle 6.5: Definitionen von Materialkategorien

Eine Standard-FR4-PCB-Platine verfügt über eine typische CTI von 200-250, mit Lötstopplack bis zu 400 (Klasse 111a); wenn die Platine jedoch PTFE-beschichtet ist, kann sich die CTI auf > 600 erhöhen (Klasse I). Der Verschmutzungsgrad (PD) fügt bei der Berechnung der Mindestkriechstrecken den Faktor der Oberflächenfeuchtigkeit oder von Verunreinigungen hinzu, da der Abstand erhöht werden muss, um eine Änderung der Vergleichszahl der Kriechwegbildung (CTI) in verschmutzten industriellen Umgebungen oder Außenbereichen auszugleichen.

Verschmutzungsgrad 1	Verschmutzungsgrad 2	Verschmutzungsgrad 3	Verschmutzungsgrad 4
Keine oder nur trockene nichtleitende Verschmutzung, was keine Einwirkung auf Leitfähigkeit hat.	Normalerweise tritt nur nichtleitende Verschmutzung auf. Vorübergehende Kondensation kann auftreten.	Leitende Verschmutzung sowie Kondensation treten häufig auf.	Leitende Verschmutzung und Kondensation treten dauerhaft auf.
gekapselte Komponenten	Büroumgebung	industrielle Umgebung	Außenbereiche

Tabelle 6.6: Verschmutzungsgrad

Tabelle 6.7 unten zeigt die Mindestkriechstrecken gemäß Betriebsspannung, Materialgruppe und Verschmutzungsgrad.

Mindestkriechstrecke							
Spitzen- spannung (V)	Verschmutzungsgrad						
	1	2			3		
	Alle Material- gruppen	Material- gruppe					
		I	II	III	I	II	III
		mm	mm	mm	mm	mm	mm
25	0.125	0.500	0.500	0.500	1.250	1.250	1.250
32	0.14	0.53	0.53	0.53	1.30	1.30	1.30
40	0.16	0.56	0.80	1.10	1.40	1.60	1.80
50	0.18	0.60	0.85	1.20	1.50	1.70	1.90
63	0.20	0.63	0.90	1.25	1.60	1.80	2.00
80	0.22	0.67	0.95	1.30	1.70	1.90	2.10
100	0.25	0.71	1.00	1.40	1.80	2.00	2.20
125	0.28	0.75	1.05	1.50	1.90	2.10	2.40
160	0.32	0.80	1.10	1.60	2.00	2.20	2.50
200	0.42	1.00	1.40	2.00	2.50	2.80	3.20
250	0.56	1.25	1.80	2.50	3.20	3.60	4.00
320	0.75	1.60	2.20	3.20	4.00	4.50	5.00
400	1.0	2.00	2.80	4.00	5.0	5.6	6.3
500	1.3	2.50	3.60	5.00	6.3	7.1	8.0
630	1.8	3.20	4.50	6.30	8.0	9.0	10.0
800	2.4	4.00	5.60	8.0	10.0	11.0	12.5
1000	3.2	5.00	7.10	10.0	12.5	14.0	16.0

Tabelle 6.7: Kriechstrecken

Wie sind diese Tabellen also in praktischen Fällen anzuwenden? Die oben genannte Spitzenspannung ist die höchste Spannung am Wandler unter normalen Einsatzbedingungen. Deshalb müsste ein Wandler mit einem Eingangsbereich von 2:1 und nominalen 48V Eingang und 24V Ausgang die Kriechstreckenanforderungen für eine maximale Eingangsspannung von 72V plus Ausgangsspannung von 24V = 96VDC erfüllen. Daher sollte die in Tabelle 6.4 oben angegebene nächst höhere Spannung von 100V verwendet werden (siehe hervorgehobene Zeile in Tabelle 6.7).

Ein voll gekapselter DC/DC-Wandler ist gegen Staub, Feuchtigkeit und Verschmutzungen hermetisch abgedichtet und wird deshalb unabhängig von den Einsatzbedingungen als PD1 eingestuft. Deshalb beträgt die erforderliche Mindestkriechstrecke, die die Ein- und Ausgangsleiterbahnen trennen soll, 0,25mm. Handelt es sich bei dem Wandler um eine Open-Frame-Konstruktion, dann erhöht sich die Kriechstrecke in einer Büroumgebung auf 1,4mm und in einer industriellen Umgebung auf 2,2mm. Ein Open-Frame-Wandler ist nicht für eine Anwendung im Außenbereich geeignet, weshalb keine Mindestkriechstrecken angegeben sind.

Gemäß Tabelle 6.4 würde der minimale Pin-Abstand bei diesem vakuumvergossenen Wandler 1 mm bei funktionaler Isolation, 1,6mm bei Basisisolation und 3,2mm bei verstärkter Isolation betragen. In der Praxis müssen auch Toleranzen für die minimale Trennung zwischen PCB-Pads, in die der Wandler eingelötet wird, gewährleistet sein. Deshalb benötigt ein gekapselter DC/DC-Wandler zum Einsatz in industriellen Umgebungen mit funktionaler Isolation vom Eingang zum Ausgang eine Mindestkriechstrecke und einen Mindestluftstrecke von 1mm, da die Kriechstrecke nicht kürzer als die Luftstrecke sein darf. Derselbe Wandler in einer Open-Frame-Konstruktion würde immer noch eine Luftstrecke von 1mm erfordern, die Kriechstrecke würde sich aber auf 2,2mm verlängern.

6.1.4 Schutzerdung

Zusätzlich zu einer elektrischen Isolation kann auch eine Schutzerdung (PE) als Schutz vor Stromschlägen verwendet werden. Somit erfüllt eine AC/DC-Stromversorgung mit Basisisolierung und einem mit PE verbundenen Ausgang die Sicherheitsanforderungen zweier Schutzverfahren. Wenn die Ausgangsspannung nicht geerdet ist, sind nur Stromversorgungen mit doppelter oder verstärkter Isolierung akzeptabel. Gefährliche Hochspannungen dürfen nicht freizugänglich sein, und freiliegende leitende Teile dürfen während des Normalbetriebs und durch einen Einzelfehler nicht gefährdend wirken können.

Es gelten folgende IEC-Schaltkreisklassifikationen:

Klasse-I-Geräte: Systeme, die Schutzerdung (z. B. ein geerdetes Metallgehäuse oder einen geerdeten Ausgang) und Stromabschaltung bei Ausfall (Schmelzsicherung oder Leistungsschalter) als Schutz verwenden und damit nur Basisisolierung erfordern. Keine freiliegenden gefährlichen Hochspannungen (geerdete Metallgehäuse oder nichtleitende Gehäuse). Klasse-I-Versorgungen müssen mit dem Erde-Symbol gekennzeichnet werden:



Klasse-II-Geräte: Verwendung doppelter oder verstärkter Isolierung, um die Notwendigkeit eines geerdeten Metallgehäuses auszuschließen, keine freiliegenden gefährlichen Hochspannungen (nichtleitendes Gehäuse). Keine PE-Verbindung erforderlich, eine Filtermasseverbindung ist aber zulässig (Funktionserdung statt Schutzerdung). Klasse-II-Versorgungen müssen mit dem Symbol für doppelte Isolierung gekennzeichnet werden:



Anmerkung: Wenn eine AC/DC-Stromversorgung über eine Filtermasse (FG)-Verbindung verfügt, um die EMV-Richtlinien zu erfüllen, kann sie dennoch als Stromversorgung der Klasse II eingestuft werden, wenn sie keinen Erdanschluss benötigt, um Schutz vor Stromschlag zu gewährleisten.

Klasse-III-Geräte: Gespeist von einer SELV-Quelle und ohne elektrisches Potential zur Erzeugung gefährlicher interner Hochspannungen, d. h. nur funktionale Isolation erforderlich. Funktionserdung ist möglich, eine Verbindung zu PE ist jedoch nicht zulässig (kein Rückstrompfad zur Masse durch das Netzteil). Klasse-III-Versorgungen müssen mit dem Klasse III Symbol gekennzeichnet werden:



Es ist etwas irreführend, dass die NEC-Klassifizierung zur Beschreibung verschiedener Schutzgrade auch ein ähnliches „Klassen“-System verwendet, aber zum Beschreiben des Schutzgrades gegenüber möglicherweise kritischer Energielevels (Brandgefahr) arabische Ziffern benutzt.

Die NEC-Schaltkreisklassifizierungen sind wie folgt:

- Schaltkreise der Klasse 1: Leistung auf $< 1\text{kVA}$ beschränkt und Ausgangsspannung $< 30\text{VAC}$
- Schaltkreise der Klasse 2: Leistung auf $< 100\text{VA}$ beschränkt, Eingangsspannung $< 600\text{VAC}$ und Ausgangsspannung $< 42,5\text{VAC}$
- Schaltkreise der Klasse 3: Leistung auf $< 100\text{VA}$ beschränkt, Eingangsspannung $< 600\text{VAC}$ und Ausgangsspannung $< 100\text{VAC}$; zusätzlicher Schutz vor Stromschlägen notwendig.

**Praktischer
Hinweis**

Wenn man über eine Stromversorgung der „Klasse zwei“ spricht, ist es also wichtig, zu wissen, ob es sich um die NEC- oder IEC-Definition handelt.

Für einen DC/DC-Wandler sind die oben erwähnten Klassifizierungen nur indirekt relevant. Fast alle DC/DC-Wandler für niedrige Eingangsspannungen können als Stromversorgungen der Klasse III eingestuft werden. Ausnahmen sind DC/DC-Wandler mit hohen Eingangs- und Ausgangsspannungen, z. B. solche die in Bahnanwendungen, in Photovoltaikanlagen oder in Anwendungen der Elektromobilität eingesetzt werden, da diese zusätzliche Schutzmaßnahmen erfordern, um das Risiko von Stromschlägen zu reduzieren. Da die Ausgänge von AC/DC-Stromversorgungen der Klasse I oder Klasse II isoliert sind, können sie massebezogen sein. Dasselbe gilt für einen isolierten oder nicht isolierten DC/DC-Wandler. Abb. 6.6 zeigt eine Erdungsschaltung, wie sie üblicherweise in Fernmeldeleitungen verwendet wird, in der alle DC-Spannungen massebezogen sind (Klasse-I-Eingang, Klasse-2-Ausgang).

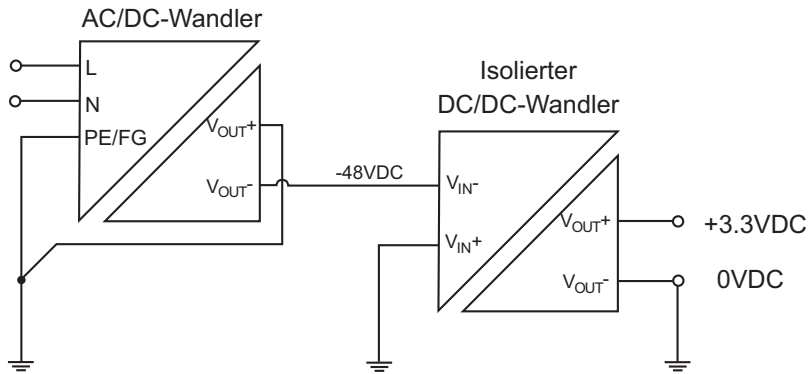


Abb. 6.6: Erdungsschaltung für Telekommunikationsnetzversorgung

6.2 Gefährliche Energielevels

Die Definition der IEC/UL 60950 von gefährlicher elektrischer Energie lautet „verfügbarer Leistungspegel von mindestens 240VA mit einer Dauer von mindestens 60 Sek. oder ein gespeicherter Energiepegel von mindestens 20J (z. B. von einem oder mehreren Kondensatoren) bei einem elektrischen Potential von mindestens 2V“.

Diese Pegel wurden ausgewählt, weil die durch eine derartige Energiefreisetzung (z. B. durch einen Kurzschluss oder Berühren spannungsführender Teile) freigesetzte Energie ausreichen könnte, um Verletzungen zu verursachen oder einen Brand auszulösen.

Im weiteren Verlauf dieses Kapitels wird das Risikomanagement erörtert. Die wichtigsten Verfahren zur Reduktion der Verletzungswahrscheinlichkeit oder einer Entzündung infolge dieser gefährlichen Energie sind die folgenden:

- physikalischer Schutz (z. B. Gehäuse, Abschirmverbinder, Einkapselung)
- Energieabbau (z. B. Entladekreise von Kondensatoren)
- Funkenlöschung (z. B. Snubber-Netzwerke)
- Eigensicherung (z. B. Energie ist konstruktionsbedingt eingeschränkt)
- Überstrombegrenzung (z. B. Schmelzsicherung)

Diese Schutzverfahren können im Hinblick auf zusätzliche Sicherheit kombiniert werden. Sandgefüllte Sicherungen werden beispielsweise eingesetzt, um das Netzteil während des Ausfallzustands abzutrennen und abzuschalten und zusätzlich über ein physikalisches Mittel (den Sand) dafür zu sorgen, dass der heiße Schmelzdrahtes nicht zu einer Zündquelle wird.

6.2.1 Schmelzsicherungen

Um die verfügbare Energie in einer Kurzschluss- oder Überstrom (OC)-Situation zu begrenzen, werden gefährliche Energien, durch Einsetzen einer Schmelzsicherung oder eines Leitungs- oder Lasttrenners in die Versorgung, unterbrochen. Der Unterbrechungs-Nennwert einer OC-Schutzvorrichtung wird sowohl als Strom wie auch als Spannung angegeben, da es eine Ansprechzeit gibt, bevor der Ausfall beseitigt werden kann, die von beiden abhängig ist.

Für eine Schmelzsicherung ist die Ansprechzeit wie folgt: $t_{\text{clear}} = t_{\text{melt}} + t_{\text{quench}}$, wobei t_{melt} die Schmelzzeit ist, die vom Schmelzintegral I_{2t} , der Umgebungstemperatur, der Vorbelastung und der Sicherungskonstruktion abhängt und wobei t_{quench} die Bogenbrennzeit ist, die von der Spannung an der Schmelzsicherung und der Sicherungskonstruktion abhängt. Während der Ansprechzeit t_{clear} fließt der Strom immer noch durch die Schmelzsicherung. Dieser Stromfluss, während dem Ansprechen der Schmelzsicherung, wird als Durchgangsleistung bezeichnet.

Es ist wichtig, dass die Nennspannung einer Schmelzsicherung nicht überschritten wird, da die Lichtbogenenergie anderenfalls nicht sicher innerhalb des Sicherungshalters verbleibt und der sich ergebende schnelle Temperaturanstieg die Schmelzsicherung zum Explodieren bringen und so diese selber zu einer Zündquelle machen könnte. Der Abschaltstrom (der maximale Strom, den die Schmelzsicherung sicher unterbrechen kann, nicht zu verwechseln mit dem Nennstrom) muss höher sein als der maximal verfügbare Strom aus der Stromversorgung (manchmal wird er als zu erwartender Kurzschlussstrom bezeichnet).

Die Auswahl des am besten geeigneten Durchlassstromes hängt sehr von der Anwendung und Umgebung, in der sie sich befindet, ab. Der Mechanismus, der den Strom unterbricht, ist ein Schmelzdraht, der bei ausreichender Energie schmilzt. Ist der Draht schon heiß, weil die Umgebungstemperatur hoch, die Kühlung schlecht oder der durch die Schmelzsicherung fließende Ruhestrom zu hoch sind, ist die Schmelzzeit kürzer, als wenn der Schmelzdraht vollständig kalt wäre.

Der Nennstrom von Schmelzsicherungen wird in der Regel bei 25°C Umgebungstemperatur spezifiziert.

Je nach Schmelzsicherungstyp muss der Nennstrom um eine Größe zwischen 5% und 40% für Einsatz bei Temperaturen von 85°C (Abb. 6.7) herabgesetzt werden.

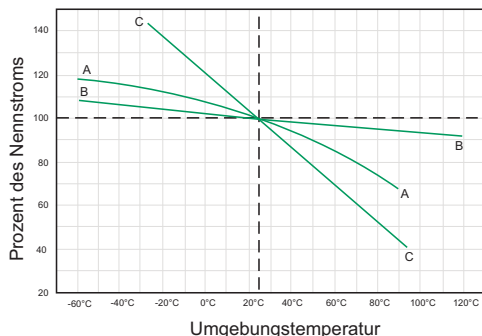


Abb. 6.7: Temperatur-Derating-Kurve für verschiedene Schmelzsicherungstypen; bei A handelt es sich um eine Patronensicherung, bei B um einen Schmelzdraht, bei C um eine PPTC-rückstellbare Schmelzsicherung

Praktischer Hinweis

Der Nennstrom hängt auch von der Meereshöhe ab. In großen Höhen ist die Luft dünner und die Wärme, die infolge des Ruhestroms innerhalb der Schmelzsicherung erzeugt wird, wird weniger gut durch Konvektion abgegeben. Der Sicherungsnennwert muss für Höhen über 2000m pro 100m um zusätzliche 0,5% für verringert werden. Eine auf dem Meeresspiegel auf 1,5A ausgelegte Schmelzsicherung sollte beispielsweise bei 4000m wie eine 1,35A Sicherung angesehen werden.

Außerdem fügt bei hochfrequenten Strömen die Induktivität des Schmelzdrahtes einen zusätzlichen reaktiven Widerstand hinzu, was auch den Wärmeverlust im Schmelzdraht erhöht. Der Durchlassstrom sollte ebenso herabgesetzt werden, wenn die Frequenz des Restwelligkeitsstroms 1kHz überschreitet:

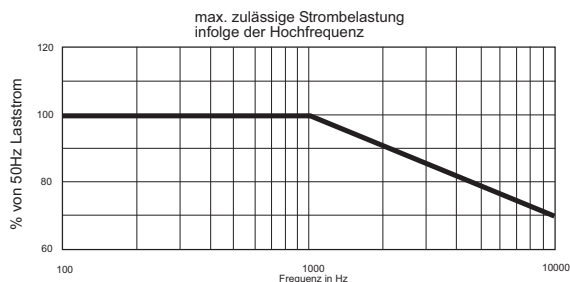


Abb. 6.8: Typische Frequenz-Derating-Kurve für eine Schmelzsicherung

Schließlich gibt es Alterungseffekte, die vorzeitiges Durchbrennen oder frühzeitige Verschmörung der Schmelzsicherung verursachen können. Die meisten dieser Alterungseffekte ergeben sich aus der Wärmeausdehnung und Kontraktion des Schmelzdrahtes durch wiederholtes Ein- und Ausschalten der Anwendung. Deshalb zeigt ein Gerät, das immer eingeschaltet ist, weniger Alterungseffekte der Schmelzsicherungen als eine Anwendung, die regelmäßig ein- und ausgeschaltet wird. Bei einer Anwendung, die häufig ein- und ausgeschaltet wird, verursachen die zyklischen thermischen Beanspruchungen Kaltverfestigungen und Mikrorisse im Schmelzdrahtmaterial, die zu internen Kontaktfehlern führen und zu einer schlechten mechanischen Verbindung zwischen der Schmelzsicherung und dem Sicherungsträger beitragen, was Kontaktprobleme verursachen kann.

Praktischer Hinweis

Angeichts all dieser Derating-Faktoren erscheint es vorteilhaft, einen Sicherungsnennwert zu wählen, der wesentlich höher ist als der Ruhestrom, um unerwünschtes Durchbrennen zu vermeiden. Es ist jedoch von erheblicher Wichtigkeit, dass die Schmelzsicherung auf einen Ausfallzustand reagiert, ohne zuerst eine gefährliche Energiemenge durchzulassen. In der Praxis ist ein Sicherungsnennwert, der im ungünstigsten Fall auf Ströme zwischen dem 1,3- und 1,5-Fachen des Ruhestroms ausgelegt ist und über eine träge Sicherungskennlinie verfügt, um kurzzeitige Stromspitzen und Reaktionen auf schnelle Lastwechsel standzuhalten, ein guter Kompromiss.

6.2.1.1 Ansprechzeiten von Schmelzsicherungen und Einschaltspitzenströme

Die prinzipielle Konstruktion einer Schmelzsicherung bestimmt deren Ansprechzeit. Es kann z.B. erwünscht sein, dass ein Gerät einem sehr hohen Überstromzustand standhält, um unerwünschtes Durchbrennen infolge von Einschaltspitzenströmen, Stromspitzen aufgrund von Lastwechseln oder Impulsbelastung zu vermeiden. Das ist bei DC/DC-Wandlern besonders wichtig, weil Wandler infolge der Eingangsfilterkondensatoren, die sich zur gleichen Zeit aufladen, in der sich auch das Magnetfeld im Transformator bildet, einen sehr hohen Anlaufstrom ziehen. Daher kann auch der Einschaltspitzenstrom eines DC/DC-Wandlers einige Ampere betragen (siehe Abb. 4.16).

In diesem speziellen Beispiel hat ein DC/DC-Wandler mit 2W, der bei der Eingangs-Nennspannung und Volllast normalerweise 80mA zieht, einen Einschaltspitzenstrom von 7,9A. Daher ist eine Sicherungskennlinie erforderlich, die den Wandler im ungünstigsten Fall mit einem Normaleingangsstrom von 100mA sicher schützen kann, aber eine achtzig Mal so hohe Stromspitze zulassen muss, wenn auch nur für die Dauer von 8µs. Obwohl ca. 8A als enorm hoher Strom erscheinen, bedeutet die nur kurze Dauer dieses Stromflusses, dass das Schmelzintegral ($I^2 t$) nur 0,000512 Ampere² Sekunden beträgt. Diese Energie reicht nicht aus, um den Schmelzdraht einer für 100mA ausgelegten Schmelzsicherung zum Schmelzen zu bringen. Zusätzlich führt jegliche Streuinduktivität des Schmelzdrahtes, des Sicherungsträgers und der Leiterbahnen zu einer wesentlichen Verringerung des Spitzenstroms durch die Schmelzsicherung. Um einen langfristig zuverlässigen Betrieb zu gewährleisten, würden wir für dieses Anwendungsbeispiel eine auf 150mA ausgelegte träge Schmelzsicherung empfehlen.

Einige träge Schmelzsicherungen verwenden spiralförmig gewickelten Schmelzdraht, um die Selbstinduktivität zu erhöhen und somit die Widerstandsfähigkeit gegenüber Stromspitzen zu steigern, im stationären Zustand aber ein Ansprechen beim festgelegten Unterbrechungszustand sicherzustellen. Ein weiteres Konstruktionsverfahren besteht darin, ein Metalltröpfchen auf den Schmelzdraht zu setzen, damit er wie ein Kühlkörper zu einer Verzögerung der Sicherungsansprechzeit fungiert, (Abb. 6.9).

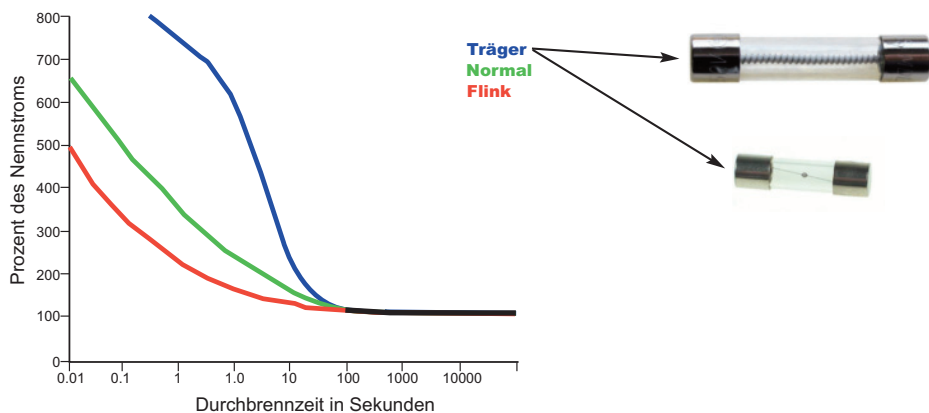


Abb. 6.9: Strom-Zeit-Verhältnis von flinken, Standard und trägen Sicherungen

Generell wird eine träge Eingangssicherung mit einem Nennstrom von ca. 150% des maximalen Eingangsstroms (gemessen bei minimaler V_{IN} und Volllast) empfohlen. Es ist nicht üblich, Ausgangsströmsicherungen an DC/DC-Wandler anzubringen, da Ausgänge meistens über eine Strombegrenzung und einen Kurzschlussschutz verfügen. Es gibt jedoch einige Anwendungen, bei denen eine Ausgangsströmsicherung aus anderen Gründen erforderlich sein kann. Ein solcher Grund besteht darin, dass ein einzelner DC/DC-Ausgang viele Schaltkreise speist, wobei jeder einzelne nicht in der Lage ist, genügend Strom zu ziehen, um den Wandlerausgang bei einem Ausfall zum sicheren Abschalten zu veranlassen. Einzelne Schmelzsicherungen können den Strom zu jedem Schaltkreis selektiv begrenzen. Ein anderer Grund könnte darin bestehen, dass die zulässige Grenze der gefährlichen Energie niedriger ist, als die Energiemenge die der Wandler im Fehlerfall, insbesondere infolge der in den Ausgangsfilterkondensatoren im Wandlerinneren gespeicherten Energie, liefern könnte.

Für Anwendungen, die eine Ausgangsströmsicherung erfordern, ist die Auswahl größer und hängt stärker von den Lasttypen (Last eher nur ohmsch, eher induktiv oder eher kapazitiv), ob sie nun statisch oder dynamisch sind, bzw. von der Definition gefährlicher Energie in der jeweiligen Anwendung ab. Während des Einschaltens der Stromversorgung steigt die Ausgangsspannung des DC/DC-Wandlers relativ langsam linear an, da jeder Schaltzyklus nur ein gewisses Energiepaket vom Eingang zum Ausgang überträgt. Bei näherer Betrachtung können die einzelnen Schritte des Ausgangsspannungsanstiegs manchmal am Oszilloskop beobachtet werden. Eine Rampenspannung generiert einen viel niedrigeren „Ausschaltstoß“-Strom als ein typischer kurzer und hoher Einschaltspitzenstrom. Daher sollten die Ausgangsströmsicherungen flink sein, um so einen sicheren Schutz der Anwendung vor gefährlicher Energie zu bieten.

6.2.2 rückstellbare Überstromschutzschalter (circuit breaker)

Um den Strom zu unterbrechen, benutzen thermisch-magnetische Überstromschutzschalter (engl.: miniature circuit breakers - MCB) zwei separate Auslösemechanismen. Der Magnetauslöser reagiert schnell auf einen starken Kurzschlussstrom (typisch innerhalb $\approx 5\text{ms}$) und der Wärmeauslöser auf eine mehrere Sekunden lang andauernde Überlast.

Das Gleichgewicht zwischen den zwei Auslösemechanismen ist durch den Aufbau des Schutzschalters festgelegt; z. B. ist es möglich, einen Miniaturschutzschalter mit einer langen Ansprechzeit für Überlast, aber immer noch schnell reagierend auf Kurzschlussfehler, zu finden.

Es wird generell angenommen, dass Schmelzsicherungen schneller als Überstromschutzschalter reagieren, aber dies ist nicht immer zwingend der Fall. Insbesondere können für DC/DC-Wandler mit hohen Leistungen mit Eingangs- oder Ausgangsströmen von einigen Ampere die auslösenden Einstellwerte von Überstromschutzschalter viel näher am Ruhestrom dimensioniert werden, als dies bei einer Schmelzsicherung der Fall wäre. Fehlauslösungen wären dadurch deutlich unwahrscheinlicher als beim Einsatz von Schmelzsicherungen. Ein Überstromschutzschalter reagiert auch auf einen länger andauernden Überlastbetrieb mit nur geringfügig über dem Nennwert liegendem Überlastwert, wenn dieses Energie-Zeit-Integral als gefährliche Energie eingestuft werden kann.

Andere Vorteile der Überstromschutzschalter bestehen darin, dass sie rückstellbar sind, eine sichtbare Anzeige besitzen (die anzeigt ob sie ausgelöst wurden), über einen sehr hohen Abschaltstromnennwert verfügen, eine interne physikalische Lichtbogenunterdrückung enthalten (z. B. Lichtbogenleitbahnen und Lichtbogenlöschkammer) und Mehrfachverbindungen (Pole) gleichzeitig ausschalten können. Die Hauptnachteile bestehen darin, dass sie wesentlich teurer und größer sind als Schmelzsicherungen und niedrigere Nennströme bieten als diese.

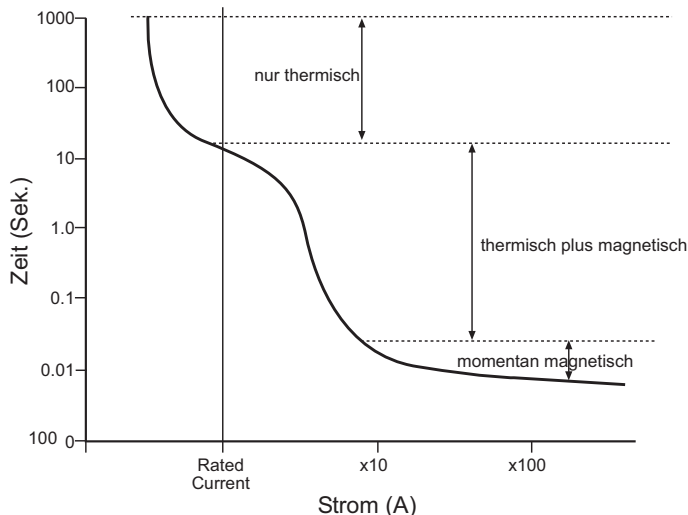


Abb. 6.10: Typisches Strom-Zeit-Verhältnis eines Überstromschutzschalters (MCB)

6.3 Inhärente Sicherheit

Im Unterschied zur intrinsichen oder Eigensicherheit (siehe folgenden Abschnitt) versucht die inhärente Sicherheit, Gefahren durch Verringern des Energieniveaus im Schaltkreis auszuschließen, sodass Störungen kein Risiko darstellen. Das Prinzip entstand in der chemischen Industrie, als verstanden wurde, dass man die mit der Herstellung chemischer Stoffe verbundenen Gefahren wesentlich verringern kann, indem man die Stoffe statt in sehr großen Mengen nur in kleineren Mengen erzeugt.

Die vier Prinzipien der inhärenten Sicherheit sind Minimieren, Ersetzen, Verringern und Vereinfachen. Bei der Anwendung auf Stromversorgungen besteht das Konzept darin, Stromversorgungen mit minimaler Menge an gespeicherter Energie zu entwickeln; eine einzelne zentrale Hochleistungs-Stromversorgung durch mehrere kleinere zu ersetzen; die Ströme, die sowohl intern als auch extern fließen können, zu verringern; und die Anwendung der Stromversorgung zu vereinfachen (z. B. durch Verwendung gepolter Stecker, sodass eine Querverbindung unmöglich wird, sowie durch Einbau von Indikatoranzeigen wie z.B. „Power ok“-Leuchten). Um die Störungsanfälligkeit zu verringern, müssen alle Netzversorgungen zusätzlich getrennt und in entsprechend sicherer Umgebungen betrieben werden, damit sie nicht durch die Umgebungsbedingungen zusätzlich belastet werden.

Das Beispiel in Abb. 6.11 zeigt einen Teil eines inhärent sicheren Stromversorgungskonzeptes. Der Primär-DC/DC-Wandler liefert eine leistungsbegrenzte Versorgung aus einem 24-V-Industriestandardbus. In diesem Beispiel wird eine Versorgungsspannung von 24V verwendet, um einen auf 5W begrenzten isolierten 24-V-Ausgang zu gewährleisten. Ein linearer Serienspannungsregler wird als zusätzlicher Strombegrenzer verwendet (Anmerkung: der Masseanschluss GND ist potentialfrei, weshalb der Regler als eine Stromquelle fungiert). Beim Einsatz eines typischen 5-V-Reglers mit einem Widerstand von 25Ω beträgt der maximale Strom, der fließen kann, $5/25 = 200\text{mA}$. Bei 24V Versorgungsspannung begrenzt dies die maximale Ausgangsleistung auf 4,8W. Der zweite DC/DC-Wandler setzt die 24V auf die Spannung auf Leiterplattebene auf 3,3V herab.

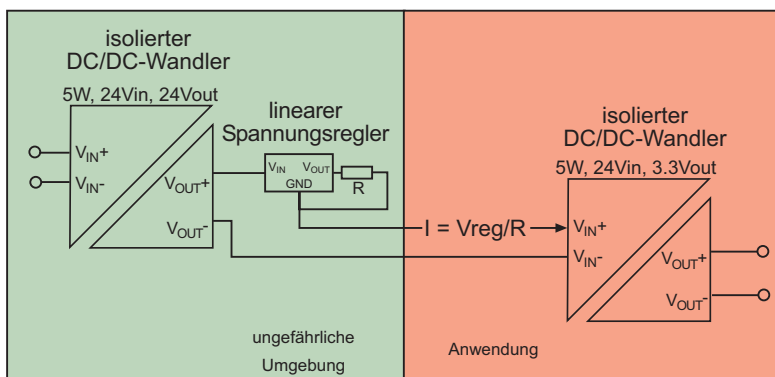


Abb. 6.11: Inhärent sichere Stromversorgung

In der Praxis gilt es bei diesen strombegrenzten Versorgungsketten Vorsicht walten zu lassen. Bei dem zweiten DC/DC-Wandler wird während des Einschaltens der Eingangsstrom begrenzt, und er wird daher höchstwahrscheinlich Anlaufschwierigkeiten zeigen. Es ist nicht möglich, als Hilfe eine zusätzliche Eingangskapazität hinzuzufügen, da dies dem Prinzip der inhärenten Sicherheit widersprechen würde.

Daher ist es hilfreich, die Last erst zuzuschalten, nachdem sich die Ausgangsspannung stabilisiert hat (Abb. 6.11). Eine andere Option besteht darin, die internen Komponenten des DC/DC-Wandlers zu ändern, z. B. durch Reduktion der Eingangsfilterkapazität, um einen Wandler mit geringerer gespeicherter Energie zu erhalten.

Das Konzept der inhärenten Sicherheit steht oft im Widerspruch zu industriellen Anforderungen und der Forderung nach preiswerten Hochleistungsnetzgeräten, die starke Anlassströme liefern und schnell auf Belastungstransienten reagieren können.

6.4 Eigensicherheit

Eigensichere Schaltkreise werden so entwickelt, dass in der Stromversorgung nicht ausreichend Energie vorhanden ist, um eine lokale Überhitzung oder Funkenbildung, die ein brennbares Gas entzünden würde, überhaupt erst entstehen zu lassen. Es gibt zwei Hauptschutzgrade: einzelner Fehler und doppelter Fehler. Darüber hinaus muss die Umgebung oder „Zone“, in der die Stromversorgung eingesetzt wird, berücksichtigt werden.

Zone	Beschreibung
0	explosionsfähige Atmosphäre dauerhaft vorhanden
1	explosionsfähige Atmosphäre kann im Normalbetrieb vorkommen (<10 Std/Jahr)
2	explosionsfähige Atmosphäre unwahrscheinlich (< 10 Stunden/Jahr)

Tabelle 6.8: Eigensicherheitszonen

Stromversorgungen der Zone 0 müssen vor doppelten Fehlern (Klasse ia), Stromversorgungen der Zone 1 vor einzelnen Fehlern (Klasse ib) geschützt werden. Stromversorgungen der Zone 2 können ohne Fehlerabsicherung auskommen, wenn sie geschlossen oder eingekapselt sind, um eine mögliche explosionsfähige Atmosphäre auszuschließen; in der Praxis sind sie aber häufig auch vor einzelnen Fehlern geschützt. Da Flugstaub brennen kann (Staubexplosion), schließt die Definition der explosionsfähigen Atmosphäre sowohl Staub als auch Gase ein.

Zone 0 (Gas) = Zone 20 (Staub), Zone 1 (Gas) = Zone 21 (Staub) und Zone 2 (Gas) = Zone 22 (Staub). Der Hauptunterschied besteht im Gehäusetyp (luftdicht gegenüber staubdicht).

Für die Stromversorgung fordert die Eigensicherheitskompatibilität eine FMEA (Failure Mode Effect Analysis = Fehlermöglichkeits- und Fehlereinflussanalyse, oder Auswirkungsanalyse) für alle verwendeten Komponenten sowie eine Analyse der Oberflächentemperaturen im ungünstigsten Fall und der Zündquellenenergie innerhalb der Versorgung, wie Spitzenstrom, Spitzenspannung und die MOSDE (Maximum Output Short-circuit Discharged Energy = maximale Ausgangskurzschluss-Entladungsenergie).

DC/DC-Wandler sind ein wichtiges Element in eigensicheren Schaltkreisen, da sie zur Gewährleistung galvanischer Trennung verwendet werden, die Menge an verfügbarer Energie begrenzen und minimale Trennungen, Luft- und Kriechstrecke gewährleisten können.

Ein Beispiel einer eigensicheren Stromversorgung ist unten gezeigt. Zwei DC/DC-Wandler werden sowohl für die Sicherheitsisolierung als auch zur Reduzierung der DC-Versorgungsspannung verwendet, um die Leerlaufzündspannung der explosions-fähigen Atmosphäre zu senken.

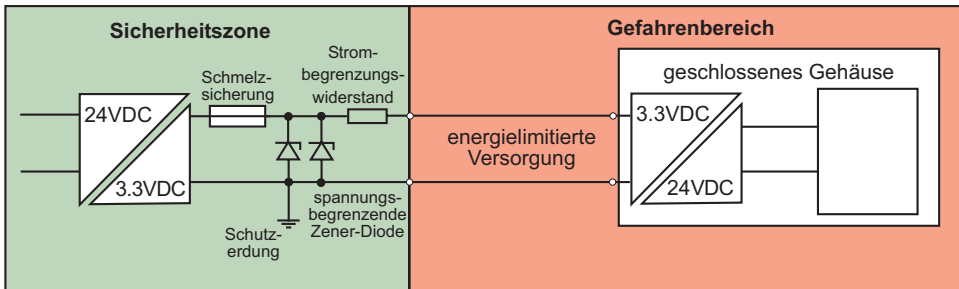


Abb. 6.12: Eigensichere Stromversorgung

Typische Einsatzbereiche für eigensichere Stromversorgungen sind Minen, petrochemische Werke, Mahlbetriebe und Verarbeitungsanlagen für Staubfutter. DC/DC-Wandler mit doppelter oder verstärkter Isolierung werden infolge der erforderlichen Doppelfehlertoleranz normalerweise verwendet, um (ia)-Schutzgrade zu erfüllen. Wenn DC/DC-Wandler in luftdichten Gehäusen eingesetzt werden, muss darauf geachtet werden, dass gewährleistet ist, dass sie sich infolge des beschränkten Luftstroms nicht überhitzen. Ein Verfahren besteht darin, sie unter Verwendung von wärmeleitfähigen Lückenfüllungen (Wärme-Pads) dicht an den Gehäusewänden oder internen Kühlkörpern zu montieren.

6.4.1 Brennbare Materialien

Als Energiequellen könnten Stromversorgungen auch zu brennen beginnen oder so heiß werden, dass sich angrenzende brennbare Materialien unter Fehler- oder nicht bestimmungsgemäßen Einsatzbedingungen entzünden können. Die meist allgemeine Standardanforderung für brennbare Stromversorgungsbauteile (PCB, Gehäuse, Vergussmasse usw.) ist, dass sie UL94-V0-konform sein müssen. Dieser Norm liegt eine Reihe von standardisierten Tests zugrunde, mit denen festgestellt wird, ob die in einer Stromversorgung verwendeten leicht entzündlichen Materialien nach Entzündung durch eine Flamme selbstlöschend sind und ob sich die Flamme ausbreiten kann. Zusätzlich kontrolliert der Standard auch, ob die verwendeten Materialien durch Funken oder elektrische Lichtbögen entzündet werden können. Dies ist wichtig, da innerhalb der Stromversorgungen, zur elektrischen Isolierung, häufig Kunststoffe verwendet werden.

Das UL94-Prüfprotokoll kategorisiert brennbare Materialien gemäß ihrer Tendenz zur Ausbreitung einer Flamme oder zum Glimmen in der Horizontalebene (HB) und zur senkrechten Ausbreitung (V-0, V-1, V-2 oder VTM-0, VTM-1 und VTM-2 für Dünnschichtmaterialien). Die Gefahr der Fremdzündung wird überprüft, indem ein Stückchen Baumwolle unter das Prüfobjekt gelegt wird, das durch geschmolzene Kunststoffropfen nicht entzündet werden darf.

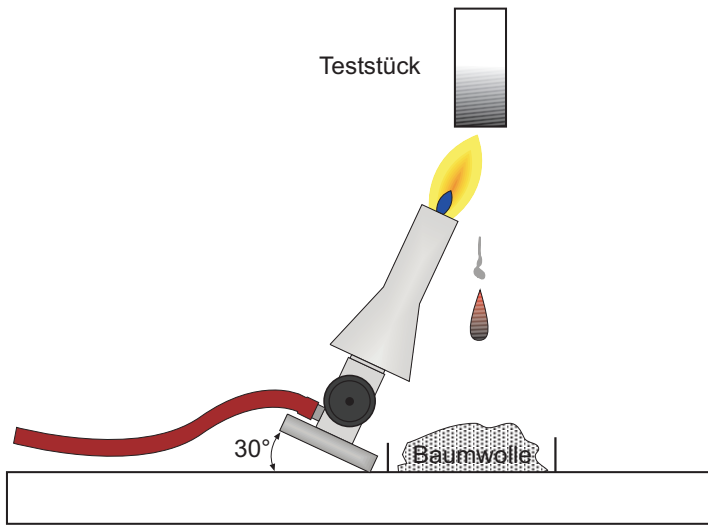


Abb. 6.13: Brennbarkeitsprüfung UL 94 durch Tropfenverbrennung

Andere brennbare Materialnormen stufen Stromversorgungen laut der Masse an verwendetem brennbarem Material, der Position der Stromversorgung (z. B. an einem bemannten oder unbemannten Ort, innerhalb eines Möbelstücks oder in einem Schlafabteil eines Zuges) und laut der physikalischen Trennung zwischen der Stromversorgung und anderen potentiellen Zündquellen oder brennbaren Materialien ein.

Um alle auf Brandgefahren anwendbaren Normen zu erfüllen, müssen Hersteller einer Stromversorgung daher die Anwendung des Endverbrauchers bis ins kleinste Detail kennen – einschließlich der Orientierung der Stromversorgung, wenn sie in der Anwendung installiert ist, des senkrechten und horizontalen Abstands, der Anforderungen an die Feuerwiderstandsfähigkeit und des beabsichtigten Gebrauchs. In vielen Situationen, in denen Stromversorgungen durch Drittanbieter oder Katalogvertriebsfirmen verkauft werden, ist dies nicht durchführbar. Daher werden in den Datenblättern die grundlegenden Standards für Brandgefahren der UL94-V0 verwendet, und weitere Konformitätsprüfungen oder -kontrollen können nur in enger Zusammenarbeit mit dem Kunden oder durch den Endverbraucher selbst durchgeführt werden.

6.4.2 Rauch

Mehr Brandtodesfälle werden durch die Inhalation von Rauch, als infolge von Brandwunden, verursacht. Rauch ist definiert als alle in der Luft enthaltenen Verbrennungsprodukte einschließlich Teilchen wie Ruß, Gase wie Kohlenmonoxid, flüchtige Stoffe wie organische Moleküle und Sprays wie Dampf. Die sichtbaren Teile von Rauch, die hauptsächlich durch Ruß und Dampf verursacht werden, sind oftmals selbst nicht toxisch, sofern sie keine Säure, die durch das Feuer freigesetzt werden kann, enthalten; durch die Verschlechterung der Sicht und Behinderung von Fluchtwegen tragen sie jedoch zur Brandgefahr bei. Gase wie Kohlenmonoxid verursachen Hyperventilation, die die Lunge dem Rauch noch stärker aussetzt, und verbinden sich mit Hämoglobin im Blut, was die reguläre Sauerstoffaufnahme im Körper blockiert.

Stromversorgungen enthalten in der Regel keine großen Mengen an brennbarem Material, abgesehen von Vergussmaterialien, die bei einem Brand eine nicht unwesentliche Rauchentwicklung auslösen würden. Vergussmasse wird eingesetzt zur Isolierung der Stromversorgung gegenüber Kontamination durch Feuchtigkeit, Staub und korrodierende Stoffe, um einen besseren thermischen Übergang zwischen wärmeerzeugenden Komponenten und dem Gehäuse zu gewährleisten und um die internen Komponenten vor den Einwirkungen von mechanischem Schock und Vibration zu schützen. Um ökologische und thermische Anforderungen an die Anwendung zu erfüllen, ist die Vergussmasse in vielen Anwendungen wesentlich. Die meist verbreiteten Werkstoffe sind Zweikomponenten-Epoxidharz, Silikonkautschuk oder Polyurethan. Davon erfüllt nur Silikonkautschuk keine strikte Anforderungen an Rauch, aber auch raucharme Vergussmaterialien sind verfügbar (jedoch leider zu einem hohen Preis).

Eine Norm zur Bewertung des Rauchrisikos ist der NF-F-Brand- und Rauchtest. Die Normen NF F 16-101/102 bedienen sich einer Risikobewertungstabelle, um die erforderlichen Spezifikationen für die französische Staatsbahn zu bestimmen. Die Risiken der Flammbarkeit und Rauchentwicklung können dann nach der Anwendung oder Zugstreckenführung bewertet werden (wenn die Strecke des Zuges durch Tunnel führt, ist Rauch ein größeres Risiko, als wenn er nur oberirdisch eingesetzt wird). Die Charakteristik des Rauchs hängt sowohl von der Rauchtrübung als auch von der Giftigkeit des Rauchs ab. Im Beispiel unten wäre die Stromversorgung für einen U-Bahnzug geeignet (Anforderungen in Weiß dargestellt), da Prüfungen zeigten, dass Silikonkautschuk-Vergussmaterial eine niedrigere Flammbarkeit aufweist, jedoch zu viel Rauch (F3/I3) abgibt als akzeptabel wäre (F1/I2).

Flammability			Smoke emission	
I0	for I.O. ≥ 70	no inflammation at 960°C	F0	for I.F. ≥ 5
I1	for I.O. 45-69	no inflammation at 960°C	F1	for I.F. 6-20
I2	for I.O. 32-44	no inflammation at 850°C	F2	for I.F. 21-40
I3	for I.O. 28-31	no afterburning at 850°C	F3	for I.F. 41-80
I4	for I.O. ≥ 20		F4	for I.F. 81-120
NC	not classified		F5	for I.F. ≥ 120

	I0	I1	I2	I3	I4	NC
F0						
F1			X			
F2						
F3				X		
F4						
F5						

X = Silicone Rubber
 X = Epoxy

Tabelle 6.9: Brand/Rauch-Risikobewertung

6.5 Verletzungsrisiko

Während die meisten DC/DC Stromversorgungen auf Platinen montiert sind und für den Kunden nicht zugänglich sind, müssen sie zu Service- und Reparaturzwecken für technisches Personal zugänglich sein. Eine Verletzungsgefahr durch hohe Temperaturen oder scharfe Kanten sollte ausgeschlossen sein.

6.5.1 Heiße Oberflächen

Die Hautverbrennungstemperatur beträgt ca. 48°C bei einem Kontakt von mehr als 10 Minuten. Aufgrund unserer Reflexhandlung, von einer heißen Oberfläche zurückzutreten oder einen heißen Gegenstand fallen zu lassen, sind in Stromversorgungen auch höhere Temperaturen zulässig. Die internationale Norm IEC60950 spezifiziert je nach Berührungszeit und Material die folgenden Temperaturgrenzen für berührbare Oberflächen oder Gegenstände:

Gegenstand	Berührungszeit	Material	Temperaturgrenze
Handgriffe, Drehknöpfe usw.	ununterbrochen	alle	55°C
	kurz (<10s)	metallisch	60°C
		nichtmetallisch	70°C
heiße Oberflächen	berührbar (<1s)	metallisch	70°C
		nichtmetallisch	80°C

Tabelle 6.10: Temperaturgrenzen für berührbare Gegenstände

In vielen Stromversorgungen können metallische Teile wie Kühlkörper so hohe Temperaturen erreichen, dass sie unter bestimmten Betriebsbedingungen als Gefahrenquellen für Hautverbrennungen einzustufen sind. In diesen Fällen muss die Stromversorgung zum Schutz gekapselt und möglicherweise mit einem Warnsymbol markiert sein, sodass erkennbar ist, dass sie heiße Oberflächen enthält und das Reparatur-oder Bedienpersonal angemessen über das Risiko von Hautverbrennungen informiert wird.



Abb. 6.14: Warnsymbol für heiße Oberflächen

6.5.2 Scharfe Kanten

Stromversorgungen enthalten keine beweglichen Teile, abgesehen von Lüftern bei Spannungsversorgungen für besonders hohe Leistungen. Deshalb trifft die Gefahr der durch bewegliche mechanische Teile verursachten kinetischen Energie in den meisten Fällen nicht zu. Alle scharfen Kanten und Ecken an den mechanischen Teilen können jedoch sowohl für Montage- als auch für Reparaturpersonal eine Gefahr darstellen. Der Standardtest für scharfe Kanten wird durchgeführt mit einem Weichgummistück, das umwickelt ist mit Spezialbändern, die die Schnittfestigkeit der Haut wiedergeben. Das Probestück ist mit einer Springfeder versehen, um einen gleichmäßigen Druck zu gewährleisten, und kann entlang jeglicher mechanischer Teile bewegt werden. Wenn scharfe Kanten vorhanden sind, spaltet sich das Band und legt den darunterliegenden Gummi, der eine andere Farbe hat, frei. So kann eine einfache Aussonderungsprüfung durchgeführt werden.

Praktischer Hinweis

Die meisten Kleinleistungs-DC/DC-Stromversorgungen haben ein durch Spritzguss hergestelltes Kunststoffgehäuse, wobei durch geeignete Spritzgussformen die Ausbildung zu scharfer Kanten vermieden werden kann. Wenn das Material jedoch entlang einer Trennfuge einen dünnen Streifen überschüssigen Kunststoffes, genannt „Spritzgrat“, ausbildet kann dieser eine gefährliche, scharfe Kante bilden. Diese muss durch Entgratung entweder manuell oder durch Einsatz entsprechender Schleifmitteln entfernt werden. Einige Hochleistungswandler verfügen über ein Metallgehäuse oder enthalten Kühlkörper. Sie werden in der Regel aus Weichmetallen wie Aluminium oder Kupfer hergestellt, und es ist relativ einfach, sämtliche Verarbeitungsgrate vor der Montage mit Handwerkzeug zu entfernen.

6.6 Auf Sicherheit ausgerichtete Konstruktionsweise

Der erste Schritt in einer risikobewussten (HB) Vorgehensweise im Rahmen einer auf die Sicherheit ausgerichteten Konstruktionsweise besteht darin, sämtliche Gefahren, die in der Anwendung vorhanden sind, zu bewerten. Im folgenden Schritt wird durch Übertragung auf den menschlichen Körper untersucht, wie diese Gefahren zu Risiken werden könnten. Und der letzte Schritt besteht darin, wirksame Schutzmaßnahmen zu entwerfen. Alle diese Schritte können während des Konstruktionsbewertungsverfahrens durchgeführt werden, also noch bevor die Geräte gebaut werden. Eine erfolgreiche Vorgehensweise basiert auf der Ermittlung bereits vorhandener Erkenntnisse über potentielle Gefahren, Übertragungsmechanismen und menschlicher, körperlicher Anfälligkeit gegenüber Verletzungen und darin, fundierte Grundsätze der Elektrotechnik anzuwenden.

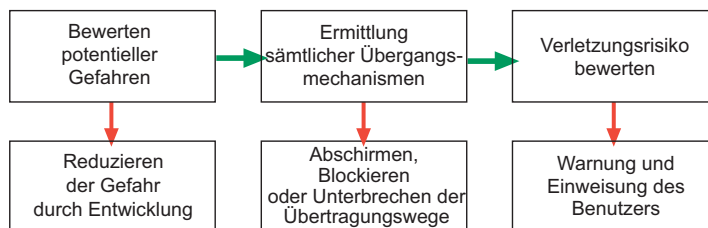


Abb. 6.15: Blockdiagramm einer auf Sicherheit ausgerichteten Konstruktionsweise

Das obige Blockdiagramm veranschaulicht eine auf Sicherheit ausgerichtete Konstruktionsweise. Die verschiedenen Sicherheitsstandards, Richtlinien und Zertifikate wurden auf Basis der Erforschung vorausgegangener Unfälle entwickelt, um elektrische und elektronische Geräte betriebssicher zu machen. Jedoch ist es ein Unterschied, ob Stromversorgungen so konstruiert werden, dass sie Sicherheitsstandards gerade so erfüllen, oder ob der Entwicklungsprozess von vornherein auf die Sicherheit ausgerichtet ist und die Gefahren, Übertragungsmechanismen und Verletzungsrisiken von Anfang an betrachtet und geeignete Schutzmechanismen gleich mit entworfen werden. Ein Beispiel für Minimierung von Gefahren durch entsprechende Konstruktion könnte darin bestehen, für Geräte, die nicht ständig gespeist werden müssen, eine batteriebetriebene statt einer Netzstromversorgung zu wählen.

Für die Batteriespannung kann dann eine Niederspannung verwendet werden; somit basiert die Sicherheit auf einer entsprechenden Entwicklung. Ein weiteres Beispiel sind die gelb gefärbten Transformatoren, die üblicherweise auf Baustellen zur Versorgung von Elektrowerkzeugen verwendet werden. Durch Herabsetzung der Netzstromversorgung von 230VAC auf 110VAC und Verwendung eines Trafos, mit einer an die Masse angeschlossenen Mittelanzapfung, beträgt die maximale AC-Spannung zur Masse 55V und ist somit sicher(Abb. 6.16).

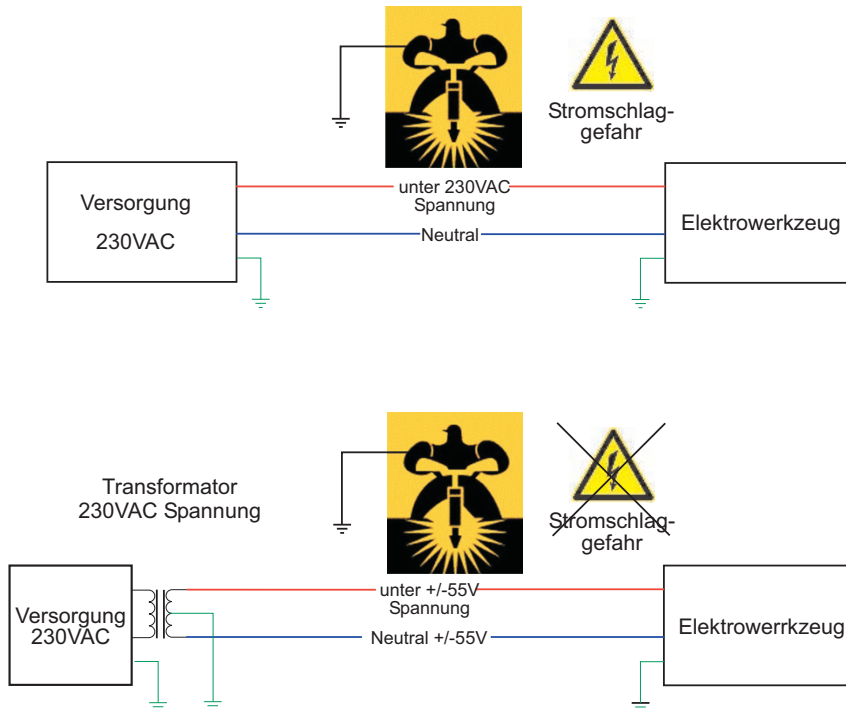


Abb. 6.16: Beispiel von Sicherheit durch Konstruktion

Kann eine Gefahr nicht durch die Konstruktion ausgeschlossen werden, besteht die nächste Stufe des Prozesses darin, die Übertragungsmechanismen zu untersuchen, die zu einer Berührung durch Menschen und somit zu potentieller Verletzung führen könnten. Die einfachste Methode besteht darin, jeglichen zufälligen Kontakt physikalisch durch den Einsatz von Abschirmung, Gehäuse oder Isolierung zu blockieren. Der Code für den Schutz gegen Eindringen (IP-Code, IP = Ingress Protection) bestimmt verschiedene Grade von physikalischem Schutz unter Verwendung eines standardisierten Probestücks („Prüffinger“), das mit einem Durchmesser von 12,5mm und einer Länge von 80mm einen menschlichen Finger simuliert.c

Grad	Gegenstands- größe	Effektiver Schutz
0		Kein effektiver Schutz gegen Kontakt oder Eindringen von Gegenständen
1	>50mm	Geschützt gegen zufälligen Kontakt großer Körperteile (z. B. Hände), aber kein Schutz gegen absichtlichen Kontakt
2	>12.5mm	Geschützt gegen zufälliges Berühren mit dem Finger
3	>2.5mm	Geschützt gegen zufällige Berührung durch Instrumente, dicke Drähte usw.
4	>1mm	Geschützt gegen zufälliges Berühren durch dünne Drähte, kleine Teile usw.
5	staubgeschützt	Eindringen von Staub wird nicht vollständig verhindert, hat aber keinen Einfluss auf die Funktion. Vollschutz gegen zufälliges Berühren
6	staubdicht	Abgedichtet gegen Staub; Vollschutz gegen zufälliges Berühren

Tabelle 6.11: Schutzgrade gegen das Eindringen von Feststoffpartikeln

Die zweite Ziffer im IP-Code stellt den Schutz gegen das Eindringen von Flüssigkeiten dar und variiert von Grad 1 (kein Schutz vor Eindringen von Wasser) bis Grad 8 (vollständiges Eintauchen).

Grad	Test	Effektiver Schutz
0	kein Schutz	keiner
1	Tropfwasser	geschützt vor leichtem Regen
2	Tropfwasser bis zu 15°	geschützt vor starkem Regen
3	Sprühwasser	geschützt gegen Sprühwasser bis zu einem Winkel von 60°
4	Wasserspritzer	geschützt gegen Wasser von einem beliebigen Winkel
5	Wasserstrahl	geschützt gegen Wasser aus einer Düse
6	starker Wasserstrahl	geschützt gegen Wasser aus einer großen Düse
6K	starker Wasserstrahl mit Druck	geschützt vor starken Wasserstrahlen
7	Eintauchen bis zu 1 m	geschützt gegen Eintauchen in Wasser für die Dauer von 30 Minuten
8	Eintauchen über 1 m	geschützt gegen die Wirkungen beim dauernden Untertauchen in Wasser

Tabelle 6.12: Schutzgrade gegen das Eindringen von Flüssigkeiten

Für den Einsatz im Innenbereich wird eine IP-Einstufung von IP20 für die Anwendung im Trockenbereich als ausreichend erachtet.

Für den Einsatz in Feuchträumen (z. B. in Baderäumen) ist mindestens IP41 (Schutz gegen kleine herabfallende Gegenstände und Tropfwasser) und für den Außeneinsatz ein IP-Wert von mindestens IP54 (staubdicht und Wasserspritzer haben keine schädliche Wirkung) erforderlich, häufiger jedoch ein Wert von IP67 (Vollschutz gegen das Eindringen von Staub und Schutz beim Eintauchen in Wasser) erforderlich.

Die ultimative Sicherheitsmaßnahme im Hinblick auf eine durch die Konstruktion bestimmte Sicherheit besteht darin, den Benutzer durch die Anbringung von Warnhinweisen an den Geräten und durch das Beifügen von Merkblättern mit Sicherheitsinformationen und Bedienungsvorschriften zu warnen, damit er nicht in eine gefährliche Situation gerät. Die Anforderungen an die Sicherheitsinformationen werden in allen relevanten Hygiene- und Sicherheitsvorschriften dokumentiert und Hinweise sind – je nach Schwere der Gefahr – mit „Warnung“ oder „Vorsicht“ oder „Gefahr“ klar markiert. Für DC/DC-Stromversorgungen ist es aber nur selten notwendig Gefahrenwarnhinweise anzubringen, wenn keine gefährlichen Hochspannungen erzeugt werden.

6.6.1 FMEA

Auch die Fehlermöglichkeits- und –einflussanalyse (FMEA) ist ein wichtiges Verfahren im Rahmen einer auf die Sicherheit gerichteten Konstruktionsweise. Auf der einfachsten Ebene kann angenommen werden, dass jede separate Komponente oder jedes separate Bauteil in der Konstruktion unter normalen Betriebsbedingungen, entweder als Kurzschluss oder als offener Stromkreis (Nenneingangsspannung, Volllast und Raumtemperatur) ausfällt. Das Ergebnis dieses Fehlers kann im Hinblick auf die Folgen für die Sicherheit und die Performance sowohl für die Stromversorgung als auch für alle anderen Geräte oder Systeme, die daran angeschlossen werden sollen, analysiert werden. Die FMEA kann verwendet werden, um die Schwere eines Fehlers laut folgender Definitionen zu bewerten:

Schwere	Definition
katastrophal	Führt zu mehreren Todesfällen und/oder Vollverlust von mehrfachen Systemfunktionen
gefährlich	Beeinträchtigt die Systemsicherheit oder schafft gefährliche Bedingungen: - starke Verringerung der Sicherheit oder Funktionsfähigkeit - schwere und tödliche Körperverletzung
risikobehaftet	Verringert die Systemsicherheit wesentlich und schafft gefährliche Bedingungen: - wesentliche Verringerung der Sicherheit oder Funktionsfähigkeit - Unwohlsein einschließlich Verletzungen - schwerer Umweltschaden und/oder schwerer Sachschaden
geringfügig	Verringert die Systemsicherheit nicht wesentlich: - gewisses physisches Unwohlsein - geringfügiger Umweltschaden, und/oder geringfügiger Sachschaden
keine Wirkung	Hat keine Wirkung auf Sicherheit oder Performance

Tabelle 6.13: Einstufung der Schwere eines Fehlers

Die FMEA kann als grundlegendes Verfahren bei einer auf die Sicherheit gerichteten Konstruktionsweise verwendet werden, ohne die Wahrscheinlichkeit des Auftretens eines Fehlers zu betrachten. Wird sie jedoch als Teil einer Risikobewertung verwendet, dann wird das Risiko wie folgt berechnet: Risiko = Schwere × Wahrscheinlichkeit

Daher muss die Wahrscheinlichkeit des Auftretens eines Ausfalls bewertet werden:

Wahrscheinlichkeit	Definition
wahrscheinlich	• Auftreten ein- oder mehrmals während der Gesamtsystemlebens-/Gebrauchsdauer eines Bauteils erwartet • Quantitativ: Wahrscheinlichkeit des Auftretens pro Betriebsstunde ist größer als 1×10^{-5}
entfernt vorhanden	• Qualitativ: Auftreten für jedes Bauteil während dessen Gesamtlebensdauer unwahrscheinlich. Kann während der Lebensdauer eines Gesamtsystems oder Wagenparks mehrmals auftreten. • Quantitativ: Wahrscheinlichkeit des Auftretens pro Betriebsstunde ist geringer als 1×10^{-5}, aber größer als 1×10^{-7}
sehr entfernt vorhanden	• Qualitativ: Auftreten für jedes Bauteil während dessen Gesamtlebensdauer nicht erwartet. Kann während der Lebensdauer einen Gesamtsystems oder Wagenparks mehrmals auftreten • Quantitativ: Wahrscheinlichkeit des Auftretens pro Betriebsstunde ist geringer als 1×10^{-5}, aber größer als 1×10^{-7}
extrem unwahrscheinlich	• Qualitativ: so unwahrscheinlich, dass das Auftreten während der Gesamtgebrauchsdauer eines Gesamtsystems oder Wagenparks nicht zu erwarten ist • Quantitativ: Wahrscheinlichkeit des Auftretens pro Betriebsstunde ist geringer als 1×10^{-9}

Tabelle 6.14: Einstufung der Ausfallwahrscheinlichkeit

6.7 Medizintechnische Sicherheit

Im Vergleich zu Betriebssicherheitsanforderungen müssen DC/DC-Stromversorgungen wesentlich strikteren Spezifikationen entsprechen, wenn sie als geeignet für den Einsatz in medizinischen Geräten zertifiziert werden. Eine der bedeutendsten Anforderungen besteht darin, dass ein Unterschied zwischen Patienten- und Bedienerschutz gemacht wird. In der MOOP-Kategorie (MOOP = means of operator protection, dt. Mittel zum Schutz des Bedieners) folgen die Anforderungen im Wesentlichen den Standardschutzmaßnahmen, die für Industrieenanrichtungen verwendet werden, in der MOPP-Kategorie (MOPP = means of patient protection, dt. Mittel zum Patientenschutz) jedoch wurden die Anforderungen erhöht – insbesondere in Bezug auf Kriech- und Luftstrecken. Tabelle 6.15 zeigt einige entsprechende Anforderungen der Isolierung für beide Schutzkategorien des Schutzverfahrens (MOP).

Isolations- anforderung	MOOP			MOPP		
	Luft- strecke	Kriech- strecke	Isolations- spannung	Luft- strecke	Kriech- strecke	Isolations- spannung
<250VAC	2.0mm	3.2mm	1500VAC	2.5mm	4.0mm	1500VAC
	4.0mm	6.4mm	3000VAC	5.0mm	8.0mm	4000VAC
<43VDC	1.0mm	2.0mm	1000VAC	1.0mm	2.0mm	1500VAC
<30VAC	2.0mm	4.0mm	2000VAC	2.0mm	4.0mm	3000VAC

Tabelle 6.15: Schutzanforderungen an medizinische Geräte

Eines der Hauptsicherheitsprinzipien besteht darin, zwei Schutzgrade zu gewährleisten, falls ein Grad versagt. Daher sind zwei Schutzmaßnahmen (means of protection - MOP, Schutzmittel) sowohl für die Betreiber- als auch für die Patientensicherheit (2×MOPP, 2×MOOP) erforderlich. Das MOP kann zwischen einigen Komponenten aufgespalten werden, sodass z. B. eine AC/DC-Stromversorgung beim Einsatz 2×MOOP gewährleisten kann und ein darauf folgender isolierter DC/DC-Wandler 2×MOPP gewährleisten kann. Es liegt aber im Wesen von Konstruktionen medizintechnischer Qualität, eine „belt-and-braces“-Methode zu verwenden und 2×MOPP und 2×MOOP sowohl für die AC/DC- als auch für die DC/DC-Stromversorgung zu fordern.

Dank der MOOP/MOPP-Konzeption wurden jedoch die Anforderungen für Patienten-Leckströme im Vergleich zu früheren medizinischen Normen, je nach Klassifikation der Anwendung im Hinblick auf Patientenkontakt, herabgesetzt. Die drei Patientenkontaktklassifikationen sind in Typ B (Body – kein direkter Patientenkontakt), Typ BF (Body Float – physikalischer Kontakt mit dem Patienten) und CF (Direktkontakt mit dem menschlichen Herzen) unterteilt. Je näher der Kontakt mit dem Herzen des Patienten ist, desto niedriger sind die zulässigen Leckströme. Tabelle 6.16 zeigt die entsprechenden Grenzen für Normalbetrieb (NC – normale Bedingungen) sowie für den Fall einer Störung (SFC – Einzelfehlerbedingung).

Ableit- strom	Typ B		Typ BF		Typ CF	
	NC	SFC	NC	SFC	NC	SFC
Erde	500µA	1mA	500µA	1mA	500µA	1mA
Gehäuse	100µA	500µA	100µA	500µA	100µA	500µA
Patient	100µA	500µA	100µA	500µA	10µA	50µA

Tabelle 6.16: Leckstrom-Grenzwerte für medizinische Geräte

Der Vorteil von neuen, höheren Stromgrenzwerten besteht darin, dass es für medizintechnisch geeignete DC/DC-Stromversorgungen viel einfacher ist, auch die EMV-Anforderungen zu erfüllen. Die früheren Grenzen verursachten wesentliche Probleme, da Gleichaktfilterkondensatoren zwischen Ausgang und Eingang gänzlich unzulässig waren. Der risikobasierte Ansatz für medizinische Sicherheit stellt Stromversorgungshersteller und Hersteller von medizinischen Geräten vor eine neue Herausforderung, da die Zertifizierung ein formales Risikomanagementverfahren (RM) nach ISO 14971 beinhaltet.

Das Risikomanagement deckt nicht nur die Risikoanalyse in der Planungs- und Produktionsphase ab, sondern beinhaltet auch Anforderungen, das Risiko im Laufe langer Nutzungszeit zu überwachen, um die Auswirkungen von Alterung, bestimmungsgemäßer Nutzung und Missbrauch, die die Gerätesicherheit über die Gesamtlebensdauer des Produktes gefährden könnten, nachzuweisen.

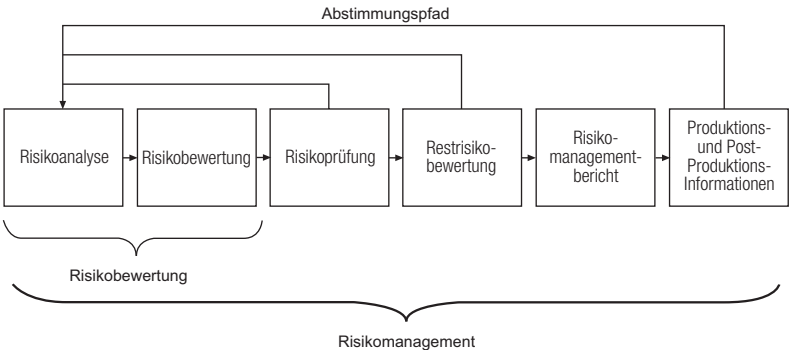


Abb. 6.17: Flussdiagramm Risikomanagement (nach ISO 14971)

Der Ablauf der Risikobewertung basiert auf einer Risikoindex-Matrix, welche Risiken analysiert und gewichtet, die durch die Stromversorgung im Normalbetrieb und unter Ausfallbedingungen entstehen können. In der Risikoindex-Matrix wird die Eintrittswahrscheinlichkeit (von unwahrscheinlich bis häufig) gegenüber der Schwere (von unbedeutend bis katastrophal), in jeweils fünf Grade unterteilt, gegenübergestellt. Das Gesamtrisiko ergibt sich als Produkt aus Wahrscheinlichkeit und Schweregrad. Liegt der Wert ≤ 6 , kann das Risiko in der Regel als annehmbar eingestuft werden. Die Akzeptanz des Grenzwertes liegt jedoch beim Benutzer. Im folgenden Beispiel ist R1 ein annehmbares Risiko, R2 ein angemessenes Risiko für die Anwendung, und R3 ist unannehmbar.

			Schweregrad				
			Unbedeutend	Geringfügig	Schwer	Kritisch	Katastrophal
		Gewichtung	1	2	3	4	5
Wahr- schein- lichkeits- grad	Häufig	5	R1	R3	R3	R3	R3
	Wahrscheinlich	4	R1	R2	R3	R3	R3
	Gelegentlich	3	R1	R1	R2	R3	R3
	Entfernt	2	R1	R1	R1	R2	R3
	Unwahrscheinlich	1	R1	R1	R1	R1	R2

Tabelle 6.17: Beispiel einer Risiko-Matrix

Der Risikomanagementplan ermittelt auch Verifizierungsverfahren, um sowohl zu gewährleisten, dass die Risikoprüfmaßnahmen im Gerät verwirklicht sind, als auch um sicherzustellen, dass die verwirklichten Risikoprüfmaßnahmen das Risiko tatsächlich verringern. In vielen Fällen profitieren die Qualitätssicherungssysteme sowohl des Herstellers als auch des Benutzers davon, dass das Risikomanagementverfahren in die Dokumentations- und Prüfverfahren aufgenommen wird, insbesondere weil die Aufgaben der Risikoprüfung und Überwachung über die Gesamtlebensdauer des Produktlebenszyklus, einschließlich Herstellung, Montage, Wartung und Überalterung, fortgesetzt werden müssen.

7. Betriebszuverlässigkeit

7.1 Vorhersage der Betriebszuverlässigkeit

Fast seit Beginn der Nutzung der Elektrotechnik ist es für den Anwender notwendig, zu wissen, wie lange die Geräte ordnungsgemäß arbeiten werden. Da niemand die Zukunft vorhersagen kann, wurden statistische Verfahren entwickelt, um die Betriebszuverlässigkeit von Bauteilen, Komponenten, Baueinheiten oder Geräten vorauszusagen.

Einer der frühesten systematischen Ansätze zur Zuverlässigkeit elektronischer Bauteile und -gruppen war das „Militärhandbuch – Zuverlässigkeitsvorhersage von elektronischen Geräten“ („Military Handbook – Reliability Prediction of Electronic Equipment“) der amerikanischen Armee, allgemein bekannt als MIL-HDBK-217, das hauptsächlich aus in einer großen Datenbank zusammengefassten, gemessener Ausfallraten verschiedener elektronischer Bauteile besteht und auf der empirischen Analyse einer großen Anzahl von Betriebsausfällen elektrischer, elektronischer und elektro-mechanischer Komponenten basiert, durchgeführt durch die Universität Maryland.

Das Handbuch wurde bis 1995 kontinuierlich aktualisiert und verbessert, wodurch die endgültige Version, MIL-HDBK 217, Revision F, Anmerkung 2, entstand. Obwohl dieses Werk nicht mehr aktualisiert wird, werden die Daten und Verfahren noch heute am häufigsten verwendet.

Das Handbuch enthält zwei Verfahren der Zuverlässigkeitsvorhersage, die Belastungsanalyse der Teile (PSA) und die Zählanalyse der Teile (PCA). Das PSA-Verfahren erfordert eine größere Menge an ausführlichen Informationen und wird normalerweise in einer späteren Planungsphase besser anwendbar, wenn Messdaten und das vorläufige Ergebnis in die Betriebszuverlässigkeitsmodelle einfließen können, während das PCA-Verfahren nur die minimalen Informationen, wie die Bauteilanzahl, Qualitätsgrad und Einsatzbedingungen, erfordert. Der größte Vorteil der Methodologie des MIL HDBK 217 besteht darin, dass das PCA-Verfahren eine nur auf der Materialliste (BOM = Bill of Materials) und der erwarteten Nutzung basierende Zuverlässigkeitsvorhersage ergibt, weshalb eine Zuverlässigkeitsvorhersage sogar für ein Produkt, das noch nicht gebaut wurde, gemacht werden kann:

$$\lambda_P = (\sum N_C \lambda_C) (1 + 0.2 \pi_E) \pi_F \pi_Q \pi_L$$

Gleichung 7.1: Berechnung der Ausfallrat

Wo:

- N_C Anzahl der Teile (pro Komponententyp)
- λ_C Ausfallrate für jedes Teil (Bezugswert aus der Datenbank)
- π_E Umweltbelastungsfaktor (anwendungsspezifisch)
- π_F Hybridfunktionsbelastung (Zusatzbelastung, hervorgerufen durch Wechselwirkung der Komponenten)
- π_Q Sortierungsgradfaktor (Normteiltoleranzen oder vorsortiert)
- π_L Reifegradfaktor (bekannte und geprüfte Konstruktion oder neues Konzept)

Die Berechnung ergibt einen Wert für jede verwendete Komponente.

Die Gesamtbetriebszuverlässigkeit kann dann durch Addierung aller individuellen Ergebnisse ermittelt werden:

Nummer	Teile	Qualität	π_p Ausfallrate	π_p Ausfallrate
			$[10^{-6}/h] T_{AMB} = 25^{\circ}C$	$[10^{-6}/h] T_{AMB} = 85^{\circ}C$
1	Transistoren	2	0.0203	0.0609
2	Dioden	2	0.1089	0.5443
3	Widerstände	3	0.0370	0.1716
4	Kondensatoren	5	0.1699	1.7000
5	Transformatoren	1	0.2256	1.9200
6	PCB, PIN	2	0.0092	0.0092
π_p Gesamtausfallrate $10^{-6}/H$			0.5708	4.4060
Stunden mittleren Ausfallabstandes (MIL-HDBK-217F)			1.751.927	226.963
Kondition	Eingang	Nominaler Eingang		Nominaler Eingang
	Ausgang	Volllast		Volllast

Tabelle 7.1: Beispiel für die Berechnung eines mittleren Ausfallabstandes durch Teilezählung für einen einfachen DC/DC-Wandler

Die Ausfallraten sind entweder als Zeitintervall zwischen zwei Ausfällen – in Stunden –, mittlere Betriebsdauer zwischen Ausfällen/Mean Time Between Failures (MTBF) oder als Zeit bis zum ersten Ausfall – mittlere Lebensdauer/Mean Time To Failure (MTTF) – definiert. Das Verhalten der Standardausfallrate wird durch die allgemein bekannte „Ausfallratenkurve“ beschrieben. Abbildung 7.1 zeigt die Form der Kurve. Die Form der Kurve ist für alle Komponenten und Systeme ungefähr dieselbe – nur die Skalierung der Zeitachse ist unterschiedlich. Sie ist in drei Hauptbereiche unterteilt: Frühausfall (I), Ausfall während der Nutzungszeit (II) und Ausfall zum Ende der Nutzungszeit (III). MTTF beinhaltet die Zonen I und II, während MTBF (mittlerer Ausfallabstand) nur Zone II einschließt.

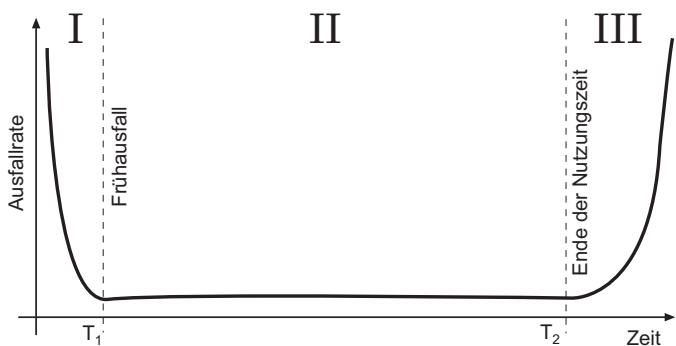


Abb. 7.1: Ausfallratenkurve

Abschnitt I beschreibt den Bereich von Frühausfällen, die in der Regel durch versteckte Werkstoff- oder Produktionsfehler verursacht werden, die die Qualitätskontrolle des Herstellers passieren konnten. Der Frühausfall erfolgt üblicherweise nach einer relativ kurzen Dauer – sogar bei komplizierten Systeme gibt es selten Frühausfälle nach 200 Einsatzstunden; im Falle von DC/DC-Wandlern treten die meisten Frühausfälle innerhalb der ersten 24 Betriebsstunden auf.

Dies mag bei einem Wandler mit einer garantierten Lebensdauer von 3 Jahren sehr kurz erscheinen, aber in einem DC/DC-Wandler, der bei 100kHz läuft, werden Schalttransistoren und Transformator schon am ersten Betriebstag mehr als 140 Millionen mal genutzt, und jeder Ausfall infolge von Fehlern in den eingesetzten Bauteilen ist dann mit hoher Wahrscheinlichkeit bereits eingetreten.

Da Wärmebeanspruchung einer der den Ausfall beschleunigenden Faktoren ist, kann die Übergangszeit (T_1) vom Frühausfall in die Nutzungszeit durch einen Thermo-trainingsprozess (engl.: burn-in-process) in einer Umgebung mit erhöhter Temperatur wesentlich verkürzt werden (Abb. 7.2). Wenn Wandler durch den Betrieb unter Volllast bei höheren Temperaturen beansprucht werden, reicht eine Burn-in-Zeit von ca. 4 Stunden aus, um den Großteil der Frühausfälle zu entdecken. Treten immer noch Frühausfälle in der Endanwendung auf, kann die Burn-in-Zeit entsprechend verlängert werden. Für Anwendungen, die eine hohe Zuverlässigkeit erfordern, wie beispielsweise Bahnapplikationen, ist eine Burn-in-Zeit von 24 Stunden üblich.



Abb. 7.2: DC/DC-Wandler während des Tests in einer Burn-in-Kammer ($T_{AMB} = 40^\circ\text{C}$)

Während der durch Bereich II charakterisierten Nutzungsdauer ist die Ausfallrate auf einem niedrigen Niveau konstant. Die zweite Übergangszeit (T_2) von der Nutzungsdauer in das Ende des Produktlebenszykluses wird durch viele Faktoren, wie der Entwicklungsqualität und der Qualität der verwendeten Bauteile, der Produktionsqualität und von Umweltbeanspruchungen in der Anwendung, beeinflusst. Bereich III stellt die Endphase des Produktlebenszyklus dar. Während dieser Phase kann eine reduzierte Performance, verursacht durch Verschleiß, chemischem Abbau der verwendeten Materialien und plötzliche Ausfälle, erwartet werden.

Da die meisten DC/DC-Hersteller einen Burn-in-Prozess einsetzen, um Frühausfälle zu erkennen, wird in den Datenblättern häufiger MTBF als MTTF angegeben.

Einige Hersteller bevorzugen den Reziprokwert der MTBF-Ausfallrate. Dieser ist auf 109 Stunden normiert und wird als „Failures In Time“ (FIT) bezeichnet.

$$FIT = \frac{10^9}{MTBF}$$

Gleichung 7.2: Verhältnis von FIT gegenüber MTBF (mittlerer Ausfallabstand)

7.2 Umweltbelastungsfaktoren

MIL-HDBK-217 enthält Betriebszuverlässigkeitsmodelle, die auf allgemeinen Militäranwendungen basieren. Die Art der Anwendung, in der ein DC/DC-Wandler verwendet werden soll, hat einen starken Einfluss auf dessen Betriebszuverlässigkeit. Wenn der Wandler beispielsweise in einem Schiff zum Einsatz kommen soll, werden die Korrosionseinwirkungen der salzhaltigen Luft dessen Lebensdauer verringern, selbst wenn er in einem Trockenbereich eingesetzt wird.

Umgebung	π_E -Symbol	Beschreibung MIL-HDBK-217F	Kommerzielle Interpretation oder Beispiele
Boden kritisch	GB	Nicht mobil, temperatur- und feuchtigkeitsgeregelte Umgebungen, mit einfachem Zugang zur Wartung	Laboraüstung, Testgeräte, Desktop-PCs, feste Fernmeldetechnik
Boden mobil	GM	Geräte in Rad- oder Kettenfahrzeugen und manuell transportierte Geräte	Fahrzeuginstrumentierung, mobile Funk- und Telekommunikationstechnik, tragbare PCs
See. geschützt	NS	Geschützte oder Unterdeck-Geräte in Überseeschiffen oder U-Booten	Navigations-, Funk- und Messgeräte von Schiffen unter Deck
Frachtflugzeug	AIC	Typische Bedingungen in Frachträumen, die durch Flugzeugbesatzung besetzt werden können	Überdruckkabinen und Cockpits, Anwendungen zur Unterhaltung während des Fluges (inflight entertainment) und nicht sicherheitskritische Anwendungen
Raumflug	SF	Erdumlaufbahn; Fahrzeug weder im aktiven Flug noch im Wiedereintritt in die Atmosphäre	Umlauf-Kommunikationssatellit, Geräte zum nur einmaligen Einsatz in-situ
Raketenstart	ML	Schwerbetrieb bezüglich Raketenstart	Schwervibrationsschlag und sehr hohe Beschleunigungskräfte, Satellitenstartbedingungen

Tabelle 7.2: Anwendungsklassen nach MIL-HDBK-217

Wenn die Endanwendung bekannt ist, kann ein Korrekturfaktor für die Berechnungen des mittleren Ausfallabstandes basierend auf „Boden unkritisch“ (Ground Benign; GB) als Referenzumweltbelastung mit einem Koeffizienten von 1 einbezogen werden:

Umgebung	π_E Symbol	π_E -Wert	Divisor
Ground Benign	GB	0.5	1.00
Ground Mobile	GM	4.0	1.64
Naval Sheltered	GNS	4.0	1.64
Aircraft Inhabited Cargo	AIC	4.0	1.64
Space Flight	SF	0.5	1.00
Missile Launch	ML	12.0	3.09

Tabelle 7.3: MTBF Korrekturfaktor je nach Umgebung

Ein DC/DC-Wandler mit einer MTBF-Zahl von 1 Million Stunden gemäß Datenblatt (basierend auf GB-Bedingungen) müsste beim Einsatz in tragbaren Geräten z. B. bis auf ca. 610.000 Stunden „reduziert“ werden, um den zusätzlichen Umweltbelastungen infolge der mit tragbaren Geräten verbundenen Schlägen, Zusammenstößen, schnellen Temperaturveränderungen usw. Rechnung zu tragen.

Eines der möglicherweise merkwürdigen Ergebnisse der Analyse des MIL-HDBK-217 ist, dass die Raumflugumgebung also genauso günstig wie eine terrestrische Umgebung gilt. An Bord eines Satelliten oder Raumfahrzeugs werden die Umweltbedingungen geregelt, und es gibt keine Vibration oder Luftverschmutzung, sodass elektronische Geräte theoretisch eine sehr hohe Lebensdauer haben. In der Praxis kann jedoch z.B. die Höhenstrahlung Löcher durch Halbleiterträger schneiden und Ausfälle hervorrufen.

Es ist möglich, DC/DC-Wandler mit „strahlungsbeständigen“ Bauteilen aufzubauen, die einen zusätzlichem Schutz vor Hochenergiestrahlung besitzen. Es kann aber auch zuverlässiger sein, einfachere Schaltungen ohne jegliche ICs zu verwenden. Ein FET kann beträchtlichen durch die Höhenstrahlung verursachten Schäden widerstehen, weil die Trägermaterialfläche relativ groß und Punktdefekten gegenüber toleranter ist. Somit ist ein einfacher Gegentakt-DC/DC-Wandler, der nur aus diskreten Bauteilen aufgebaut ist, häufig für Weltraumfahrtanwendungen geeignet.

7.3 Verwendung von MTBF-Angaben

MTBF-Angaben stiften manchmal große Verwirrung, weil sie häufig falsch verstanden und gelegentlich auch durch die Hersteller absichtlich verzerrt werden. Eine MTBF-Zahl von 1 Million Stunden bedeutet nicht, dass das Produkt eine Lebensdauer von:

$$\frac{1,000,000}{24 \times 365} = 114 \text{ Jahren hat!}$$

MTBF ist einfach als inverser Wert der tatsächlichen Ausfallrate definiert. Dieser Wert sagt lediglich, dass statistisch gesehen ein DC/DC-Wandler aus einer Gesamtheit von 100 Exemplaren nach 10.000 Betriebsstunden ausfällt:

$$MTBF = \frac{10,000}{1/100} = 1 \text{ Millionen Stunden}$$

Wenn die Ausfallrate für eine bestimmte Anzahl unter Einsatzbedingungen weniger als 1 % jährlich betragen soll, sollte der erforderliche MTBF-Nennwert der Stromversorgungen wie folgt aussehen:

$$\text{erforderlicher MTBF} = \frac{365 \times 24}{1\%} = 876 \text{ tsd. Stunden}$$

Wenn die MTBF-Zahlen korrekt verwendet werden, können sie helfen, die zu erwartenden Wartungs- und Reparaturkosten unter Einsatzbedingungen vorherzusagen. Die in Tausenden oder Millionen Stunden angegebenen MTBF-Werte rufen Verwirrung bei all denen hervor, die nicht damit vertraut sind. Nehmen wir das erste oben aufgeführte Beispiel; die Wandler haben einen MTBF-Wert von einer Million Stunden (äquivalent zu 114 Jahren), aber ein einzelner Wandler fiel schon nach 13 Monaten Nutzung aus. Vielleicht kann ein geläufigeres Beispiel helfen, dieses scheinbare „Verrechnen“ zu erklären; und zwar das menschliche Leben. Die durchschnittliche „Ausfallrate“ eines 25-jährigen Menschen beträgt 0,1 %; das heißt, wir können erwarten, dass ein 25-Jähriger aus einer Menge von 1000 25-Jährigen sterben wird. Die Berechnung ergibt einen menschlichen MTBF-Wert von 800 Jahren! Der Grund, warum die MTBF-Zahlen so hoch (und so variabel) sind, liegt darin, dass die Ausfallrate im flachen mittleren Abschnitt der Lebensdauerkurve (der Nutzungszeit) sehr, sehr niedrig ist.

Multipliziert über eine lange Zeit, bedeutet dies, dass die winzigen Änderungen im Ausfallratendelta (Änderungsrate der Ausfallrate im Laufe der Zeit) große Änderungen in der berechneten MTBF hervorrufen. Das erklärt auch, warum wir alle nicht 800 Jahren leben. Mit 25 sind die meisten Menschen in ihrem gesündesten Alter, und die primäre Todesursache sind Unfälle. Wenn wir nicht altern oder an Krankheiten leiden würden, könnten wir alle 800 Jahre leben, wenn der einzige Grund des Todes der Zufall wäre. Würde andererseits ein anderes Alter gewählt, z. B. 45 Jahre, entstünde eine ganz andere menschliche MTBF-Zahl, weil wir Menschen in einem relativ frühen Alter anfangen zu „verschleissen“.

Da die Ausfallrate während der Nutzungsendphase einem exponentiellen Gesetz folgt, kann die Betriebszuverlässigkeit aus MTBF unter Verwendung der folgenden Formel berechnet werden:

$$Reliability = e^{-\frac{MTBF}{T}}$$

Gleichung 7.3: Verhältnis der Betriebszuverlässigkeit gegenüber MTBF

Ist die Zeit (T) gleich MTBF, reduziert sich die Gleichung zu e^{-1} oder 37%. Dies kann so gedeutet werden, dass bei $T = MTBF$ nur 37% der Wandler immer noch arbeiten werden, oder alternativ, dass die statistische Wahrscheinlichkeit, dass bei $T = MTBF$ immer noch alle Wandler arbeiten werden, nur 37% beträgt.

7.4 Demonstrierte MTBF

Die meisten Stromversorgungshersteller können nicht jahrelang warten, um die tatsächliche Ausfallrate und Delta-Ausfallrate ihrer Produkte zu messen, geschweige denn über 50 Jahre, um ausreichend Versuchsdaten für hochzuverlässige Produkte zu sammeln. Der praktischste Weg, Zuverlässigkeitskenngrößen zu ermitteln, besteht darin, die empirischen Ergebnisse für die einzelnen Komponenten laut solchen Datenbanken wie MIL HDBK 217 zu verwenden und anzunehmen, dass die Ausfallraten konstant bleiben. Das Ergebnis ist nicht ideal, aber es ist sehr viel besser als eine einfache Annahme oder jahrzehntelanges Warten auf aussagekräftige, zuverlässige Daten. Wurden jedoch ausreichend viele Produkte auf den Markt gebracht oder in langfristigen Testreihen geprüft, ist ein genaueres Maß für die Betriebszuverlässigkeit verfügbar, die sogenannte demonstrierte MTBF, die auf tatsächlich registrierten Ausfällen basiert. Da der Stichprobenumfang und die Beobachtungszeit aus praktischen Gründen begrenzt sind, kann die Anzahl der tatsächlichen Ausfälle niedrig sein. Daher sind statistische Analysemittel wie die X²- (Chi hoch zwei) Verteilung notwendig, um die demonstrierte MTBF-Zahl mit einem vernünftigen Maß an statistischer Sicherheit (z. B. 95%) berechnen zu können:

$$Demonstrated\ MTBF = \frac{2\ T}{\chi^2(0.05, v)}$$

wobei T = Zeit (in Stunden), 0,05 ist gleich der 95%-Konfidenzgrenze und v ist der Freiheitsgrad der Chi-Quadrat-Funktion.

Gleichung 7.4: Berechnung von demonstriertem MTBF

Es gibt auch noch andere Datenbanken und statistische Modelle, die zur Ermittlung von Ausfallraten verwendet werden können. Die meist verbreiteten neben MIL HDBK 217F sind Bellcore/Telcordia TR-NWT-332 und IEC61709. Die Ergebnisse variieren stark zwischen den Methodologien, weil verschiedene Annahmen gemacht und verschiedene Arbeitsbelastungen in den Berechnungen (z. B. benutzt MIL HDBK 100% Last, Bellcore/Telcordia nur 50% Last) eingesetzt werden. Für denselben DC/DC-Wandler mit 30 W gibt MIL HDBK 217F, Notice 2, eine MTBF von 435.000 Stunden, Bellcore/Telcordia TR-332 eine MTBF von mehr als 3 Millionen Stunden und IEC61709 eine MTBF von ca. 80 Millionen Stunden. Unabhängig davon, welche Methodologie verwendet wird – wenn zwei Produkte ähnliche Performancekriterien, aber verschiedene MTBF-Werte aufweisen, wird das Produkt mit der höheren MTBF-Zahl auch in der Praxis am zuverlässigsten sein – solange dasselbe Modell und dieselben Belastungsfaktoren in den MTBF-Berechnungen verwendet wurden.

7.5 MTBF und Temperatur

Die Betriebszuverlässigkeit sinkt mit zunehmender Betriebstemperatur, deshalb ist der MTBF-Wert im Datenblatt in der Regel nur bei Raumtemperatur gültig und sollte entsprechend ausgewiesen werden. Der Grund, warum die Betriebszuverlässigkeit so temperaturabhängig ist, basiert auf der Aktivierungsenergie für chemische Prozesse. Im Jahre 1898 hat der schwedische Chemiker Arrhenius einen Beweis veröffentlicht und gezeigt, dass die chemische Reaktionsgeschwindigkeit temperaturabhängig ist und sich die Reaktionsgeschwindigkeit bei jedem Temperaturanstieg um 10°K ungefähr verdoppelt.

$$k = A \exp\left(\frac{-E_A}{k_B T}\right)$$

wobei k = Reaktionsgeschwindigkeit, A = Vorfaktor, E_A = Aktivierungsenergie, k_B = Boltzmannsche Konstante und T = Temperatur in °K.

Gleichung 7.5: Arrheniussche Gleichung

Die Arrheniussche Gleichung hat viele Anwendungen außerhalb der reinen Chemie gefunden und kann ebenso zur Lebensdauerbeurteilung elektronischer Bauteile, wo viele der Alterungseffekte chemischer Natur sind (z. B. Korrosionsbelastungen, Abbau von Materialien, Dislokationen in Halbleiterkristallgittern usw.) angewandt werden. Die Gleichung kann auch umgeordnet werden, um einen Beschleunigungsfaktor zu ergeben, der temperaturabhängig ist. Für elektronische Bauteile beträgt die Aktivierungsenergie 0,6 Elektronenvolt, was den folgenden Beschleunigungsfaktor ergibt:

$$\text{Beschleunigungsfaktor} = \exp\left(\frac{0.6\text{eV}}{k_B} \left(\frac{1}{T_{REF}} - \frac{1}{T_{AMB}}\right)\right)$$

Gleichung 7.6: Beschleunigungsfaktor von Arrhenius

Für die meisten Datenblattspezifikationen wird die Bezugstemperatur als Nennraumtemperatur oder 25°C angenommen. Das ergibt folgende Beschleunigungsfaktoren:

T(amb)	Beschleunigungsfaktor
25°C	1
30°C	1.5
40°C	3
50°C	6
60°C	12
70°C	22
80°C	40

Tabelle 7.4: Beschleunigungsfaktoren für verschiedene Umgebungstemperaturen

Aus diesem einfachen Verhältnis ist ersichtlich, dass eine Verdoppelung der Umgebungstemperatur von 25°C auf 50°C den Alterungseffekt um den Faktor 6 erhöht. Würde die Temperatur um weitere 25°C auf 75°C erhöht, würde der Alterungseffekt um ca. Faktor 30 steigen.

Dasselbe Verhältnis gilt auch umgekehrt. Reduzieren der Temperatur erhöht die Betriebszuverlässigkeit elektronischer Bauteile. Bei sehr niedrigen Temperaturen (unter -20°C) können jedoch andere Faktoren, wie mechanische Beanspruchungen infolge der unterschiedlicher Kontraktionskoeffizienten der verschiedenen Materialien oder Lötverbindungen, die dadurch z.B. brüchig werden können, eine wiederum steigende Ausfallrate hervorrufen. Daher kann die Arrheniussche Gleichung nicht ohne Grenze extrapoliert werden.

Für die MTBF-Berechnung gibt es neben diversen Alterungseffekten auch andere Belastungsfaktoren, aber die Berechnungen zeigen einen klaren Zusammenhang der Zuverlässigkeitsminderung bei Temperaturanstieg:

Umgebungstemperatur	MTBF (MIL-HDBK-217F) (Vollast)
25°C	1.368.813 hours
50°C	711.033 hours
85°C	226.072 hours

Tabelle 7.5: Beispiel der MTBF-Änderung bei unterschiedlicher Temperaturbeanspruchung (für einen RECOM 2W DC/DC-Wandler)

7.6 Auf Betriebszuverlässigkeit ausgerichtete Konstruktionsweise

Die Betriebszuverlässigkeit kann in Stromversorgungen mitprojektiert werden, indem bei der Auswahl der entsprechenden Bauteilkennwerte Topologien und Spezifikationen mit langer Lebensdauer Berücksichtigung finden.

Die Hauptkriterien in der Entwurfsphase sind die Auswahl der Bauteilspezifikation, Verwendung von bewährten und getesteten Schaltkreis-Layouts und Berücksichtigung von zu erwartenden elektrischen, thermischen sowie sonstigen einwirkenden Umweltbelastungen. Darüber hinaus besteht eine der wichtigsten Methoden einer auf die Betriebszuverlässigkeit gerichteten Konstruktionsweise darin, zu gewährleisten, dass sich die Stromversorgung in der Produktion einfach überprüfen lässt. Das bedeutet, den physischen Zugriff in die Konstruktion zur Prüfung von Signalformen, Spannungen sowie Temperaturen zu erlauben, um sicherzustellen, dass die Stromversorgung innerhalb der Normalbetriebsgrenzen arbeitet. Jegliche Werte, die nahe an den Toleranzgrenzen liegen, oder Signalformen, die „nicht richtig aussehen“, könnten darauf hindeuten, dass selbst, wenn der Wandler die Datenblattspezifikationen erfüllt, die Auswirkung dieser Abweichungen auf die Lebensdauer nicht optimal sein könnte.

Wie im vorhergehenden Abschnitt erwähnt, sind hohe Temperaturen der Feind einer langen Lebensdauer und hohen Betriebszuverlässigkeit. Jede Komponente hat einen durch den Hersteller bestimmten maximalen Betriebstemperaturbereich, aber eine gut strukturierte auf die Betriebszuverlässigkeit gerichtete Konstruktionsweise wird keine Komponente bis ans Limit der jeweiligen Betriebstemperaturgrenze beanspruchen. Die nächste Tabelle gibt typische maximale Betriebstemperaturen, sowie empfohlene maximale Komponenten-De-rating-Temperaturen für einige allgemeine DC/DC-Wandlerkomponenten an:

Komponente	Max. Einsatztemperatur (Datenblätter des Herstellers)	Empfohlene max. Berechnungs- temperatur (ungünstigster Fall)
SMD-Widerstand	125°C	115°C
SMD-Kondensator	125°C	115°C
SMD Diode	125°C	115°C
FET (Junktionstemperature)	155°C	140°C
Transformator	130°C	120°C
Optokoppler	110°C	100°C
PCB (FR4)	140°C	130°C

Tabelle 7.6: Maximale Berechnungs-Betriebstemperaturen

7.7 Empfehlungen zum PCB-Layoutentwurf zum Erzielen einer hohen Zuverlässigkeit

Wenn eine Komponente die berechnete Temperatur überschreitet sollte, kann entweder zusätzliches Kupfer auf der PCB hinzugefügt werden, um die Wärme abzuleiten, oder mehrere einzelne Komponenten parallel geschaltet werden, um den Strom aufzuteilen. Auch DC/DC-Wandlerpins können eingesetzt werden, um dazu beizutragen, einen Teil der Wärme vom Wandler zur Haupt-PCB abzuleiten. Bei DC/DC-Wandlern mit Metallgehäuse hilft die Platzierung von heißen Bauteilen im Wandlerinneren so nah wie möglich am Gehäuse oder das Hinzufügen von wärmeübertragenden Kupferblöcken zur gezielten Lenkung der Wärmeübertragung.

Wärmeübertragung von einem Bauteil zum anderen kann vermieden werden, indem die beiden Komponenten nicht in unmittelbarer Nähe zueinander platziert werden.

Ein klassischer Fehler besteht darin, den Optokoppler fast in Berührung mit dem Transformator zu platzieren, um Platz zu sparen. Während der Transformator inneren Temperaturen um 120°C leicht standhält, hat der Optokoppler eine Betriebstemperaturgrenze von 100°C. Also obwohl fast keine Wärme innerhalb des Optokopplers entsteht, kann er durch die angrenzende heiße Komponente überhitzt und so zum begrenzenden Faktor für die Ermittlung der maximalen Umgebungstemperatur des Wandlers werden. Ähnliche Bedingungen können bei an Dioden angrenzenden Komponenten entstehen. Dioden erzeugen normalerweise aufgrund des Vorwärtsspannungsabfalls und des Durchflusstromes eine große Menge an Abwärme. Wird ein passives Element zwischen zwei Dioden (z. B. der Ausgangskondensator zwischen zwei Ausgangsgleichrichterdioden) platziert, wird der Kondensator, der von zwei Seiten erwärmt wird, dann womöglich außerhalb seiner Temperaturgrenzen bleiben, obwohl die Dioden durchaus innerhalb ihrer Temperaturgrenzen arbeiten.

Chipwiderstände gehören zu den zuverlässigsten in DC/DC-Wandlern verwendeten Komponenten. Dennoch muss auch das PCB-Layout der Anwendung einer genauen Betrachtung unterzogen werden, da der Keramikträger spröde ist:

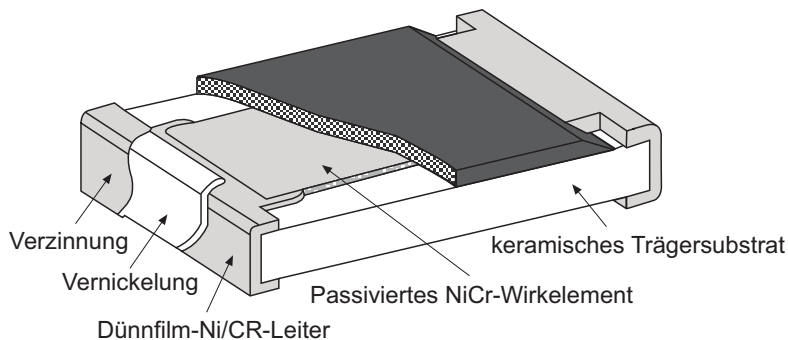


Abb. 7.3: Chipwiderstandkonstruktion

Ein typisches PCB-Panel kann viele gleiche Teilplatten beinhalten, die in einem Trennprozess in einzelne Platinen geschnitten werden müssen. Dieser mechanische Trennprozess (der unterschiedlich genannt wird: de-paneling, dicing oder singulation) kann mit Schneidplättchen (V-Schnitt), Schneidstempeln, Fräs- oder Laserschneidmaschinen ausgeführt werden.

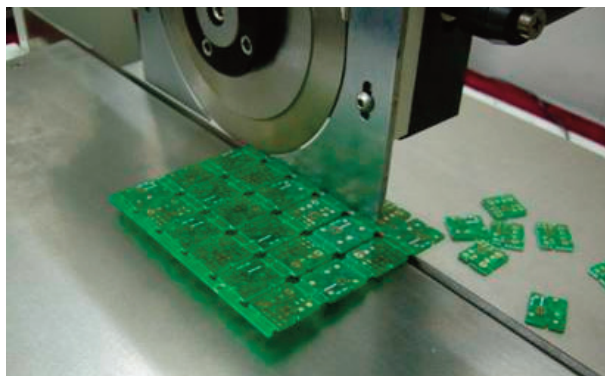


Abb. 7.4: Zum Teilen von PCB-Platinen verwendete V-Schnitt-Maschine

Die alle Einzelplatine zur besseren Verarbeitbarkeit tragende Mutterplatine kann auch mit einer Stanze oder einem Bohrer perforiert werden, sodass die Boards später zur Endverarbeitung der einzelnen Produkte manuell abgetrennt werden können. Jede Entfernung von überschüssigem PCB-Material, Schneiden oder Vereinzeln (dicing) beansprucht die PCB mechanisch und kann Durchbiegungen und mechanische Spannungen auslösen. Wenn Chipwiderstände und andere keramische Komponenten zu nahe an den Kanten platziert sind und so Durchbiegung, Verwölbung oder Krümmung ausgesetzt werden, können sie reißen und ausfallen.

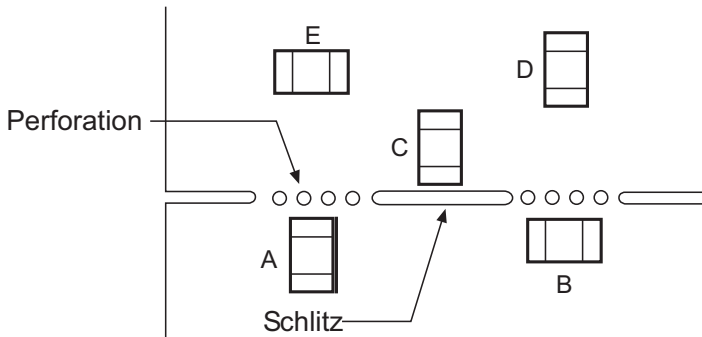


Abb. 7.5: Positionierung des Chipwiderstands

In Abb. 7.5 wird Bauteil A mechanisch stärker beansprucht als Bauteil B, wenn die PCB vereinzelt wird. Auf ähnliche Weise wird Bauteil C durch die Slot-Stanze oder Fräswerkzeug mehr beansprucht als Bauteil D. Bauteil E befindet sich in einer idealen Lage: parallel zum Schnitt und um mehr als die eigene Länge von der Kante entfernt. Die Perforationen sollten durch Durchbiegung der PCB in der Richtung der Komponenten gebrochen werden, sodass die maximale mechanische Beanspruchung auf der anderen Seite der PCB auftritt (Abb. 7.6).

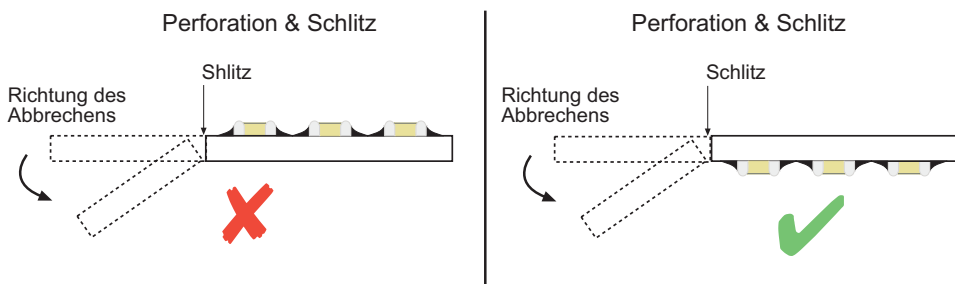


Abb. 7.6 Korrektes und unkorrektes Abbrechen entlang der Perforation

Die Beanspruchungen der PCB, und folglich der eingelöteten Bauteile, kann durch Optimierung des Layouts der einzelnen PCBs reduziert werden. Dies ist ein wichtiger Punkt in der Planung des Fertigungsprozesses und hat einen maßgeblichen Einfluss auf die Betriebszuverlässigkeit des Endproduktes. Das folgende Beispiel hilft, dies zu erklären.

Abb. 7.7 zeigt das originale Layout einzelner PCBA- (printed circuit board assemblies) Layouts auf einer Platte. Die schraffierten Bereiche bezeichnen die Bereiche, an denen die einzelnen Slots auseinandergeschnitten werden, die aus konstruktionsbedingten Gründen nah an den PCB-Komponenten platziert sind. Die ursprüngliche Idee bestand darin, die PCBAs um 180° versetzt abzuwechseln, um so eine gleichmäßigere Gewichtsverteilung zu erzielen (jede Teilplatine trägt eine große und schwere Induktivität). Das Gesamtgewicht der Bauteile wurde zu groß für die Mutterplatine, Ihre Durchbiegung im Löt- und weiteren Verarbeitungsprozess war zu hoch, ebenso die Gesamt-Ausfallrate in weiterer Folge.

Abbildung 7.8 zeigt das überarbeitete Layout der einzelnen PCBA-Layouts. Die Anzahl der PCBAs je Mutterplatine wurde von 50 auf 40 reduziert, und 2 mm breite Abstände zwischen den PCBAs wurden zur Erhöhung der Steifigkeit vorgesehen. Der breitere Abstand erlaubt es dem Bediener der V-Schnitt-Maschine, den Vereinzelungsvorgang besser einzurichten und so sicherzustellen, dass die Beanspruchungen der Bauteile auf den Einzelplatten gering gehalten werden. Als Resultat dieser Maßnahme eines neuen Layouts sank die Gesamtausfallrate deutlich.

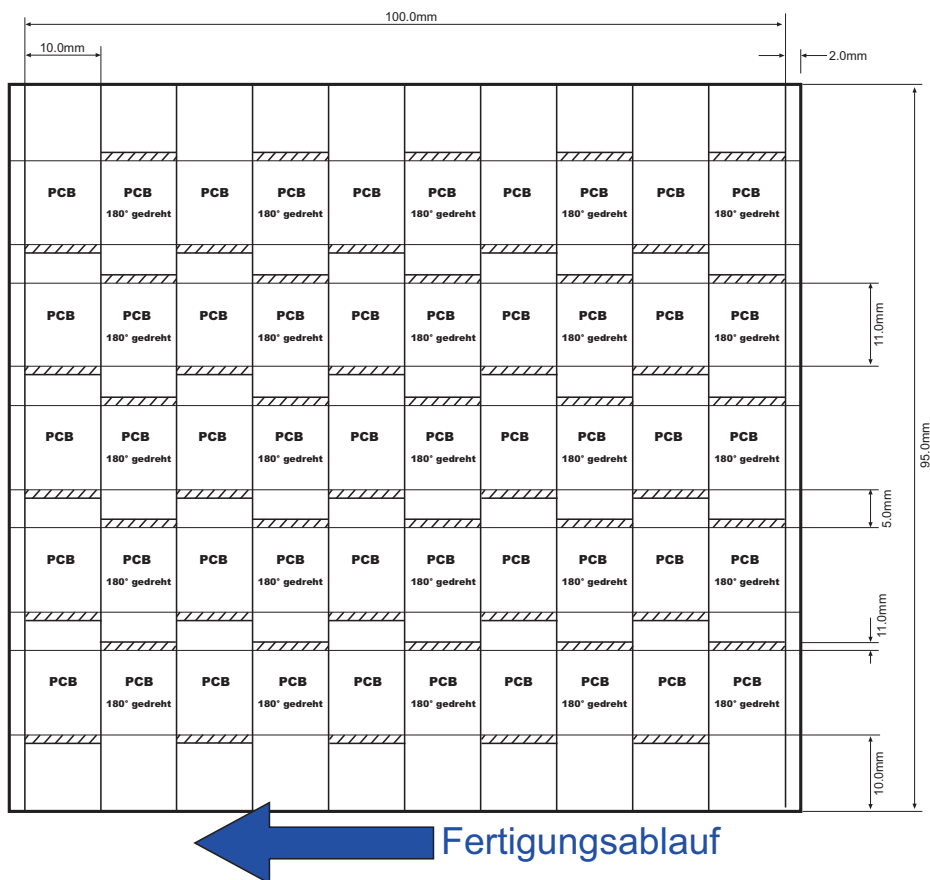


Abb. 7.7: Ursprüngliches PCBA-Layout mit unannehmbare Ausfallrate

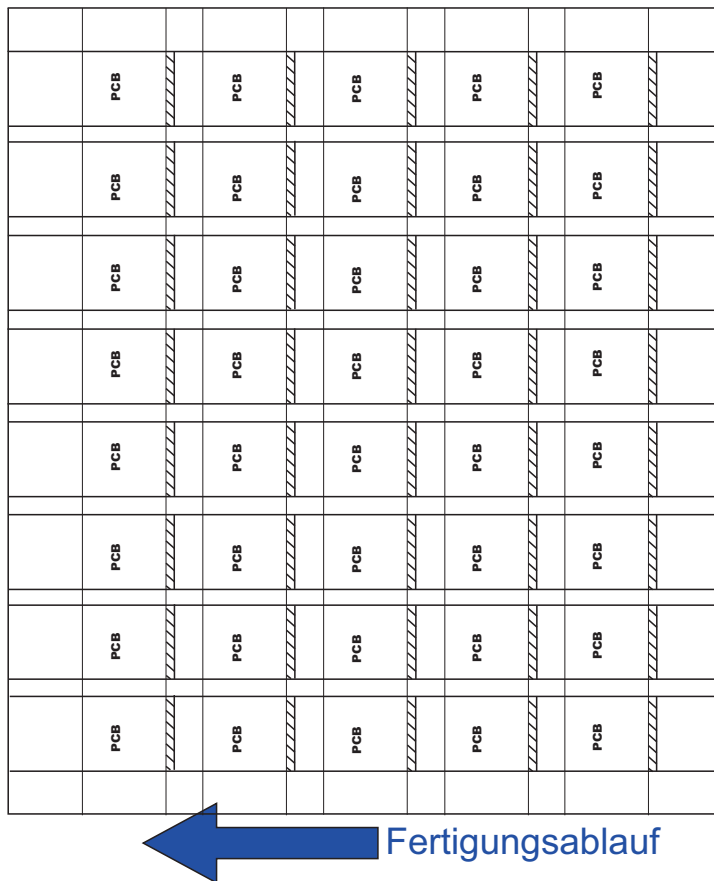


Abb. 7.8: Überarbeitetes PCBA-Layout mit zusätzlichem Abstand zur Erhöhung der Steifigkeit und Optimierung für Zuschnitt via V-Schnitt-Maschine.

7.8 Kondensatorzuverlässigkeit

DC/DC-Wandler enthalten normalerweise Mehrschicht-Keramikkondensatoren (MLCC), Tantal- oder Elektrolytkondensatoren. Man findet auch SMD-Polyesterfolien-Kondensatoren, die aber meist unerwünscht groß sind und daher nicht in die von der Industrielektronik geforderten, kompakten Gehäusegrößen passen.

7.8.1 MLCC

Mehrschicht-Keramikkondensatoren (MLCC) sind der am häufigsten verwendete Kondensatortyp, der in DC/DC-Stromversorgungen eingesetzt wird. Sie bieten eine hohe volumetrische Kapazität, sind nicht gepolt, verfügen über einen sehr niedrigen ESR- und ESI-Wert und bieten eine stabile Nennkapazität über einen breiten Frequenz- und Temperaturbereich, wodurch sie für Einsätze sowohl in Filtern wie auch als Ladekondensatoren bestens geeignet sind.

MLCCs können jedoch leicht ausfallen, wenn sie oberhalb ihrer Maximalgrenzspannungen betrieben werden. Diese maximal zulässige Spannung ist begründet in ihrem Aufbau in vielen metallischen Schichten, abgetrennt durch einen dünnen keramischen Isolator. Tritt ein Funkenüberschlag zwischen den Schichten auf, gibt es im Gegensatz zu Elektrolytkondensatoren keinen Selbstheilungsmechanismus, um den Schaden zu reparieren, und der MLCC wird früher oder später ausfallen.

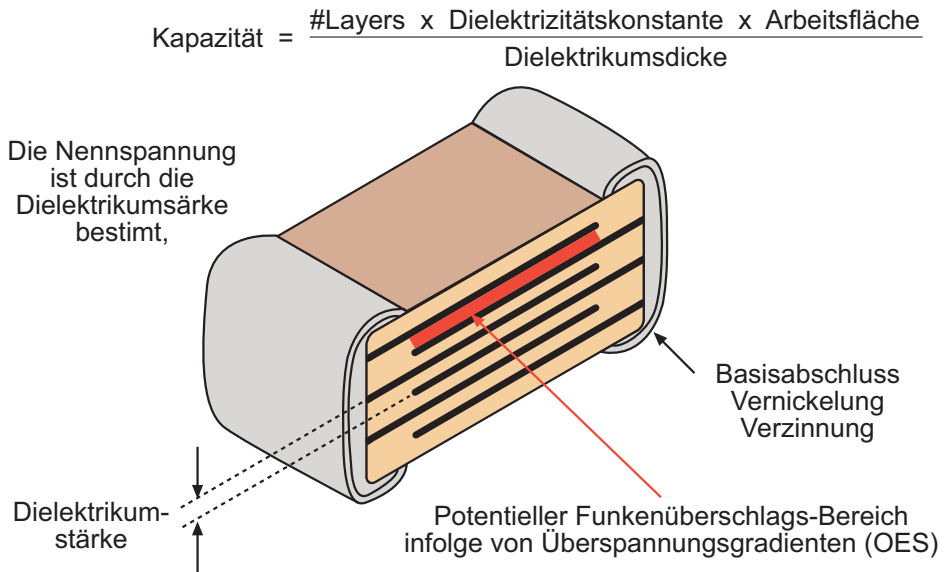
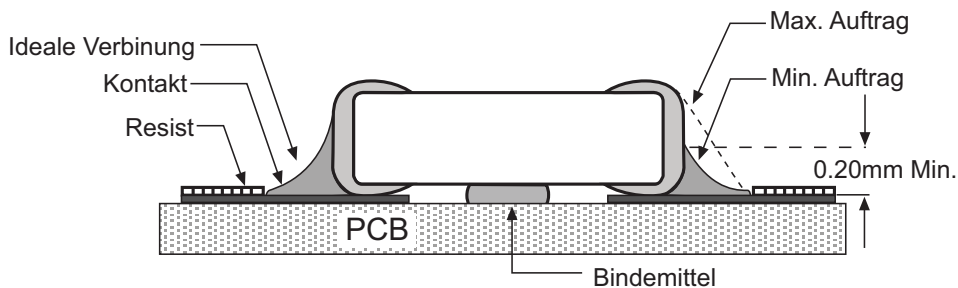


Abb. 7.9: MLCC Funkenüberschlag

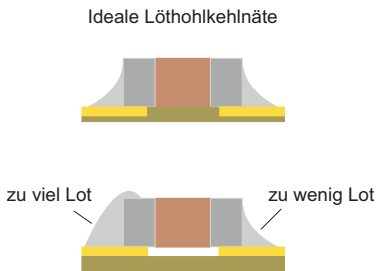
Die Entwurfskriterien müssen deshalb gewährleisten, dass der MLCC nicht einmal einer kurzzeitigen elektrischen Überspannung (EOS) unterzogen wird. Da die meisten Hersteller Kondensatoren während der Herstellung mit 100% ihrer Nennspannung testen, können sie bis zu ihrer Nennspannung ohne jegliche Verminderung der Nennspannung verwendet werden. Da EOS jedoch dem kubischen Gesetz folgen wenn der Nennwert überschritten wird, würde eine zweckmäßige Auslegungsgrenze bei DC- oder Niederfrequenz-AC-Spannungen 90% der Nennspannung des Herstellers betragen. Es muss darauf geachtet werden, ob das Signal am Kondensator nah an der Resonanzfrequenz liegt, weil dies eine lokale Überhitzung und einen Frühausfall verursachen könnte. Der interne Temperaturanstieg infolge von Eigenerwärmung darf höchstens 20°C betragen.

Ein anderer Resonanzeffekt, der bei MLCCs auftreten kann, sind durch einen Piezoeffekt an den keramischen Schichten hervorgerufene, mitunter sogar akustisch vernehmbare Störgeräusche. Wenn die akustische Resonanzfrequenz die AC-Wellenform am Kondensator anpasst, kann der Kondensator zu „singen“ oder „winseln“ beginnen. Die Lösung besteht darin, ein anderes MLCC-Format mit einer höheren oder niedrigeren akustischen Resonanz zu verwenden. Laut den Herstellern beeinträchtigen Sing-Kondensatoren ihre Betriebszuverlässigkeit nicht, aber als Anwender mag man mitunter laute Stromversorgungen nicht besonders gerne.

Die Menge an verwendeter Lötpaste und das Layout der Pads (Kontakte) sind ebenfalls wichtig, um die Belastung auf den Kondensator infolge einer ungleichmäßigen Menge von Lötzinn an den Anschlüssen zu reduzieren. Eine gut ausgeführte Fließlötverbindung verfügt über eine Löthohlkehlnaht, die mit dem Auftragsmeniskus konkav ist und sollte bis zu 75% der Kondensatorhöhe ansteigen.

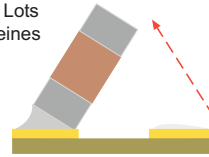


Lötmetallvolumen und Kehlhaftform



Manhattan-Effekt (auch Grabsteineffekt genannt)

Unterschied in der Oberflächenspannung des flüssigen Lots verursacht den Anstieg eines Endes des Chips.



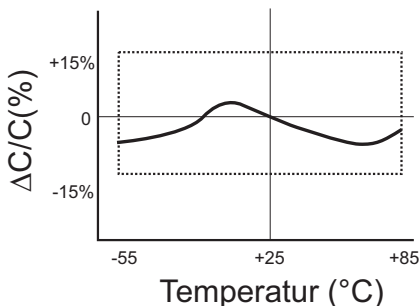
Während des Erhitzens im Reflow-Ofen steigt das Chipbauelement an und steht aufrecht. Dies kann durch zu viel Lötpaste, ungleichmäßige Erwärmung, Veränderung der Montageposition usw. hervorgerufen werden.

Abb.7.11: Form der MLCC-Löthohlkehlnaht

Handlötten ist besonders bei MLCC-Kondensatoren gefährlich. Nicht nur die Kondensatoroberfläche kann durch die Lötkolbenspitze leicht angekratzt werden, auch die asymmetrische Erhitzung des Kondensators, da jedes Ende des Abschlusses separat gelötet wird, verursacht beträchtliche mechanische Beanspruchungen innerhalb der Struktur, was zur Rissbildung führen kann. Wenn ein MLCC-Kondensator zufällig fallen gelassen wird, sollte er aussortiert und nicht weiter verwendet werden.

Obwohl Keramikkondensatoren eine sehr stabile Kapazität vs. Temperaturverlauf zeigen können, hängt die Stabilität stark von der vorgegebenen Toleranz ab. Die preiswerteren Typen können je nach Temperatur stark variieren; +22% bis -82% über den ganzen Betriebstemperaturbereich. Eine für einen korrekten Betrieb bei 25°C entwickelte Schaltung, kann sich deshalb bei niedrigen oder hohen Umgebungstemperaturen ganz anders verhalten und unzuverlässig werden.

Typischer Temperaturverlauf (ex.X5R)



Typischer Temperaturverlauf (ex.Y5V)

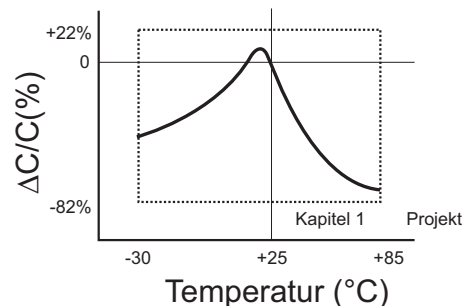


Abb. 7.12: Kapazitätsänderung vs. Temperatur für verschiedene MLCC-Klassen

Buchstabencode niedrige Temperatur	Zifferncode max. Temperatur	Buchstabencodeänderung der Kapazitäten über dem Temperaturbereich
X = -55°C (-67°F)	4 = +65°C	P = ±10%
Y = -30°C (-22°F)	5 = +85°C	R = ±15%
Z = +10°C (+50°F)	6 = +105°C	S = ±20%
	7 = +125°C	T = +22/-33%
	8 = +150°C	U = +22/-56%
	9 = +200°C	V = +22/-82%

Tabelle 7.7: MLCC-Klassifizierungscodes (nach EIA RS-198)

Ein X7R-Kondensator arbeitet beispielsweise von -55°C bis +125°C mit einer Kapazitätsvarianz von ±15%. Ein letzter Punkt zur Nutzung von MLCC ist, dass dessen Kapazität auch DC-spannungsabhängig ist. Um eine hohe volumetrische Effizienz zu erreichen, ist das Isolierdielektrikum zwischen den Schichten sehr dünn (nur wenige Mikrometer stark) und weshalb Belastung durch zu hohe elektrische Felder einen Rückgang der Speicherwirkung zeigt. Ein auf 10µF bei 6,3V ausgelegter MLCC 0806 zeigt mit angelegter Spannung einen wesentlichen Kapazitätsrückgang, der in der Konstruktionsberechnung berücksichtigt werden muss:

Kapazitätswert	angelegte Spannung
10µF	0V
8.8µF	2V
7.2µF	4V
5.7µF	6V

Tabelle 7.8: Änderung der gemessenen Kapazität bei unterschiedlichen angelegten Spannungen für einen 10µF-MLCC spezifiziert für 6,3VDC.

7.8.2 Tantal- und Elektrolytkondensatoren

Elektrolytkondensatoren werden selten in DC/DC-Wandlermodulen verwendet, sind aber häufig als externe Komponenten zu finden. Entweder als Teil eines EMV-Filters oder um die Eingangs- oder Ausgangsspannung zu stabilisieren oder um kurzzeitigen Spitzenstrombedarf zu puffern. Ein Beispiel der letzteren Kategorie ist in IGBT-Treiberschaltungen zu finden. Die durch die IGBT-Treiber gezogene Durchschnittsleistung beträgt in der Regel nur ca. 2W, aber die Gate-Ansteuerungsstromspitzen können mehrere Ampere betragen. Ein großer Elektrolytkondensator mit niedrigem ESR am Ausgang des DC/DC-Wandlers kann den Treiberspitzenstrom liefern, was ein Kleinleistungs-DC/DC-Wandler nicht kann. Tantal- und Aluminiumelektrolytkondensatoren verwenden eine ähnliche Konstruktion von durch einen Isolator abgetrennten und entweder mit Flüssigkeit oder Gelelektrolyt getränkten Leiterschichten (Tantal oder Aluminium). Die Nutzung des Elektrolyts führt zu einer starken Erhöhung der volumetrischen Kapazität, weshalb ein Elektrolytkondensator bei gleichem Volumen mehr als die doppelte Kapazität als ein MLCC bieten kann. Ein anderer Vorteil des Elektrolyts besteht darin, dass es den Kondensatoren erlaubt, kleine Schäden zwischen den Platten selbst wiederherzustellen (selbstheilend).

Selbstheilung tritt auf, weil jeglicher Funkenüberschlag zwischen den Schichten das Wasser im Elektrolyt elektrolysiert, wobei mittels Hydrolyse Wasserstoff und Sauerstoff erzeugt werden. Der Sauerstoff verbindet sich mit der Anoden-Schicht, verursacht eine Oxid-Schichtbildung und heilt somit das Leck. Der freigesetzte Wasserstoff entweicht entweder aus dem Kondensator oder wird chemisch absorbiert.

Tantal steht in der Liste von Konfliktmineralien, deren Verwendung in den USA durch das Dodd-Frank-Gesetz beschränkt ist. Daher flaut die Verwendung von Tantal-kondensatoren ab, obwohl sie etwaige Vorteile gegenüber Aluminium-Elektrolyt-kondensatoren bieten, wie etwa eine höhere Nennkapazität bei gleicher Größe und eine stabilere Kapazität über Temperatur und Zeit. Ihre Tendenz jedoch in Thermostabilität überzugehen, wenn sie durch transiente Überspannung beschädigt werden, hat Tantalkondensatoren einen schlechten Ruf als unzuverlässig eingebracht. Tritt ein Ausfall ein, kann die durch den Funkenüberschlag erzeugte Wärme das Manganoxid-Kathodenmaterial entzünden und den Kondensator zum Brennen bringen. RECOM verwendet keinerlei Tantalkondensatoren in den Produkten.

Elektrolytkondensatoren (Elkos) haben einen schlechten Ruf in Bezug auf Zuverlässigkeit. Dies ist zum größten Teil unberechtigt. Wenn die Kondensatoren korrekt spezifiziert sind und innerhalb ihrer Spezifikationen eingesetzt werden, können sie tatsächlich sehr zuverlässige Komponenten sein. RECOM hat ein AC/DC-Produkt, in dem Elektrolyt-kondensatoren Verwendung finden, deren Lebenserwartung bei 25°C mehr als 22 Jahre beträgt. Der Druck des Stromversorgungsmarkts, die Produkte zu immer niedrigeren Preise anzubieten, bedeutet jedoch, dass häufig Kompromisse in der Konstruktion, Auswahl der Komponenten und Betriebsbedingungen von Elektrolytkondensatoren gemacht werden – so sehr, dass sie häufig das schwächste Element in der Zuverlässigkeitskette sind und am wahrscheinlichsten ausfallen. Eine Frage, die häufig gestellt wird, ist welche Lebensdauer Elkos haben und wie viele von ihnen ausfallen werden. Faktisch sind dies zwei Fragen und sie erfordern zwei unterschiedliche Antworten. Die Betriebslebensdauer von Kondensatoren ist normalerweise die Zeit, die es dauert, bis sich die ESR-Zahl verdoppelt, oder 10% der Kondensatoren durch Kurzschluss oder offenen Stromkreis ausgefallen sind. Einer der primären Alterungseffekte der Elkos vernichtet ist, dass der Elektrolyt allmählich austrocknet. Dies verursacht mit der Zeit einen allmählichen Anstieg des ESR-Wertes, bis die Performance als unannehmbar gilt:

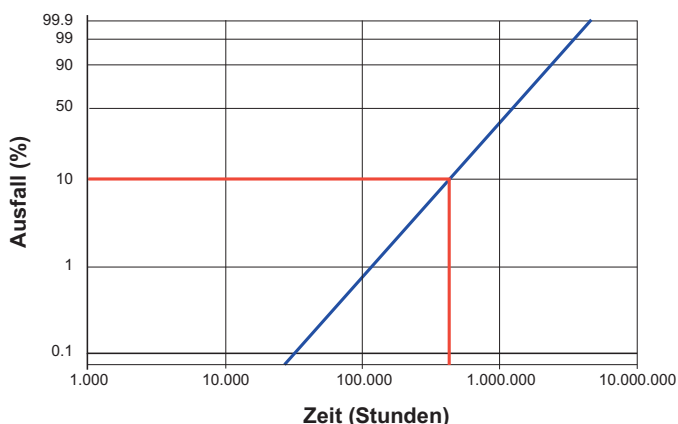


Abb. 7.13: Elko-Ausfallrate über längere Zeit infolge von zunehmenden ESR

Da die Lebensdauer von Elkos sehr lang sein kann, werden sie bei maximaler Belastungsgrenze getestet (HASS oder Highly Accelerated Stress Screening), um Alterungsprozesse zu beschleunigen und dies wird in den Datenblättern entsprechend spezifiziert. Dies bedeutet, dass eine typische Lebenserwartung nur einige Tausend Stunden betragen könnte. In der Praxis kann die Lebensdauer durch den Betrieb der Komponente unterhalb ihrer absoluten maximalen Belastungsgrenzen wesentlich länger sein.

Das folgende Gleichungssystem gibt eine vereinfachte Rechenregel, um die tatsächliche Lebensdauer eines Elkos zu berechnen:

$$Life_{working} = Life_{rated} V_{stress} 2^{(T_{rated} - T_{actual})/10}$$

wobei Spannungsgradient:

$$V_{stress} = 4.3 - 3.3 \frac{V_{actual}}{V_{rated}}$$

Gleichung 7.7: Berechnung der Elkolbensdauer (vereinfacht)

Aus diesem Gleichungssystem können drei wichtige Schlussfolgerungen gezogen werden:

- 1: Kapazität spielt keine Rolle. Die Lebensdauer hängt nicht von der Nennkapazität ab; daher ändert die Verwendung von Kondensatoren höherer Kapazität in der Konstruktion die Zuverlässigkeit nicht.
- 2: Die Lebensdauer hängt nicht von der absoluten Spannung ab, nur das Verhältnis von V_{actual} zu V_{rated} . Ein Hochspannungskondensator ist deshalb nicht inhärent mehr oder weniger zuverlässig als ein Niederspannungskondensator.
- 3: Die Temperatur ist wichtig, weil sie dem quadratischen Gesetz folgt.

Wenn ein Kondensator beispielsweise auf 2000 Stunden/85°C ausgelegt und in einer Stromversorgung mit einer DC-Spannung von 70% der Nennspannungsgrenze bei einer Betriebstemperatur von 15°C über der Umgebungstemperatur verwendet wird, beträgt die Erhöhung der Lebensdauer laut den Berechnungen:

$$V_{stress} = 4.3 - 3.3 \frac{V_{actual}}{V_{rated}} = 4.3 - (3.3 \times 0.7) = 2V$$

$$Life_{working} = Life_{rated} V_{stress} 2^{(T_{rated} - T_{actual})/10} = 2000 \times 2 \times 2^{(45)/10} = 90,500 \text{ Stunden}$$

Deshalb erhöht sich die Kondensatorlebensdauer durch die Kombination aus reduzierter Spannung und geringerer Temperaturbeanspruchung auf mehr als 90.500 Stunden (> 10 Jahre).

Dieses Beispiel zeigt auch, dass eine auf die Betriebszuverlässigkeit gerichtete Konstruktionsweise durch die Reduzierung des Spannungsgradienten an den Kondensatoren (als Faustregel sollte die Mittelspannung 70% der nominellen Kondensatorspannung nicht überschreiten), sowie durch das Betreiben von Elkos bei möglichst niedrigen Temperaturen vorgenommen werden kann. Gutes Wärmedesign bei Elkos hat am meisten Einfluss auf die Lebensdauer und folgt zwei Grundregeln. Erstens trägt die innere Erwärmung wesentlich zur Wärmebeanspruchung bei: je höher der Welligkeitsstrom, desto höher ist die innere Verlustleistung, die am eigenen ESR des Kondensators entsteht.

Da der Kondensator altert, nimmt der ESR zu und trägt somit zum Anstieg der Innentemperatur bei und beschleunigt die Reduzierung der Lebensdauer weiter. Da die Zuverlässigkeit nicht vom Kapazitätswert abhängt, weist ein Kondensator mit größerer Kapazität eine niedrigere elektrolytische Erwärmung und eine längere Dauerhaftigkeit auf. Er besitzt auch eine größere Fläche zum Abtransport der Innenwärme. Zweitens sollten Elkos durch das Layout weit entfernt von Kühlkörpern, Transformatoren oder heißen Halbleiterbauteilen positioniert werden, sodass die Umgebungstemperatur möglichst niedrig bleibt. Vorsicht ist auch bei ungeschirmten induktiven Komponenten geboten, die am elektromagnetischen Feld, das Wirbelströme innerhalb der Kondensatorschichten generiert und somit lokale Erwärmung hervorruft, ausstrahlen können. Kondensatoren dürfen nicht mit Induktivitäten in Kontakt kommen. Angetrieben vom Wunsch immerkleinere, immer preiswertere Stromversorgungs-lösungen zu bauen, werden diese einfachen Entwurfsregeln jedoch häufig ignoriert. Daher der schlechte Ruf von Elkos in Bezug auf die Zuverlässigkeit.

7.9 Zuverlässigkeit von Halbleitern

Leistungshalbleiterelemente werden in DC/DC-Wandlern zum Schalten und Regeln von Strömen und Spannungen verwendet und sind deshalb hohen elektrischen und thermischen Beanspruchungen ausgesetzt, was ihre Lebensdauer reduziert. Die Zuverlässigkeit von Halbleitern wird jedoch am stärksten durch die Herstellungsqualität beeinflusst. Die Arbeitscharakteristiken von Halbleiterübergängen sind sehr empfindlich gegenüber jeglichen Verunreinigungen in den verwendeten Materialien und jeglichen Verunreinigungssubstanzteilchen in Dünnschicht-Metallschichten. Das Problem mit Halbleitern besteht darin, dass eine nicht dem Standard entsprechende Herstellung, wie schlechtes Drahtbonds, Einschließung von Staubteilchen auf dem Kristall und unvollständige hermetische Abdichtung des Gehäuses zu den Pins, nicht mit bloßem Auge sichtbar sind und keine direkte Einwirkung auf die Performance haben. Erst später verursachen solche Fehler Frühaussagen. Wie bei vielen Komponenten ist die Suche nach und das Beibehalten eines Qualitätslieferanten in Bezug auf die Gesamtqualität des Endproduktes wichtig.

Elektrische und thermische Beanspruchungen auf den Halbleitern können durch eine entsprechend ausgelegte Konstruktion verringert werden. Alle Halbleiter sind beispielsweise für Schäden durch transiente Überspannungen anfällig. Eine Eintaktschaltungstopologie generiert am Schalt-FET die doppelte Eingangsspannung, während die FETs in einer Gegentakt-Topologie nur die Eingangsspannung schalten müssen. Wenn ein Spannungsstoß am Eingang auftritt, ist es viel wahrscheinlicher, dass die Eintaktschaltungstopologie ihre Nennspannung überschreitet und ausfällt, als das Gegentakt-Äquivalent.

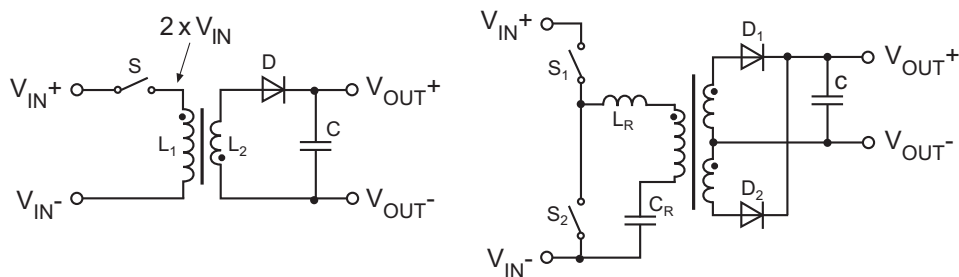


Abb. 7.14: Die Eintaktschaltungstopologie links legt den doppelt so großen Spannungsgradienten an die Schaltelemente an wie die äquivalente Gegentaktschaltungstopologie rechts.

Die heißesten Halbleiterkomponenten in einem beliebigen Leistungswandler sind Schalt-FETs und Gleichrichterioden, da sie beide in der Haupthochstrombahn zwischen Eingang und Ausgang liegen. Passive Elemente tendieren zu relativ homogener Überhitzung, während Halbleiter zu ungleichmäßiger Überhitzung neigen. Die Wärmebeanspruchung beginnt sich an einer lokalen Schwachstelle oder Grenze innerhalb des Geräts zu konzentrieren und ruft schnell weiteren Schaden und schließlich Thermoinstabilität hervor. Generell beinhalten sogar große Halbleitergehäuse nur sehr kleine Einzel-Chip-Schaltkreise. Daher herrscht zu wenig Wärmeträgheit um dazu beizutragen beliebige transiente Überhitzungsereignisse zu absorbieren.

Ein Kühlkörper für einen FET oder eine Diode kann womöglich hilfreich sein um die, innerhalb des Halbleiterbauelements, erzeugte mittlere Wärme, nach außen an die Umgebung abzuführen. Er ist aber gegen sich plötzlich entwickelnde lokale Überhitzungen oder andere Wärmeunregelmäßigkeiten infolge des Wärmewiderstands zwischen Sperrschicht und Gehäuse nicht wirksam. Um jegliche Einwirkung transients Überhitzung zu berücksichtigen, besteht die sicherste Methode bei einer auf die Betriebszuverlässigkeit gerichteten Konstruktionsweise darin, den Nennstrom der Halbleiter entsprechend sicher zu wählen.

Komponente	Parameter	Derating-Faktor
Diskrete Halbleiterbauelemente (Dioden, FETs usw.)	Nennspitzleistung	70% max.
	Nennspitzstrom	50% max.
	durchschnittlicher Nennstrom	50% max.
lineare Spannungsregler	Nennstrom	50% max.
Sperrwandler	Nennstrom	80% max.
Signalkreisioden	Nennstrom	85% max.

Tabelle 7.9: Vorgeschlagener Halbleiter-Derating-Faktor

7.10 ESD

Schaden durch elektrostatische Entladung (ESD) kann auftreten, wenn statische Elektrizität durch elektronische Bauteile zur Masse entladen wird. Die verbreitetste Quelle ist Reibungs- oder Triboelektrizität – Differenzen der statischen Aufladung, die auftreten, wenn zwei verschiedene Isolierstoffe aneinander gerieben werden. Wenn Bediener elektronische Bauteile oder PCBs ohne ESD-Schutzmittel manipulieren, kann die statische Elektrizität durch Bewegungen von Kleidung, das Aufstehen vom Stuhl oder sogar durch das Abziehen von Plastikverpackungen Spannungen in Höhe von Zehntausenden Volt erzeugen. Wenn dieser Bediener dann eine geerdete PCB berührt oder die PCB einem geerdeten Kollegen übergibt, können buchstäblich Funken fliegen, die irreparablen Schaden an Halbleitern oder anderen ESD-empfindlichen Komponenten verursachen. Außer der Sicherstellung, dass alle Bediener, Geräte, Stühle, Fußböden und Labortische geerdet sind, um einen ESD-geschützten Bereich zu bilden, ist es manchmal auch nützlich Luftbefeuchter zu verwenden, um die statische Aufladung zu verringern:

Quelle	niedrige Feuchtigkeit	hohe Feuchtigkeit
Laufen auf Teppich	35.000V	1.500V
Laufen auf Vinylboden	12.000V	250V
Bediener am Labortisch	6.000V	100V
Entfernen einer Plastikverpackung	20.000V	1.200V
Aufstehen vom Stuhl	18.000V	1.500V

Tabelle 7.10: Beispiele statischer Elektrizitäts-Spannungen bei niedriger und hoher Feuchtigkeit

Der Grund, warum die Halbleiter so empfänglich gegenüber ESD-Schäden sind, ist dessen Dünnschicht-Konstruktion, bei der Hochspannungen einen Defekt in der Metalloxydisolierung und lokales Schmelzen verursachen können. Dasselbe trifft auf MLCCs zu, infolge der nur wenige Mikrometer dünnen dielektrischen Trennung zwischen den Schichten, die durch eine transiente Überspannung leicht gefährdet wird. Ein gängiger Irrtum besteht darin anzunehmen, dass Komponenten, sobald sie auf eine PCB gelötet wurden, gegen ESD-Schäden unempfindlich oder durch die Ein- und Ausgangsfilterungskomponenten geschützt sind.

Obwohl es sein kann, dass der ESD-Strompfad in einer Baugruppe direkt zur Masse und nicht durch die elektrostatisch gefährdeten Bauelemente fließt, kann das induzierte elektrostatische Feld immer noch stark genug sein, um gewissen Schaden anzurichten.

Ein Transistor oder eine Diode, die durch ein ESD-Ereignis beschädigt wurde, fällt möglicherweise nicht sofort aus. Elektronenmikroskopbilder könnten lokales Schmelzen und kleine Löcher, die durch Schichten ausgestanzt werden zeigen. Die Komponente arbeitet zwar mit erhöhten Leckströmen, aber ansonsten normal weiter. Diese Art von verstecktem Schaden ist jedoch eine Zeitbombe und es ist ungewiss wann diese hochgehen kann. Zu einem gewissen Zeitpunkt tritt ein elektrischer Fehler auf und die Komponente fällt plötzlich aus. Ein ESD-Schaden ist die häufigste Ursache von „unerklärten“ Frühausfällen.

Die Empfindlichkeit von Komponenten und Untergruppen gegen ESD-Schäden kann geprüft werden. Das allgemeinste Verfahren ist das Modell des menschlichen Körpers (HBM = Human Body Model), das die Energie, die durch menschliche Bewegungen durch die Aufladung eines Kondensators mit 100pF auf Hochspannungen und die anschließende Entladung in das Prüfobjekt über einen Widerstand von 1,5k Ω erzeugt werden kann, simuliert.

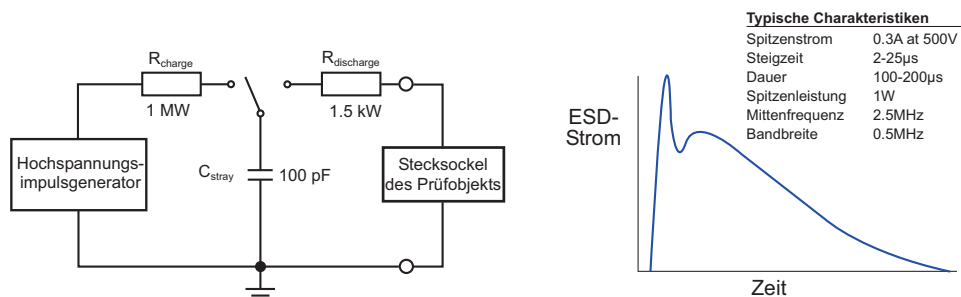


Abb.7.15 Prüfschaltung und Gesamt-Wellenform des Modells des menschlichen Körpers (HBM = Human Body Model)

Um den ESD-Nennwert des Teils zu bestimmen, wird die ESD-Prüfung bei immer ansteigenden Hochspannungen wiederholt:

Klasse	HBM-Prüfspannung
0	250VDC
1A	500VDC
1B	1kVDC
1C	2kVDC
2	4kVDC
3	8kVDC

Tabelle 7.11: ESD-Klassifikationen

Durch das Hinzufügen von ultraschnell schaltenden Dioden zur Begrenzung der Eingänge und durch den Einbau von Funkenstrecken in die PCB zur Ableitung von Energie von empfindlicheren Komponenten kann ein ESD-Schutz realisiert werden. Aber der Kostendruck des Marktes schließt fast alle, außer den preiswertesten Lösungen, aus. Genau so effektiv für die Gesamtzuverlässigkeit ist es, DC/DC-Wandler in ESD-geschützten Bereichen herzustellen, zusammenzubauen und zu verpacken und dabei antistatische Verpackung zu verwenden, sodass auch die Vibration beim Transport keine wesentliche Reibungselektrizität erzeugt. Es obliegt dann klar dem Endverbraucher, die Kette nicht zu unterbrechen und in seiner Produktion ebenfalls ESD-Schutzmaßnahmen zutreffen.

7.11 Induktivitäten

Induktivitäten finden sich in fast jedem DC/DC-Wandler. Der Transformator ist das Herzstück jeder Wandlerkonstruktion und zweifellos die kritischste Komponente bei der Bestimmung der Gesamtperformance. Die üblichsten Transformatortypen sind der Ferritringkern- oder Spulenkörpertyp (bobbin type), da diese bei Hochfrequenzen gut arbeiten und mit geschlossenen Induktionslinien gebaut werden können. Die Ferritkerne bestehen aus Eisenoxid (Fe_2O_4), kombiniert mit anderen Metallen wie Mangan-Zink (MnZn) und Nickel-Zink (NiZn) und einem Binder, dann in Form gepresst und eingebrannt, um eine leicht magnetisierbare Kristallstruktur zu erzeugen.

Die Kerne sind spröde und müssen sorgfältig gehandhabt werden. Subminiaturtoroide werden häufig – um eine glatte, gleitfähige Oberfläche zu erzeugen – zusätzlich mit Nylon oder Epoxidharzlack beschichtet, um die Wahrscheinlichkeit von Transportschäden zu verringern und die Wicklung von Hand zu vereinfachen. Es mag überraschend klingen, dass die meisten auf dem Markt befindlichen serienmäßigen Kleinleistungs-DC/DC-Wandler handgewickelte Ringkerntransformatoren verwenden. Das Problem der Entwicklung eines automatischen Prozesses zur Herstellung von Transformatoren mit einem Kerndurchmesser von nur ca. 6mm und einem 3mm großen Mittelloch, die bis zu sechs separate Wicklungen erfordern, konnte bislang noch nicht vollständig gelöst werden. RECOM verfügt über zwei automatische Wickelmaschinen, die solche Ringkerntransformatoren herstellen können; diese sind eine Eigenentwicklung und kommerziell nicht verfügbar. Das Problem existiert nicht bei Transformatoren mit Spulenkörpern, die maschinengewickelte Wicklungen auf einem Kunststoffspulenkörper verwenden. Um den Transformator zu bilden, werden zwei Hälften des Ferritkerns ringsum die Spule aufeinander geklebt.

Von den beiden Konstruktionstypen sind die Ringkerntypen am zuverlässigsten. Sie sind auch selbstschirmend, da der Magnetfluss innerhalb des Ringkerns eng gebunden ist. Transformatoren mit Spulenkörpern können ausfallen, wenn sich bei der Montage oder während des Einsatzes Risse im Kern bilden, was einen unerwünschten Luftspalt darstellen und somit die Arbeitsparameter des Wandlers ändern würde. Dasselbe betrifft das Zusammenkleben der zwei Ferrithälften; die Hälften müssen unbeschädigt, korrekt geebnet und gut poliert sein, damit ein enger Kontakt entsteht. Bei der auf die Betriebszuverlässigkeit gerichteten Konstruktionsweise besteht die Hauptüberlegung darin, die Kerntemperatur unter dem Curie-Punkt – der Temperatur, bei der der Kern beginnt, seine magnetischen Eigenschaften zu verlieren – zu halten. Verschiedene Ferritgemische weisen unterschiedliche Curie-Punkte auf.

Außer dem Transformator können auch Induktivitäten für Sperrwandler oder RF-Drosseln auf der PCB als SMD-Komponenten verwendet werden. SMD-Induktivitäten sind üblicherweise auf einem Keramik- oder Ferritspulenträger gewickelt, während RF-Drosseln, ähnlich wie MLCCs, auch mit einer mehrschichtigen Konstruktion hergestellt werden können.

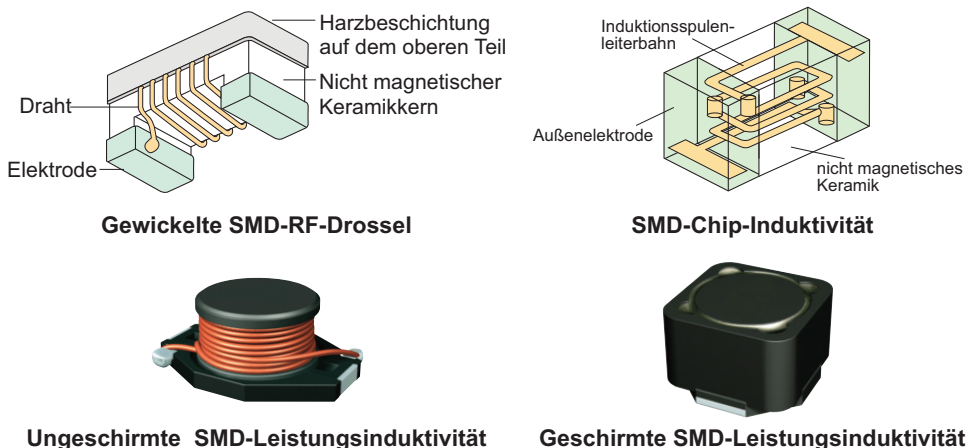


Abb. 7.16: Konstruktionsunterschiede zwischen SMD-Induktivitäten

Es wird dringend empfohlen, geschirmte Induktivitäten zu verwenden. Die magnetische Abschirmung blockiert nicht nur Störungen zwischen angrenzenden Komponenten, sondern wenn zwei ungeschirmte Induktivitäten dicht nebeneinander platziert werden, verringert deren Wechselwirkung auch deren effektive Induktivität:

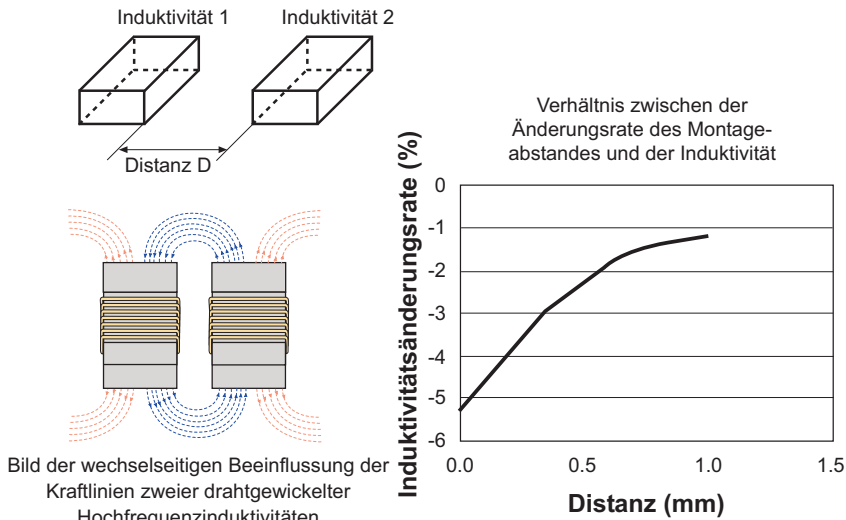


Abb. 7.17: Wechselwirkung zwischen angrenzenden ungeschirmten Induktivitäten

Ungeschirmte Induktivitäten können auch die Gesamtzuverlässigkeit einer Anwendung außerhalb des DC/DC-Wandlers beeinträchtigen. Streumagnetfelder induzieren Ströme, die in jedem angrenzenden Leiter, ob es sich um eine PCB-Leiterbahn, Kondensatorschicht oder einen Kabelbaum handelt, fließen. Die meisten Ferrite verbessern ihre ursprüngliche Performance, wenn sie sich „setzen“. Das zyklische Erwärmen und Abkühlen beim Ein- und Ausschalten des Wandlers, sowie die Erregung des schnell wechselnden Magnetflusses, verursachen eine Selbstausrichtung der magnetischen Grenzen, was eine langsame Erhöhung der Permeabilität hervorruft. Obwohl dies aus der Perspektive der Performance des DC/DC-Wandlers positiv ist, bedeutet es, dass ungeschirmte Induktivitäten während der ersten Betriebswochen ihre projektierten Magnetfelder so lange langsam erhöhen können, bis dies plötzlich zu einem Problem führen kann. Andererseits verdichten geschirmte Induktivitäten jegliche magnetische Kraftlinienstreuungen langsam. Dieser Effekt nimmt nach ca. 50-60 Stunden der Nutzung ab und stabilisiert sich dann.



7.18: Röntgenvergleich zwischen einer geschirmten Induktivität, verwendet in RECOMs Sperrwandler R-78 (Abbildung links), und einer ungeschirmten Induktivität, verwendet in dem kopierten Produkt eines Mitbewerbers (Abbildung rechts)

8. LED-Charakteristiken

Im Krieg lautet die erste Soldatenpflicht "Du musst wissen, wer Dein Feind ist". Dasselbe Prinzip gilt für die Festkörperbeleuchtung (SSL = Solid State Lighting) – wenn man nicht verstehen, wie sich die LED (Leuchtdiode) verhält, darf man sich nicht wundern, wenn die Anwendung nicht gelingt.

LEDs sind nichtlineare Bauteile. Wenn an die LED eine niedrige Spannung angelegt wird, leitet sie nicht. Steigt nun die Spannung an, passiert sie einen Schwellenwert, ab dem der Strom durch die LED steil ansteigt und die LED plötzlich Licht auszustrahlen beginnt. Steigt die Spannung weiter an, überhitzt sich die LED schnell und brennt durch. Der Trick besteht darin, die LED in dem schmalen Band zwischen voll Aus und voll Ein zu betreiben.

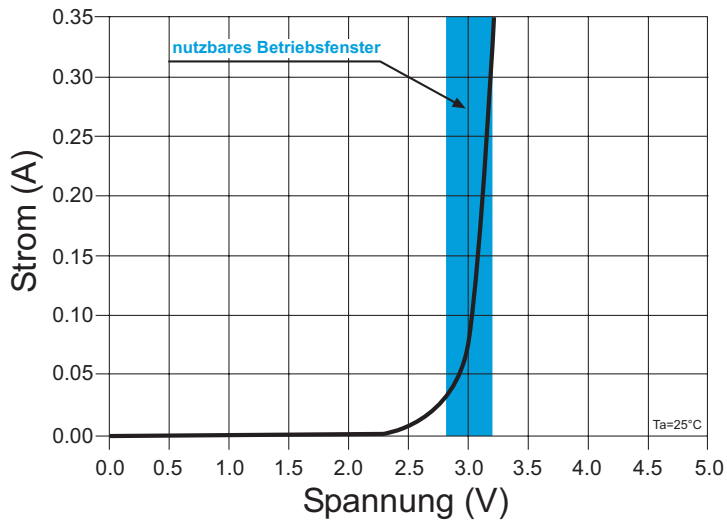


Abb. 8.1: Nutzfläche für Hochleistungs-LEDs [$T_{AMB} = 25^\circ\text{C}$]

Es gibt jedoch eine zusätzliche Schwierigkeit. Das Fenster der nutzbaren Spannung ist bei verschiedenen Hochleistungs-LEDs unterschiedlich (sogar innerhalb der LEDs ein und derselben Charge und ein und desselben Lieferanten), und dieser Spannungsbereich ändert sich mit der Umgebungstemperatur und dem Alter der LED.

Abb. 8.2 zeigt die Nutzfläche detaillierter. Dieses Beispiel zeigt vier identische LEDs, die gemäß ihrem Datenblatt über dieselbe Spezifikation verfügen. Alle LED-Hersteller sortieren LEDs nach der Lichtfarbe, die sie ausstrahlen (dies wird "binning" = "Klasseneinteilung" genannt – die LEDs werden während der Herstellung geprüft und laut ihrer Farbtemperatur sortiert).

Die Folge ist, dass eine Lieferung mehrere verschiedene Fertigungslose beinhalten kann. Daher ist eine breite Variation an Schwellenwerten bzw. Durchlassspannungen (V_F) zu erwarten. Die meisten Datenblätter von Hochleistungs-LEDs spezifizieren eine V_F -Toleranz von ca. 20%, weshalb die breiten Variationen, die in Abb. 8.2 gezeigt werden, nicht übertrieben sind.

Wenn man, wie in diesem Beispiel, als Versorgungsspannung z. B. 3V wählt, wird LED 1 übersteuert, LED 2 zieht 300mA, LED 3 zieht 250mA, und LED 4 zieht nur 125mA.

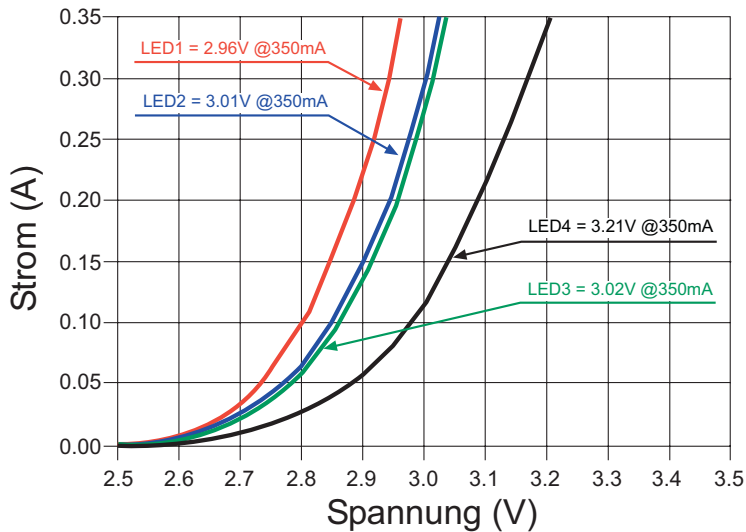


Abb. 8.2: LED-Charakteristik im Detail

Darüber hinaus sind diese Kurven dynamisch. Da sich die LEDs bis zu ihrem Arbeitstemperaturbereich erwärmen, driften alle Kurven nach links (die Durchlassspannung V_F verringert sich mit dem Temperaturanstieg).

Die Lichtleistung der LED ist jedoch direkt proportional zum Strom, der durch sie fließt (Abb. 8.3). Daher wird in obigem Beispiel bei einer Versorgungsspannung von 3V LED 1 (mitunter nur kurzzeitig) wie eine Supernova, LED 2 ein wenig heller als LED 3 und LED 4 sehr dunkel leuchtend erscheinen.

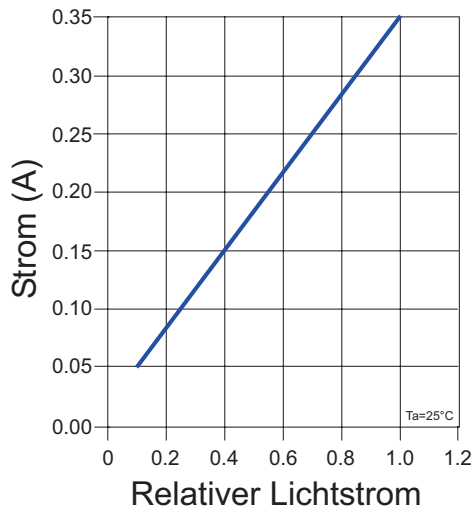


Abb. 8.3: Verhältnis der Lichtleistung gegenüber dem LED-Strom

8.1 Betreiben von LEDs mit Dauerstrom

Die Lösung dieses Problems der Veränderlichkeit der Durchlassspannung V_t besteht darin, einen Konstantstrom statt einer Konstantspannung zu verwenden, um die LEDs zu betreiben.

Der LED-Treiber regelt die Ausgangsspannung automatisch, um den Ausgangsstrom und somit die Lichtleistung konstant zu halten. Das funktioniert mit einer einzelnen LED oder mit einer Kette von in Reihe geschlossenen LEDs. Solange der Strom durch alle LEDs derselbe ist, weisen sie dieselbe Helligkeit auf, selbst wenn V_F an jeder LED verschieden ist (siehe Abb. 8.4).

Während sich die LEDs bis zu ihrer Betriebstemperatur erwärmen, verringert der Konstantstromtreiber die Erregungsspannung automatisch, um den Strom durch die LEDs konstant zu halten. Daher ist die Helligkeit der LEDs auch von der Betriebstemperatur unabhängig.

Ein weiterer wichtiger Vorteil besteht darin, dass der Konstantstromtreiber es nicht zulässt, dass auch nur eine einzige LED in einer Kette übersteuert wird und somit sicherstellt, dass alle LEDs über eine lange Betriebsdauer verfügen. Fällt eine LED durch Kurzschluss aus, arbeiten die restlichen LEDs immer noch mit dem korrekten Strom.

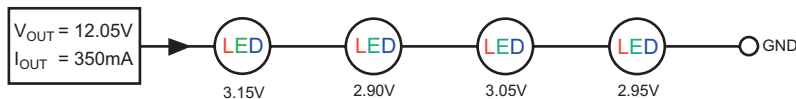


Abb. 8.4: LED-Kette

8.2 Einige DC-Konstantstromquellen

Die einfachste Konstantstromquelle ist eine Konstantspannungsversorgung, die LEDs über einen Widerstand betreibt. Wenn der Spannungsabfall am Widerstand ungefähr gleich der Durchlassspannung einer LED ist, dann ruft eine 10%ige Änderung von V_F eine ähnliche Änderung des LED-Stroms hervor (vergleichen Sie dies mit den in Abb. 8.2 gezeigten Kurven, bei denen eine 10%ige Änderung von V_F eine etwa 50%ige Änderung des LED-Stroms verursacht). Diese Lösung ist zwar preiswert, besitzt jedoch eine schlechte Stromregelung und ist wenig energieeffizient. Viele preiswerte LED-Cluster-Birnen, die als Ersatz für Niederspannungs-Halogenlampen angeboten werden, verwenden dieses Verfahren. Es versteht sich von selbst, dass, wenn eine der LEDs durch Kurzschluss ausfällt, der Widerstand überlastet wird und normalerweise nach einer relativ kurzen Zeit durchbrennt, weshalb die zu erwartende Lebensdauer dieser LED-Cluster-Lampen relativ kurz ist.

Eine weitere einfache Konstantstromquelle ist ein Linearstromregler. Auf dem Markt sind preiswerte LED-Treiber erhältlich, die dieses Verfahren verwenden, oder man kann einen Standard-Linearspannungsregler im Konstantstrombetrieb verwenden. Der interne Feedbackkreis hält den Strom innerhalb von ca. $\pm 5\%$ geregelt, die Verlustleistung am Regler muss aber als Wärme abgeführt werden, weshalb eine gute Kühlung des Reglers erforderlich ist. Der Nachteil ist der niedrige Wirkungsgrad dieser Lösung, was eher gegen die Verwendung von hocheffizienten SSL-Geräten spricht.

Die beste Konstantstromquelle ist ein Sperrwandler. Der Preis des Treibers ist höher als bei den anderen Lösungen, aber die Genauigkeit des Ausgangsstroms kann bei $\pm 3\%$ über einem breiten LED-Lastbereich liegen und Wandlungswirkungsgrade können bis zu 96% betragen, was bedeutet, dass nur 4% Energie als Wärme verloren geht und die Treiber bei hohen Umgebungstemperaturen verwendet werden können.

Beispiel von Konstantstromquellen für LEDs:

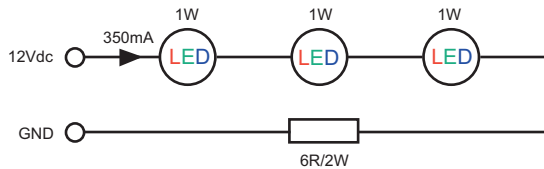


Abb. 8.5: Einfacher Widerstand: preiswert, aber ungenau und unwirtschaftlich

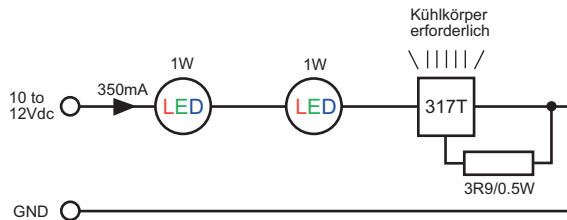


Abb. 8.6: Linearer Spannungsregler: preiswert und genau, aber unwirtschaftlich

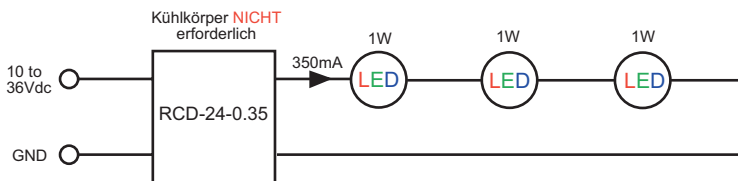


Abb. 8.7: Sperrwandler: höherer Preis, aber genau und effizient

Ein wichtiger Unterschied zwischen den oben gezeigten Varianten sind die Eingangs- und Ausgangsspannungsbereiche.

Ein DC/DC-Sperrwandler verfügt über einen breiten Eingangs- und Ausgangsspannungsbereich, über welchem die Konstantstromregelung gut arbeitet (z. B. arbeitet der RCD-24.0.35 von 5V bis 36VDC und hat einen Ausgangsspannungsbereich von 2 bis 34VDC). Ein breiter Ausgangsspannungsbereich ermöglicht nicht nur viele verschiedene Kombinationen von LED-Kettenlängen, sondern erlaubt auch einen breiten Dimmungsbereich.

Zwei andere oben gezeigte Varianten haben Verlustleistungsprobleme, wenn nur eine LED benötigt wird, da der Widerstand oder der Linearspannungsregler einen größeren Spannungsabfall an ihnen hat, was die Verlustleistung noch weiter erhöht. Der Eingangsspannungsbereich muss aus demselben Grund beschränkt werden.

8.3 Verbindung von LEDs in Ketten

Die meisten weißen Leistungs-LEDs werden zum Betrieb bei 350mA Dauerstrom entwickelt. Der Grund dafür ist, dass die Chemie der zur Herstellung einer Weißlicht-LED verwendeten Materialien eine Durchlassspannung von ca. 3V und $3,0V \times 0,35A \sim 1 \text{ Watt}$ festlegt, was einer geeigneten LED-Leistung entspricht.

Die meisten DC/DC-Konstantstrom-LED-Treiber sind Buck- oder Abwärtswandler. Das bedeutet, dass die Maximalausgangsspannung niedriger als die Eingangsspannung ist. Also hängt die Anzahl der LEDs, die angetrieben werden können, von der Eingangsspannung ab.

Eingangsspannung	5Vdc	12Vdc	24Vdc	36Vdc	5Vdc
Typische Anzahl an LEDs in Kette	1	3	7*	10*	15

Tabelle 8.1: Anzahl von LEDs, die man in einer Kette antreiben kann, gegenüber der Eingangsspannung

Wenn die Eingangsspannung nicht geregelt ist (z. B. bei einer Batterie), müssen die meisten LEDs entsprechend der minimal verfügbaren Eingangsspannung verringert werden.

Beispiel:

Wie viele LEDs mit 1W können von einer 12-V-Blei-Säure-Batterie angetrieben werden?

Batteriespannungsbereich	9 ~ 14 VDC
DC/DC-Treiber-Headroom	1 V
LED-Treiber-Ausgangsspannungsbereich daher	8 ~ 13 VDC,
wenn die LED-Durchlassspannung V_F	typischerweise 3,3 V*
dann beträgt die maximale Anzahl an LEDs, die angetrieben werden können,	2

Zwei LEDs sind nicht sehr viel! Eine Möglichkeit, dieses Problem zu umgehen, besteht darin, entweder einen Booster-Wandler, bei dem die Ausgangsspannung höher ist als die Eingangsspannung, oder zwei oder mehr parallel geschaltete Ketten von LEDs zu verwenden. Für jede verwendete 350-mA-Kette von LEDs muss der Treiberstrom erhöht werden, um einen korrekten Gesamtstrom zu liefern. Eine einzelne Kette benötigt also einen Treiber von 350mA, zwei parallelgeschaltete Ketten benötigen einen Treiber 700mA, drei parallelgeschaltete Ketten würden eine Quelle von 1,05A benötigen usw.

Deshalb hängt die Auswahl von LED-Treibern von der verfügbaren Eingangsspannung und der Anzahl der Ketten von LEDs, die angetrieben werden müssen, ab.

Abb. 8.8, 8.9 und 8.10 zeigen einige mögliche Kombinationen für eine feste Versorgung von 12VDC unter Verwendung von typischen weißen 1-W-LEDs. Mit einer geregelten 12-V-Versorgung können bis zu drei LEDs in einer Kette angetrieben werden ($3 \times 3,3V = 9,9V$, was 2,1V Headroom für die Regelung des Dauerstromtreibers ergibt).

*** Anmerkung:** Es ist ein gängiger Irrtum, dass die Anzahl an LEDs, die angetrieben werden können, von der maximalen V_F , die im LED-Datenblatt angegeben wird, abhängt. In der Praxis ist dies nicht der Fall, weil V_F wesentlich abfällt, wenn die LEDs ihren jeweiligen Arbeitstemperaturbereich erreichen. Somit kann die im Datenblatt angegebene typische V_F betriebssicher verwendet werden. Ein typisches Datenblatt kann festlegen, dass V_F bei 25°C Umgebungstemperatur mindestens 3,3V, typischerweise 3,6V und maximal 3,9V beträgt. Die Zahlen würden jedoch bei 50°C näher an mindestens 3,0V, typischerweise 3,3V und maximal 3,6V liegen. Deshalb kann eine feste 24-V-Versorgung sieben LEDs sicher antreiben und eine 36-V-Versorgung zehn LEDs, selbst wenn die Spannung am LED-Treiber geringfügig abfällt.

8.4 Parallelschaltung von LED-Ketten

3 LEDs in einer einzelnen Kette mit 350-mA-Treiber:

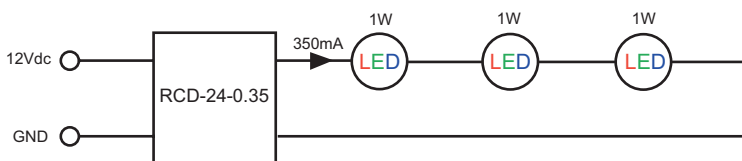


Abb. 8.8: Vorteile: genauer LED-Strom, ausfallsicher / **Nachteile:** geringe Anzahl an LEDs je Treiber

6 LEDs in zwei Ketten mit 700-mA-Treiber:

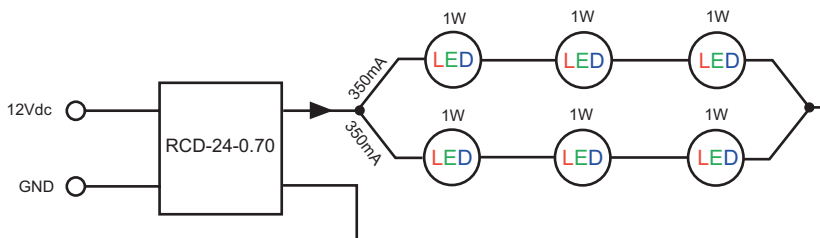


Abb. 8.9: Vorteile: doppelte Anzahl von LEDs je Treiber / **Nachteile:** nicht ausfallsicher, unsymmetrische Ströme in den Ketten

9 LEDs in drei Ketten mit 1050-mA-Treiber:

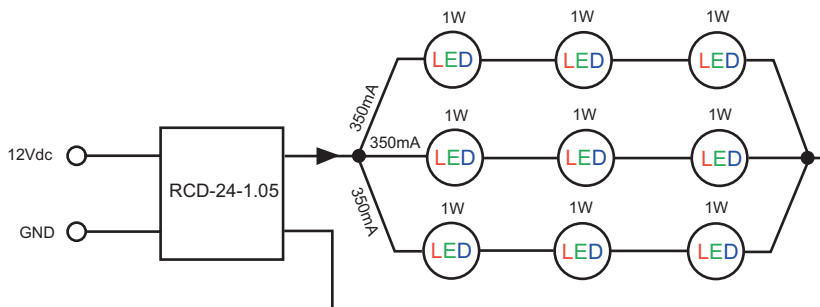


Abb. 8.10: Vorteile: dreifache Anzahl von LEDs je Treiber / **Nachteile:** nicht ausfallsicher, unsymmetrische Ströme in den Ketten

Die sicherste und zuverlässigste Methode zum Antrieb von LEDs besteht darin, eine einzelne Kette von LEDs an den LED-Treiber anzuschließen. Fällt eine LED durch einen offenen Stromkreis aus, wird der Strom zu den restlichen LEDs in der Kette unterbrochen. Fällt eine LED durch Kurzschluss aus, bleibt der Strom in den restlichen LEDs unverändert.

Das Betreiben von Mehrfachketten durch einen einzelnen LED-Treiber hat den Vorteil, dass mehr LEDs betrieben werden können, aber es besteht die Gefahr, dass eine LED ausfällt. Wenn bei zwei parallelgeschalteten Ketten eine LED durch einen offenen Stromkreis ausfällt, fließt der Dauerstrom von 700mA durch die restliche LED-Kette und ruft nach sehr kurzer Zeit auch deren Ausfall hervor. Wenn bei drei parallelgeschalteten Ketten eine einzelne LED ausfällt, wird der Antriebsstrom von 1A auf die restlichen beiden Ketten aufgeteilt.

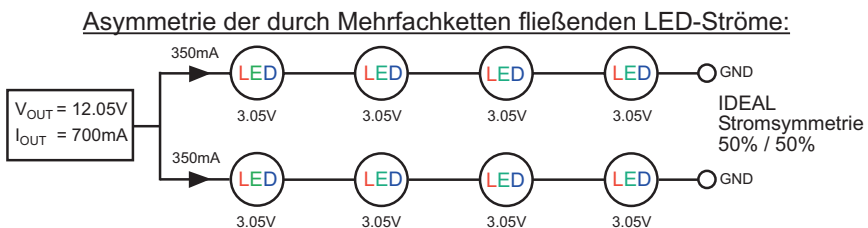
Beide Ketten werden mit jeweils 500mA pro Kette überlastet. Je nach dem, wie gut die LEDs Wärme ableiten können, halten die LEDs dem wahrscheinlich für eine gewisse Zeit stand, aber schließlich wird der Überstrom den Ausfall einer anderen LED hervorrufen, wonach die dritte Kette den gesamten Strom von 1A übernimmt und fast sofort ebenso ausfällt.

Fällt eine einzelne LED durch Kurzschluss aus, werden die in den Ketten fließenden Ströme sehr unsymmetrisch und der größte Teil des Stromes fließt durch die Kette mit der kurzgeschlossenen LED. Dies ruft schließlich den Ausfall der Kette mit demselben katastrophalen Dominoeffekt auf die restlichen Ketten hervor, wie oben beschrieben.

Hochleistungs-LEDs sind sehr betriebssicher, weshalb die oben beschriebenen Ausfälle nicht sehr häufig vorkommen. Viele Entwickler von LED-Beleuchtungen entscheiden sich für die einfache Lösung und niedrigen Kosten durch den Einsatz von Mehrfachketten, die durch einen einzelnen Treiber betrieben werden, und gehen dabei das Risiko ein, dass Mehrfach-LEDs ausfallen, wenn eine einzelne LED ausfällt.

8.5 Symmetrierung des LED-Stroms in parallelgeschalteten Ketten

Ein anderes wichtiges Problem ist die Symmetrie der Ströme, die in Mehrfachketten fließen. Wir wissen, dass zwei oder mehrere Ketten von LEDs verschiedene kombinierte Durchlassspannungen haben. Der LED-Treiber liefert einen Dauerstrom bei einer Spannung, die dem Mittelwert der kombinierten Durchlassspannungen jeder Kette entspricht. Diese Spannung ist bei einigen Ketten zu hoch und bei anderen zu niedrig, weshalb die Ströme nicht gleich aufgeteilt werden.



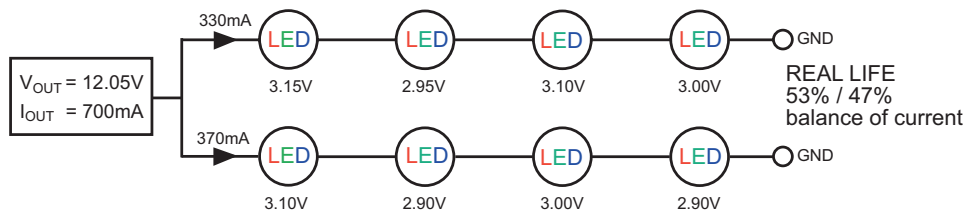


Abb. 8.12: Reale Symmetrie

In obigem Beispiel reicht die Stromasymmetrie nicht aus, um den Ausfall der überlasteten Kette hervorzurufen, weshalb die beiden LED-Ketten betriebssicher arbeiten. Es existiert jedoch eine Differenz von 6% in der Lichtleistung zwischen den zwei Ketten.

Die Lösung des Problems der unausgebalancierten Ketten besteht darin, entweder einen Treiber je Kette zu verwenden oder für den Stromausgleich einen externen Stromkreis hinzuzufügen. Solch eine Schaltung ist ein Stromspiegel.

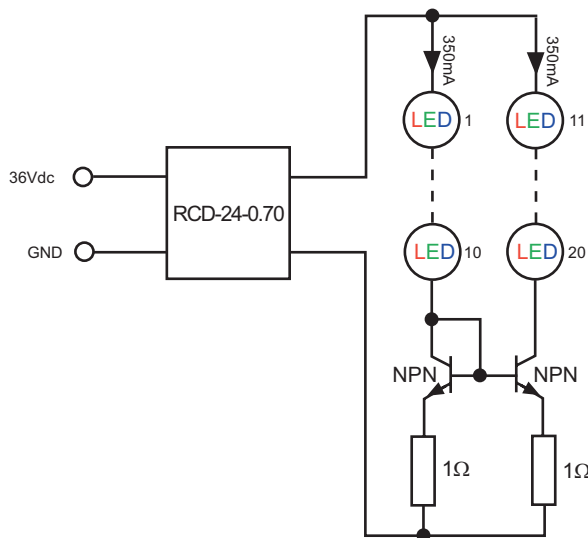


Abb. 8.13: Symmetrierung von LED-Strömen unter Verwendung eines Stromspiegels

Der erste NPN-Transistor wirkt als Referenz. Der zweite NPN-Transistor "spiegelt" diesen Strom wider. Somit werden die Ströme in der Kette automatisch gleichmäßig geteilt. Theoretisch sind Emitterleitungswiderstände mit 1Ω für den Stromspiegel nicht erforderlich, in der Praxis aber helfen sie, die Differenzen in V_{be} zwischen den Transistoren auszugleichen und eine genauere Stromsymmetrie zu erzielen.

Ein Stromspiegel hilft auch als Schutz gegen LED-Ausfälle. Fällt eine LED in der ersten Kette durch einen offenen Stromkreis aus, wird die zweite Kette geschützt (der Referenzstrom ist Null, daher fällt der Strom in den anderen Ketten auch auf Null). Auch wenn eine LED durch Kurzschluss ausfällt, bleiben die Ströme weiterhin ausgeglichen.

Einige Hersteller von LED-Treibern behaupten, dass LEDs den Strom automatisch gleichmäßig teilen würden und solche externen Stromspiegelkreise unnötig wären. Dies ist nicht der Fall. Eine Asymmetrie gibt es nur dann nicht, wenn die kombinierten Durchlassspannungen der LED-Ketten absolut identisch sind.

Wenn beispielsweise zwei parallelgeschaltete Ketten auf einem gemeinsamen Kühlkörper montiert sind und eine Kette mehr Strom zieht als die andere, wird sie heller und heißer laufen. Die Kühlkörpertemperatur steigt langsam an, was zum Abfallen der V_F der zweiten Kette und auch dazu führt, dass sie versucht, mehr Strom zu ziehen. In der Theorie sollten die beiden Ketten dann wegen der thermischen negativen Rückführung ihre jeweiligen Ströme ausgleichen. In der Praxis kann dieser Effekt zwar gemessen werden; er ist jedoch ungenügend, um einen genauen Stromausgleich zu garantieren.

Wenn es sich bei den beiden Ketten um zwei separate LED-Lampen handelt, findet außerdem keine Rückführung zur Temperaturkompensation statt. Die Lampe mit der niedrigsten kombinierten V_F zieht den meisten Strom, läuft am heißesten, und die V_F fällt immer noch weiter. Dies verschlechtert die Asymmetrie und kann zu Thermoinstabilität und LED-Ausfall führen.

Als der Schaltkreis in Abb. 8.13 veröffentlicht wurde, gab es einige Kritik im Internet, dass ein Stromspiegel keine ideale Lösung wäre und dass schon das Hinzufügen von Widerständen mit 1Ω helfen würde, den Strom auszugleichen. Dies ist bis zu einem gewissen Grad richtig, wenn man jedoch eine genaue Stromsymmetrie benötigt, dann ist der Stromspiegel – abgesehen vom Antreiben jeder einzelnen Lampe durch einen eigenen Treiber – immer noch die einfachste und beste Lösung.

8.6 Parallele Ketten oder Grid-Array – was ist besser?

In Abschnitt 8.4 wurden die Folgen von einzelnen LED-Ausfällen durch einen offenen Stromkreis oder Kurzschluss besprochen. Je größer die Anzahl von parallelgeschalteten Ketten ist, desto niedriger ist die Gefahr, dass ein Einzelfehler in einer Kette zum Ausfall der restlichen Ketten führen würde. Wenn also fünf Ketten parallel geschaltet wären, würden bei Ausfall durch einen offenen Stromkreis einer LED-Kette die restlichen vier Ketten nur mit 125 % des Nennstroms übersteuert. Die LEDs würden zwar deutlich heller leuchten, aber sie würden wohl kaum ausfallen, solange eine gute Wärmeabfuhr gewährleistet ist.

Der Nachteil einer Parallelschaltung vieler Ketten besteht darin, dass ein Treiber, der mehrere Ampere liefern kann, erforderlich ist und dieser teuer oder schwer erhältlich sein könnte. Außerdem ist bei der Auswahl von LED-Treibern, die viele Ampere an Strom liefern können, eine gewisse Sorgfalt erforderlich. Wenn die LED-Last zu niedrig ist, weil z. B. ein Steckverbinder an einigen der Ketten einen fehlerhaften Anschluss hat, brennt der Strom die restlichen LEDs sofort durch. Bevor der LED-Treiber eingeschaltet wird muss genau darauf geachtet werden, dass alle Verbindungen einwandfrei sind. Viele teure LED-Beleuchtungsinstallationen wurden schon aufgrund fehlerhafter Verdrahtung und Einsatz von Hochstrom-LED-Treibern beschädigt!

In der Praxis ist es sicherer, die Anzahl von parallelgeschalteten Ketten je Treiber auf fünf oder weniger zu begrenzen und mehrere Niederstrom-Treiber statt eines einzelnen Hochstromtreibers zu verwenden, wenn entsprechend viele LEDs angetrieben werden müssen.

Die Verwendung langer Ketten ist auch keine schlechte Idee, denn wenn eine LED durch Kurzschluss ausfällt, ist der Stromanstieg in dieser Kette proportional geringer je länger die Ketten sind.

Die nächste Frage ist, ob die LEDs in individuellen Ketten geschaltet oder die Ketten kreuzweise verbunden werden sollen, um ein LED-Array anzufertigen. Das folgende Beispiel mit 15 LEDs veranschaulicht diese beiden Varianten (in beiden Fällen ist der Treiber derselbe). Man könnte die 15 LEDs in fünf Säulen von 3 LEDs anschließen, aber aus dem oben erwähnten Grund sind drei Säulen aus 5 LEDs eine sicherere Anordnung.

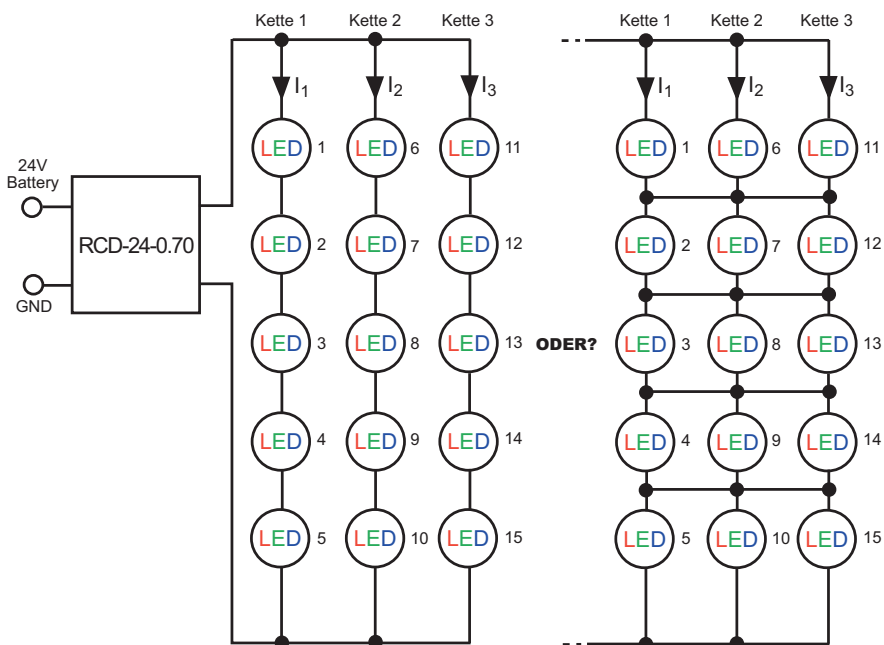


Abb. 8.14: Verbindung von LEDs in parallelen Ketten oder ein Grid-Array

Der Vorteil eines Grid-Arrays besteht darin, dass, wenn eine LED ausfällt, nicht die ganze Säule von LEDs ausfällt, sondern nur die LEDs in derselben Reihe wie die ausgefallene LED überlastet werden. Fällt eine LED durch Kurzschluss aus, werden die LEDs in derselben Reihe nicht mehr leuchten, aber der durch die restlichen LEDs fließende Strom bleibt immer noch korrekt.

Wenn es darauf ankommt, dass eine Lampe mit 15 LEDs zuverlässig arbeitet und dass sie weiterhin Licht ausstrahlt, selbst wenn einzelne LEDs durch einen offenen Kreis oder Kurzschluss ausfallen, dann ist die Grid-Array-Lösung der beste Weg, die LEDs zu verschalten.

Der Nachteil eines Grid-Arrays besteht darin, dass es sich bei der VF in jeder Reihe um einen Durchschnittswert handelt und die Toleranz von $\pm 20\%$ in einzelnen LED-Durchlassspannungen bedeuten kann, dass nicht alle LEDs dieselbe Helligkeit aufweisen. Dies kann zu lokalen Überhitzungsstellen und einer verringerten LED-Lebensdauer einiger LEDs und auch den ästhetischen Anspruch verletzen (z.B. man stelle sich eine Designerlampe vor wo manche LEDs deutlich dunkler leuchten als andere).

Wenn es wichtig ist, dass eine Lampe mit 15 LEDs eine sehr gleichmäßige Lichtleistung ohne lokale Überhitzungsstellen aufweist, dann ist eine Schaltung der LEDs in parallelgeschalteten Ketten die beste Lösung.

Wenn sowohl die Ausfallsunempfindlichkeit als auch eine gleichmäßige Lichtleistung wesentlich sind, dann ist es am besten, drei Ketten und drei 350-mA-Treiber zu verwenden!

8.7 Dimmen von LEDs

Wenn LEDs gedimmt werden sollen – sei es durch 1-10V Analogspannung, -Phasenanschnittsteuerung, Stromnetz (power-line), Digitaleingänge wie DALI oder eine WLAN-Verbindung – dann gibt es tatsächlich nur zwei Möglichkeiten, den Ausgang einer LED zu dimmen: entweder durch lineares Reduzieren des Stroms durch die LED (Analogdimmung) oder durch deren sehr schnelles Ein- und Ausschalten mit verschiedenen Impuls-Pausen-Verhältnissen (PWM-Dimmung). Obwohl beide Verfahren denselben Effekt erreichen, gibt es wichtige Unterschiede in der Art und Weise, wie sie in der Praxis arbeiten. Daher ist die richtige Wahl des Verfahrens zur Dimmung für viele Anwendungen von entscheidender Bedeutung.

8.7.1 Analoges Dimmen gegenüber PWM-Dimmen

Eine LED arbeitet in einem sehr engen Durchlassspannungsbereich. Eine typische, lichtstarke Einzelchip-LED beginnt bei ca. 2,5V zu leuchten, erreicht 10% Helligkeit bei 2,7V und volle Helligkeit bei 3,1V. Die Aufgabe eines Konstantstrom-LED-Treibers besteht darin, die an der LED angelegte Spannung kontinuierlich abzugleichen, um den Strom durch sie konstant zu erhalten, obwohl die LED mit der Temperatur und Zeit driftet. Um die Stabilität zu erhöhen, wird der Strom in der Regel durch Messung der Spannung an einem niederohmigen Serienwiderstand und Zuführung des Ergebnisses in eine analoge Rückführungsschleife mit einer relativ langsamen Ansprechzeit überwacht. Analoges Dimmen lässt sich dann leicht einführen, indem man eine Komparatorstufe in die Rückführungsschleife einfügt.

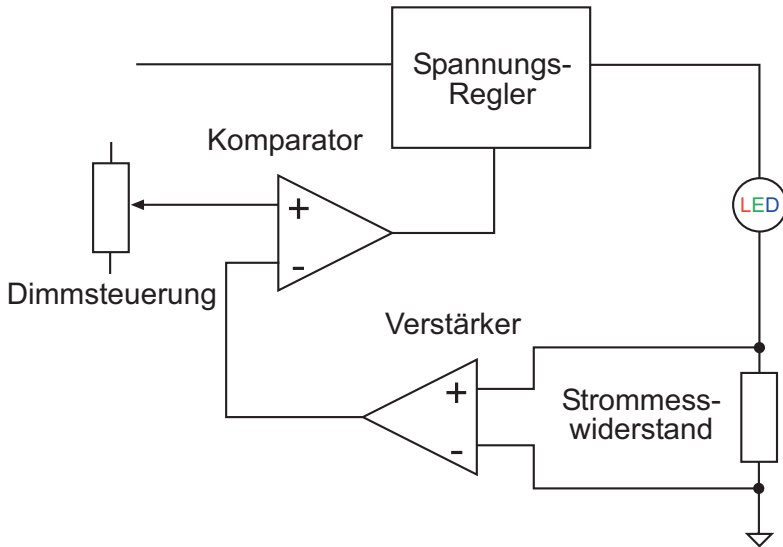


Abb. 8.15 Analoge Dimmsteuerung

Analoges Dimmen kann – außer bei den extremen Einstellwerten bei fast voller Helligkeit oder fast voller Dunkelheit – sehr lineare Dimmkurven liefern. Auf den hellsten Dimm-Niveaus können Sättigungseffekte im Komparator nichtlineares Ansprechen erzeugen, während auf den dunkelsten Lichtniveaus der Strom durch den Shuntwiderstand so niedrig ist, dass die Eingangsoffsetspannungen im Messverstärker zu einer wesentlichen Fehlerquelle werden. Das Gesamtergebnis ist unvermeidliches nichtlineares Dimmen in den unteren und den oberen 3% des Dimmungsbereichs, und dies gilt sogar für einen mit Sachverstand entwickelten Analogdimmung-Schaltkreis.

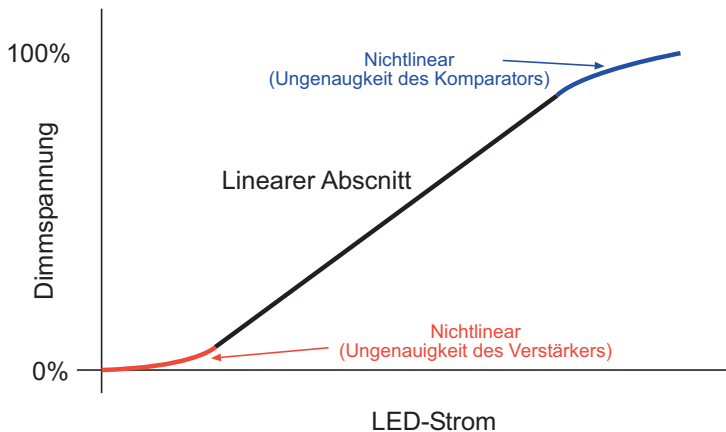


Abb. 8.16: Nichtlinearitäten von analogem Dimmen

Eine Alternative zu analogem Dimmen ist PWM-Dimmen. Hierbei sind zwar auch ein Serienwiderstand und ein Strommessverstärker erforderlich, um den maximalen durch die LED fließenden Strom zu überwachen, aber die angelegte LED-Spannungsquelle wird mit einem PWM-Signal ein- und ausgeschaltet. Diese Vorgehensweise wird aufgrund ihrer Einfachheit sehr häufig bei Einzelchip-LED-Treibern wendet.

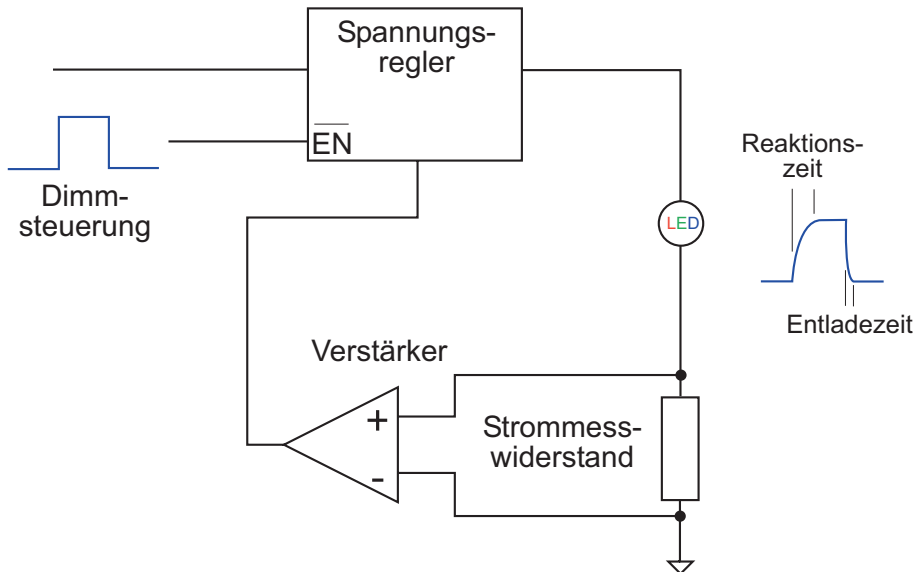


Abb. 8.17: PWM-Dimmungssteuerung

PWM-Dimmen ist nicht so linear wie analoges Dimmen. Wenn der PWM-Steuereingang nach unten geht, schaltet die Ausgangsspannung nicht sofort aus, da sich die Ausgangskapazität durch die LED-Last entladen muss. Wenn der PWM-Eingang nach oben geht, hat der Spannungsregler eine verzögerte Ansprechzeit auf den Freigabeeingang, da er sich zuerst einschalten muss. Diese Ein- und Ausschaltverzögerungen bedeuten, dass relativ niederfrequente PWM-Signale (einige Hundert Hz) verwendet werden müssen und das Ansprechverhalten des Dimmens nichtlinear ist. In vielen Konstruktionen bedeuten diese Verzögerungen, dass PWM-Dimmen unter 10% nicht möglich ist, weil der Treiber nicht rechtzeitig auf das kurze Eingangssignal reagieren kann.

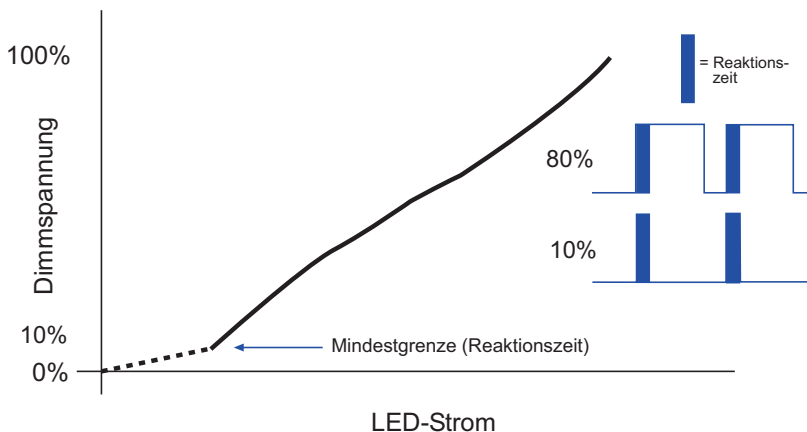


Abb. 8.18: Nichtlinearitäten beim PWM-Dimmen

Eine Alternative zur Ansteuerung des Freigabeeingangs mit dem PWM-Dimmsignal besteht darin, die Masseverbindung an der LED-Kette durch einen FET zu unterbrechen.

Da der FET viel schneller als die Stromregler-Rückführungsschleife reagiert, ist tiefes Dimmen unter 5 % möglich. Wenn die LED-Last jedoch abgeschaltet wird, schwebt der Ausgang bis zum maximalen Grenzwert, weshalb die Stromregelung einige Zeit zur erneuten Stabilisierung benötigt, wenn die LED wiedereingeschaltet wird. Dieser Stromüberschusseffekt tritt in jedem PWM Zyklus auf. Daher muss die Langzeitwirkung auf die Lebensdauer der LED beachtet werden.

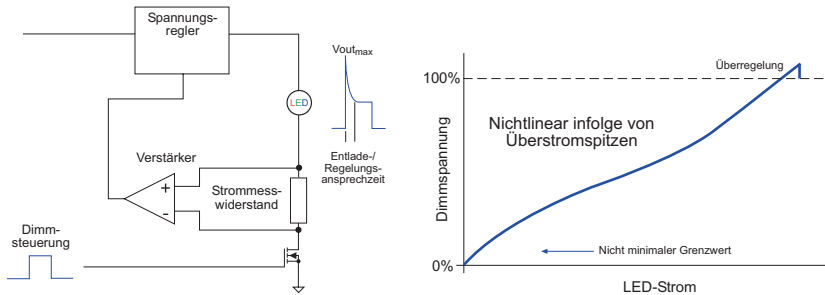


Abb. 8.19: Schalt-Dimmen

8.7.2 Wahrgenommene Helligkeit

Nachdem wir erkennen mussten, dass es keine ideale Methode für das Dimmen von LEDs gibt, stoßen wir auf das nächste Problem: unsere Augen. Die menschliche Wahrnehmung von Helligkeit ist nichtlinear. Bei schwachen Lichtpegeln öffnet sich unsere Regenbogenhaut automatisch, um mehr Licht einzulassen – daher nehmen wir LEDs heller wahr, als ein einfacher Lichtmesser es anzeigen würde. Um das Verhältnis zwischen wahrgenommener und gemessener Helligkeit auszurechnen, nimmt man die Quadratwurzel aus dem gemessenen Normlicht, z. B. würde eine auf ein Viertel (0,25) des LED-Nennstroms gedimmte LED für unsere Augen $\sqrt{0,25} = 0,5$ oder halb so hell erscheinen.

Obwohl sich fast alle Hersteller von LED-Treibern sehr darum bemühen, ihre Dimmer möglichst linear und mathematisch genau zu dimmen, bevorzugen unsere Augen die nichtlineare Kurve der natürlichen Glühlampe, da sie unserer Helligkeitswahrnehmung sehr viel näher kommt als das lineare Ansprechen der LEDs. Auf dem LED-Markt ist zurzeit die Nachfrage nach Linearität größer als nach Natürlichkeit, da dies die Abstimmung von verschiedenen Lichtarten einfacher macht. In Zukunft könnte sich dies jedoch ändern, da der Markt reifer und die Nachfrage nach natürlicherem Dimmen größer werden wird.

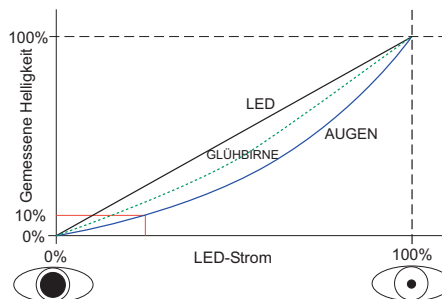


Abb. 8.20: Menschliche visuelle Wahrnehmung

8.7.3 Abschlusswort zum Dimmen

Das LED-Dimmen mag von vielen Lieferanten von Vorschaltgeräten als "eine beschlossene Sache" präsentiert werden; selbstbewusst schreiben sie Spezifikationen wie Dimm-Verhältnisse von 1:1000 in ihre Datenblätter, obwohl ihre Ausgangsgenauigkeit nur $\pm 5\%$ (1:20) beträgt, aber diese kurze Erörterung zeigt, dass genaues, lineares und flimmerfreies LED-Dimmen immer noch keine Selbstverständlichkeit ist, ungeachtet vieler Tausender von verschiedenen dimmbaren LED-Treibern auf dem Markt. Da sich die LED-Technik jedoch ständig weiterentwickelt, um mehr Licht für weniger Strom anzubieten, interessieren sich die Anwender weniger für verfügbare Lichtleistung und mehr für die Steuerung dieser Leistung. Somit werden dimmbare LEDs immer mehr zum Standard – umso mehr, weil neue Faktoren, wie z. B. die Energieeffizienzrichtlinien, zunehmend dimmbare Beleuchtung fordern, um den Stromverbrauch zu verringern.

8.8 Thermische Überlegungen

Hochleistungs-LEDs benötigen gute Wärmeableitung, wenn sich ihre tatsächliche Lebensdauer der im Datenblatt angegebenen nähern soll. Die erste Frage könnte sein, warum LEDs mit hohem Wirkungsgrad heiß werden? Es scheint unlogisch, dass eine LED mit einer Lichtausbeute von ca. 50 Lumen-per-Watt ein sorgfältigeres Wärmedesign als beispielsweise ein Scheinwerfer mit einem Bruchteil des Wirkungsgrades benötigt.

Das folgende Beispiel kann helfen: Ein Halogenstrahler mit 100W liefert 5W nützliches Licht. Aus den restlichen 95W aufgenommener Leistung werden ca. 80W im Infrarotbereich ausgestrahlt und nur 15W als Wärme an den Beleuchtungskörper geleitet. Eine 50-W-LED-Beleuchtung liefert ebenfalls 5W nützliches Licht. Aber die restliche Leistung von 45W wird als Wärme an das Gehäuse abgeleitet. Obwohl die Lichtausbeute der LED-Beleuchtung also doppelt so hoch ist, wie die der Glühlampe, muss ein Gehäuse entwickelt werden, das fast dreimal soviel Wärmeabfuhr sicherstellt.

Ein anderer wichtiger Unterschied zwischen Glühlampen- und LED-Lichtquellen besteht darin, dass sich die Glühlampe, um zu arbeiten, auf hohe Temperaturen verlässt (der Glühfaden leuchtet schließlich in Weißglut), während die LED-Lebensdauer stark beeinträchtigt wird, wenn die Sperrschichttemperatur im Halbleiterkristall der LED über 100°C ansteigt.

Sperrschichttemperatur	<100°C	100-115°C	115-125°C	>125°C
LED-Lebensdauer B50/50% der Überlebensrate	1	3	7*	10*

Tabelle 8.2: LED-Lebensdauer gegenüber Sperrschichttemperatur

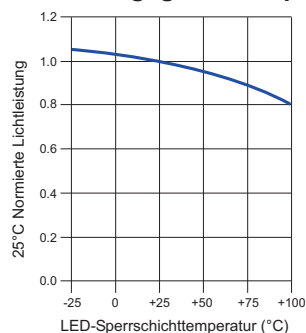


Abb. 8.21: LED-Lichtstrom gegenüber LED- Sperrschichttemperatur

Auch Hochleistungs-LEDs verlieren Lichtausbeute mit ansteigender Sperrschichttemperatur. Die im Datenblatt angegebenen Lichtleistungszahlen werden normalerweise nur für 25°C angegeben.

Bei einer Sperrschichttemperatur von 65°C fällt die Lichtleistung normalerweise um 10%, und bei 100°C gehen 20% Helligkeit verloren (Abb. 8.21).

Somit läuft eine mit Sachverstand entwickelte LED-Lampe bei einer maximalen Temperatur des LED-Grundkörpers von ca. 65°C. Eine Möglichkeit sicherzustellen, dass die LED-Temperatur nicht auf einen zu hohen Wert ansteigt, besteht darin bei ansteigender Temperatur ein Derating anzuwenden. Das nächste Kapitel gibt einige praktische Beispiele.

8.9 Temperatur-Derating

Eine LED kann nur konsequent unter Volleistung laufen, wenn die Wärmeableitung adäquat ist und die Umgebungstemperatur innerhalb der angemessenen Grenzen bleibt. Steigt die Temperatur des LED-Grundkörpers auf einen zu hohen Wert an, müssen Vorkehrungen getroffen werden, um die interne Verlustleistung zu reduzieren.

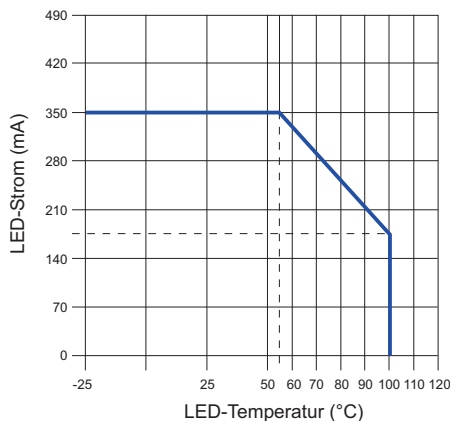


Abb. 8.22: Typische LED-Temperatur-Derating-Kurve

Abb. 8.22 zeigt ein ideales LED-Strom vs. Temperatur-Verhältnis. Bis zur vom Hersteller angegebenen maximalen Arbeitstemperatur bleibt der LED-Strom konstant. Wenn die LED-Temperatur den Grenzwert überschreitet, wird der Strom und somit die Leistung reduziert, die LED wird gedimmt um sie vor Überhitzung zu schützen. Diese Kurve nennt man "Derating-Kurve"; sie hält die LED im Betrieb innerhalb ihrer sicheren Verlustleistungsgrenzen. Die 55°C "Schwellen"-Temperatur in obigem Diagramm ist die Grundkörper- oder Kühlkörpertemperatur – die LED selbst wird normalerweise 15°C wärmer (d. h. 70°C) und die interne Sperrschichttemperatur nahezu 35°C wärmer (d. h. 90°C). Somit ist 55°C eine sichere Leistungsgrenze, obwohl sie auf ein Maximum von 65°C für leistungsfähige LED-Lampen erhöht werden könnte.

8.9.1 Hinzufügen von automatischem Temperatur-Derating zu einem LED-Treiber

Verfügt ein LED-Treiber über einen Dimmeingang, dann kann man leicht einen externen Temperatursensor und einige externe Stromkreise hinzufügen, um die gewünschte Derating-Kennlinie wiederherzustellen, wie in Abb. 8.22 gezeigt.

Der LED-Treiber der Serie RCD-24 von Recom verfügt über zwei verschiedene Dimmeingänge und ist somit ein idealer Kandidat, um die verschiedenen Möglichkeiten zu erklären, wie Übertemperaturschutz in einen LED-Treiberkreis integriert werden kann.

8.9.2 Übertemperaturschutz durch Verwendung eines PTC-Widerstands

Ein Thermistor ist ein Widerstand, der seine Größe mit der Temperatur ändert. Steigt der Widerstand mit ansteigender Temperatur an, weist er einen positiven Temperaturkoeffizienten (PTC) auf. Es gibt PTC-Thermistoren mit einer sehr nichtlinearen Kennlinie (Abb. 8.23).

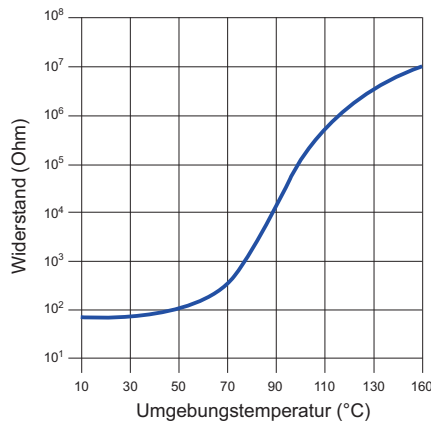


Abb. 8.23: Typische Widerstand-Temperatur-Kurve von PTC-Thermistoren

Solange die Temperatur unter einem gegebenen Schwellenwert bleibt, in diesem Fall 70°C, weist der PTC-Thermistor einen relativ stabilen Widerstand in der Größenordnung von ca. hundert Ohm auf. Über diesem Schwellenwert steigt der Widerstand sehr schnell an: bei 80°C beträgt der Widerstand 1kΩ; bei 90°C 10kΩ und bei 100°C 100kΩ.

Praktischerweise sind viele PTC-Thermistoren auch mit einer mechanischen (Schraub)Befestigung erhältlich, die zur Temperaturüberwachung sehr leicht am Kühlkörpergehäuse der LED-Beleuchtung angebracht werden kann.

Dieses Ansprechverhalten kann verwendet werden, um eine sehr einfache, preiswerte und zuverlässige Übertemperaturschutzschaltung unter Verwendung eines analogen Dimmeinganges der LED-Treiber der Serie RCD-24 herzustellen (Abb. 8.24 und Abb. 8.25).

Der analoge Dimmeingang wird durch eine externe Spannung angesteuert. Wenn die Eingangsspannung festgelegt ist, sind lediglich ein PTC-Thermistor und zwei Spannungsteilerwiderstände als einzige zusätzliche Komponenten erforderlich, um eine automatische Temperatur-Derating-Funktion zu realisieren.

Wenn andere Derating-Temperaturpunkte erforderlich sind, sind PTC-Thermoresistoren mit verschiedenen Temperaturschwellen in 10°C-Schritten von 60°C bis zu 130°C erhältlich. Daher ist es einfach eine Frage der Auswahl des richtigen Teils, um den Spezifikationen der LED zu entsprechen. Ist die Eingangsspannung variabel, dann kann eine Zenerdiode oder ein Linearspannungsregler hinzugefügt werden, um eine stabile Vergleichsspannung zu gewährleisten.

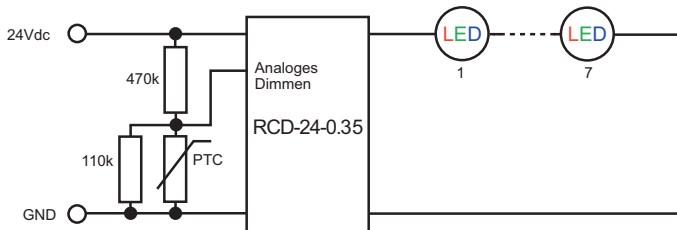


Abb. 8.24: PTC-Thermistorkreis

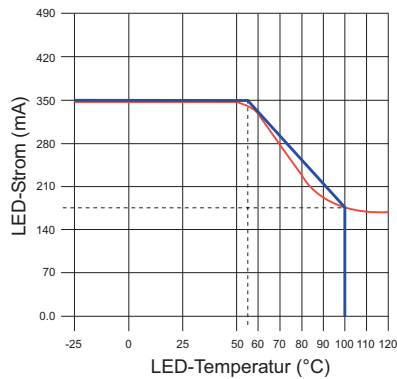


Abb. 8.25: Resultierende LED-Derating-Kurve (rote Linie)

8.9.3 Übertemperaturschutz unter Verwendung von analogem Temperatursensor-IC

Es sind viele IC-Temperatursensoren, die über einen temperaturlinearen Ausgang verfügen, erhältlich. Sie kosten nicht viel mehr als PTC-Thermistoren und haben den Vorteil, dass ihre Linearität und Offsets sehr genau sind, weshalb eine Temperaturüberwachung mit <1°C Auflösung möglich ist. Um eine nützliche Steuersignalspannung zu erzeugen, muss der Ausgang verstärkt werden. Daher werden sie am häufigsten mit einer Operationsverstärkerstufe verwendet.

Die unten vorgeschlagene Schaltung (Abb. 8.26) verwendet einen gebräuchlichen Temperatursensor IC und einen Operationsverstärker. Ähnliche Produkte sind von vielen anderen Herstellern erhältlich. Der Ausgang der Schaltung wird an den analogen Spannungsdimmingeingang der RCD-Treiber-Serie angeschlossen. Dieser Steuereingang dimmt die LED-Helligkeit je nach der am Pin vorhandenen Spannung linear.

Im Schaltkreis unten versorgt der Temperatursensor die lineare Ausgangsspannung je nach Umgebungstemperatur. Der Ausgang ist auf $10\text{mV}/^{\circ}\text{C} + 600\text{mV}$ vorkalibriert, sodass die Ausgangsspannung bei 55°C $1,15\text{V}$ beträgt. Der Operationsverstärkerblock enthält zwei Kleinleistungs-Operationsverstärker und eine präzise Spannungsreferenz von 200mV . Die Offset-Voreinstellung gleicht den Offset auf $1,15\text{V}$ ab, und die Verstärkung ist so eingestellt, dass die LED bei 100°C mit 50% Nennstrom läuft. Der Vorteil dieser Schaltung besteht darin, dass nur ein Aufbau benötigt wird, um verschiedene LED-Charakteristiken von unterschiedlichen Herstellern zu kompensieren, da der Maximalwert der Derating-Kurve einstellbar ist.

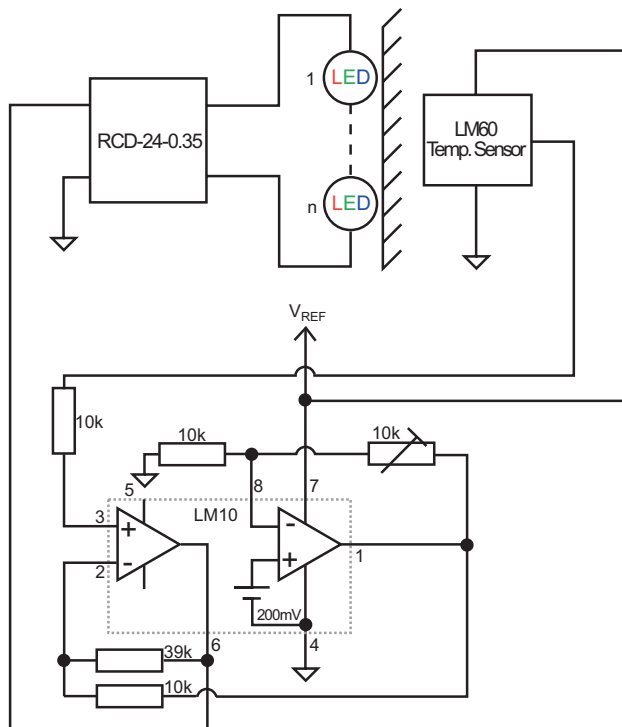


Abb. 8.26: Analoge Übertemperaturschutzschaltung

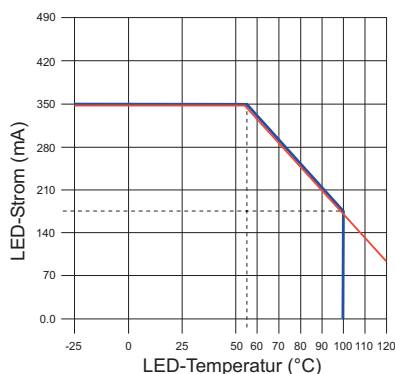


Abb. 8.27: Resultierende LED-Derating-Kurve (rote Linie)

8.9.4 Übertemperaturschutz unter Verwendung eines Mikrocontrollers

Die zweite Möglichkeit des Dimmeingangs der RCD-Serie ist der PWM-Eingang. Impulsbreitenmodulation verwendet ein digitales Signal zur Steuerung der Helligkeit der LEDs, das die LEDs so schnell ein- und ausschaltet, dass dies für das menschliche Auge nicht sichtbar ist. Wenn die LED längere Zeit AUS bleibt als EIN, erscheint sie gedimmt. Wenn die LED längere Zeit EIN als AUS bleibt, dann erscheint sie hell. Der PWM-Eingang antwortet auf den Logikeingang und ist daher für die Anbindung eines Digitalreglers ideal.

Es gibt einige ICs, die eine Temperatur direkt in ein PWM-Signal umwandeln (z. B. einige Lüfterregler, MAX6673, TMP05 usw.). Jedoch ist etwas eingebaute Intelligenz normalerweise erforderlich, um die Temperaturschwelle einzustellen und das PWM-Signal auf die Derating-Kurve der LED abzustimmen. Deshalb ist es häufig einfacher, einen Mikrocontroller zu verwenden.

Der Schaltungsvorschlag unten (Abb. 8.28) verwendet einen Mikrocontroller, um bis zu acht LED-Treiber zu überwachen und zu steuern. Da nur sechs Ein- und Ausgangsanschlüsse verwendet werden, kann die Schaltung leicht erweitert werden, um mehr LED-Treiber zu steuern, oder es könnte durch Verwendung der freien Ports ein Fernüberhitzungsalarm hinzugefügt werden.

In diesem Beispiel wird die Temperaturmessung mittels MAX6575L/H ICs, die Kleinleistungs-Temperatursensoren sind, realisiert. Bis zu acht Temperatursensoren können eine Dreidraht-Schnittstelle teilen. Temperaturen werden mittels Messung von Zeitverzögerung zwischen dem durch den Mikroprozessor initiierten Triggerimpuls und der Abfallflanke der nachfolgenden Impulsverzögerungen abgelesen, wie von den Geräten ermittelt. Verschiedene Sensoren an derselben Ein- und Ausgabeleitung verwenden verschiedene Multiplikatoren für die Wartezeit, um eine Überlappung der Signale zu vermeiden. Eine ähnliche Konstruktion könnte ebenso leicht mit anderen Temperatursensoren von verschiedenen Herstellern eingebaut werden – z. B. der TPM05 im Verkettungsmodus.

Der ansprechbare Kleinleistungs-Latchflipflop 74HC259 kann mit einem Rückstellimpuls zurückgesetzt werden, womit alle LED-Treiber eingeschaltet werden. Der Mikroprozessor kann dann jeden Ausgang nach einer angemessenen Zeitverzögerung individuell einstellen, um acht PWM-Signale zu erzeugen und jeden LED-Treiber unabhängig zu steuern.

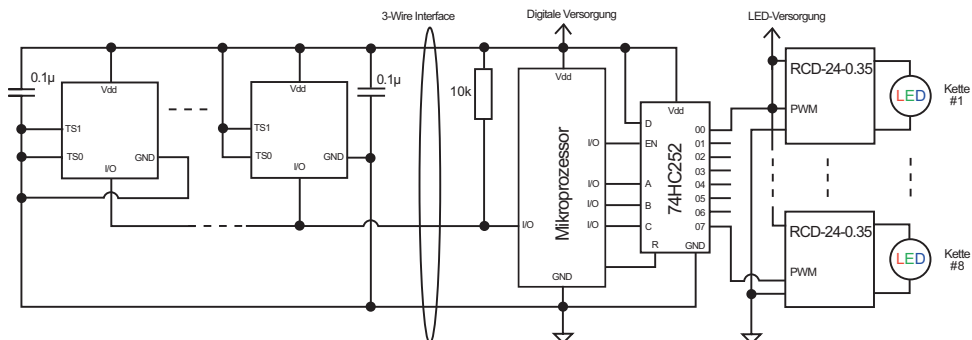


Abb. 8.28: Mikroprozessorgesteuerter PWM-Regler für bis zu acht LED-Treiber

8.10 Korrektur der Helligkeit

So wie die Temperaturmessung in einem Regelkreis verwendet werden kann, um die LED-Temperatur konstant zu halten, kann ein Lichtaufnehmer verwendet werden, um die Lichtleistung der LED konstant zu halten.

Alle LEDs verlieren Lichtausbeute im Laufe der Zeit:

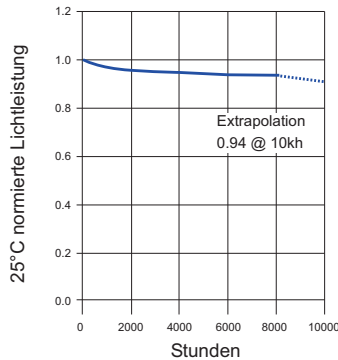


Abb. 8.29: Verlust von Lichtleistung im Laufe der Zeit

Wenn in einem Zimmer eine LED-Lampe angebracht wird und dann zwei Monate später eine identische Lampe hinzugefügt wird, wird die neue Lampe fast 5% heller sein.

Eine Lösung dieses Problems besteht darin, die Nennlichtleistung durch Verwendung eines Lichtaufnehmers auf 95% herabzusetzen, z. B. eine Fotodiode wie in Abb. 8.30 gezeigtem Beispiel. Die Anschlussleitungen der Fotodiode müssen kurz gehalten werden, um das Eindringen von zu viel Rauschen im Schaltkreis zu vermeiden. R_1 sollte so gewählt werden, dass die Ausgangsspannung des Rail-to-Rail-Operationsverstärkers ICL7611 ca. 200mV beträgt, wenn die LED-Lampe neu ist.

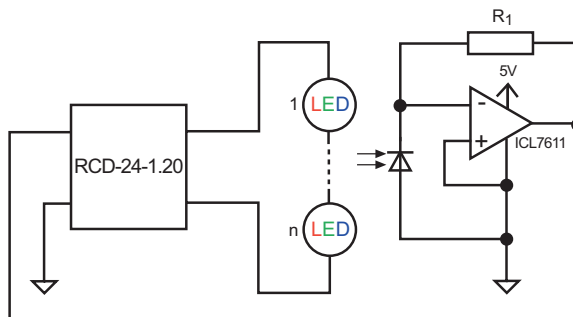


Abb. 8.30: Lichtaufnehmer-Rückkopplungskreis

Da die LED die Lichtstromeffizienz mit der Zeit verliert, erhöht der Rückkopplungskreis den LED-Strom automatisch, um dies zu kompensieren. Das in Abb. 8.30 gezeigte Schaltungskonzept kann modifiziert werden, wenn sowohl eine stabile maximale Lichtleistung als auch Dimmen erforderlich sind.

Der LED-Treiber der RCD-Serie ist nahezu einzigartig, da er über zwei Dimmeingänge verfügt, die beide gleichzeitig verwendet werden können.

Somit kann der analoge Dimmeingang für die Korrektur der LED-Helligkeit verwendet werden, während der PWM-Eingang zum unabhängigen Dimmen von LEDs eingesetzt werden kann.

Abb. 8.31 zeigt ein Schaltkreiskonzept, das eine Track-and-Hold-Technologie verwendet, um den Rückkopplungsspannungspegel der Helligkeitskorrektur zu speichern, während die LED eingeschaltet ist, den Pegel aber ignoriert, wenn die LED AUS ist. Somit ist die an den RCD übertragene Rückkopplungsspannung unabhängig vom PWM-Dimmeingang.

Eine durch den Widerstand mit $10\text{k}\Omega$ und den Kondensator mit 10nF gebildete geringe Verzögerung stellt sicher, dass die Ansprechzeit des LED-Treiberausgangs beachtet wird, bevor die Operationsverstärker-Ausgangsspannung gemessen und auf dem Kondensator mit 470nF gespeichert wird.

Die genauen Bauteilwerte müssen für einzelne Anwendungen eventuell optimiert werden.

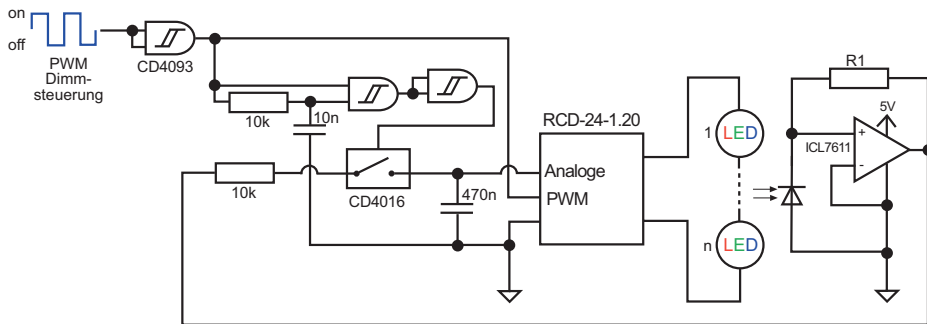


Abb. 8.31: Dimmbarer Lichtaufnehmerrückkopplungskreis

Eine andere verbreitete Anwendung für optische Rückführung ist ein Umgebungslichtsensor. Die Idee besteht darin, keine konstante LED-Lichtleistung zu haben, sondern die Umgebungslichtpegel zu messen und die LEDs tagsüber zu dimmen und die Helligkeit der LEDs allmählich zu erhöhen, wenn es dunkelt wird, um den Lichtstrom konstant zu erhalten.

Ein verbreiteter, preiswerter Lichtaufnehmer ist ein LDR oder lichtempfindlicher Widerstand. Dieser weist ein lineares Ansprechen auf den natürlichen Logarithmus des Helligkeitspegels auf ($R = \text{Lux} \times e^{-b}$) und kann mit einigen Vorwiderständen leicht eingesetzt werden, um das gewünschte Umgebungslichtniveau einzustellen.

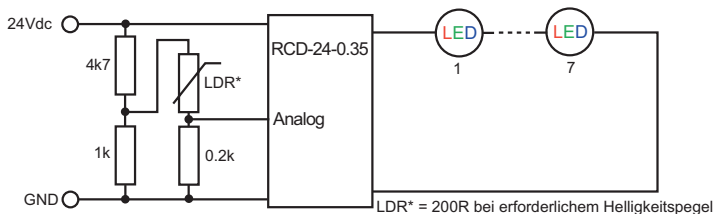


Abb. 8.32: Umgebungslichtsensor-Rückkopplungskreis

8.11 Einige Schaltungskonzepte unter Verwendung eines RCD-Treibers

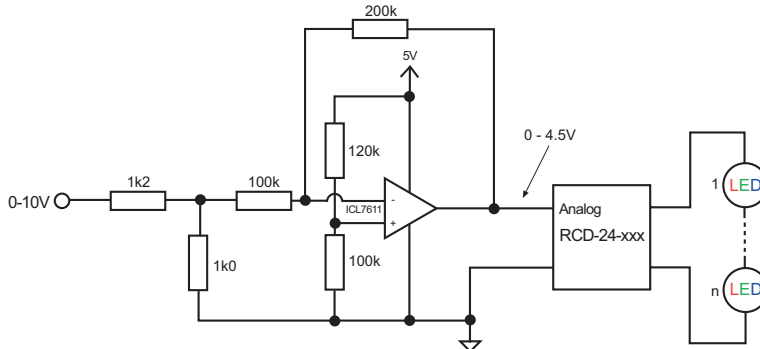


Abb. 8.33: 0 – 10V Dimmsteuerung (0V = 0 %, 10V = 100 %)

So funktioniert es:

Der Rail-to-Rail-Operationsverstärker wird als invertierender Verstärker konfiguriert. Der nichtinvertierende Eingang wird bei einer "virtuelle Masse"-Spannung von 2,25V durch die 120k- und 100k-Widerstandsteilerkette gehalten. Beträgt die Eingangsspannung 0V, dann muss die Operationsverstärkerspannung 4,5V betragen, damit die Spannung am invertierenden Eingang des Opamp auch 2,25V beträgt. Beträgt die Eingangsspannung 10V, dann senkt der Eingangsspannungsteiler von 1k2 und 1k0 die Eingangsspannung der Verstärkerschaltung auf 4,5V. Nur wenn die Ausgangsspannung 0V beträgt, sind die Eingänge am Operationsverstärker ausgeglichen.

Eine geringfügige Modifikation dieser Schaltung erlaubt Steuerspannungen von 1 bis 10V.

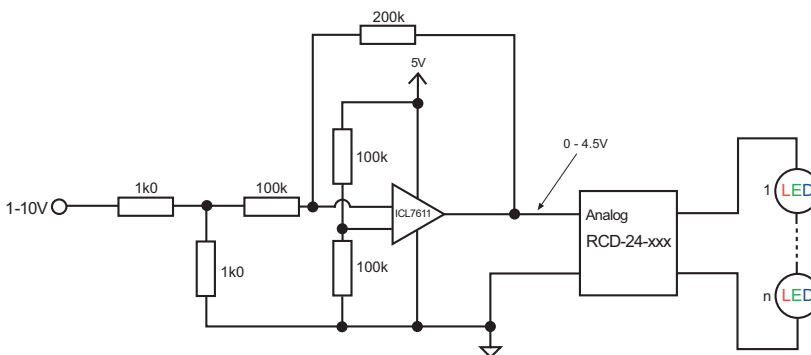


Abb. 8.34: 1 – 10 V - Dimmsteuerung (1 V = 0 %, 10 V = 100 %)

So funktioniert es:

Der Rail-to-Rail-Operationsverstärker wird als invertierender Verstärker konfiguriert. Der nichtinvertierende Eingang wird bei einem virtuellen Erdpotential von 2,5V durch die 100k-Widerstandsteilerkette gehalten. Beträgt die Eingangsspannung 1V, dann senkt der 1k 0-Eingangsspannungsteiler die Spannung am invertierenden Eingang des Opamp auf 0,5V. Nur wenn die Ausgangsspannung 4,5V beträgt, sind die Eingänge am Operationsverstärker ausgeglichen.

Beträgt die Eingangsspannung 10V, dann senkt der Eingangsspannungsteiler die *Spannung am invertierenden Eingang des Opamp auf 5V. Nur wenn die Ausgangsspannung 0V beträgt, sind die Eingänge am Operationsverstärker ausgeglichen.

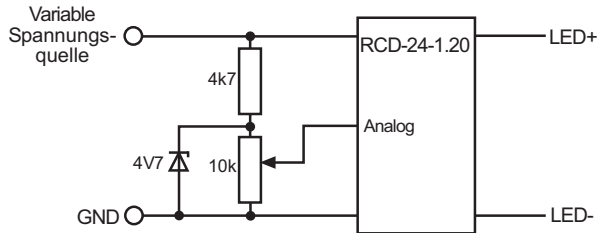


Abb. 8.35: Einfacher Potentiometer-Dimmer

So funktioniert es:

Wenn die Eingangsspannung (z. B. eine Batterie) nicht stabilisiert ist, dann muss die Eingangsdimmspannung geregelt werden. Eine einfache Zenerdiode ist alles, was benötigt wird, obwohl auch ein 5-V-Regler verwendet werden könnte, wenn eine sehr genaue Dimmsteuerung benötigt wird.

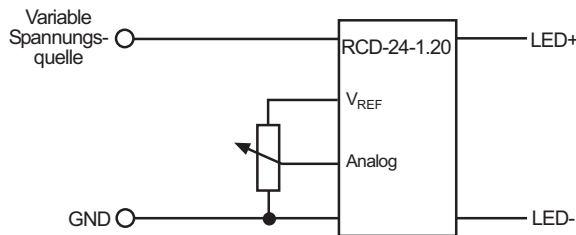


Abb. 8.36: Genauer Potentiometer-Dimmer

PWM zu Analog:

Der analoge Dimmeingang kann auch mit einem PWM-Signal angesteuert werden. Dies vermeidet die maximalzulässige Frequenz am PWM-Eingang und ist für Mikrocontroller, die auf internen Timern basierende PWM-Ausgänge besitzen und nicht einfach niederfrequente PWM-Signale ausgeben können, nützlich.

Der Nachteil dieses Verfahrens besteht darin, dass die Ansprechzeit des LED-Ausgangs auf eine Änderung des Dimmniveaus länger ist, weil der Kondensator bis auf den Mittelwert des Eingangsspannungspegels geladen oder entladen werden muss.

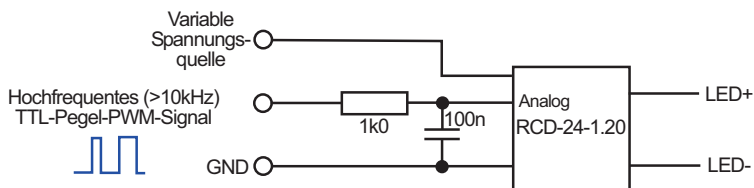


Abb. 8.37: PWM auf Analog-Dimmsteuerung

Handgesteuerte PWM-Generatoren:

PWM-Signale verfügen über den Vorteil, dass die Signale über sehr lange Entfernungen ohne Verlust gesendet werden können und dass sie größtenteils nicht für interne Störungen anfällig sind.

Es ist manchmal besser das PWM-Impuls-Pausen-Verhältnisses über eine Handsteuerung (z. B. Potentiometer) einzustellen, als das Signal in digitaler Form zu erzeugen. Die folgenden beiden Schaltungen sind Beispiele für einfache, universelle PWM-Generatoren für die RCD-Serie:

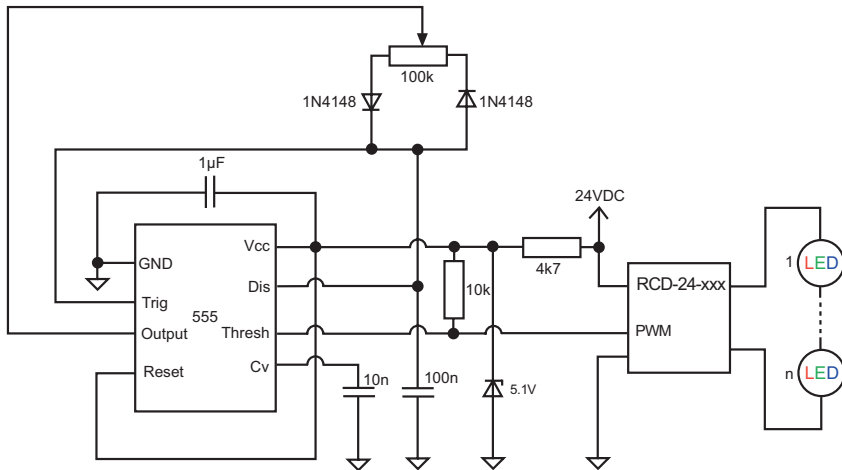


Abb. 8.38: 555-gestützter PWM-Generator (Potentiometerregelung)

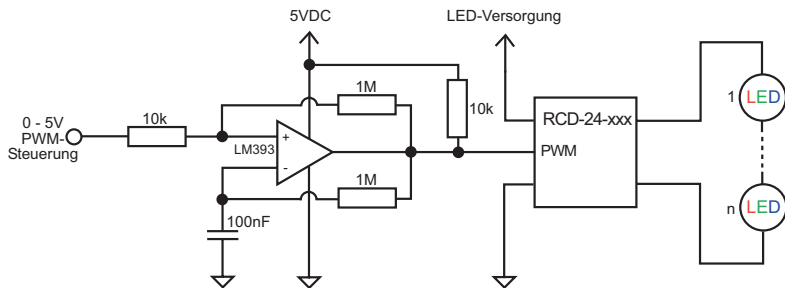


Abb. 8.39: Komparatorgestützter PWM-Generator (Potentiometer- oder Spannungsregelung)

Schalten von LED-Ketten:

Dieses Anwendungskonzept ermöglicht das ein- bzw. ausschalten von vier LED-Ketten mit einem 4-Bit-Steuersignal, ohne die aktiven Stränge zu übersteuern.

Die Ketten können auch unabhängig voneinander gedimmt werden, ggf. unter Verwendung eines PWM-Dimmeingangs.

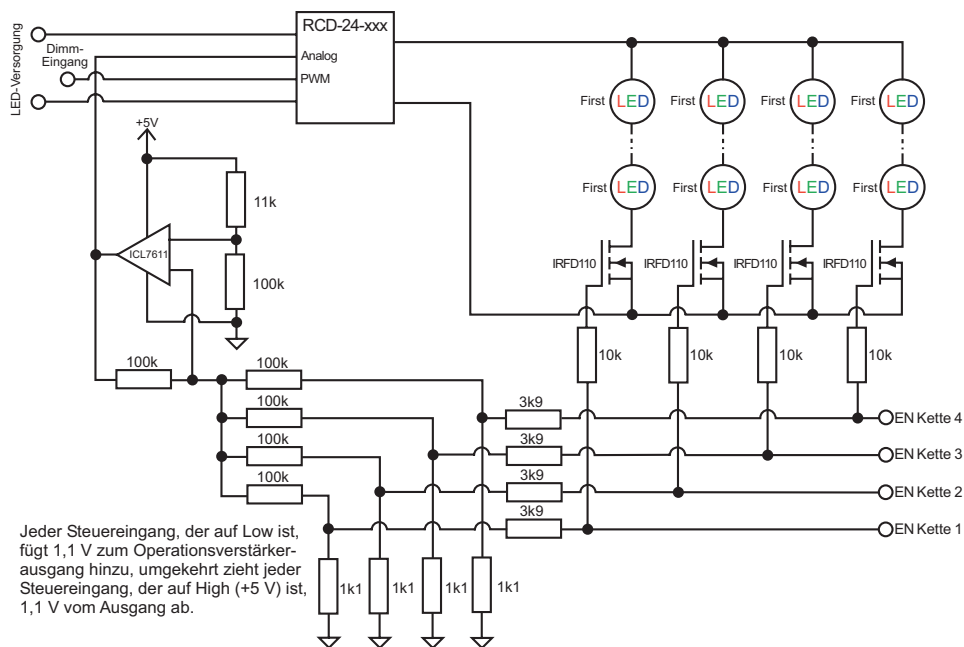


Abb. 8.40: Schaltbare LED-Ketten mit LED-Stromausgleich

S1	S2	S3	S4	Operations- verstärker Vout	LED-Strom
0	0	0	0	4.4V	0mA
0	0	0	1	3.3V	250Ma
0	0	1	0	3.3V	250mA
0	0	1	1	2.2V	500mA
0	1	0	0	3.3V	250mA
0	1	0	1	2.2V	500mA
0	1	1	0	2.2V	500mA
0	1	1	1	1.1V	750mA
1	0	0	0	3.3V	250mA
1	0	0	1	2.2V	500mA
1	0	1	1	1.1V	750mA
1	1	0	0	2.2V	500mA
1	1	0	1	1.1V	750mA
1	1	1	0	1.1V	750mA
1	1	1	1	0V	1000mA

Tabelle 8.3: 4-Bit-Steuerelemente

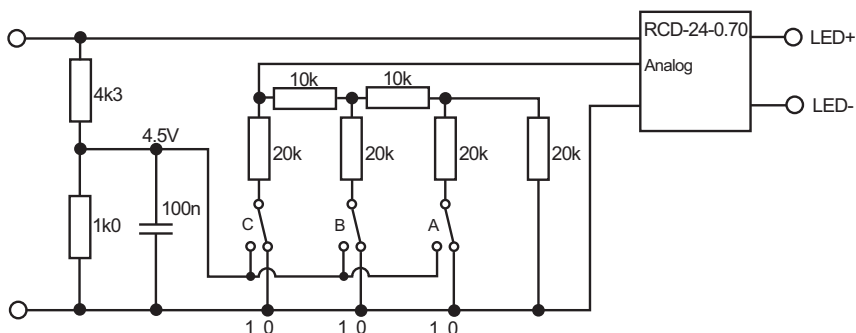


Abb. 8.41: Schaltkreis LED-Hintergrundbeleuchtung

C	B	A	Analogeingang	LED-Strom
0	0	0	0.00V	700mA
0	0	1	0.64V	600mA
0	1	0	1.28V	500mA
0	1	1	1.93V	400mA
1	0	0	2.25V	300mA
1	0	1	3.21V	200mA
1	1	0	3.86V	100mA
1	1	1	4.50V	000mA

Tabelle 8.4: 3-Bit-Zweipunkteingang

So funktioniert es:

Das R2R-Widerstandsnetzwerk wandelt den 3-Bit-Binäreingang in eine 8-stufige Steuerungspannung um.

Der Vorteil dieser Schaltung gegenüber der vorhergehenden besteht darin, dass keine aktiven Komponenten benötigt und die R2R-Widerstandskette auf eine beliebige Anzahl von Bits erweitert werden kann, wenn eine höhere Auflösung erforderlich ist. R2R-Netzwerke sind als fertige Widerstandsnetzwerke in einem kompakten SIP-Format erhältlich.

Dieser Typus eines Schaltkreises wird häufig als Regler für Hintergrundbeleuchtung verwendet, da acht Helligkeitsniveaus für die meisten Anwendungen in diesem Bereich ausreichend sind.

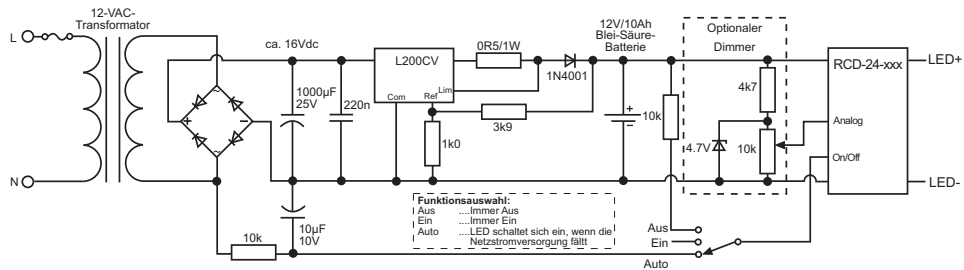


Abb. 8.42: Diese Schaltung verwendet einen Niederspannungs-Netztransformator zum Laden einer Blei-Säure-Reservebatterie

Eine lineare Vorregelung begrenzt sowohl die Ladespannung der Batterie als auch den maximalen Ladestrom, um es demselben Schaltkreis zu ermöglichen, sowohl eine entladene Batterie wiederaufzuladen als auch die laufende Nachladung einer vollständig geladenen Batterie durchzuführen. Der LED-Treiber kann umgeschaltet werden, um die LED-Beleuchtung automatisch einzuschalten, wenn ein Netzeingang ausfällt.

So funktioniert es:

Der 12-VAC-Ausgang aus dem Transformator ist gleichgerichtet und geglättet, um den Linearregler mit ca. 16 VDC zu versorgen. Der Regler wird so eingestellt, dass er 13,8 V bei 1A maximalem Strom liefert, um eine 12-V-Blei-Säure-Batterie zu laden. Die Diode am Ausgang des L200 sperrt den Strom aus der Batterie in den Regler, wenn der Netzeingang abgeschaltet ist; da das Reglerreferenzeingangssignal jedoch nach der Diode abgenommen wird, hat es keine Auswirkung auf die Ausgangsspannung.

Der Freigabeeingang des LED-Treibers des RCD kann auf drei Positionen geschaltet werden:

AUS Der EIN/AUS-Eingang wird durch einen hochohmhigen Widerstand auf 12V gezogen. Dies erlaubt es einem 12-V-Signal, einen 5-V-Eingang zu steuern. Diese Methode wird anstelle eines Spannungsteilers gewählt, weil der Teiler die Batterie mit der Zeit entladen würde.

EIN Der Steuereingang wird offen gelassen, daher sind die LEDs EIN voreingestellt.

AUTO Solange der Netzeingang aktiv ist, wird der 12-VAC-Ausgang durch den Widerstand mit 10kΩ und den Kondensator mit 10µF geglättet, um der Mittelspannung ca. 6VDC zu liefern, was den LED-Treiber blockiert. Wenn das Versorgungsnetz ausfällt oder ausgeschaltet wird, fällt diese Mittelspannung auf Null, und der Treiber wird aktiviert.

Einfacher RGBW-Mischer:

Die RGB-Mischer-Anwendungsschaltung, wie im Recom-Datenblatt angegeben, kann durch Einführung von RGBW-LEDs leicht erweitert werden.

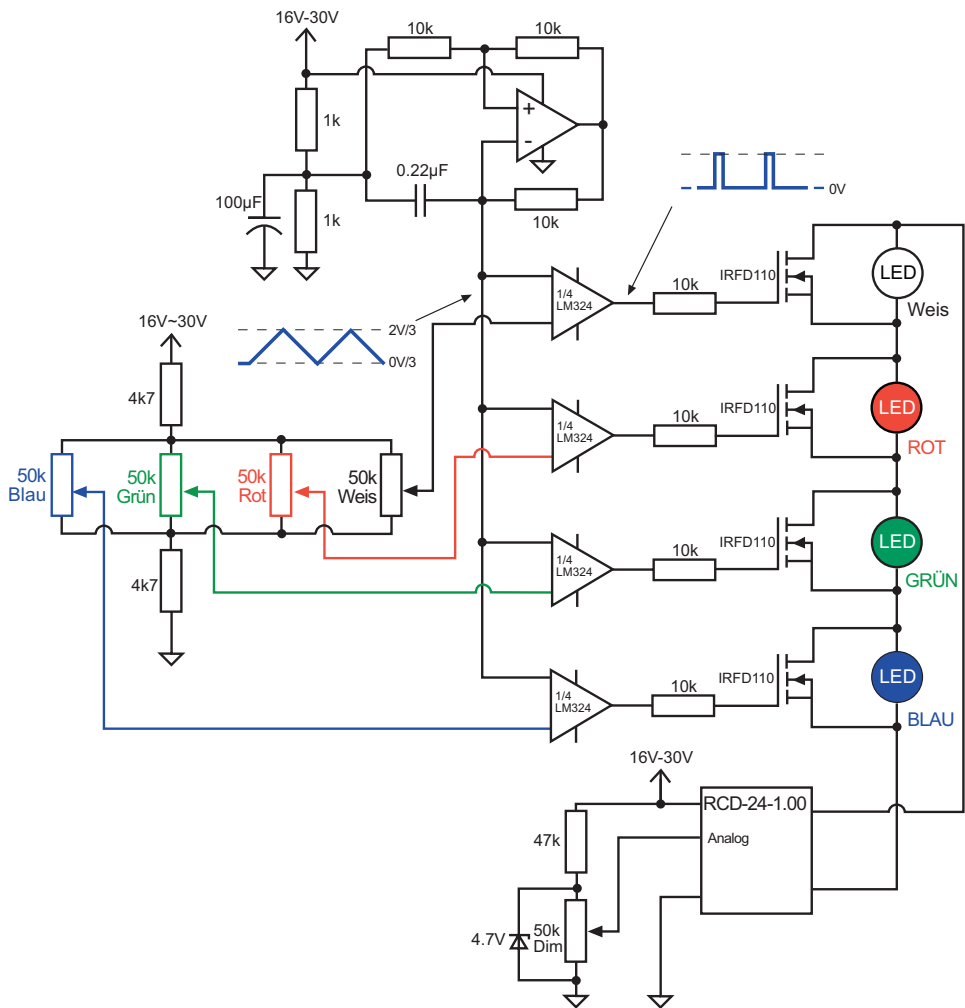


Abb. 8.43: RGBW-Mischer

9. DC/DC-Anwendungsideen

9.1 Einleitung

Viele Anwendungen erfordern eine DC/DC-Spannungswandlung. So viele, dass nach Einschätzungen der Weltmarkt bis zum Jahre 2020 die 35-Milliarden-Dollar-Marke überschritten haben wird. Für viele Anwendungs-Entwickler ist aber der DC/DC-Wandler eine "Blackbox"; eine Komponente zur Erfüllung einer Funktion - genau wie andere Komponenten auch, wie z. B. eine Induktivität oder ein Transistor. Als Allzweck-Funktionsblock kann man DC/DC-Wandler überall bzw. da, wo man sie gerade benötigt, verwenden; es gibt keinen "typischen" Einsatzbereich. Dieses abschließende Kapitel untersucht einige ungewöhnlichere Möglichkeiten, wie DC/DC-Wandler verwendet werden können und soll veranschaulichen, wie breit der Anwendungsbereich für DC/DC-Wandler ist.

9.2 Polaritätswechselung

Ein isolierter DC/DC-Wandler besitzt einen potentialfreien Ausgang. Gleichermäßen legitim ist es, sich vorzustellen, dass er einen potentialfreien Eingang hat. Deshalb kann jeder isolierte DC/DC-Wandler verwendet werden, um die Polarität einer Versorgungsspannung zu invertieren. Wenn keine Isolation benötigt wird, aber ein gemeinsamer Referenzpunkt, dann kann jeder Ausgang mit jedem Eingang sowie mit jeder gewünschten Referenzspannung verbunden werden. Die nächste Abbildung zeigt einige in Frage kommende Konfigurationen, um aus positiven Eingangsspannungen negative Ausgangsspannungen zu erzeugen und umgekehrt.

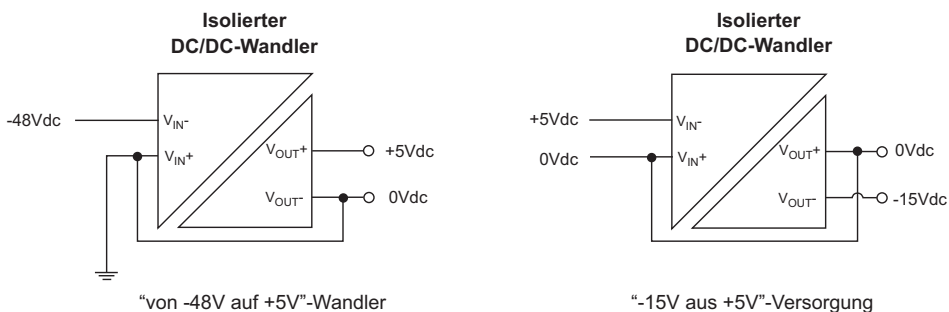


Abb. 9.1a + 9.1b: Beispiele der Umpolung unter Verwendung isolierter Wandler

Eine Anwendung für „von -48V auf +5V“-Wandler könnte ein GSM-Sendermodul sein, das aus einer Telefonleitung gespeist wird (zusätzliche Begrenzungsschaltungen für die Eingangsspannung könnten benötigt werden, um gegenüber dem Anrufsignal von 90V abzusichern). Eine Anwendung für die Versorgung +5 auf -15V kann die negative Spannungsversorgung für Analogschaltungen wie Operationsverstärker oder Analog-Digital-Wandler gewährleisten.

Es ist auch möglich, einen Schaltregler zu verwenden, um aus einer positiven Eingangsspannung eine negative Ausgangsspannung bereitzustellen. Dies funktioniert, weil der Abwärtsschaltregler die Ausgangsspannung auf den gemeinsamen Pin bezieht, aber nur die Spannungsdifferenz betrachtet, nicht jedoch die absolute Spannung. Wenn der Ausgang geerdet ist und der gemeinsame Pin unverbunden gelassen wird, wird der

gemeinsame Pin negativ, um die korrekte Spannungsdifferenz aufrechtzuerhalten. Abb. 9.1a zeigt ein Anwendungsbeispiel, in dem ein positiver 5-V-Schaltregler verwendet werden kann, um aus dem positiven 5-V-Eingang einen negativen 5-V-Ausgang zu erzeugen: eine Schaltungsvariante die mit einem linearen Spannungsregler unmöglich ist.

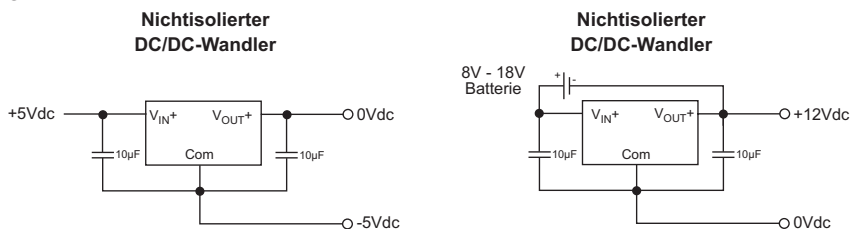


Abb. 9.2a und 9.2b: Beispiele nicht isolierter Positiv-zu-Negativ-Wandler

Eine weitere Anwendung für den nichtinvertierenden Positiv-zu-Negativ-Wandler ist als Ausgangsspannungsregler der 12-V-Batterie (Abb. 9.2b). Da der Schaltregler die Eingangsspannung wie die Differenz zwischen einer positiven Eingangsspannung und der erzeugten negativen Ausgangsspannung betrachtet, kann der Schaltkreis eine stabile 12-V-Ausgangsspannung gewährleisten, selbst wenn die Batteriespannung nur 8V beträgt (weil die Eingangsspannung, soweit sie den Wandler betrifft, $8V + |12V| = 20V$ beträgt, was immer noch ausreichend ist, um einen 12-V-Ausgang zu regeln). Die obere Eingangsspannungsgrenze wird durch die sichere maximale Eingangsspannung des Schaltreglers minus der Ausgangsspannung festgelegt, das heißt $V_{MAX} - 2V$ (2V ist die Sicherheitsreserve) - $|V_{OUT}|$, was für einen RECOM R-7812-Regler $(32V - 2V) - 12V = 18V$ ergibt. Deshalb liefert dieser Schaltkreis einen stabilen 12-V-Ausgang aus einem Eingang von 8V bis 18V. Man beachte bitte, dass der Batterie-Minuspol der Anschluss für den +12V-Ausgang ist. Dieser ungewöhnliche Schaltkreis funktioniert nur, weil die Batterie in Bezug auf Masse schwebend ist. Wird ein Batterieladegerät während dem Betrieb dieser Schaltung angeschlossen, dann muss der Ausgang des Ladegeräts auch in Bezug auf Masse schwebend sein, um Kurzschlüsse zu vermeiden.

9.3 Spannungsverdoppler

Es gibt DC/DC-Anwendungen, bei denen keine Isolation erforderlich ist, aber eine höhere Ausgangsspannung als die Eingangsspannung benötigt wird. Das folgende Beispiel in Abb. 9.3 zeigt einen Spannungsverdoppler, der bei doppelter Leistung des DC/DC-Wandlers die doppelte Eingangsspannung erzeugt:

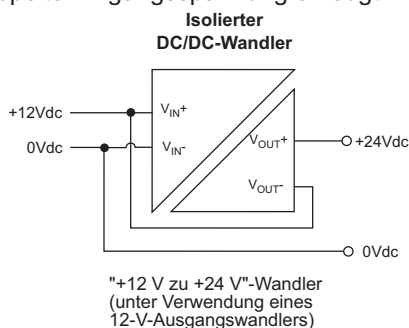


Abb. 9.3: Leistungsverdoppler

Wurde der DC/DC-Wandler auf 15 W (z. B. RP15-1212S) ausgelegt, dann liefert sein Ausgang bei 1,25A 12V. Diese Ausgangsspannung kommt jedoch zur Eingangsspannung von 12V hinzu. Deshalb betrachtet die Last 24V bei 1,25A bzw. 30W. Eine typische Anwendung besteht in solchen Fällen, in denen aus einer stabilen 12-V-Versorgung zum Speisen einer Pumpe oder Leistungsspule 24VDC erzeugt werden müssen, der zur Verfügung stehende Platz oder das Entwicklungsbudget jedoch einen größeren und teureren Leistungs-DC/DC-Wandler ausschließen. Da ein RP15 nur 1" x 1" groß ist, kann mit diesem Schaltkreis eine sehr kompakte 30-W-Stromversorgung realisiert werden.

9.4 Kombination von Schaltreglern und DC/DC-Wandlern

DC/DC-Wandler können kombiniert werden, um von den Vorteilen jeden Typs zu profitieren, oder in Reihe geschaltet werden, um die Ausgangsspannung oder Isolation zu erhöhen. Den Kombinationsmöglichkeiten sind fast keine Grenzen gesetzt. Daher zeigen wir hier nur einige wenige Beispiele.

Beispiel 1: Dieses Beispiel bezieht sich auf eine 5-V-Versorgung, die in einem Gabelstapler verwendet wird. Die Batteriespannung beträgt nominell 48V, und die Stromversorgung sollte in einer Anzeigetafel mit sehr begrenztem Platz eingebaut werden. RECOM stellt die Schaltregler-Serie R-78HB für hohe Eingangsspannungen her, wobei der Ausgangsstrom jedoch auf 0,5A begrenzt und damit viel zu niedrig für den für die Anzeigetafel erforderlichen Strom (1A) ist. Was benötigt wird, ist die Kombination einer hohen Eingangsspannung mit einem hohen Ausgangsstrom:

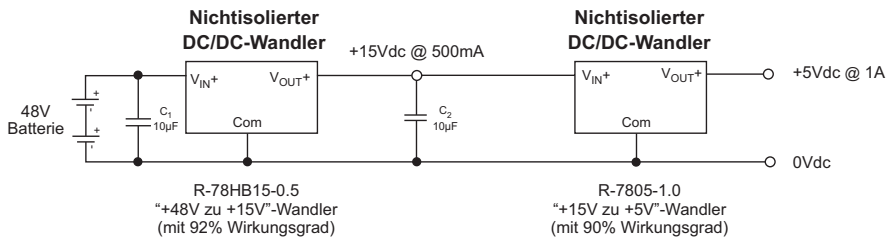


Abb. 9.4: Kaskaden-Schaltregler

Die Batteriespannung kann variieren von 65V beim Laden bis runter auf 20V wenn sie entladen ist. Dieser Spannungsbereich liegt noch innerhalb des Bereiches von 4:1 eines isolierten 5W-Wandlers (z. B. REC5-4805SRWZ in DIP24). Da bei dieser Anwendung aber keine Isolation benötigt wird, könnte die Lösung auch in der Verwendung zweier nichtisolierter Wandler in Serie liegen. Dies erfüllt dieselbe Funktion, ist aber preiswerter und benötigt weniger Grundfläche.

Der vordere Regler R-78HB15 senkt die Nennspannung der 48-V-Batterie wirksam auf 15V, begrenzt den Strom aber auf 500mA. Der zweite R-7805-Regler liefert 5V bei 1A, zieht aber nur 370mA vom ersten Regler. Der durchschnittliche Eingangsstrom für einen Schaltregler kann unter Verwendung des folgenden Verhältnisses berechnet werden:

$$I_{IN} = \frac{V_{OUT}}{V_{IN}} \times \frac{I_{OUT}}{\eta}, \text{ where } \eta = \text{efficiency}$$

Gleichung 9.1: Berechnung des Schaltreglereingangsstroms

Für dieses Beispiel zieht der R-7805 $I_{IN} = 5/15 \times 1/0,9 = 0,37 \text{ A}$ vom ersten Regler. Der Gesamtwirkungsgrad beträgt $> 82\%$, und die Stromversorgung benötigt nur $21 \times 12 \text{ mm}$ Platinenfläche (ca. $1/3$ der Größe eines DIP24-Gehäuses). Alternativ können diese beiden Regler flach liegend auf der PCB montiert werden, um eine Gesamthöhe von weniger als 8mm zu erreichen. Da die Leerlaufstromaufnahme unter 5mA liegt, ist kein Ein/Aus-Schalter notwendig.

Obwohl die Schaltregler der Serie R-78 für den Normalbetrieb keine externe Beschaltung erfordern, ist eine entsprechende Beschaltung, wie folgend, für diese Anwendung zu empfehlen. Der Eingangskondensator C_1 hilft, den Regler R-78HB15-0.5 vor schnellwirkenden Einschwingspannungen (Transienten) und Spitzen, hervorgerufen durch Laständerungen an der Batterie, zu schützen (sog. Load-Dump-Schutz). Ein geringwertiger Elektrolytkondensator mit hohem ESR arbeitet am besten, weil sein Innenwiderstand als Snubber-Netzwerk wirkt. Der zweite $10\text{-}\mu\text{F}$ -Kondensator C_2 ist erforderlich, weil der Regler R-7805-1.0 für Anwendungen mit hohen Anlaufströmen Spitzenströme von bis zu 3A liefern kann. C_2 hilft, diesen vorübergehenden Zusatzstrom zu liefern. Ein MLCC mit niedriger ESR wäre eine gute Wahl.

Beispiel 2: Wenn Isolation benötigt wird, kann der zweite Abwärtsschaltregler durch einen isolierten DC/DC-Wandler ersetzt werden. In diesem zweiten Anwendungsbeispiel sind die Anforderungen etwas anders. Die Anwendung ist eine Batterie-Ladezustandsüberwachung. Die Eingangsspannung sollte eine universelle Gleichspannung sein, sodass 12-V-, 24-V-, 28-V-, 36-V- oder 48-V-Batterien überwacht werden können. Der Ausgang muss ein isoliertes Stromschleifensignal mit 4 bis 20mA sein, dass über lange Entfernungen gesendet werden kann.

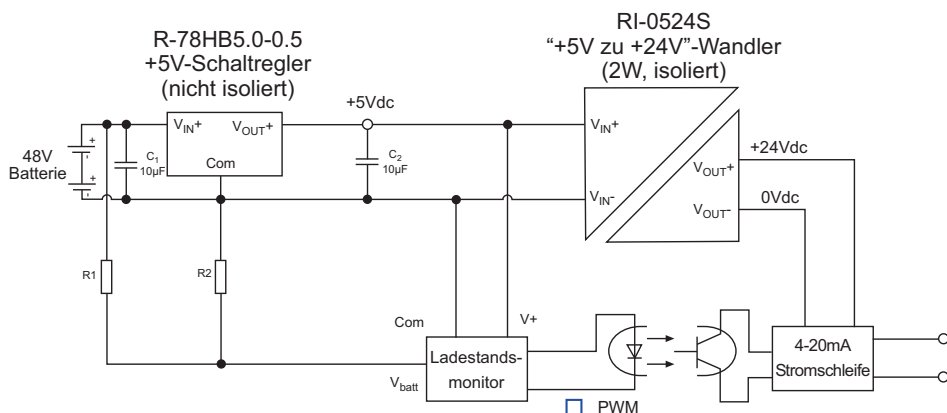


Abb. 9.5: Im Entladeanzeiger verwendeter Kaskaden-Wandler

In diesem Beispiel lässt der breite Eingangsspannungsbereich des Schaltreglers R-78HB5.0 (9-72VDC) die Anwendung mit vielen unterschiedlichen Batteriespannungen zu. Der stabile 5-V-Ausgang wird zum Speisen eines Niederspannungs-Ladestandsüberwachungs-ICs, der ein PWM-Signal proportional zur Batteriespannung erzeugt, verwendet. Das PWM-Signal ist optoisoliert und wird zur Steuerung eines Ausgangssignals von 4-20mA verwendet, das per Kabel über viele Kilometer gesendet werden kann. Die Versorgung der Stromschleife für 4-20mA wird durch einen isolierten 2-W-DC/DC-Klein-Wandler gespeist, der die 5V auf 24V anhebt und einen Strom von 83mA liefert.

Die gesamte Schaltung kann auf einer PCB realisiert werden, die kleiner ist als eine Streichholzschachtel ist.

Bis jetzt haben wir die Kaskadierung von DC/DC-Wandlern und Schaltreglern in Reihe und die Nutzung eines Schaltreglers als Frontregler behandelt. Es bleibt die Kombination, einen Schaltregler als Nachregler zu verwenden. Nachregelung ist eine sehr gängige Anforderung in Stromversorgungen. Der größte Vorteil eines Schaltreglers als Nachregler besteht darin, dass er ein Leistungsumformer ist, weshalb ein Niederspannungsausgang bei einem hohen Strom weniger Strom aus der höheren Gesamtversorgung zieht, nämlich direkt proportional zu der Spannungsdifferenz (siehe Gleichung 9.1).

Dieser Vorteil tritt nur bei Schaltreglern und nicht bei Linearreglern auf. Linearregler ziehen immer so viel Strom, wie sie liefern.

Beispiel 3: In diesem einfachen Beispiel werden Schaltregler verwendet, um Zwischenspannungen aus einem isolierten DC/DC-Wandlers Ausgang zu versorgen. Die Anforderungen des Schaltkreises sind:

+12V @ 0,4A

5V @ 1,5A und

-9V @ -0,2A,

alle separat geregelt und isoliert gegenüber der Versorgungsspannung von 24VDC.

Diese Anforderung kann durch die Verwendung von drei separaten DC/DC-Wandlern realisiert werden: einen 5-W-DC/DC für die +12-V-Versorgung, einen 7,5-W-DC/DC für die +5-V-Versorgung und einen 2-W-DC/DC für -9-V-Versorgung (Abb. 9.6). Die Stromversorgung benötigt 42 mm² Platinenfläche.

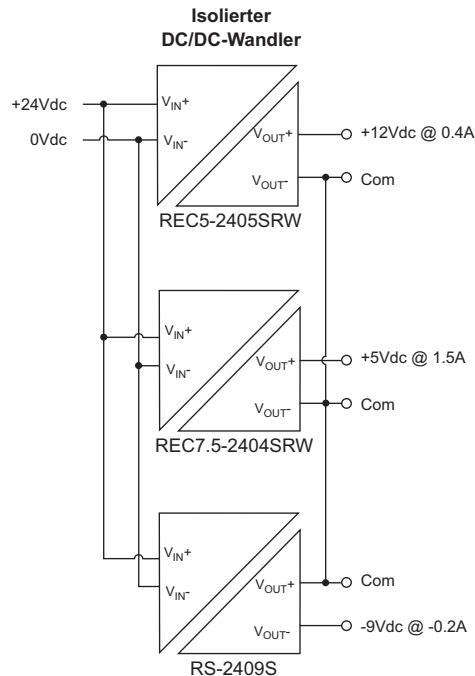


Abb. 9.6: Dreifache Ausgangsstromversorgung unter Verwendung von DC/DC-Wandlern

Obwohl die obige Lösung alle Stromversorgungsanforderungen erfüllt, benötigt der unten gezeigte alternative Schaltkreis unter Anwendung eines einzelnen DC/DC-Wandlers und Nachregelungs-Schaltregler jedoch weniger Platinenfläche (30 mm²) und ist preiswerter.

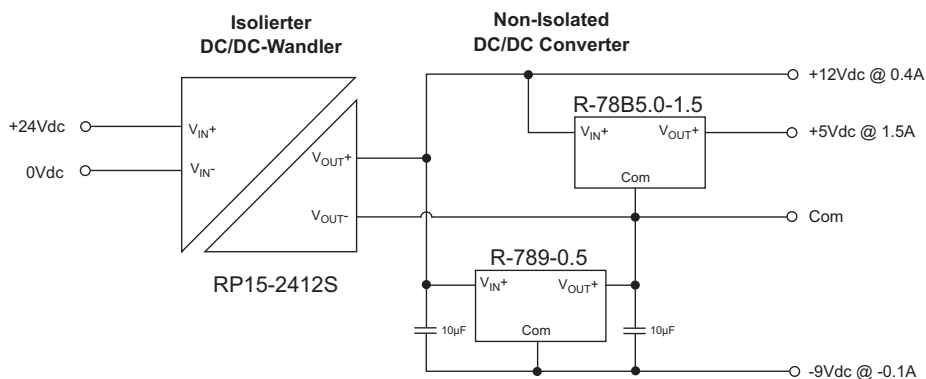


Abb. 9.7: Dreifache Ausgangsstromversorgung realisiert mit Schaltreglern

9.5 Schaltung von Wandlern in Serie

Generell sollten die Ausgänge von DC/DC-Wandler nicht parallel geschaltet werden um den Ausgangsstrom zu erhöhen. Die einzige Ausnahme ist, wenn sie über einen Lastverteilungseingang verfügen oder wenn ein separater Lastverteilungs-Regler (engl. load share controller) verwendet wird, um die Lastströme symmetrisch auf die einzelnen Wandler zu verteilen. DC/DC-Wandler können jedoch in Serie geschaltet werden, um die Ausgangsspannung und somit die Ausgangsleistung zu erhöhen. Abb. 9.8 zeigt, wie zwei DC/DC-Wandler in Serie geschaltet sind. Die Dioden werden benötigt, um entsprechende Kurzschlussicherheit zu gewährleisten. Es ist offensichtlich, dass ein asymmetrischer +/- -Ausgang unter Verwendung von zwei Wandlern mit unterschiedlichen Ausgangsspannungen erzeugt werden kann, wenn der Mittelanschluss als gemeinsamer Pin verwendet wird. Die Dioden werden immer noch benötigt, um Kurzschlüsse zwischen der positiven und negativen Versorgungsspannung zu bewältigen.

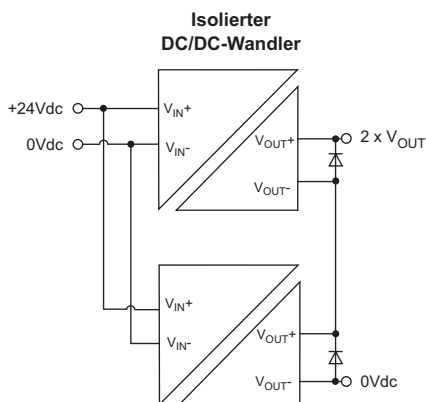


Abb. 9.8: Schaltung von DC/DC-Wandlern in Serie

Isolierte Wandler können „gestapelt“ (engl.: stacked) betrieben werden, um Hochspannungen zu erzeugen. Das folgende Anwendungsbeispiel bezieht sich auf eine Hochspannungs-Kleinleistungsversorgung für einen Ionisator. Jeder DC/DC-Wandler erzeugt 150V aus einer 12-V-Versorgung.

Der oberste Wandler hat einen ständigen Spannungsgradienten von 600 V an der Isolationsbarriere, deshalb sollte er mit einer Isolationsspannung von mindestens 2kVdc/1s ausgelegt werden.

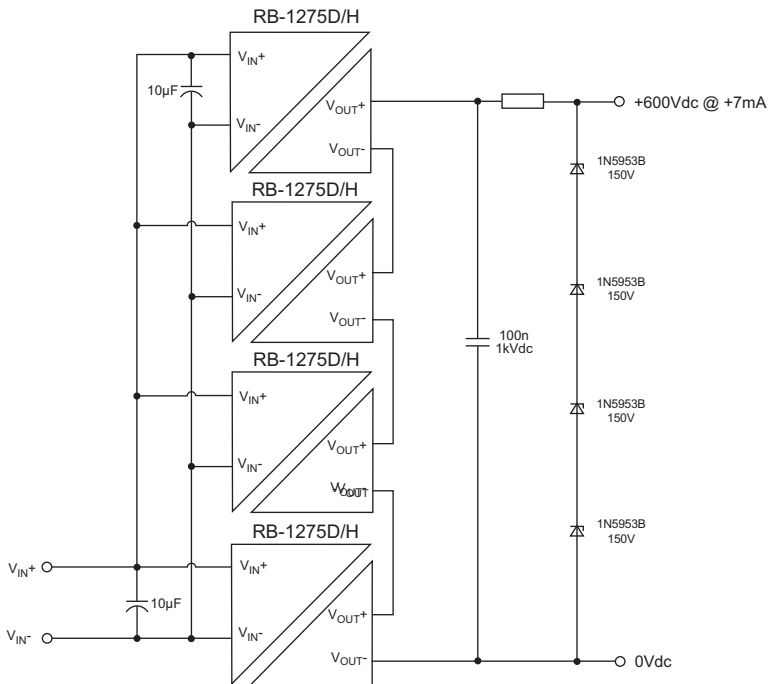


Abb. 9.9: Stromversorgung 600VDC

9.6 Erhöhung der Isolation

DC/DC-Wandler können auch kaskadiert werden, um die Isolation zu erhöhen. Abgleichwiderstände werden benötigt, um sicherzustellen, dass der Spannungsgradient an den Wandlern gleichmäßig verteilt ist, was diese Schaltungsvariante auf nicht-sicherheitskritische Anwendungen beschränkt.

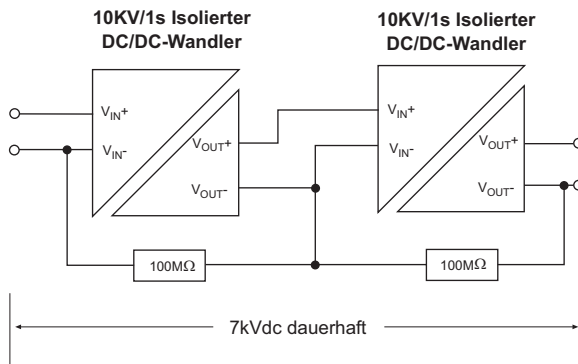


Abb. 9.10: Kaskadierung von Wandlern zur Erhöhung der Isolation

Ein Beispiel einer Anwendung, die sehr hohe Isolation erfordert, ist eine Überwachungsschaltung für eine Hochvakuumpumpe. Hochvakuumwerte können nur durch die Verwendung von Ionenpumpen, die Arbeitsspannungen von bis zu 7kVDC verwenden, erreicht werden. Ein DC/DC-Wandler mit 10kDC Isolationsnennspannung ist nicht geeignet, da diese Isolationsspannung nur für die Dauer von 1 Sekunde zulässig ist. Für eine Dauerspannung von 7kVDC an der Isolationsbarriere des Wandlers wäre ein Nennwert von 14kVDC pro 1s erforderlich. Solche Wandler sind nicht nur schwer erhältlich, sondern auch sehr teuer. Zwei preiswerte Wandler, die auf 10kVDC/1s ausgelegt sind, können jedoch, wie in Abb. 9.10 gezeigt, kaskadiert hintereinander geschaltet werden. Der Isolations-Widerstand dieser Wandler beträgt jeweils 10 Gigaohm, weshalb die Spannungsabgleichwiderstände allerhöchstens bei 1/100 dieses Wertes liegen sollten.

9.7 5-V-Bereinigung

Wenn sich analoge und digitale Schaltungen eine übliche 5-V-Versorgung teilen, kann es Probleme mit hochfrequenten Störungen von digitalen zu analogen ICs geben. Dies ist besonders störend in Audio- oder Videoanwendungen, bei denen die Überlagerung von digitalem Rauschen auf die Analogsignale das Erscheinen von Balken im Bild oder das Hören von unerwünschten Zischlauten im Audiosignal hervorrufen kann.

Abb. 9.11 zeigt einen auf den ersten Blick zwecklos erscheinenden Schaltkreis: einen nicht isolierten 5-Vin-zu-5-Vout-Wandler. Der Grund, warum diese Schaltung tatsächlich einen Sinn hat, liegt in den Spezifikationen des DC/DC-Wandlers.

Der Eingangsspannungsbereich des Wandlers beträgt $5V \pm 10\%$, während sein Ausgang $5V \pm 5\%$ beträgt. Daher bereinigt er alle kleinen Spannungsschwankungen am 5-V-Eingang.

Der gemeinsame Masseanschluss wird häufig dort benötigt, wo Audio- und Videoschaltkreis dieselbe Versorgung teilen.

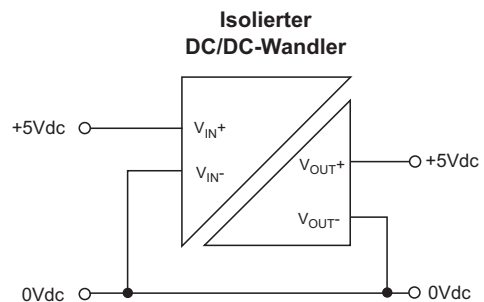


Abb. 9.11: Nicht isolierter 5-V-zu-5-V-Wandler

Eine andere Variante zu diesem Thema ist die in Abb. 9.12 gezeigte, sehr saubere 5-V-Stromversorgung, in der ein Linearspannungsregler verwendet wird, um eine sehr rauscharme Ausgangsspannung zu gewährleisten. Der Linearspannungsregler kann mit dem 5-V-Eingang nicht einfach in Reihe geschaltet werden, da sogar ein LDO (Low Drop Out)-Spannungsregler immer noch einige hundert Millivolt Headroom benötigt. In diesem Anwendungsbeispiel wird ein Dual-Ausgangs-3,3-V-SMD-DC/DC-Wandler verwendet, um 6,6V zu liefern, die dann durch jeden geeigneten rauscharmen Linearspannungsregler auf 5V heruntergeregelt werden kann. Der gemeinsame Pin eines Zweifachausgangs-DC/DC-Wandlers kann offen gelassen werden, wenn er nicht benötigt wird (NC). Mit einer geeigneten BauteilAuswahl kann mit dieser Schaltung ein Ausgangsrauschpegel von höchstens 5µVp-p erzielt werden.

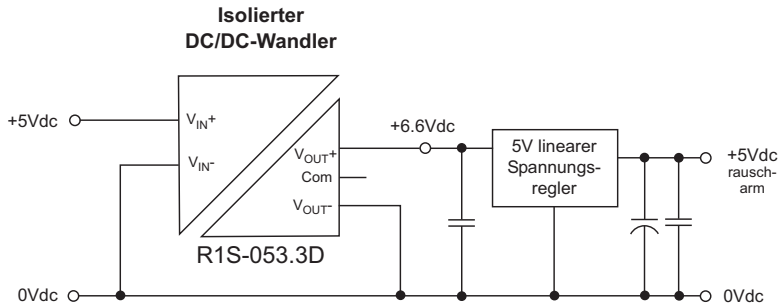


Abb. 9.12: Sehr rauscharme 5-V-Versorgung

9.8 Verwendung des CTRL-Steuerungspins

Einige DC/DC-Wandler verfügen über einen ON/OFF-Steuerungspins, der zwar die Ausgangsleistungsstufe ausschaltet, den Anschluss des Wandlers an die Versorgung aber immer noch aufrechterhält. Der Vorteil dieser Anordnung besteht darin, dass die Ansprechzeit auf ein Freigabesignal sehr kurz ist und das Wiederanlaufen des Wandlers keine große Einschaltspitzenströme verursacht, da die Eingangsfilterkondensatoren schon aufgeladen sind. Dieser letztere Punkt ist besonders für batteriebetriebene Systeme wichtig, bei denen der Einschaltspitzenstrombedarf in einer entladenen Batterie ein Abfallen der Batteriespannung unterhalb des minimalen Grenzwertes hervorrufen könnte.

Das folgende Beispiel zeigt eine Ausführung mit extrem niedriger durchschnittlicher Stromaufnahme für ein batteriebetriebenes Gerät, das nur fallweise aktiv sein muss. Typische Anwendungsfälle könnten ein Fernüberwachungssystem für ein netzunabhängiges System, wie eine mit Sonnenenergie betriebene Pumpenanlage oder Hochgebirgs-Wetterstation, sein. In voreingestellten Intervallen wird der Mikrocontroller angeschaltet, um Messungen wie Pumpendruck oder Umgebungstemperatur, die dann in einem nichtflüchtigen RAM gespeichert werden, auszuführen. Nach einigen Messzyklen könnte der Mikrocontroller so programmiert werden, dass er die GSM-Verbindung aktiviert und die gespeicherten Daten überträgt. Wenn entweder die Messungen oder die Übertragung abgeschlossen sind, löst der Mikrocontroller den 555-Timer aus und schaltet dabei seine eigene Versorgung aus. Während der Ruhezeiten wäre das System, mit Ausnahme der Timerschaltung, vollständig abgeschaltet. Der aus der Batterie im Bereitschaftsbetrieb gezogene Strom beträgt nur ca. 120µA (100µA für den Zeitgeberschaltkreis und 20 µA für den Regler R-78AA im Bereitschaftsbetrieb). Somit könnte die Lebensdauer der Batterie problemlos ein ganzes Jahr oder länger betragen. Eine Option dabei ist die Möglichkeit für die Versorgung des Mikrocontrollers eine Spannung von 5V oder 3,3V per Jumper auszuwählen, um eine Verwendung unterschiedlicher Mikrocontroller-Familien zuzulassen.

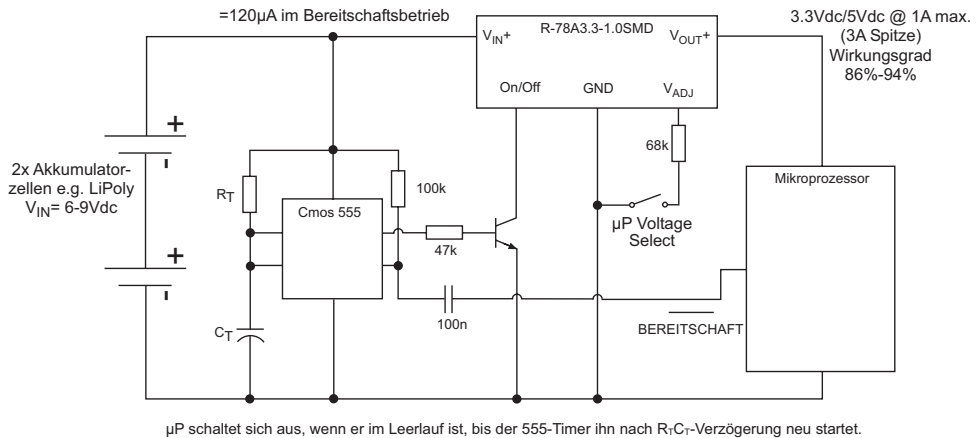


Abb. 9.13: Selbstabschaltkreis

9.9 Verwendung des V_{ADJ} -Pins

Viele DC/DC-Wandler verfügen zur Einstellung der Ausgangsspannung über einen Abgleichpin. Die Ausgangsspannung wird intern über einen festen Widerstandsteiler mit einer Spannungsreferenz verglichen (Abb. 9.14). Der Abgleichpin gewährleistet einen externen Zugangspunkt zur Verbindung des Spannungsteilers im Wandlerinneren und ermöglicht es so dem Benutzer, die Ausgangsspannung nach oben oder nach unten abzugleichen. Ein externer Widerstand vom Abgleichpin zur Masse setzt die Spannung V_{REF} herab und zwingt den Wandler, zur Kompensierung die Ausgangsspannung zu erhöhen. Ähnlich erhöht ein Widerstand vom Abgleichpin zu V_{OUT+} die V_{REF} -Spannung und veranlasst den Wandler, die Ausgangsspannung herabzusetzen.

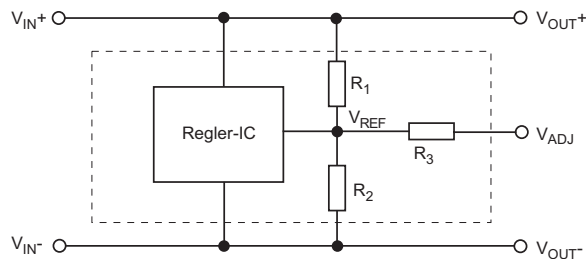


Abb. 9.14: Typischer Schaltkreis zum externen Abgleich der Ausgangsspannung

Im vorgezeichneten Schaltkreis begrenzt R_3 den Abgleichbereich bis zur Sicherheitsgrenze für die Stabilität der Ausgangsspannung. Der typische 'Abgleichbereich beträgt $\pm 10\%$.

Um Kabeldämpfung oder ohmsche Leiterbahnverluste zu kompensieren, ist es üblich, die Ausgangsspannung nachzustellen. Wenn die Last einigermaßen konstant ist, wird kein Messeingang mit Feedback benötigt. Trimming nach unten wird weniger häufig verwendet, hauptsächlich zum Sicherstellen, dass die Ausgangsspannung den absoluten Höchstwert der kritischen Bauteile nicht überschreitet.

Ein Schaltregler ist über einen viel breiteren Bereich von Ausgangsspannungen mit denselben internen Bauteilen stabil, weshalb Abgleichbereiche von mindestens 50% sehr typisch sind. Im folgenden Anwendungsbeispiel wird ein Schaltreglermodul mit 1A verwendet, das eine einfache mit Sonnenenergie betriebene Ladestation bildet. Die Nennausgangsspannung des Reglers R-6112P beträgt 12V, aber der Wahlschalter lässt auch 13,8V zu um eine Blei-Säure-Batterie zu laden, 6V um eine Nickel-Kadmium-Batterie zu laden und 5V für einen USB-Mobiltelefon-Ladeanschluss. Die Abgleichwiderstände sind so ausgewählt, dass sie den Spannungsabfall an der Ausgangsdiode D_2 kompensieren, was erforderlich ist, um den Wandler vor Rückströmen zu schützen, wenn das Solarpanel keinen Strom liefert. Das Widerstandsnetzwerk am USB-Anschluss ist erforderlich, um dem Telefon mitzuteilen, wie viel Strom es sicher aus der 5-V-Versorgung ziehen kann.

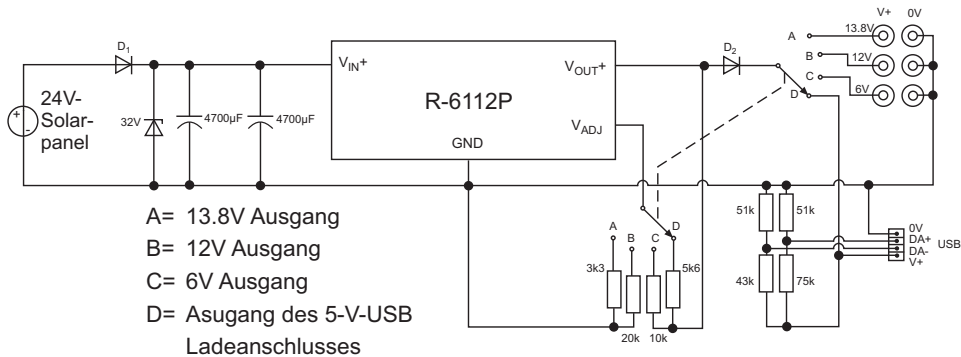


Abb. 9.15: Einfache mit Sonnenenergie betriebene Aufladestation

Anmerkung: Garantien für Mobiltelefone können verfallen, wenn ein nicht zertifiziertes Ladegerät verwendet wird. RECOM übernimmt keine Verantwortung für diesen Anwendungsvorschlag oder dessen Nutzung.

Quellen und Referenzmaterialien

Bell, B (2003) "Introduction to Push-Pull... Technologies", National Semiconductor

Dixon, L (2002) "Control Loop Cookbook", Texas Instruments Inc.

Ducard, G 2013 "Introduction to Digital Control of Dynamic systems", Lecture Notes

Falin, J (2008) "Designing DC/DC Converters based on SEPIC Technology", Texas Instruments Inc.

Gerstl, D (2003), "Power Supply Reliability", C&D Technologies

Kankanala, K (2011) "AN1369: Full-Bridge DC/DC Converter Reference Design", Microchip Technology.

Kessler, M (2010) "Synchronous Inverse SEPIC Technology", Analog Dialogue, Analog Devices

Liang, R (2012) "Design Considerations for System-Level ESD Circuit Protection", Texas Instruments Inc.

Maimone, G (2010) "AN4067: Selecting L and C Components", Freescale Semiconductor.

Mammano, B et al. (2005) "Safety Considerations in Power Supply Design", Texas Instruments

Peiravi, A (2009), "Testing and Reliability Improvement", Dept. of Electrical Engineering, University of Mashhad

Saliva, A (2013) "DN 2013-01: Design Guide for QR Flyback Converter", Infineon INFA) Corp.

Schramm, C (2004) "DC/DC-Wandler im Einsatz", DATEL GmbH

Sheehan, R (2007) "Understanding and Applying Current-Mode Control Theory", National Semi.

Walding, C (2008) "LLC Resonant Topology lowers switching losses", Fairchild Semiconductor

Anwendungshinweise:

ANP29 (2007) Zeta Converter Basics, Sipex Corporation

SLUP084 (2001) The Right-Half-Plane Zero – A Simplified Explanation, Texas Instruments Inc.

AN-1889 (2013) "How to Measure the Loop Transfer Function", Texas Instruments Inc.

DN022 (2010) "Output Ripple and Noise Measurement Methods", Ericsson Power

"Top 10 Circuit Protection Considerations" (2013) , Littelfuse

"Lighting and Electrical Equipment for use in Hazardous Atmospheres", IEC Hazardous Fundamentals,

TP018 (2011): "Safety Requirements for Board-Mounted DC/DC Converters", Ericsson Power

Fuse Technology, Overcurrent Protection Group, Cooper Bussmann

AN9003 "A Users Guide to Intrinsic Safety", Cooper Crouse-Hinds

Application Guide (2009) "Miniature Circuit Breakers", ABB

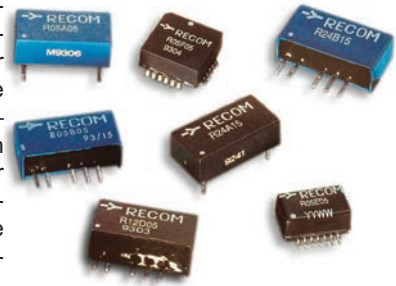
"MLCC Application Guide" (2013), Walsin Technology Corp.

"Safety Application Guide for Multilayer Ceramic Capacitors", Kyocera Corp

Über RECOM

25 Jahre ist es her, dass wir einem führenden deutschen Mobiltelefonhersteller unseren ersten, von Hand gefertigten DC/DC-Wandler vorstellten. Das Produkt bestand alle Tests mit Bravour und wenige Wochen später erhielten wir unsere erste Bestellung über 8.000 Stück. Eine neue Produktkategorie war geboren: DC/DC-Wandler in Modulbauform.

Während des Übergangs von analoger zu digitaler Elektronik erhöhte sich die Nachfrage nach DC/DC-Modulen rasant, überall wurden standardisierte Onboard-Schaltnetze benötigt. Zudem mussten I/O-Ports oder Verstärkerkanäle galvanisch isoliert werden, um die Sicherheit zu erhöhen und Erdschleifen zu vermeiden – weitere Anforderungen, die DC/DC-Wandler erfüllen konnten. RECOM liefert seit den ersten Tagen der Entwicklung von DC/DC-Wandlern zuverlässige, effiziente und modulare Lösungen für Kunden, die weder die Zeit noch die Mühen aufwenden können, Eigenentwicklungen zur Spannungsversorgung zu betreiben.



In den vergangenen zehn Jahren sind DC/DC-Wandler ein vielfältig eingesetztes Massenmarktprodukt geworden, das zu attraktiven Preisen abgeboten wird. Dadurch ist es heute für die meisten Anwender günstiger, zur Versorgung ihrer Onboard-Elektronik auf zertifizierte Wandlermodule zurückzugreifen als eine eigene Versorgung zu entwickeln. Das beschleunigt nicht nur den Entwicklungsprozess, sondern erleichtert auch die Einhaltung der EMV-Normen. Da in vielen Entwicklungsabteilungen Erfahrung mit elektromagnetischen Komponenten und Materialien fehlt, sind DC/DC-Wandler nach wie vor eine bewährte Alternative zur Stromversorgung der Marke Eigenbau.

In den vergangenen Jahren hat RECOM eine Vielzahl innovativer Produkte entwickelt, darunter den R78-Schaltregler, einen pinkompatiblen Ersatz für Linearspannungsregler, der viele Nachahmer-Produkte anderer Hersteller nach sich gezogen hat. Der Erfolg führte dazu, dass das Unternehmen schneller wuchs als der Strommarkt im Allgemeinen. Inzwischen haben wir weitere Märkte mit neuen Produkten erschlossen, wie z. B. hocheffizienten AC/DC-Netzteilen mit geringem Standby-Verbrauch, einem umfassenden Sortiment an DC- und AC-LED-Treibern und Spezialanwendungen aus den Bereichen Medizintechnik und Bahntechnologie.

Unser neues, hochmodern eingerichtetes, campusähnliches RECOM Headquarter in Gmunden/Österreich bietet den nötigen Raum für die Erweiterung unseres Geschäfts und die Umsetzung unserer ehrgeizigen Ziele. Ein Team von Ingenieuren aus aller Welt arbeitet in erstklassig ausgestatteten, modernen Labors kontinuierlich an neuen Produkten und Lösungen für unsere Kunden. Durch umfangreiche Leistungs-, Belastungs- und Umwelttests schon im frühesten Stadium der Produktentwicklung gewährleisten wir, dass unsere Hauptanliegen – Zuverlässigkeit und Qualität – stets erfüllt werden. Ebenso haben wir in ein eigenes EMV-Testlabor mit Testkammer investiert, um unabhängig von externen Einrichtungen arbeiten zu können.

An unserem neuen Hauptsitz befindet sich neben alledem auch ein vollautomatisiertes Logistikzentrum für einen Bestand von über 30.000 verschiedenen Produkten. Die Produktion und die Fertigungstechnik sind in zwei RECOM Werken in Kaohsiung/Taiwan angesiedelt. Die RECOM Manufacturing Ltd. ist für die Fertigung, Überprüfung und Versendung der fertigen Wandler zuständig, und die RECOM Technology Ltd. produziert als vollautomatisierte SMT-Fabrik mit sechs SMD-Fertigungslinien sogenannte „High Runner“, Produkte in sehr hoher Stückzahl. In dieser Aufstellung kann RECOM eine beachtliche Produkttiefe zu wettbewerbsfähigen Preisen bieten und sich immer wieder als globaler Marktführer behaupten.

Danksagung

Ich wollte dieses Buch schon vor Jahren schreiben, hatte jedoch aus Zeitgründen nie die Möglichkeit dazu, bis ich nach einem schweren Skiunfall eine mehrmonatige Auszeit nehmen musste. Da ich mehr oder weniger ans Bett gefesselt war, packte ich die Gelegenheit beim Schopfe und brachte das Werk endlich zu Papier. Zuallererst möchte ich dem örtlichen Krankenhaus, dem LKH Gmunden, dafür danken, dass ich mein Krankenzimmer vorübergehend in ein Arbeitszimmer umfunktionieren durfte. Außerdem danke ich allen Mitarbeitern von RECOM, die mich bei der Umsetzung dieses Projekts unterstützt haben.

Mein besonderer Dank gilt RECOM Test Lab Consultant Carl Schramm, der mir erlaubt hat, seine Publikation „DC/DC Wandler im Einsatz“ zu nutzen. Auf seinen Ausführungen beruht ein Großteil der Kapitel 1 und 3. Ich habe seinen Text um eigene Interpretationen und Ergänzungen erweitert, sodass etwaige Fehler ganz allein mir anzulasten sind.

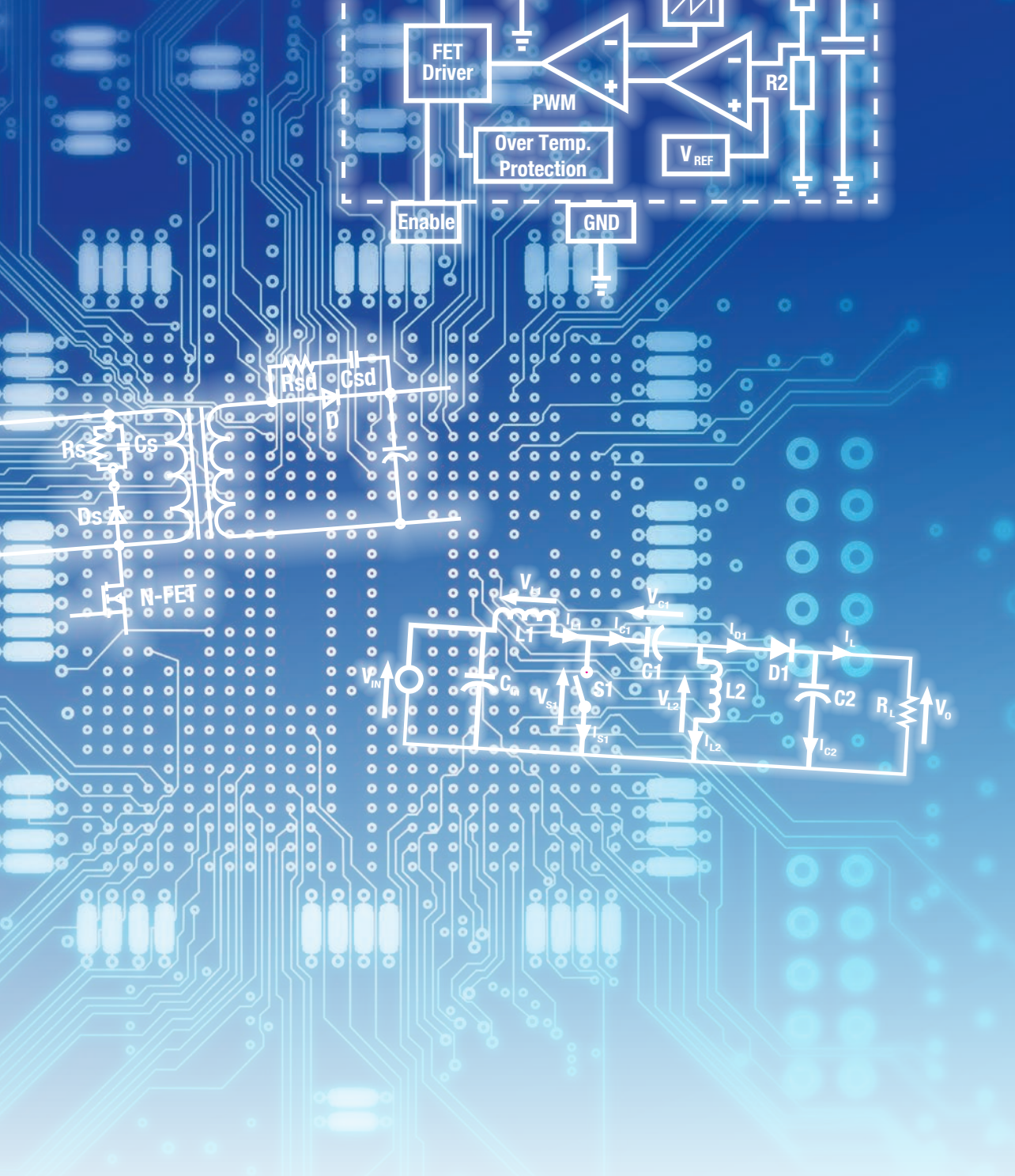
Quellen und Referenzmaterialien

Um meine Wissenslücken zu füllen, habe ich auf verschiedene Quellen zurückgegriffen. Besonders empfehlen kann ich die von Texas Instruments Incorporated, ON Semiconductor, Maxim Integrated und Fairchild Semiconductor herausgegebenen Anwendungshinweise. Zur weiteren Lektüre eignen sich Fachmagazine wie Electronic Product Design and Test, Power Systems Design und Electronic Design, die regelmäßig aufschlussreiche „How to“-Artikel veröffentlichen. Die wichtigste Informationsquelle überhaupt ist natürlich das Internet. Detaillierte Informationen finden sich häufig unter den Links im Quellenverzeichnis einschlägiger Wikipedia-Seiten.

Über den Verfasser



Steve Roberts wurde in England geboren. Nach seinem Abschluss (B.Sc.) in Physik und Elektronik an der Brunel University in London, arbeitete er für das University College Hospital. Danach war er 12 Jahre lang als Head of Interactives im Science Museum tätig und absolvierte in dieser Zeit ein Masterstudium am Londoner University College. Später zog er nach Österreich und entwickelte bei RECOM zunächst als technischer Kundenbetreuer Wandler und beantwortete als Applikationsingenieur die unterschiedlichsten Kundenanfragen. Heute ist er am neuen Hauptsitz Gmunden als technischer Direktor der RECOM Gruppe tätig.



WE POWER YOUR PRODUCTS
www.recom-electronic.com

**RECOM Power GmbH und
 RECOM Engineering GmbH und Co KG**
 Münzfeld 35, 4810 Gmunden, Österreich