

Ein Fenster zur gleichzeitigen Messung der
Übertragungsfunktion eines realen Systems
und des Leistungsdichtespektrums des
überlagerten Rauschens am Systemausgang
(Teil 2)

Helmut Repp

Erlangen - 2025

Dipl.-Ing. Helmut Repp

Liegnitzer Straße 1

D-91058 Erlangen

Germany

Telefon: +49-9131-3 66 41

Mobile: +49-173-57 19 350

E-mail: Helmut.Repp@gmx.de

Übersicht

Mit Hilfe des in [1] verwendeten Rauschklimrmessverfahrens gelingt es in einer Messung sowohl die Übertragungsfunktion als auch das Rauschleistungsdichtespektrum eines gestörten Systems zu messen. Dort wird für zeitinvariante Systeme, die von stationären und mittelwertfreien Prozessen erregt und gestört werden, gezeigt wie sich das Rauschklimrmessverfahren um eine Fensterung erweitern lässt. Durch die Einführung einer Fensterung wird es unter anderem möglich, die spektralen Eigenschaften der Störung wesentlich frequenzselektiver und dabei dennoch leistungsrichtig zu ermitteln.

Im zweiten Teil wird nun gezeigt, wie man mit dem um die Fensterung erweiterten Rauschklimrmessverfahren Systeme vermessen kann, die von Prozessen erregt und gestört werden, die zeitabhängige erste Momente aufweisen. Das erste und die zweiten zentralen Momente der Störung werden dabei getrennt ermittelt. Zudem werden hier drei Möglichkeiten untersucht, wie dieses Messverfahren zu erweitern ist, so dass man damit auch

- alle Korrelationen, die zwischen Ein- und Ausgangssignal bestehen, bei einem
- periodisch zeitvarianten System, das von einem
- zyklstationären Rauschprozess

gestört wird, messen kann. Die dazu notwendigen Modifikationen des RKM können dabei in beliebiger Kombination der drei möglichen Erweiterungen eingesetzt werden. Wenn man alle drei Modifikationen berücksichtigt, erhält man ein sehr allgemeines, und daher vielseitig einsetzbares Messverfahren für komplexwertige Systeme.

Desweiteren wird gezeigt, wie und unter welchen Umständen man Messwerte für die periodisch zeitvariante, komplexe Impulsantwort und die beiden Autokorrelationsfolgen eines komplexen, zyklstationären Approximationsfehlerprozesses berechnen kann.

Desweiteren enthält der zweite Teil auch einige zusätzliche praxisorientierte Betrachtungen zu den in [1] und hier theoretisch behandelten Varianten des Messverfahrens. Es wird eine programmflussartige Ablaufübersicht für die Variante des RKM angegeben, mit der man das lineare Systemverhalten messen kann, das alle Korrelationen des Systemein- und -ausgangs beschreibt, und bei der das erste zeitabhängige und die zweiten zentralen zeitunabhängigen Momente der Störung getrennt gemessen werden. Für die anderen Varianten

wird eine Aufwandsabschätzung durchgeführt.

Zwei Simulationsbeispiele runden die Beschreibung des RKM ab. Eines beschäftigt sich mit der Varianz und Kovarianz der Messwerte des zeitabhängigen ersten Moments der Störung. Das andere demonstriert die Messung eines von einem zyklstationären Prozess gestörten, periodisch zeitvarianten, komplexwertigen Systems mit Hilfe der umfangreichsten RKM-Variante mit Fensterung.

Auch zur Konstruktion der Fenster werden einige Ergänzungen angegeben. So wird hier eine erweiterte Variante des in [1] vorgestellten Algorithmus angegeben, bei der einige dort nicht genutzte Freiheitsgrade dazu verwendet werden können, die Eigenschaften der Fensterfolge je nach Applikation geeignet zu beeinflussen. Auch wird ein Algorithmus zur Konstruktion einer kontinuierlichen Fensterfunktion endlicher Länge mit analogen Eigenschaften angegeben. Einige MATLAB Programme zeigen, wie sich die Fenster, deren Autokorrelationsfolgen bzw. -funktionen, Fourierreihenkoeffizienten und Spektren berechnen lassen.

Es handelt sich beim zweiten Teil nicht um eine eigenständige Abhandlung. Vielmehr wird hier auf der in [1] dargestellten Theorie aufgebaut, und es wird immer wieder auf die dort gewonnenen Ergebnisse Bezug genommen. Die Referenzierung erfolgt dabei, indem der entsprechenden Kapitel-, Gleichungs-, Bild-, Tabellenummer oder Seitenzahl die Literaturstelle [1] vorangestellt wird (z. B. ... Gleichung ([1]:2.20) ...).

Inhaltsverzeichnis

1	Einleitung	1
2	Das Systemmodell	7
2.1	Interpretation der Rauschprozesse im Systemmodell	7
2.2	Theoretische Werte der Übertragungsfunktionen und der deterministischen Störung	26
2.3	Leistungsdichtespektren und Fensterung	39
3	Das Rauschklirrmessverfahren mit Fensterung	55
3.1	Messung der Übertragungsfunktionen und der deterministischen Störung .	56
3.2	Erwartungstreue der Messwerte der Übertragungsfunktionen und der deterministischen Störung	61
3.3	Messung der beiden Leistungsdichtespektren	65
3.4	Varianzen und Kovarianzen der Messwerte	73
4	Weitere Messwerte	95
4.1	Messwerte der zeitvarianten Impulsantworten	95
4.2	Messwerte der Auto- und Kreuzkorrelationsfolgen	97
5	Spektralschätzung mit zeitabhängigem ersten Moment	101
5.1	Spektralschätzung komplexer Prozesse	101
5.2	Spektralschätzung reeller Prozesse	106
6	Reellwertige Systeme	109
6.1	Erste Variante des RKM zur Messung reellwertiger Systeme	109
6.2	Weitere Varianten zur Messung reellwertiger Systeme	115

7	Ablaufsübersicht für eine Variante des RKM	123
8	Aufwandsabschätzung für andere Varianten des RKM	139
9	Beispiele für RKM Messergebnisse	153
9.1	Varianz und Kovarianz der Messwerte der deterministischen Störung	153
9.2	Messung eines komplexen, periodisch zeitvarianten Systems mit zyklota- tionärer Störung	155
10	Ergänzungen zum Fenster	167
10.1	Algorithmus für Fenster mit frei wählbaren Nullstellen	167
10.2	Berechnung des Spektrums der Fensterfolge	178
10.3	Berechnung der AKF der Fensterfolge	179
10.4	Fenster mit weiteren Nullstellen am Einheitskreis im $2\pi/F$ Raster	181
10.5	Basisfenster mit sechs zusätzlichen Nullstellen am Einheitskreis	185
10.6	Halbbandfilter	187
10.7	Berechnung einer kontinuierlichen Fensterfunktion	189
10.8	Die kontinuierlichen Fensterfunktion als Grenzwertlösung	201
10.9	Augendiagramm des AKF der kontinuierlichen Fensterfunktion	202
11	MATLAB-Programmauszüge	205
11.1	Die diskreten Fensterfolgen	206
11.2	Die Spektren der diskreten Fensterfolgen	213
11.3	Die diskreten Fensterautokorrelationsfolgen	220
11.4	Die Fourierreihenoeffizienten der kontinuierlichen Fensterfunktionen . . .	221
11.5	Die kontinuierlichen Fensterfunktionen	226
11.6	Die Spektren der kontinuierlichen Fensterfunktionen	227
11.7	Die kontinuierlichen Fensterautokorrelationsfunktionen	230
	Literaturverzeichnis	231

Anhang	233
A.1 Identität der Lösungsräume der Gleichungssysteme (2.21) und (2.22) . . .	233
A.2 Zur Konditionierung der empirischen Kovarianzmatrix	237
A.3 Zur Wahrscheinlichkeit der empirischen Varianz Null	242
A.4 Grenzen für die Hauptdiagonalelemente der inversen Kovarianzmatrix . . .	245
A.5 Zur Berechnung des Erwartungswertes einer bilinearen Form	246
A.6 Zur Berechnung der Messwerte des LDS und des KLDS	249
A.7 Zu Spur und Rang des Produktes Idempotenter Matrizen	253
A.8 Zur Berechnung der Varianzen und Kovarianzen des LDS und KLDS . . .	254
A.9 Zum Vergleich der Beträge der empirischen Varianz und Kovarianz	258
A.9.1 Messwerte des LDS und des KLDS	259
A.9.2 Messwerte der deterministischen Störung	262
A.10 Anmerkung zur Gauß-Verbundverteilung	267
 Liste der verwendeten Abkürzungen und Formelzeichen	 273

Denksportaufgabe:

Zeige, dass $\prod_{\rho=1}^{M-1} \sin\left(\frac{\pi}{M} \cdot \rho\right) = M \cdot 2^{1-M}$ gilt.

Lösung auf Seite 204.

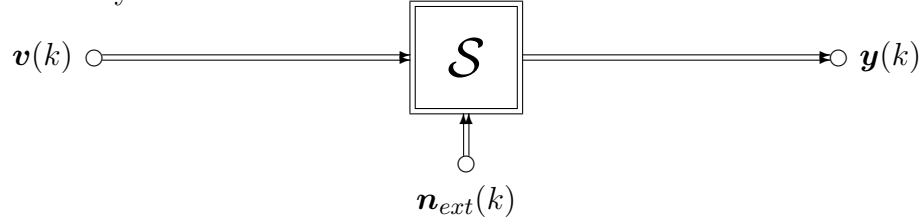
1 Einleitung

Von vielen digitalen Systemen¹ wünscht man ein lineares, zeitinvariantes oder wenigstens periodisch zeitvariantes und stabiles Verhalten. Reale Systeme können solches Verhalten nur näherungsweise realisieren. So ist z. B. in der Realität der Bereich der zulässigen Aussteuerung immer begrenzt, bei digitalen Systemen kann nur mit endlicher Wortlänge gearbeitet werden, und bei kontinuierlichen Systemen zeigen nichtlineare Bauteilkennlinien ihre Wirkung. Auch externe Störungen, die von außen in das System einstreuen und die im oberen Teil von Bild 1.1 mit $\mathbf{n}_{ext}(k)$ bezeichnet werden, verursachen Abweichungen des realen Systemverhaltens vom gewünschten Systemverhalten. In aller Regel sind solche Systeme nicht nur für die Erregung mit einer konkreten Signalfolge entworfen. Man wird daher die Erregung des Systems als einen Zufallsprozess $\mathbf{v}(k)$ betrachten, von dem die wichtigsten stochastischen Eigenschaften — wie z. B. die Varianz — bekannt sind. Auch im zweiten Teil werden wir uns auf den Fall der Erregung mit zufälligen periodischen Eingangssignalen beschränken, wobei die wesentlichen stochastischen Eigenschaften dieselben sein sollen, wie bei dem im normalen Betrieb anliegenden Zufallsprozess. Man ist nun daran interessiert, den sich am Ausgang des realen Systems $\mathcal{S}$ ergebenden Prozess $\mathbf{y}(k)$ des realen Systems in vier Anteile $\mathbf{x}(k)$, $\mathbf{x}_*(k)$, $u(k)$ und $\mathbf{n}(k)$ aufzuspalten, wie dies in Bild 1.1 im unteren Teilbild dargestellt ist. Der eine Anteil $\mathbf{x}(k)$ ist die Reaktion eines idealen linearen periodisch zeitvarianten Modellsystems $\mathcal{S}_{lin}$, das durch seine bifrequente Übertragungsfunktion $H(\mu_1, \mu_2)$ beschrieben wird, auf den erregenden Prozess $\mathbf{v}(k)$. Der zweite Anteil $\mathbf{x}_*(k)$ ist die Reaktion eines weiteren idealen linearen periodisch zeitvarianten Modellsystems $\mathcal{S}_{*,lin}$, das durch die bifrequente Übertragungsfunktion $H_*(\mu_1, \mu_2)$ beschrieben wird, auf denselben, nun aber konjugierten, erregenden Prozess $\mathbf{v}(k)^*$. Das im allgemeinen zeitabhängige erste Moment des verbleibenden Differenzprozesses $\mathbf{y}(k) - \mathbf{x}(k) - \mathbf{x}_*(k)$ ist nicht zufällig und wird in weiteren als deterministische Störung $u(k)$ bezeichnet. Der vierte und letzte Anteil des Prozess $\mathbf{y}(k)$ ist der Rauschprozess $\mathbf{n}(k)$, der somit mittelwertfrei ist. Wir wollen uns hier auf den Fall beschränken, dass der Prozess $\mathbf{n}(k)$ zyklstationär ist.

Eine Ensemblemittelung — d. h. eine Mittelung über mehrere Messungen im gleichen Zeitintervall, bei denen mit unterschiedlichen Musterfolgen des eingangsseitigen Zufallsprozesses erregt wird — kann zur Gewinnung der das System beschreibenden Größen an

¹Für die Behandlung zeitkontinuierlicher Systeme gilt das in [1] auf Seite [1]:1 gesagte.

a) Reales System:



b) Systemmodell:

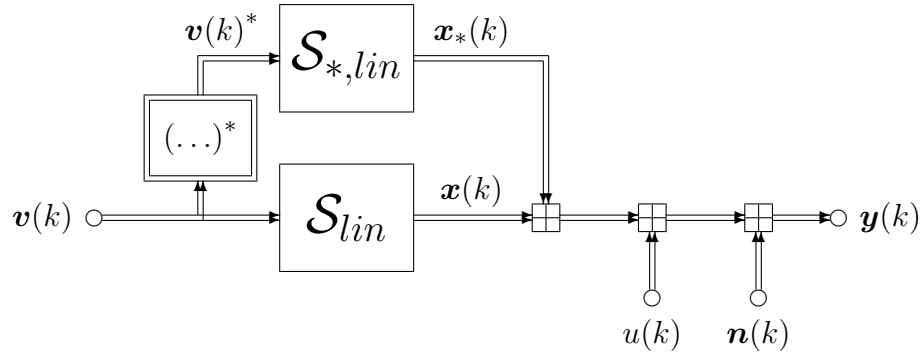


Bild 1.1: Erweitertes Modell eines gestörten nichtlinearen realen Systems

einem einzelnen System i. Allg. nicht vorgenommen werden. Daher muss die Messung dieser Größen mit Hilfe einer Zeitmittelung erfolgen. Deshalb werden bei der Messung zeitlich begrenzte Musterfolgen verwendet, die in unterschiedlichen Zeitintervallen das zu messende System erregen. Die am System auftretenden Zufallsprozesse bedürfen dann aber einer Interpretation, die an die Art des typischen Betriebs angepasst ist, und die Messung muss mit dieser Betriebsart konform durchgeführt werden. Eine Forderung an die so interpretierten Zufallsprozesse, die sich daraus ergibt, ist die Ergodizität dieser Prozesse in einem sehr weiten Sinn. Mit dieser Interpretation der Zufallsprozesse und mit der Ergodizitätsforderung unter besonderer Berücksichtigung einer deterministischen Störung wird bei der Beschreibung des Systemmodells im nächsten Kapitel begonnen.

Wie in [1] werden wir uns bei der Beschreibung des Rauschprozesses $\mathbf{n}(k)$ auf die Momente zweiter Ordnung² beschränken. Die Autokorrelationsfolge (AKF) dieses Prozesses ist der Teil der Momente zweiter Ordnung, der die Korrelationen der Zufallsgrößen des Prozesses zu einem Zeitpunkt mit den Zufallsgrößen zu einem weiteren Zeitpunkt beinhaltet. Bei den hier betrachteten Fall eines zyklstationären Prozesses $\mathbf{n}(k)$ ist die AKF

$$\phi_{\mathbf{n}}(k_1, k_2) = \text{E}\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^*\} \quad (1.1)$$

zweidimensional. Über eine zweidimensionale diskrete Fouriertransformation bezüglich

²Gemäß der Definition der deterministischen Störung $u(k)$ ist das Moment erster Ordnung des Prozesses $\mathbf{n}(k)$ immer null.

beider Zeitvariablen erhält man das bifrequente Leistungsdichtespektrum (LDS)

$$\Phi_{\mathbf{n}}(\Omega_1, \Omega_2) = \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^*\} \cdot e^{-j \cdot \Omega_1 k_1} \cdot e^{j \cdot \Omega_2 k_2}, \quad (1.2)$$

bei der das Vorzeichen bei der zweiten Frequenzvariable Ω_2 umgedreht ist. Bei komplexen Rauschprozessen ist die Beschreibung der zweiten Momente nur vollständig, wenn neben der AKF auch noch die diskrete Folge der Kovarianzen der Zufallsgrößen des Prozesses zu einem Zeitpunkt mit den konjugierten Zufallsgrößen zu einem weiteren Zeitpunkt angegeben wird. Weil bei der Definition der Kovarianz die eine daran beteiligte Zufallsgröße bereits konjugiert wird, ist also auch noch die zweidimensionale Folge der Erwartungswerte der Produkte zweier *nicht* konjugierter Zufallsgrößen anzugeben.

$$\psi_{\mathbf{n}}(k_1, k_2) = E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)\} \quad (1.3)$$

Diese sei im weiteren als Kreuzkorrelationsfolge bezeichnet, weil diese die kreuzweisen Korrelationen des Ein- und Ausgangs eines konjugierenden — im übrigen nichtlinearen — Systems enthält. Durch eine zweidimensionale diskrete Fouriertransformation bezüglich beider Zeitvariablen — wieder mit Vorzeicheninvertierung bei der zweiten Frequenzvariable — gewinnt man daraus das bifrequente Kreuzleistungsdichtespektrum (KLDS)

$$\Psi_{\mathbf{n}}(\Omega_1, \Omega_2) = \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \psi_{\mathbf{n}}(k_1, k_2) \cdot e^{-j \cdot \Omega_1 k_1} \cdot e^{j \cdot \Omega_2 k_2}. \quad (1.4)$$

Im Anschluss an die Interpretation der Rauschprozesse wird in Kapitel 2.2 auf die Voraussetzungen eingegangen, die ein reales System erfüllen muss, damit eine Beschreibung durch ein Modell nach Bild 1.1 sinnvoll ist. Es wird gezeigt, welche Übertragungsfunktionen $H(\mu_1, \mu_2)$ und $H_*(\mu_1, \mu_2)$ und welche deterministische Störung $u(k)$ sich im Systemmodell theoretisch ergeben, wenn das Modellsystem das lineare Verhalten des realen Systems in dem Sinne möglichst gut approximieren soll, dass das zweite Moment des Modellrauschprozesses $\mathbf{n}(k)$ möglichst klein wird. Die Überführung des realen Systems in das Systemmodell wird somit als die Lösung einer linearen Regression dargestellt. Der Modellrauschprozess ergibt sich daher als der Prozess des Fehlers der Approximation des realen Systems durch die beiden linearen periodisch zeitvarianten Systeme $\mathcal{S}_{lin}$ und $\mathcal{S}_{*,lin}$ und die deterministische Störung $u(k)$. Der Rauschprozess $\mathbf{n}(k)$ wird daher im weiteren auch oft als Approximationsfehlerprozess bezeichnet.

Da die Approximation des realen Systems durch das Systemmodell nach Bild 1.1 von den i. Allg. unbekannten stochastischen Eigenschaften der Störungen am realen System abhängt, kann man die zu bestimmenden Folgen und Funktionen in der Regel nicht theoretisch in geschlossener Form angeben. Man muss daher zunächst endlich viele Werte finden,

mit deren Hilfe die das Systemmodell beschreibenden Folgen und Funktionen möglichst aussagekräftig charakterisiert werden können. Desweiteren wird ein Verfahren benötigt, das es erlaubt, mit Hilfe einer Messung am realen System diese endlich vielen Werte abzuschätzen. Ein endlich langer Abschnitt der Folge der deterministischen Störung $u(k)$ besteht aus endlich vielen Werten, und ist daher zur Beschreibung des Systemmodells unverändert geeignet. Es wird gezeigt, dass bei der hier gewählten periodischen zufälligen Erregung die endlich vielen Abtastwerte $H(\mu_1, \mu_2)$ bzw. $H_*(\mu_1, \mu_2)$ der beiden bifrequenten Übertragungsfunktionen die beiden linearen Systeme $\mathcal{S}_{lin}$ und $\mathcal{S}_{*,lin}$ vollständig beschreiben. Analog zu [1] wird hier für den Fall des zyklstationären Approximationsfehlerprozesses eine dem bifrequenten LDS eng verwandte bifrequente spektrale Folge zur Beschreibung der zweiten Momente des Modellrauschprozesses angegeben. Es handelt sich dabei um die zweidimensionale Verallgemeinerung des Periodogramms eines gefensterten Signalausschnittes. Die dabei an die Fensterfolge zu stellenden Forderungen sind dabei dieselben wie in [1].

In Kapitel 3 wird dann das Rauschklimrmessverfahren in der Art modifiziert, dass man damit alle im Systemmodell vorkommenden Größen, die in Kapitel 2 vorgestellt und theoretisch hergeleitet worden sind, messtechnisch bestimmen kann. Dass diese Messwerte die entsprechenden theoretischen Größen erwartungstreu und konsistent abschätzen, wird dadurch gezeigt, dass die Erwartungswerte und Varianzen der Messwerte berechnet werden. Zusätzlich zu den Messwertvarianzen werden auch deren Kovarianzen hergeleitet und es wird angegeben wie diese abgeschätzt werden können. Damit kann man dann Konfidenzintervalle bzw. -gebiete für die Messwerte angeben, wie dies in [1] beschrieben ist.

Aus den bis dahin bestimmten Messwerten lassen sich weitere Messwerte ableiten. Dies sind zum einen die beiden periodisch zeitvarianten Impulsantworten der beiden linearen Modellsysteme in Bild 1.1 und zum anderen die Messwerte für die Auto- und Kreuzkorrelationsfolgen des zyklstationären Approximationsfehlerprozesses. Die Berechnung dieser Messwerte und die Abschätzung ihrer Messwert(ko)varianzen wird in Kapitel 4 beschrieben.

Die Spektralschätzung komplexer, stationärer und zyklstationärer, mittelwertbehafteter Prozesse wird als ein Sonderfall des des RKM in Kapitel 5 behandelt.

In Kapitel 6 werden mehrere Varianten zu Messung reellwertiger Systeme untersucht, die von Prozessen gestört werden, deren erstes Moment zeitabhängig ist.

Für den Fall, dass sich ein stationärer Approximationsfehlerprozess $\mathbf{n}(k)$ ergibt, wird eine implementierungsnahe Auflistung der bei diesem Messverfahren durchzuführenden Einzelschritte in Kapitel 7 angegeben.

Kapitel 8 enthält für einige weitere wichtige Varianten des RKM bzw. der daraus abgeleiteten Spektralschätzverfahren Abschätzungen des benötigten Speicherbedarfs.

Anhand zweier am Rechner simulierter Beispielsysteme wird in Kapitel 9 die Einsetzbarkeit der Fensterfolge sowie die neu eingeführte Schätzung der deterministischen Störung beim RKM verdeutlicht. Diese Beispiele dienen auch der Verdeutlichung einiger im theoretischen Teil der Abhandlung beschriebener Sachverhalte.

Der Rest des zweiten Teils beschäftigt sich mit dem in Kapitel [1]:6 vorgestellten Algorithmus zur Konstruktion einer geeigneten Fensterfolge. Im Kapitel 10 wird zunächst angegeben, wie man die in [1] nicht genutzten Freiheitsgrade bei der Berechnung einer Fensterfolge mit einbeziehen kann, um die Eigenschaften des Spektrums der Fensterfolge günstig beeinflussen zu können. Hierfür werden anschließend drei Beispiele angegeben. Desweiteren wird eine Methode vorgestellt, wie man das Fensterspektrum vor allem im Sperrbereich hochgenau berechnen kann. Eine weitere Methode, gibt an wie sich die Fenster AKF einzig aus den Fourierreihenkoeffizienten des Fensters berechnen lässt. Wie sich der in Kapitel [1]:6 vorgestellte Algorithmus abändern lässt, um eine kontinuierliche Fensterfunktion zu berechnen, wird ebenfalls angegeben. Dass dieser Algorithmus eine Grenzwertlösung für $M \rightarrow \infty$ darstellt, wird gezeigt. Auch wird ein Augendiagramm der AKF des kontinuierlichen Fensters dargestellt. In Kapitel 11 werden sieben MATLAB Programme aufgelistet, die der Berechnung der Fenster, sowie ihrer Spektren und AKFs dienen.

2 Das Systemmodell

In diesem Kapitel werde ich zunächst auf die Interpretation der im Systemmodell nach Bild 1.1 vorkommenden Zufallsprozesse eingehen. Dann werden die beiden bifrequenten Übertragungsfunktionen der beiden linearen Modellsysteme und das deterministische Störsignal in Abhängigkeit der stochastischen Eigenschaften der Prozesse am Eingang und am Ausgang des realen Systems berechnet. Im letzten Teil dieses Kapitels wird die in [1] angestellte Überlegung, statt des LDS des Rauschprozesses den Erwartungswert des Periodogramms des gefensterten Rauschprozesses zu dessen Beschreibung zu verwenden, für die beiden bifrequenten Leistungsdichtespektren bestätigt. Es wird gezeigt, dass die in [1] aufgestellten Forderungen für die dabei verwendete Fensterfolge weiterhin gültig sind.

2.1 Interpretation der Rauschprozesse im Systemmodell

Zunächst eine Vorbemerkung: In der Literatur findet man zwei verschiedene Definitionen für die Stationarität eines Zufallsprozesses im engen Sinne. Die erste Definition besagt folgendes: Greift man sich aus einem Zufallsprozess einen Satz von Zufallsgrößen heraus, indem man für den Parameter — bei uns meist die diskrete Zeit k — des Zufallsprozesses eine beliebige Anzahl beliebiger Werte einsetzt, so weisen diese eine gemeinsame Verbundverteilung auf. Nimmt man nun statt dieser Zufallsgrößen, diejenigen Zufallsgrößen, bei denen die Werte des Parameters um einen beliebigen aber bei allen Zufallsgrößen konstanten Wert — z. B. die gleiche Zeit κ — verschoben sind, so weisen diese ebenfalls eine gemeinsame Verbundverteilung auf. Ist die Verbundverteilung der Zufallsgrößen ohne Parameterverschiebung gleich der Verbundverteilung der Zufallsgrößen mit den verschobenen Parameterwerten, und zwar unabhängig davon, wieviele und welche Parameterwerte für den Satz der Zufallsgrößen ausgewählt wurden, und wie groß der Verschiebungswert des Parameters ist, so nennt man den Zufallsprozess im engen Sinne stationär. Die zweite Definition ist da weniger streng. Sie besagt lediglich, dass alle Momente, die man aus den Verbundverteilungen der so herausgegriffenen Zufallsgrößen berechnet, nicht von der Parameterverschiebung der beiden Sätze der Zufallsgrößen abhängen dürfen. Da es Verteilungen gibt, bei denen keine Momente oder nur solche Momente bis zu einer endlichen Ordnung existieren, lässt diese Definition bei solchen Zufallsprozessen keine Aussage über

deren Stationarität im engen Sinne zu. Beispielsweise besitzt die Cauchy-Verteilung [2] kein einziges Moment. Einen Zufallsprozess, dessen Zufallsgrößen alle unabhängig sind und dieselbe Cauchy-Verteilung aufweisen, würde man ohne weiteres als stationär bezeichnen, obwohl nach der zweiten Definition keine Aussage darüber möglich ist. In weiteren wird bei den theoretischen Herleitungen mit Stationarität im engen Sinne immer die Stationarität im engen Sinne gemäß der ersten Definition gemeint sein. Unter der Stationarität im weiten Sinne ist — wie üblich — die Definition gemäß der zweiten Variante für die Momente bis zur zweiten Ordnung gemeint. Bei realen Zufallsprozessen kann man davon ausgehen, dass die beiden Definitionen identisch sind, da dann der Bereich der möglichen Werte der Zufallsvariablen als begrenzt angesehen werden kann, so dass die Momente beliebiger Ordnung existieren. Desöfteren wird von bereichsweiser Stationarität oder Stationarität eines Zufallsprozessausschnittes die Rede sein. Da diese Begriffe nicht definiert sind sei hiermit erklärt, dass mit bereichsweiser Stationarität gemeint ist, dass bei den Definitionen der Stationarität nur solche Sätze von Zufallsgrößen betrachtet werden, bei denen die Parameterwerte aller Zufallsgrößen — also sowohl der Zufallsgrößen mit der ursprünglichen, als auch der Zufallsgrößen mit den verschobenen Parameterwerten — innerhalb des angegebenen Parameterbereichs liegen. Im weiteren wird mit einem Zufallsprozessausschnitt ein Zufallsprozess gemeint sein, bei dem der Parameterbereich auf ein angegebenes Intervall beschränkt ist. Ein Ausschnitt eines Zufallsprozesses wird ab nun als stationär bezeichnet, wenn der Zufallsprozess, aus dem er entnommen worden ist, in dem angegebenen Intervall bereichsweise stationär ist. Da bei uns der Parameter immer diskret ist — meist $k \in \mathbb{Z}$ —, wird im weiteren mit einer Zuordnung eines Ausschnittes eines Zufallsprozesses zu einem Zufallsvektor der Vorgang gemeint sein, bei dem man die Zufallsgrößen eines Zufallsprozessausschnittes als Elemente zu einem Zufallsvektor zusammenfasst, wobei man alle möglichen Werte des Parameters innerhalb des angegebenen Intervalls in streng monoton steigender Reihenfolge einsetzt. Damit kann man die Stationarität eines Zufallsvektors analog definieren, wobei eine Parameterverschiebung dann einer Indexverschiebung entspricht. Entsprechend definiert man die Zyklstationarität mit einer Periode, indem man bei der Definition die Identität der beiden Verbundverteilungen, bzw. die Identität der Momente nur für solche Werte der Parameterverschiebung fordert, die sich jeweils um ganzzahlige Vielfache der Periode unterscheiden.

Jedes am realen System auftretende determinierte Signal kann man sich als eine unbegrenzt lange Musterfolge eines Zufallsprozesses vorstellen. Die Musterfolge ist unbegrenzt lang, weil der Parameter des Zufallsprozesses die diskrete Zeit k ist, die ein Element der unendlichen Menge der ganzen Zahlen ist. Solch eine Musterfolge ist eine konkrete Realisierung unbegrenzter Länge, also eine konkrete Stichprobe vom Umfang 1 des Zufallsprozesses. Will man Aussagen über die stochastischen Eigenschaften des zugrundeliegenden

Zufallsprozesses empirisch gewinnen, so wird man sich eine endliche Anzahl L von endlich langen Signalausschnitten gleicher Länge F aus der zeitlich unbegrenzten Musterfolge in einem bestimmten Zeitraum — dem gesamten Zeitraum der Messung — herausausschneiden, und diese in geeigneter Weise verarbeiten. Die daraus gewonnen Aussagen können prinzipiell nur Schätzwerte für die zu bestimmenden theoretischen Größen des Zufallsprozesses sein. Außerdem kann so nur ein Teil der den Zufallsprozess bestimmenden Größen abgeschätzt werden.

Die L Signalausschnitte seien mit $\lambda = 1$ (1) L durchnummeriert. Der Signalausschnitt mit der Nummer λ stellt eine konkrete Realisierung, also eine Stichprobe vom Umfang 1 des Zufallsvektors mit derselben Nummer λ dar. Insgesamt hat man L konkrete Realisierungen von L Zufallsvektoren. Der Zufallsvektor mit der Nummer λ enthält als Elemente die F Zufallsgrößen des Zufallsprozesses der Zeitpunkte k , die dem Zeitintervall des Signalausschnittes mit der Nummer λ entsprechen. Bei jedem dieser L Zufallsvektoren kann man eine Verbundverteilung seiner F Elemente angeben. Im allgemeinen sind die L Verbundverteilungen der L Zufallsvektoren nicht identisch. Wenn nun die L Zeitintervalle der Signalausschnitte so gewählt worden sind, dass alle L Zufallsvektoren unabhängig sind, und dass alle dieselbe Verbundverteilung ihrer Elemente aufweisen, so kann man die L Signalausschnitte als eine konkrete Stichprobe vom Umfang L eines einzigen Zufallsvektors, der im weiteren als Modellzufallsvektor bezeichnet sei, mit derselben Verbundverteilung interpretieren. Unabhängig sind die L Zufallsvektoren dann, wenn sich die $F \cdot L$ -dimensionale Verbundverteilung von allen $F \cdot L$ Elementen aller L Zufallsvektoren als das Produkt der F -dimensionalen Verbundverteilungen der jeweils F Elemente der L Zufallsvektoren schreiben lässt. Die konkrete Stichprobe vom Umfang L , also die Gesamtheit aller L aus der konkreten Musterfolge des Zufallsprozesses herausgeschnittenen Signalausschnitte der Länge F , ist dann eine konkrete Realisierung einer mathematischen Stichprobe vom Umfang L des F -dimensionalen Modellzufallsvektors. Unter einer mathematischen Stichprobe vom Umfang L versteht man hier das $F \cdot L$ -dimensionale Zufallsexperiment, des Ziehens der konkreten Stichprobe. In unserem Fall ist die Zufälligkeit dieses Experiments dadurch gegeben, dass die zeitlich unbegrenzte Musterfolge, aus der alle L Signalausschnitte entnommen worden sind, selbst eine konkrete Realisierung eines Zufallsprozesses ist. Mathematisch wird eine Stichprobe nur dann genannt, wenn alle L Elemente der Stichprobe unabhängig sind, und wenn die Verbundverteilung jedes Elements der Stichprobe bei allen Elementen der Stichprobe gleich der Verbundverteilung des Zufallsvektors ist, aus dem die Stichprobe entnommen wurde. Genau diese Forderung wurde eben an die Art der Auswahl der L Signalausschnitte gestellt. Sinnvollerweise sollte die Auswahl der L Zeitintervalle der Signalausschnitte noch ein weiteres Kriterium erfüllen. Es ist nämlich zu fordern, dass die interessierenden stochastischen Eigenschaften des

Zufallsprozesses identisch sind mit den entsprechenden stochastischen Eigenschaften des Modellzufallsvektors, dessen konkrete Stichprobe die ausgewählten Signalausschnitte der Musterfolge sind. Welche Signalausschnittlänge man wählen muss, hängt dabei vor allem davon ab, welche stochastischen Eigenschaften man bestimmen will. Will man beispielsweise bei einem zyklstationären Prozess, dessen Momente mit 17 periodisch sind, alle 17 Korrelation zweier Signalwerte bestimmen, die um 7 Takte auseinanderliegen, so ist für die Signalausschnittlänge F wenigstens der Wert $F \geq 17+7 = 24$ zu wählen. Anderenfalls wären die an einigen der interessierenden Korrelationen beteiligten Zufallsgrößen gar nicht in den gewählten Ausschnitten des Zufallsprozesses vorhanden. Allgemein empfiehlt es sich, die Länge der Signalausschnitte schon deshalb eher größer zu wählen, als unbedingt erforderlich, weil dann zu erwarten ist, dass sich die gewünschten stochastischen Eigenschaften bei gleichem Stichprobenumfang L mit einer kleineren Varianz bestimmen lassen. Wir wollen uns im weiteren mit der Frage beschäftigen, was man bei der Auswahl der L Zeitintervalle der Signalausschnitte beachten sollte, damit die bisher aufgestellten Forderungen als erfüllt angesehen werden können.

Die Zufallsprozesse, die an einem realen System anliegen, sind immer instationär, da das System erstens nicht ewig existiert, und zweitens sich nicht immer in demselben Betriebszustand befindet. So kann z. B. ein und dasselbe System zur Übertragung verschiedener Signale dienen, deren stochastische Eigenschaften sich grundlegend unterscheiden, oder das System wird zu unterschiedlichen Zeiten in unterschiedlichen Umgebungen betrieben, in denen eine grundlegend andere Störsituation vorliegt. Daher wird die Forderung nach Gleichheit der L Verbundverteilungen nur erfüllbar sein, wenn man die Wahl der Zeitintervalle einschränkt. Man wird zunächst die unterschiedlichen Betriebsarten des Systems klassifizieren müssen. Dann wird man die Messung am realen System, beginnend mit der Entnahme der Signalausschnitte, innerhalb eines Zeitraums durchzuführen, in dem sich das System in diesem typischen Betriebszustand befindet, der nicht zuletzt auch durch die typischen stochastischen Eigenschaften des erregenden sowie des störenden Prozesses festgelegt ist. Innerhalb des Zeitraums der Messung wird man diesen Betriebszustand nicht verlassen. Es sei darauf hingewiesen, dass die stochastischen Eigenschaften der Zufallsprozesse innerhalb dieses typischen Betriebszustandes keinesfalls stationär sein müssen. Bei geeigneter Entnahme der Signalausschnitte wird man so eine Aussage über das System- und Störverhalten gewinnen, die dann auf andere Zeiträume übertragbar ist, in denen sich das reale System ebenfalls in diesem typischen Betriebszustand befindet. Um die Verwendbarkeit der Messergebnisse sicherzustellen ist es daher wichtig, den Betriebszustand, in dem sich das System bei der Messung befindet, genau zu dokumentieren. Wie genau der Betriebszustand der Messung festgelegt sein sollte, hängt davon ab, wie empfindlich das System auf eine Änderung des Betriebszustands reagiert. Es ist daher immer im Einzel-

fall entweder anhand heuristischer Überlegungen oder durch Messung mit geändertem Betriebszustand zu prüfen, wie genau der bei der Messung vorliegende Betriebszustand beschrieben werden muss. Wir wollen im weiteren davon ausgehen, dass der typische Betriebszustand während der Messung nicht verlassen wird, und sind uns darüber im Klaren, dass die gemessenen stochastischen Eigenschaften nicht die des für alle Zeiten anliegenden Zufallsprozesses sind, sondern lediglich für die Zeiträume gültig sind, innerhalb derer sich das System in demselben Betriebszustand wie bei der Messung befindet.

Innerhalb des Zeitraums der gesamten Messung muss es — wie gesagt — mindestens L Zeitintervalle der Länge F geben, innerhalb derer die L Verbundverteilungen der L Zufallsvektoren, die durch Ausschneiden der Intervalle aus dem Zufallsprozess entstehen, gleich sind. Diese F -dimensionale Verbundverteilung ist die des Modellzufallsvektors. Ob die Forderung, dass die Verbundverteilungen aller L Zufallsprozessausschnitte gleich sein sollen, erfüllt ist, kann man letztlich nicht überprüfen, da man für ein Zeitintervall immer nur eine konkrete Realisierung des zugrundeliegenden Zufallsprozesses messen kann. Bei manchen Systemen ist der Zugriff auf das System im realen Betrieb auf bestimmte Zeitintervalle begrenzt, und man kann annehmen, dass die stochastischen Eigenschaften aller Zufallsprozessausschnitte innerhalb der zugreifbaren Intervalle identisch sind, sofern der typische Betriebszustand vorliegt. Man kann daher beliebige zugreifbare Zeitintervalle auswählen. Bei anderen Systemen ist anzunehmen, dass es sich bei dem zugrundeliegenden Zufallsprozess um einen zyklstationären Prozess handelt, für dessen Beschreibung die Kenntnis einer Periode der stochastischen Eigenschaften ausreicht. Bei solchen Systemen wird man die Zeitintervalle, die man zur Gewinnung der Signalausschnitte heranzieht, synchron zur Periodizität der Eigenschaften des Zufallsprozesses wählen, so dass dadurch alle Zufallsprozessausschnitte dieselben stochastischen Eigenschaften aufweisen. Bei den meisten Systemen bestehen keinerlei Einschränkungen hinsichtlich der Zeit des Zugriffs, und es kann angenommen werden, dass der zugrundeliegende Zufallsprozess innerhalb des Zeitraums des typischen Betriebszustands stationär ist. Bei solchen Systemen sind die stochastischen Eigenschaften aller Zufallsprozessausschnitte identisch, egal wie man die Ausschnitte wählt.

Eine weitere wichtige Forderung war, dass die L Zufallsvektoren, die als die Ausschnitte des zugrundeliegenden Zufallsprozesses definiert sind, die den Zeitintervallen der Signalentnahme entsprechen, unabhängig sein müssen. Die L Zeitintervalle dürfen sich nicht überlappen, da sonst die L Zufallsvektoren nicht unabhängig sein könnten, weil unterschiedliche Zufallsvektoren dann gleiche Zufallsgrößen als Elemente enthalten würden. Bei vielen Systemen kann man davon ausgehen, dass die Unabhängigkeit dadurch erreicht werden kann, dass man die Zeitintervalle mit einem genügend großen zeitlichen Abstand wählt. Wenn die Unabhängigkeit der Zufallsprozessausschnitte gegeben ist, folgt aus der

Faktorisierbarkeit der gemeinsamen Verbundverteilung aller L Stichprobenelemente, dass das zweite *zentrale* Moment zweier beliebiger Zufallsgrößen, die aus unterschiedlichen Intervallen aus dem Zufallsprozesses entnommen worden sind, null ist. Die Autokovarianzfolge (nicht die Autokorrelationsfolge), die die zeitdiskrete zweidimensionale Folge der zweiten *zentralen* Momente ist und die von den beiden Zeitpunkten der Entnahme der Zufallsgrößen abhängt, ist in diesem Fall für alle Zeitpunkttupel, deren Elemente in unterschiedlichen Zeitintervallen liegen, null. Bei vielen realen Systemen kann man annehmen, dass die Autokovarianzfolge null ist, wenn die Zeitpunkte nur weit genug auseinander liegen, so dass die Autokovarianzfolge auf ein festes Intervall der Differenz der Zeitpunkte begrenzt ist. Es gibt aber auch technisch relevante Zufallsprozesse, bei denen diese zeitliche Begrenzung der Autokovarianzfolge nicht vorliegt, oder bei denen die zeitliche Begrenzung der Autokovarianzfolge so groß ist, dass das Einfügen von Pausen bei der Signalentnahme *de facto* nicht durchführbar ist. So kann es z. B. vorkommen, dass den Signalen an einem realen System ein periodischer Störer überlagert ist. Dieser Störer kann entweder deterministisch oder stochastisch sein. Ein deterministischer periodischer Störer bewirkt ein periodisch zeitabhängiges erstes Moment, das nicht in die Autokovarianzfunktion eingeht. Die Unabhängigkeit der Zufallsprozessausschnitte kann auch hier gegeben sein (muss aber nicht). Ein stochastischer periodischer Störer liegt beispielsweise vor, wenn die Phasenlage des periodischen Störsignals zufällig ist. Die Autokovarianzfolge enthält im letztgenannten Fall einen Anteil, der in ihren beiden unabhängigen Variablen — den Zeitpunkten der zur Berechnung des zweiten zentralen Moments herangezogenen Zufallsgrößen — periodisch und somit (wenigstens innerhalb des Zeitraums der Messung) zeitlich unbegrenzt ist. Daher kann die Unabhängigkeit der L Zufallsvektoren in Falle eines additiven, zufälligen periodischen Störers nicht dadurch erreicht werden, dass man die L Zeitintervalle mit einem hinreichend großen Abstand wählt. Anhand der einzigen Musterfolge, die im Zeitraum der Messung vorliegt, kann man nicht entscheiden, ob es sich bei diesem Signalanteil um einen deterministischen Störer oder um eine Musterfolge eines stochastischen Störers handelt. Man muss daher mit Hilfe heuristischer Überlegungen entscheiden, wie man diese Störung in der Musterfolge interpretieren will, und die Art der Festlegung der Zeitintervalle zur Signalentnahme eventuell entsprechend anpassen, um diese Störung entweder dem ersten oder dem zweiten zentralen Moment zuzuordnen. Zu berücksichtigen ist dabei, in welcher Art der Störer im normalen Betrieb des Systems vorliegt.

Erwartet man bei einem System, das auf einen festen Takt synchronisiert ist, dass im normalen Betrieb eine starre Phasenkopplung zwischen dem Systemtakt und dem Störer vorliegt, so wird man diese Art von Störung als deterministische Störung und somit als zeitabhängiges Moment erster Ordnung im zugrundeliegenden Zufallsprozess interpretie-

ren. Wenn man für die Messung der stochastischen Eigenschaften zur Signalentnahme nur solche Zeitintervalle auswählt, die um ganzzahlige Vielfache der Periode der Störung auseinanderliegen, werden auch bei Anwesenheit einer deterministischen Störung die L Zufallsvektoren unabhängig sein, und gleiche Verbundverteilungen aufweisen, wie dies ohne die deterministische Störung der Fall ist. Aus der so gewonnenen konkreten Stichprobe kann man dann Schätzwerte für die stochastischen Eigenschaften des Modellzufallsvektors berechnen.

Erwartet man jedoch, dass die Phasenlage der periodischen Störung in keinem festen Zusammenhang mit der Art des Zugriffs auf das System steht, so können die L Zufallsvektoren, die die Zufallsgrößen des zugrundeliegenden Zufallsprozesses der L festen Zeitintervalle enthalten, schon wegen der Periodizität der Autokovarianzfolge nicht unabhängig sein. Dann muss die Auswahl der Zeitintervalle der Signalentnahme in Bezug auf die Phasenlage der Musterfolge des periodischen Störers zufällig erfolgen. Um diesen Fall beschreiben zu können erzeugen wir uns zunächst eine $F \times L$ Zufallsmatrix, die wir im weiteren als zufällige Stichprobenmatrix bezeichnen wollen. Wie diese Matrix entsteht sei nun erläutert. In einem ersten Schritt legen wir eine hinreichend große Anzahl $\tilde{L}$ (wenigstens L , wenn möglich mehr) von Zeitintervallen fest, die zur Signalentnahme geeignet sind, was heißen soll, dass bei diesen $\tilde{L}$ Zeitintervallen die entsprechenden Zufallsprozessausschnitte immer die gleiche F -dimensionale Verbundverteilung — nämlich die des Modellzufallsvektors — aufweisen. Die $\tilde{L} \cdot F$ Zufallsgrößen des zugrundeliegenden Zufallsprozesses bilden $\tilde{L}$ Zufallsvektoren. Den Spaltenvektoren der zufälligen Stichprobenmatrix ordnet man nun zufällig maximal L der $\tilde{L}$ Zufallsvektoren zu. Wie diese zufällige Zuordnung vorgenommen wird sei zunächst völlig frei, so dass sich die Wahrscheinlichkeiten, einem Spaltenvektor einen bestimmten der $\tilde{L}$ Zufallsvektoren zuzuordnen, auch gegenseitig bedingen können. Beispielsweise kann man solche Zuordnungen ausschließen, die denselben Zufallsvektor zwei unterschiedlichen Spaltenvektoren zuordnen. Die freie Wahl der Art der zufälligen Zuordnung der Zufallsprozessausschnitte zu den L Spaltenvektoren kann dazu genutzt werden, gewisse stochastische Eigenschaften der zufälligen Stichprobenmatrix günstig zu beeinflussen. Die $F \cdot L$ -dimensionale Verbundverteilung der Elemente der zufälligen Stichprobenmatrix ist der Erwartungswert der $F \cdot L$ -dimensionalen Verbundverteilung über alle möglichen Zuordnungen. Er berechnet sich als die gewichtete Summe über alle $F \cdot L$ -dimensionalen Verbundverteilungen der Zufallsmatrizen, die sich bei allen möglichen konkreten Zuordnungen ergeben. Die Gewichtung erfolgt dabei nach der Wahrscheinlichkeit, mit der eine konkrete Zuordnung auftritt. Jede $F \cdot L$ -dimensionale Verbundverteilung jeder konkreten Zuordnung wird sich wegen der Anwesenheit des zufälligen periodischen Störers nicht als das Produkt der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren schreiben lassen. Unabhängig davon, wie die Wahrscheinlichkeiten der Zuordnungen ge-

wählt werden, sind alle F -dimensionalen Verbundverteilungen aller L Spaltenvektoren der zufälligen Stichprobenmatrix immer identisch und gleich der F -dimensionalen Verbundverteilung des Modellzufallsvektors. Die $\tilde{L}$ Zeitintervalle wurden ja so gewählt, dass die F -dimensionalen Verbundverteilungen der $\tilde{L}$ Zufallsvektoren gleich sind. Bildet man die F -dimensionale Verbundverteilung eines Spaltenvektors, indem man den Erwartungswert aller möglichen F -dimensionalen Verbundverteilungen der $\tilde{L}$ Zufallsvektoren — gewichtet gemäß der Auftrittswahrscheinlichkeit der Zuordnungen — berechnet, so kann diese F -dimensionale Verbundverteilung vor die Erwartungswertbildung gezogen werden, und der Erwartungswert ist dann nur mehr die Summe der Wahrscheinlichkeiten aller möglichen Zuordnungen, die immer eins ist. Die Wahrscheinlichkeiten aller möglichen Zuordnungen sollen nun möglichst so gewählt werden, dass die $F \cdot L$ -dimensionale Verbundverteilung aller Elemente der zufälligen Stichprobenmatrix gleich dem Produkt der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren ist. Ob dies möglich ist, hängt sicher von den stochastischen Eigenschaften des zugrundeliegenden Zufallsprozesses ab. Sollte dies nicht möglich sein, so wird man die Freiheit, die Wahrscheinlichkeiten aller möglichen Zuordnungen beliebig wählen zu können, dazu nutzen, eine möglichst gute Approximation des Produkts der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren durch die $F \cdot L$ -dimensionale Verbundverteilung der zufälligen Stichprobenmatrix zu erreichen. Wie man den Fehler dieser Approximation sinnvollerweise definiert, kann man nicht generell sagen. Dies hängt von der Wahl der zu bestimmenden stochastischen Eigenschaften des zugrundeliegenden Zufallsprozesses ab. Trotzdem kann man sicher sein, dass es wenigstens eine Wahl für die Wahrscheinlichkeiten der Zuordnungen gibt, die das Approximationsziel zumindest nicht schlechter annähert, als die beste aller bei der Erwartungswertbildung auftretenden Zuordnungen, also als jede feste Zuordnung mit der Wahrscheinlichkeit Eins. Hat man die Wahrscheinlichkeiten aller möglichen Zuordnungen geeignet gewählt, so kann man erwarten, dass die Messwerte der stochastischen Eigenschaften, die man anhand der einer konkreten Stichprobenmatrix erhält, mit den zu messenden stochastischen Eigenschaften des Modellzufallsvektors gut übereinstimmen. Das Hauptproblem der Wahl der Wahrscheinlichkeiten der Zuordnungen besteht darin, dass man die $F \cdot L$ -dimensionalen Verbundverteilungen der einzelnen Zuordnungen, deren Erwartungswert man zu bilden hat, ebensowenig kennt, wie das zu approximierende Produkt der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren, deren stochastische Eigenschaften gemessen werden sollen. Bei den meisten Systemen, die von einer periodischen Störung überlagert sind, kann man annehmen, dass die gleichwahrscheinliche Wahl aller Zuordnungen, unter Ausschluss der Zuordnungen, die zu identischen Spaltenvektoren führen, zu dem gewünschten Ergebnis führt, wenn man die Anzahl $\tilde{L}$ der Zeitintervalle groß genug wählt und diese gleichmäßig über mehrere Perioden einer zu erwartenden Störung verteilt. Gegebenfalls sollte man die Messung wiederholen und an-

hand der Abweichung der Messergebnisse entscheiden, ob man die Messung mit geänderter Stichprobenentnahme durchführen sollte. Theoretisch könnte man auch die Messung mehrfach wiederholen, und anschließend die Hypothese testen, dass alle Messergebnisse aus ein und demselben Zufallsvektor stammen. Dieses Verfahren dürfte allerdings in der Praxis wegen der dazu notwendigen immensen Messdauer kaum anwendbar sein. Heuristische Überlegungen bezüglich der Art der zu erwartenden Störung erscheinen hier angebrachter. Da bei der Messung der stochastischen Eigenschaften die Reihenfolge der Elemente der Stichprobe in der Regel — wie auch beim RKM — keinen Einfluss auf die Messergebnisse hat, sind alle Stichprobenmatrizen, die durch reine Spaltenpermutationen ineinander überführt werden können, gleich gut geeignet. Deshalb ist bei vielen Systemen anzunehmen, dass die Unabhängigkeit der Spaltenvektoren — also der Stichprobenelemente — in guter Näherung dadurch erreicht werden kann, dass man bei der Entnahme zweier zeitlich aufeinanderfolgender Stichproben eine Pause zufälliger Länge einlegt, so dass sich die Zeitintervalle der Stichproben über mehrere Perioden der zu erwartenden Störung etwa gleichmäßig verteilen.

Die Essenz der bisher gemachten Betrachtungen besteht darin, dass man bei Systemen, bei denen man nicht sicher sein kann, dass keine Störungen auftreten, die bei festgelegter Wahl der Zeitintervalle zu abhängigen Signalausschnitten führen, eine zufällige Wahl der Zeitintervalle der Stichprobenentnahme durchführen sollte, und dass man die Messergebnisse überprüfen sollte, wenn man sich nicht sicher ist, dass die zufällige Entnahme zur gewünschten Unabhängigkeit führt. Es wird nun anhand eines ausführlichen realitätsnahen Beispiels gezeigt, wie man bei den Ein- und Ausgangssignalen eines Systems durch die zufällige Entnahme von Signalausschnitten aus deren Musterfolgen innerhalb des Messzeitraums zu der gewünschten Stichprobenmatrix kommt. Dabei treten typische Vertreter aller oben beschriebenen Arten von Störungen auf.

In diesem Beispiel werden verschiedene Signale über die N logischen Kanäle des in Bild 2.1 dargestellten TDMA-Systems über einen linearen, zeitinvarianten, dispersiven und gestörten physikalischen Übertragungskanal übertragen. Bei dem TDMA (Time-Division Multiple Access) Verfahren werden den logischen Kanälen jeweils periodisch abwechselnd Zeitschlitz der Länge F zugeordnet. Nachdem jeweils ein Zeitschlitz von jedem logischen Kanal übertragen wurde, wird ein Zeitschlitz der Länge F eingefügt, in dem eine im Empfänger bekannte Signalsequenz gesendet wird. Diese dient der Synchronisation des Empfängers einerseits auf die Periode $P = F \cdot (N+1)$, mit der sich die Zeitschlitz aller logischen Kanäle wiederholen und andererseits auf den Takt, mit dem sich die $N+1$ Zeitschlitz abwechseln. In der weiteren theoretischen Interpretation der Prozesse am TDMA-System wird angenommen, dass die Synchronisation des Empfängers auf den Sendetakt fehlerfrei erfolgt. Die sich periodisch wiederholende Synchronisationssequenz ist

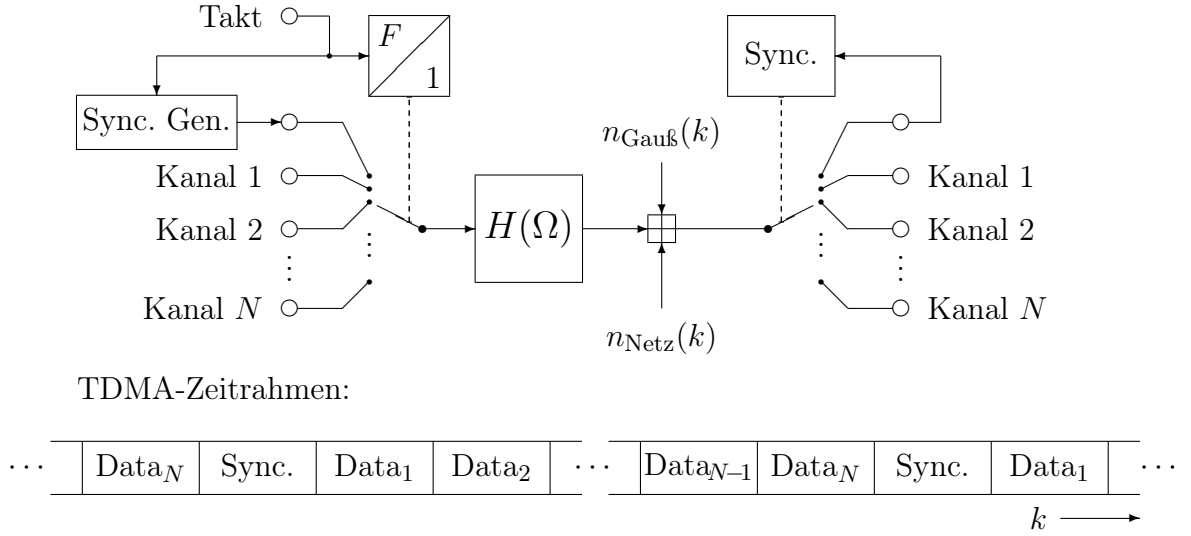


Bild 2.1: TDMA-Beispielsystem und TDMA-Zeitrahmen

ein deterministisches Signal, das in Phase und Frequenz exakt festliegt, und deren Phasenlage sich bezüglich der Zeitschlitzes der logischen Kanäle nicht ändert. Der physikalische Übertragungskanal, über den das so konstruierte TDMA-Signal übertragen wird, habe eine Impulsantwort, die auf die Länge eines Zeitschlitzes begrenzt sei. In diesem Übertragungskanal seien nun zwei Störquellen vorhanden. Die eine Störquelle ist ein additives, stationäres, weißes und gaußverteiltes Rauschen, während die andere Störquelle ein durch die Energieversorgung verursachter netzfrequenter Sinuseintonstörer der Kreisfrequenz Ω_{Netz} ist (Netzbrumm). Der typische Betriebszustand des Systems sei unter anderem dadurch beschrieben, dass der Sinuseintonstörer immer die gleiche Amplitude hat. Der Sinuseintonstörer und der Takt der Zeitschlitzes seien *nicht* synchron, d. h. ihre Frequenzen stehen nicht in einem trivialen rationalen Verhältnis zueinander. Es liegt daher nahe anzunehmen, dass die Phase des netzfrequenten Sinuseintonstörers in jedem Zeitpunkt innerhalb des Messzeitraums als im Bereich von $-\pi$ bis π gleichverteilt angenommen werden kann. Jede Verbundverteilung beliebiger Abtastwerte des Sinuseintonstörers ändert sich nicht, wenn man alle daran beteiligten Abtastwerte durch Abtastwerte ersetzt, die alle um eine beliebige aber bei allen Abtastwerten gleiche Zeit verschoben sind. Der Zufallsprozess des Sinuseintonstörers ist daher stationär. Die zweidimensionale Autokovarianzfolge zweier Abtastwerte des netzfrequenten Sinuseintonstörers ist die Folge der Abtastwerte der Kosinusfunktion mit der Netzfrequenz und dem halben Quadrat der Amplitude der Störung. Im Argument der Kosinusfunktion tritt dabei nur der zeitliche Abstand der beiden Abtastzeitpunkte der zufälligen Sinuseintonstörung auf, was aufgrund der Stationarität dieser Störung auch so sein muss. Da diese Autokovarianzfolge nicht auf eine Periode P des TDMA-Rahmens begrenzt ist, kann die Verbundverteilung der beiden Abtastwerte des Sinuseintonstörers auch dann nicht als das Produkt der beiden eindimensionalen

Verteilungen eines Sinussignals mit gleichverteilter Phase geschrieben werden, wenn die Abtastwerte zeitlich um mehr als P auseinander liegen.

In diesem Beispielsystem sollen nun die stochastischen Eigenschaften der Prozesse am Systemein- sowie -ausgang gemessen werden. Die Messung soll für alle N logischen Kanäle getrennt erfolgen. Die Wahl der $\tilde{L}$ Zeitintervalle ist nicht beliebig möglich, sondern muss mit einer festen Phasenlage bezüglich des P -fachen Systemtakts erfolgen. Alle innerhalb des Messzeitraums liegenden und für eine Signalausschnittentnahme in Frage kommenden Zeitintervalle sind also immer um ganzzahlige Vielfache von P gegeneinander verschoben. Während ein Kanal gemessen wird, werden auf den anderen Kanälen modulierte Datensymbole übertragen, die in der Art moduliert sind, dass deren Anteile am Sendesignal nur in den Zeitschlitten von null verschieden sind, die den jeweiligen logischen Kanälen entsprechen. Von den Datensymbolen, die in den gerade nicht gemessenen Kanälen übertragen werden, nehmen wir an, dass diese aus Zufallsprozessen entnommen sind, die einen Anteil am Sendesignal erzeugen, der mit der Periode P zyklstationär ist, und dessen Erwartungswert für jeden Zeitpunkt null ist. Nun erregen wir den zu vermessenden logischen Kanal mit einem zufälligen Testsignal. Zur Generierung des Testsignals wird für die Zeitschlitz des zu vermessenden logischen Kanals, die innerhalb des Zeitraums der Messung liegen, aus einem mittelwertfreien Zufallsvektor eine konkrete Stichprobe entnommen. Jedes Element der Stichprobe bildet einen Signalausschnitt. Diese Signalausschnitte werden als Sendesignale für die Zeitschlitz des zu vermessenden Kanals verwendet. Der Umfang der Stichprobe wird dabei so groß gewählt, dass alle $\tilde{L}$ Zeitschlitz des zu vermessenden Kanals, die innerhalb des Zeitraums der Messung liegen, mit Elementen der Stichprobe aufgefüllt sind. Der gesamte sendeseitige Zufallsprozess ist daher innerhalb des Zeitraums der Messung zyklstationär mit der Periode P , und die auf die Zeitschlitz des zu vermessenden Kanals begrenzten Zufallsvektoren der Elemente der mathematischen Stichprobe sind unabhängig. Wählt man sich aus allen $\tilde{L}$ Zeitschlitten des zu vermessenden logischen Kanals innerhalb des Zeitraums der Messung nun L beliebige Zeitschlitz fest aus, so stellen die L Signalausschnitte der Musterfolge des erregenden Zufallsprozesses eine konkrete Realisierung einer mathematischen Stichprobenmatrix dar, deren Spaltenvektoren alle dieselbe Verbundverteilung besitzen wie der Zufallsvektor, aus dem die Sendesignalsequenzen gewonnen wurden. Eine zufällige Auswahl der L Zeitschlitz ist daher beim ungestörten Sendesignal nicht notwendig. Eine Schätzung der stochastischen Eigenschaften des erregenden Zufallsvektors anhand einer konkreten Realisierung der zufälligen Stichprobenmatrix ist daher auch ohne eine zufällige Zuordnung der Zeitintervalle zu den Spaltenvektoren der Stichprobenmatrix möglich.

Betrachten wir nun die $\tilde{L}$ Zufallsvektoren des *ersten* logischen Kanals am Ausgang des physikalischen Übertragungskanal. Diese Zufallsvektoren enthalten die Zufallsprozess-

ausschnitte der Zeitschlitz, die den $\tilde{L}$ Synchronisationszeitschlitz im Zeitbereich der Messung unmittelbar folgen. Die stochastischen Eigenschaften dieser $\tilde{L}$ Zufallsvektoren werden durch die Eigenschaften der unterschiedlichen am Gesamtausgangsprozess beteiligten Teilprozesse bestimmt. Da ist zunächst der stationäre, weiße Gaußprozess, dessen Anteile bei allen $\tilde{L}$ Zufallsvektoren unabhängig sind. Dem überlagern sich additiv die mit der Periode P zyklstationären Prozesse, die die Sendesignale der anderen gerade nicht vermessenen logischen Kanäle erzeugen, und die durch den physikalischen Übertragungskanal linear verzerrt sind. Da wir angenommen hatten, dass die Impulsantwort des Übertragungskanals auf die Dauer eines Zeitschlitzes begrenzt ist, und da die Zeitschlitz des gerade vermessenen logischen Kanals den Synchronisationszeitschlitz folgen, liefern diese Teilprozesse keinen Beitrag zu den $\tilde{L}$ betrachteten Zufallsvektoren. Ein weiterer Teilprozess ist der mit der Periode P zyklstationäre Prozess, aus dem das Testsignal gewonnen wird, der ebenfalls durch den physikalischen Übertragungskanal linear verzerrt ist. Von diesem Teilprozess wird nur der Anteil, der nach der linearen Verzerrung in die Zeitschlitz des gerade vermessenen logischen Kanals fällt, in den $\tilde{L}$ Zufallsvektoren vorhanden sein. Da erstens die Impulsantwort des Übertragungskanals auf einen Zeitschlitz begrenzt ist, da zweitens die Zeitintervalle mit einem größeren zeitlichen Abstand entnommen werden, und da drittens die Stichprobenelemente des Testsignals unabhängig sind, sind auch die Beiträge des Testsignals zu den $\tilde{L}$ Zufallsvektoren unabhängig. Die Anteile des Testsignalprozesses und des Gaußprozesses sind ebenfalls voneinander unabhängig. Als weiterer Anteil ist noch der stationäre Zufallsprozess des netzfrequenten Sinuseintonstörers vorhanden, der ebenfalls von den anderen Teilprozessen unabhängig ist. Da alle Signalausschnitte des Sinuseintonstörers voneinander *abhängig* sind, sind auch deren Anteile an den $\tilde{L}$ Zufallsvektoren voneinander abhängig. Der Anteil des mit P periodischen deterministischen Synchronisationssignals, der durch die lineare Verzerrung des physikalischen Übertragungskanals jeweils in die den Synchronisationszeitschlitz folgenden, gerade vermessenen Zeitintervalle fällt, ist bei allen $\tilde{L}$ Zufallsvektoren gleich, und überlagert sich additiv. Er verursacht bei allen $\tilde{L}$ Zufallsvektoren dieselbe Verschiebung ihrer Verbundverteilungen. Da alle Teilprozesse voneinander unabhängig und stationär bzw. mit der Periode P zyklstationär sind, und da sich alle Teilprozesse additiv überlagern, ist die F -dimensionale Verbundverteilung aller $\tilde{L}$ Zufallsvektoren gleich. Diese Verbundverteilung ist die unseres Modellzufallsvektors. Nun greifen wir uns aus den $\tilde{L}$ Zufallsvektoren L heraus und fassen diese Vektoren als Spaltenvektoren zu der zufälligen Stichprobenmatrix zusammen. Zunächst erfolgt das Herausgreifen der L Zufallsvektoren — also die Zuordnung eines Teils der $\tilde{L}$ Zufallsvektoren zu den L Spaltenvektoren der Stichprobenmatrix — zwar fast beliebig aber *nicht* zufällig. Wir wollen nur solche Zuordnungen zulassen, bei denen die L Spaltenvektoren der zufälligen Stichprobenmatrix aus L unterschiedlichen Zeitschlitz entnommen werden, so dass es nicht vorkom-

men kann, dass irgendeiner der $\tilde{L}$ Zufallsvektoren mehrfach als Spaltenvektor auftritt. Bei dem Element der zufälligen Stichprobenmatrix in der Zeile κ und in der Spalte λ erzeugt der Teilprozess des Sinuseintonstörers den Anteil $\sin(\Omega_{\text{Netz}} \cdot \kappa + \phi_\lambda)$, wobei ϕ_λ eine zufällige Phase ist, die bei allen Elementen eines Spaltenvektors gleich ist, und die sich aus zwei Anteilen zusammensetzt. Der erste Anteil ist die zufällige Phase ϕ der im Zeitraum der Messung vorliegenden Sinusstörung. Sie ist bei allen Elementen der zufälligen Stichprobenmatrix gleich. Der zweite Anteil ist die Phasenverschiebung $\Omega_{\text{Netz}} \cdot P \cdot \tilde{\lambda}_\lambda$, die bei allen Elementen eines Spaltenvektors jeweils gleich ist. $\tilde{\lambda}_\lambda$ ist dabei die Nummer des Zeitschlitzes, dessen Zufallsprozessausschnitt bei der gewählten Zuordnung dem λ -ten Spaltenvektor zugeordnet wird. Die zufälligen Abtastwerte der Sinusstörung, die in allen Elementen der zufälligen Stichprobenmatrix additiv vorhanden sind, sind also unterschiedlich phasenverschobene nichtlineare Funktionen von ein und derselben Zufallsgröße und somit nicht unabhängig. Die $F \cdot L$ -dimensionale Verbundverteilung der zufälligen Abtastwerte der Sinusstörung ist daher nicht faktorisierbar. Weil die an der zufälligen Stichprobenmatrix beteiligten Teilprozesse alle voneinander unabhängig sind, ergibt sich die $F \cdot L$ -dimensionale Verbundverteilung der Elemente der zufälligen Stichprobenmatrix als die mehrfache $F \cdot L$ -dimensionale Faltung¹ der $F \cdot L$ -dimensionalen Verbundverteilungen der einzelnen beteiligten Teilprozesse. Weil wir identische Spaltenvektoren in der zufälligen Stichprobenmatrix ausgeschlossen haben, ergibt sich *ohne* den Zufallsprozess des netzfrequenten Sinuseintonstörers bei der Faltung eine $F \cdot L$ -dimensionale Verbundverteilung, die als Produkt der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren geschrieben werden kann. Bei der $F \cdot L$ -dimensionalen Faltung mit der $F \cdot L$ -dimensionalen Verbundverteilung des zufälligen Sinuseintonstörers ergibt sich keine faktorisierbare Verbundverteilung, weil die $F \cdot L$ -dimensionale Verbundverteilung des zu-

¹Begonnen wird mit einer Verbundverteilung eines Teilprozesses. Dann wird wiederholt die Faltung der Verbundverteilung eines noch nicht berücksichtigten Teilprozesses mit der bisher berechneten Verbundverteilung durchgeführt, bis alle Teilprozesse verarbeitet wurden. Bei der $F \cdot L$ -dimensionale Faltung handelt es sich um eine Integration über alle $F \cdot L$ Dimensionen des Raumes. In welchem Sinne die Integration hier definiert ist (Riemann, Stieltjes oder im Sinne der Distributionentheorie) sei dahingestellt. Integriert wird über das Produkt der $F \cdot L$ -dimensionalen Verbundverteilungsdichte des einen Teilprozesses und der $F \cdot L$ -dimensionalen Verbundverteilung des anderen Teilprozesses (Beim Stieltjes-Integral muss es heißen: Integral der Verbundverteilungsdichte bezüglich der Verbundverteilung). Bei der $F \cdot L$ -dimensionalen Verbundverteilung treten dabei die Veränderlichen der $F \cdot L$ -dimensionalen Verbundverteilungsdichte gespiegelt und verschoben auf. Über die $F \cdot L$ Veränderlichen der $F \cdot L$ -dimensionalen Verbundverteilungsdichte wird integriert, während die $F \cdot L$ Parameter der dabei auftretenden Verschiebung die Veränderlichen des Ergebnisses der Faltung sind. Festgehalten sei auch noch, dass sich beide Verbundverteilungen faktorisieren lassen, wenn die Spaltenvektoren der an der Faltung beteiligten Teilprozessausschnitte unabhängig sind. Dann kann die $F \cdot L$ -fache Integration als Produkt von L Mehrfachintegralen geschrieben werden, wobei jedes Mehrfachintegral nur über einem F -dimensionalen Raum zu berechnen ist. Daher ist die sich ergebende $F \cdot L$ -dimensionale Verbundverteilung ebenfalls wieder faktorisierbar.

fälligen Sinuseintonstörers nicht faktorisiert ist. Es besteht nun die Hoffnung, dass sich bei einer zufälligen Zuordnung eine faktorisierbare $F \cdot L$ -dimensionale Verbundverteilung ergibt. Wie oben hergeleitet, ergibt sich diese durch eine Erwartungswertbildung aus allen nicht faktorisierbaren Verbundverteilungen aller möglichen Zuordnungen. Wir haben also den Erwartungswert eines $F \cdot L$ -dimensionale Faltungsintegrals zu bilden, wobei an diesem Faltungsintegral einerseits die *nicht* faktorisierbare $F \cdot L$ -dimensionale Verbundverteilung des Sinuseintonstörers und andererseits die faktorisierbare $F \cdot L$ -dimensionale Verbundverteilung aller anderen Anteile beteiligt sind. Die Erwartungswertbildung kann nun in das $F \cdot L$ -dimensionale Faltungsintegral gezogen werden. *Ohne* Berücksichtigung des Sinuseintonstörers ist die $F \cdot L$ -dimensionale Verbundverteilung das Produkt der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren, das bei allen möglichen Zuordnungen gleich ist. Diese $F \cdot L$ -dimensionale Verbundverteilung kann daher vor die Erwartungswertbildung gezogen werden. Daher ergibt sich bei einer zufälligen Zuordnung die $F \cdot L$ -dimensionale Verbundverteilung der Stichprobenmatrix als die $F \cdot L$ -dimensionale Faltung des Produkts der L identischen F -dimensionalen Verbundverteilungen der Spaltenvektoren ohne Berücksichtigung des Sinuseintonstörers mit dem Erwartungswert der $F \cdot L$ -dimensionalen Verbundverteilung des zufälligen Sinuseintonstörers. Dieser Erwartungswert sollte durch die Wahl der Wahrscheinlichkeiten der möglichen Zuordnungen faktorisiert gemacht werden. Die Zuordnung $\tilde{\lambda}_\lambda$, die dem λ -ten Spaltenvektor die Nummer des Zeitschlitzes zuordnet, ist nun selbst eine Zufallsgröße, die von der zufälligen Phase ϕ der im Zeitraum der Messung vorliegenden Sinusstörung unabhängig ist. Die zufälligen Phasen ϕ_λ , setzen sich daher jeweils aus zwei unabhängigen Zufallsgrößen additiv zusammen, sind aber immer noch bei allen Elementen *eines* Spaltenvektors jeweils gleich. Weil alle zufälligen Abtastwerte eines Spaltenvektors jeweils Funktionen von ein und derselben Zufallsgröße ϕ_λ sind, genügt es, die Unabhängigkeit der L Phasen ϕ_λ zu zeigen, um die Unabhängigkeit der Spaltenvektoren der Zufallsmatrix der Abtastwerte des Sinuseintonstörers zu zeigen. Nun wollen wir annehmen, dass die Messdauer groß genug gewählt worden sei, so dass die $\tilde{L}$ Ausschnitte des netzfrequenten Sinuseintonstörers Phasenlagen $\Omega_{\text{Netz}} \cdot P \cdot \tilde{\lambda}_\lambda$ besitzen, die im gesamten Zeitraum der Messung in einen Bereich von mehreren Vielfachen von 2π fallen. Wählt man nun alle möglichen Zuordnungen, mit Ausnahme der Zuordnungen, die zu identischen Spaltenvektoren führen, gleich wahrscheinlich, so ist die Phasendifferenz zweier Stichprobenelemente des Sinuseintonstörers bei allen $\tilde{L}-1$ möglichen diskreten Phasendifferenzen gleichwahrscheinlich, und die möglichen Phasendifferenzen liegen — modulo 2π berechnet — etwa gleichmäßig im Bereich von $-\pi$ bis π verstreut. Die L -dimensionale Verbundverteilungsdichte der Phasen ϕ_λ ist als Ableitung der Verbundverteilung nach allen Veränderlichen definiert, und daher eine Funktion, die nur außerhalb eines eindimensionalen nichtlinearen Raumes (ein kontinuierlicher Freiheitsgrad der zufälligen Phase ϕ des Störers) existiert, und dort null ist.

Weil die Phase ϕ in alle Zufallsphasen ϕ_λ additiv eingeht, besteht der eindimensionale nichtlineare Raum aus $\tilde{L}/(\tilde{L}-L)!$ Geradenstücken, die alle parallel zu dem Vektor liegen, der nur Einsen enthält.

Im Teilbild a) des Bildes 2.2 ist dieser eindimensionale Raum, in dem die Verbundverteilungsdichte der Phasen ϕ_λ Impulslinien aufweist, für $L = 2$ und $\tilde{L} = 10$ dargestellt. Jedes Geradenstück entspricht dabei einer konkreten Zuordnung. Das Geradenstück einer konkreten Zuordnung ist um einen Vektor verschoben, der die bei der Zuordnung auftretenden $L = 2$ Phasenversätze $\Omega_{Netz} \cdot P \cdot \tilde{\lambda}_1$ und $\Omega_{Netz} \cdot P \cdot \tilde{\lambda}_2$ beider Spaltenvektoren enthält. Da die Elemente der Spaltenvektoren der Zufallsmatrix der Abtastwerte des Sinuseintonstörers phasenverschobene Sinusfunktionen von ϕ_λ sind, und somit mit 2π periodisch sind, kann man ebensogut die modulo 2π berechneten Zufallsgrößen betrachten. Wie man in Teilbild b) sieht, werden durch die modulo 2π Reduktion die verschobenen Geradenstücke zerteilt und durch Parallelverschiebung um Vielfache von 2π auf den Bereich des Quadrats mit der Kantenlänge 2π um den Koordinatenursprung abgebildet. Diese zerteilten Geradenstücke der ersten beiden und der letzten der insgesamt 90 möglichen konkreten Zuordnungen sind in Teilbild c) des Bildes 2.2 dargestellt. Innerhalb jedes zerteilten Geradenstückes einer konkreten Zuordnung, in dem die zweidimensionale Verbundverteilungsdichte nicht definiert ist, ist die eindimensionale bedingte² Verteilungsdichte der zufälligen Phase ϕ eine Gleichverteilung. Will man daraus für eine konkrete Zuordnung die L -dimensionale bedingte Verbundverteilung berechnen, so hat man allgemein das uneigentliche L -dimensionale Volumenintegral über die L -dimensionale bedingte Verbundverteilungsdichte zu berechnen, wobei die oberen Integrationsgrenzen die freien Veränderlichen der L -dimensionalen Verbundverteilung sind. Im ersten Bild des Teilbildes c), das die Lage der Impulslinien der zweidimensionalen Verteilungsdichte bei der ersten möglichen Zuordnung darstellt, ist das Gebiet hervorgehoben, über das zu integrieren ist, um einen bestimmten Wert der zweidimensionalen Verbundverteilung zu erhalten. In unserem Fall des Sinuseintonstörers sind ein oder zwei eindimensionale Integrale über den konstanten Wert $1/(2\pi)$ zu berechnen, wobei die Integrationsgrenzen von den beiden freien Veränderlichen der zweidimensionalen Verbundverteilung abhängen. Bei der hier vorliegenden Gleichverteilung innerhalb der Geradenstücke ist dieses Integral proportional zur Länge der Anteile der Geradenstücke, die in das hervorgehobene Gebiet fallen. Die sich bei der Integration ergebenden zweidimensionalen bedingten Verbundverteilungen der drei für die Graphik ausgewählten Zuordnungen sind jeweils unter den in Teilbild d) dargestellten Verbundverteilungsdichten als „Höhenliniengraphik“ dargestellt. Das Phasentupel, für das sich der Wert des Integrals über das hervorgehobene Gebiet ergibt, ist im linken Teilbild der ersten möglichen Zuordnung als „×“ markiert. Die Erwartungs-

²unter der Bedingung der konkreten Zuordnung

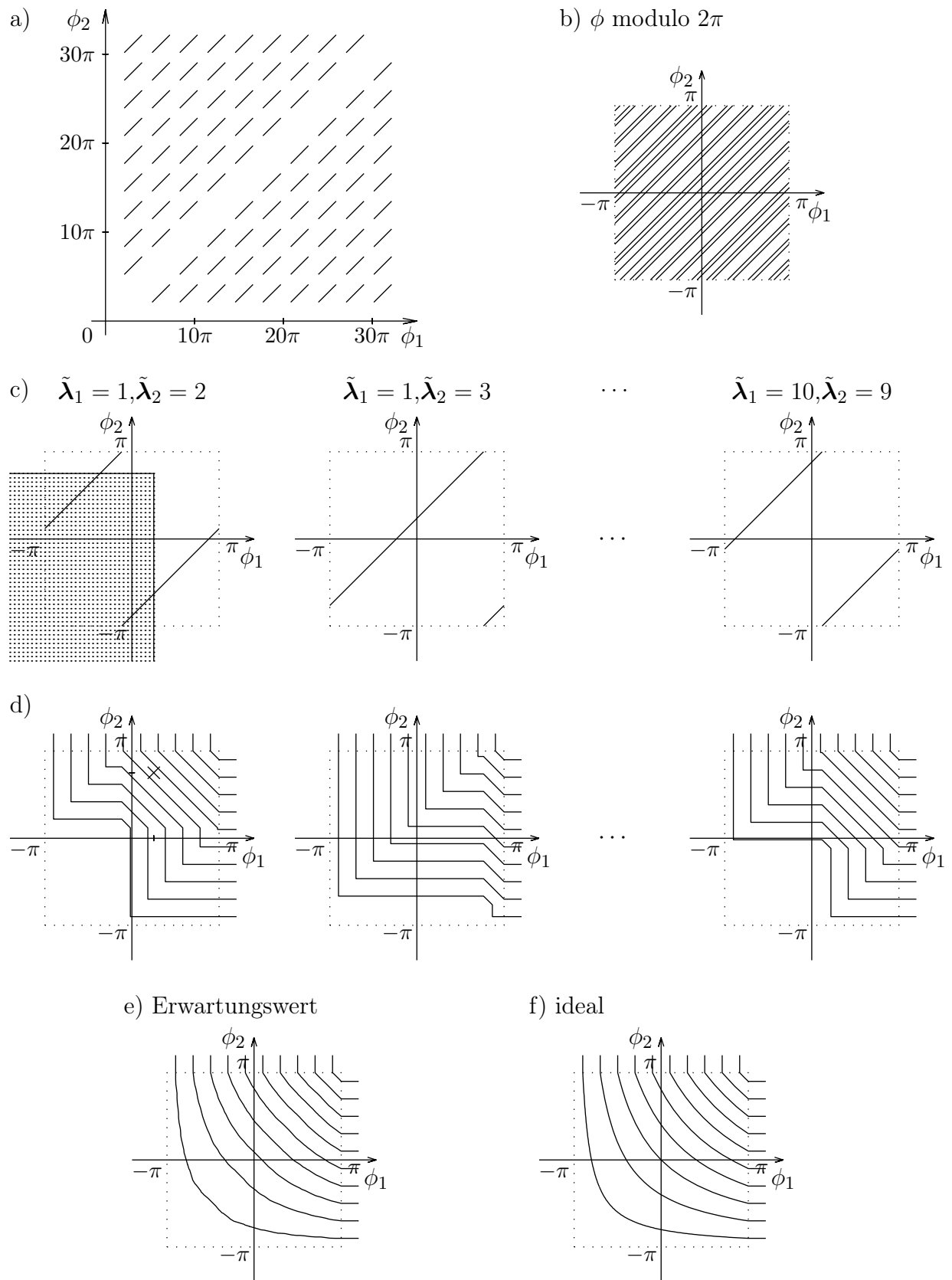


Bild 2.2: Zur Verbundverteilung der Phase des Sinusstörers

wertbildung über die zweidimensionalen bedingten Verbundverteilungen aller möglichen Zuordnungen liefert die zweidimensionale Verbundverteilung der zufälligen modulo 2π reduzierten Phasen ϕ_λ . Sie ist für unser einfaches Beispiel in Teilbild e) des Bildes 2.2 dargestellt. Wie man in Teilbild b) erkennt, sind die zerteilten Geradenstücke dadurch, dass die Phasenversätze $\Omega_{\text{Netz}} \cdot P \cdot \tilde{\lambda}_\lambda$ im Bereich von $-\pi$ bis π etwa gleichmäßig verstreut sind, ebenfalls etwa gleichmäßig auf den gesamten Bereich des Quadrats verstreut. Es wird sich — nicht nur in unserem Beispiel mit $L = 2$ — bei der Berechnung des Integrals etwa dieselbe L -dimensionale Verbundverteilung ergeben, wie bei der L -fachen Integration über eine L -dimensionale Gleichverteilung $(2\pi)^{-L}$, die sich problemlos faktorisieren lässt. Zum Vergleich ist die sich bei der Integration über die zweidimensionale Gleichverteilung ergebende Verbundverteilung in Teilbild f) zu sehen. Man kann daher annehmen, dass die modulo 2π reduzierten Phasen ϕ_λ bei einer zufälligen Zeitschlitzzuordnung in wesentlich besserer Näherung als unabhängig angesehen werden können, als im Fall einer festen Zuordnung der Zeitschlitzze, bei der sich die zweidimensionale Verbundverteilung als das Integral über ein einziges zerteiltes Geradenstück ergibt, wie dies in den Teilbildern d) zu sehen ist. Daher kann man auch die Spaltenvektoren der Zufallsmatrix der Abtastwerte des Sinuseintonstörers als näherungsweise unabhängig betrachten.

Wie wir oben gesehen haben wird somit auch die $F \cdot L$ -dimensionale Verbundverteilung des zufälligen Sinuseintonstörers näherungsweise faktorisierbar sein. Da diese bei der Berechnung der $F \cdot L$ -dimensionalen Verbundverteilung aller Anteile der zufälligen Stichprobenmatrix mit der faktorisierbaren $F \cdot L$ -dimensionalen Verbundverteilung der restlichen Anteile gefaltet wird, ist anzunehmen, dass sich auch die $F \cdot L$ -dimensionale Verbundverteilung aller Anteile der zufälligen Stichprobenmatrix am Ausgang des TDMA-Systems näherungsweise faktorisieren lässt. Somit sind die Spaltenvektoren der zufälligen Stichprobenmatrix näherungsweise unabhängig. Es sei nochmal darauf hingewiesen, dass die F -dimensionale Verbundverteilung jedes Spaltenvektors von der zufälligen Zuordnung nicht beeinflusst wird und bei allen Spaltenvektoren gleich ist. Insgesamt bilden die Spaltenvektoren der zufälligen Stichprobenmatrix daher näherungsweise eine mathematische Stichprobe des Modellzufallsvektors, der dieselbe F -dimensionale Verbundverteilung aufweist, wie jeder beliebige der $\tilde{L}$ Zufallsvektoren. Eine konkrete Stichprobe vom Umfang L — also eine konkrete Realisierung der zufälligen Stichprobenmatrix — erhält man, indem man aus der konkreten Musterfolge des Ausgangssignals aus den $\tilde{L}$ konkreten Signalausschnitten zufällig einen Satz von L Signalausschnitten auswählt. Wenn man anhand dieser konkreten Stichprobe die gesuchten stochastischen Eigenschaften empirisch bestimmt, so wird man erwarten können, dass diese die stochastischen Eigenschaften jedes beliebigen der $\tilde{L}$ Zufallsvektoren und somit auch die stochastischen Eigenschaften des Ausgangsprozesses des TDMA-Systems innerhalb des Zeitschlitzes des ersten logischen Kanals adäquat abschätzen.

Abschließend wird noch untersucht, was sich bei den eben durchgeführten Betrachtungen ändert, wenn man die Messung an einem anderen als dem logischen Kanal durchführt, der dem Zeitschlitz entspricht, der dem Synchronisationszeitschlitz folgt. In allen anderen logischen Kanälen ist die deterministische Störung durch das Synchronisationssignal aufgrund der begrenzt angenommenen Dauer der Systemimpulsantwort in dem analog gebildeten Modellzufallsvektor nicht mehr vorhanden. Stattdessen wird der mit der Periode P zyklstationäre Prozess, der durch die lineare Verzerrung des Übertragungskanals aus dem Prozess entsteht, aus dem die Sendesignale des der Messung vorangehenden Kanals erzeugt werden, als Störer mit zeitabhängigen stochastischen Eigenschaften in unserem Modellzufallsvektor auftreten. Diese Art der Störung ist auch diejenige, die man beim normalen Betrieb des TDMA-Systems in einem Kanal, der nicht dem Synchronisationszeitschlitz folgt, erwartet. Je nachdem ob die Zufallsvektoren, aus denen die Datensignale entstammen, die in dem vor dem zu messenden logischen Kanal entsprechenden Zeitschlitz übertragen werden, unabhängig sind oder nicht, sind deren Anteile an den Spaltenvektoren der zufälligen Stichprobenmatrix bei einer festen Zuordnung ebenfalls unabhängig oder nicht. Wenn diese unabhängig sind, wird sich deren Unabhängigkeit nicht ändern, wenn man die feste Zuordnung durch eine davon unabhängige zufällige Zuordnung ersetzt. Sind die Spaltenvektoranteile jedoch abhängig, so kann man auch hier erwarten, dass die Qualität³ der Abhängigkeit der Spaltenvektoren der zufälligen Stichprobenmatrix sich deutlich verringert, wenn man eine zufällige Zuordnung der Zeitschlitze zu den Spaltenvektoren der zufälligen Stichprobenmatrix einführt. Eine Messung der stochastischen Eigenschaften der Zufallsprozessausschnitte der Zeitschlitze anhand einer konkreten Stichprobenmatrix wird dann auch bei diesen logischen Kanälen erst durch die zufällige Wahl der Zuordnung ermöglicht, selbst wenn kein netzfrequenter Störer vorhanden ist.

Will man stochastische Eigenschaften bestimmen, in die sowohl Zufallsgrößen am Eingang wie auch am Ausgang unseres TDMA-Beispielsystems eingehen — um daraus z. B. die Übertragungsfunktion des physikalischen Übertragungskanals zu berechnen —, so muss man weiterhin fordern, dass auch alle $\tilde{L}$ Zufallsvektoren, die sich aus den $\tilde{L}$ Zufallsvektoren des Ein- und Ausgangsprozesses zusammensetzen, eine Verbundverteilung besitzen, die bei allen $\tilde{L}$ zusammengesetzten Zufallsvektoren identisch ist. Man muss dann beim Eingangssignal dieselbe Zuordnung verwenden, wie beim Ausgangssignal, auch wenn — wie bei unserem Beispielsystem — eine zufällige Zuordnung beim Signal am Systemeingang nicht notwendig ist, wenn man nur die stochastischen Eigenschaften des Eingangssignals bestimmt. Wenn man dann die ggf. zufällige Zuordnung der Spaltenvektoren der zufälligen Stichprobenmatrizen am Systemein- und ausgang so vornimmt, dass auch noch

³Abweichung der Approximation des Produkts der L identischen F -dimensionalen Randverbundverteilungen der Spaltenvektoren durch die $F \cdot L$ -dimensionale Verbundverteilung der Matrix

sichergestellt ist, dass die Spaltenvektoren der einen Stichprobenmatrix von den Spaltenvektoren der anderen unabhängig sind, wenn sie einen unterschiedlichen Spaltenindex aufweisen, bilden die Paare der Spaltenvektoren mit gleichen Spaltenindex der beiden Stichprobenmatrizen des Ein- und Ausgangs die Elemente einer mathematischen Stichprobe des Zufallsvektors, der sich aus den beiden Zufallsvektoren des Ein- und Ausgangs zusammensetzt.

Damit möchte ich das TDMA-Beispiel zur Interpretation der Zufallsprozesse im Systemmodell nach Bild 1.1 abschließen. Es sollte damit vor allem gezeigt werden, dass man immer darauf achten sollte, ob und wie eine geeignete Stichprobenentnahme an einem realen System möglich ist. Im weiteren werden unter den im Systemmodell auftretenden Zufallsvektoren nicht mehr die Ausschnitte der zeitverschobenen Zeitintervalle des tatsächlich am System anliegenden Zufallsprozesses gemeint sein, sondern immer die entsprechenden zeitverschiebungsunabhängigen Modellzufallsvektoren, die wegen der geforderten Wahl der Zeitintervalle dieselbe Verbundverteilung aufweisen, wie alle entsprechenden Zufallsprozessausschnitte. Es wird dabei angenommen, dass die Entnahme der Signalausschnitte — wenn notwendig — zufällig vorgenommen wurde, so dass die oben beschriebene Betrachtungsweise der L Signalausschnitte als eine konkrete Stichprobe einer mathematischen Stichprobe vom Umfang L — also als ein Ensemble — des Modellzufallsvektors ebenso zutrifft, wie die zuletzt beschriebene Interpretation des Paares der beiden Stichprobenmatrizen des Ein- und des Ausgangs des Systems. Die diskrete Zeitvariable k wird im weiteren sowohl für die tatsächliche physikalisch unaufhaltsam verrinnende Zeit — etwa beim Zugriff auf eine bestimmte Zufallsgröße eines tatsächlich am System anliegenden Zufallsprozesses — als auch als Laufindex der Elemente des Modellzufallsvektors verwendet, ohne dass die sich dabei ergebende ggf. zufällige Zeitverschiebung berücksichtigt wird. Dadurch werden mehrere konkrete Realisierungen des Modellzufallsvektors möglich, ohne dass jedesmal eine andere Zeitverschiebung zu berücksichtigen ist oder eine neue Variable für die Elemente des Modellzufallsvektors eingeführt werden muss. Ob der Zufallsprozess mit der physikalischen Zeit oder die Zufallsgrößen des Modellzufallsvektors mit der verschobenen Zeit gemeint sind, wird aus dem Zusammenhang ersichtlich. Die Forderung, dass die Verbundverteilung des Modellzufallsvektors gleich der Verbundverteilung aller Zufallsvektoren derselben Länge ist, die aus dem realen Zufallsprozess innerhalb der Zeitintervalle entnommen werden, die alle in einem Zeitraum liegen, in dem sich das System in einem typischen Betriebszustand befindet, ist eine der Ergodizität analoge Forderung für Prozesse, die nicht stationär oder zyklstationär sein müssen.

Die bisher durchgeführte Betrachtung bezüglich der Gewinnung einer Stichprobenmatrix lässt es zu, dass man beliebige stochastische Eigenschaften anhand einer konkreten Stichprobenmatrix bestimmen kann. Manchmal kann sich jedoch der Fall ergeben, dass die

geeignete Gewinnung einer Stichprobenmatrix nicht möglich oder praktisch nicht durchführbar ist. Wenn man jedoch nur ganz bestimmte stochastische Eigenschaften abschätzen will — beim RKM sind dies die ersten und zweiten zentralen Momente —, kann es genügen, dass man die Stichprobenmatrix in einer Art erzeugt, die bezüglich der zu messenden Merkmale als zufällig angesehen werden kann. Das bedeutet, dass lediglich die Schätzwerte der zu bestimmenden Merkmale keine systematischen oder methodischen Fehler aufweisen, die durch die Art der Stichprobenerhebung verursacht werden. Wie solch eine Stichprobe zu erheben ist, kann allgemein nicht gesagt werden, sondern ist immer anhand des konkret zu vermessenden Systems zu entscheiden. Die im weiteren hergeleitete Theorie geht immer davon aus, dass die Art der Erhebung der Stichprobe so gewählt wurde, dass sich eine mathematische Stichprobenmatrix ergibt.

2.2 Theoretische Werte der Übertragungsfunktionen und der deterministischen Störung

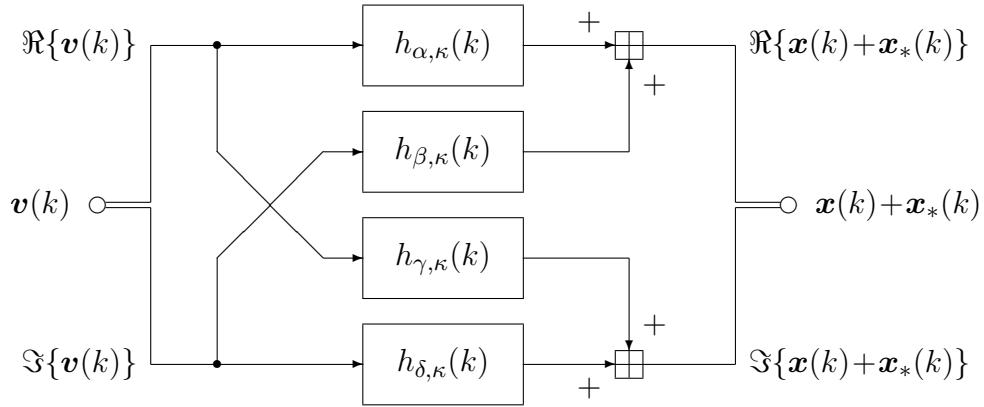
In [1] haben wir nur *ein* lineares Modellsystem angesetzt. Ein solches Systemmodell liefert jedoch nur eine unvollständige Aussage über die Korrelationen, die zwischen den Real- und Imaginärteilen des Ein- und Ausgangssignals am realen System vorhanden sein können. Das erste Teilbild in Bild 2.3 zeigt ein erweitertes Modellsystem, bei dem alle vier möglichen linearen Verknüpfungen der Real- und Imaginärteile von Ein- und Ausgang vorhanden sind. Es besteht aus den vier reellwertigen linearen Teilsystemen, die durch ihre zeitvarianten Impulsantworten $h_{\alpha,\kappa}(k)$, $h_{\beta,\kappa}(k)$, $h_{\gamma,\kappa}(k)$ und $h_{\delta,\kappa}(k)$ beschrieben werden. Wie ein Vergleich der Ein- und Ausgangssignale zeigt, lässt sich dieses i. Allg. nichtlineare, erweiterte Modellsystem in die beiden im zweiten Teilbild dargestellten linearen komplexwertigen Modellsysteme überführen, wenn man für deren zeitvariante Impulsantworten

$$h_{\kappa}(k) = \frac{h_{\alpha,\kappa}(k) + h_{\delta,\kappa}(k)}{2} + j \cdot \frac{h_{\gamma,\kappa}(k) - h_{\beta,\kappa}(k)}{2} \quad (2.1a)$$

und
$$h_{*,\kappa}(k) = \frac{h_{\alpha,\kappa}(k) - h_{\delta,\kappa}(k)}{2} + j \cdot \frac{h_{\gamma,\kappa}(k) + h_{\beta,\kappa}(k)}{2} \quad (2.1b)$$

wählt. Somit enthält das Systemmodell in Bild 1.1 nun zwei lineare, stabile, zeitdiskrete und i. allg. komplexwertige Modellsysteme $\mathcal{S}_{lin}$ und $\mathcal{S}_{*,lin}$, deren Antworten $h_{\kappa}(k)$ und $h_{*,\kappa}(k)$ auf einen Impuls $v(k) = \gamma_0(k - \kappa)$ zum Zeitpunkt κ wir zunächst als *abhängig* von dem Zeitpunkt κ ansetzen. Da reale Systeme immer kausal sind, ist es ausreichend Modellsysteme anzusetzen, deren Impulsantworten für $k < \kappa$ null sind. Desweiteren kann man bei einem realen stabilen System davon ausgehen, dass dessen Impulsantwort nach einer

a) Erweitertes, nichtlineares Modellsystem:



b) Äquivalentes Modellsystem, das zwei lineare komplexwertige Modellsysteme enthält:

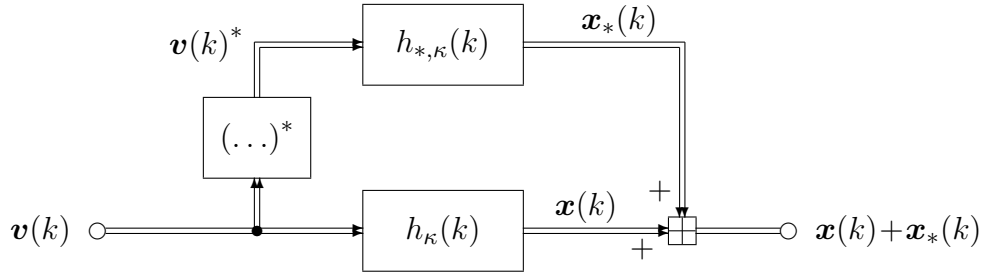


Bild 2.3: Erweitertes Modellsystem, das alle Ein- und Ausgangskovarianzen berücksichtigt

hinreichend groß gewählten Einschwingzeit E soweit abgeklungen ist, dass die Impulsantworten beider Modellsysteme in guter Näherung als zeitlich begrenzt anzusehen sind. Da die Modellsysteme von dem Zufallsprozess $\mathbf{v}(k)$, bzw. dem konjugierten Zufallsprozess $\mathbf{v}(k)^*$ erregt werden, erhalten wir an den Ausgängen der Modellsysteme die Zufallsprozesse

$$\mathbf{x}(k) = \sum_{\kappa=k-E}^k \mathbf{v}(\kappa) \cdot h_{\kappa}(k) \quad \text{und} \quad (2.2a)$$

$$\mathbf{x}_*(k) = \sum_{\kappa=k-E}^k \mathbf{v}(\kappa)^* \cdot h_{*,\kappa}(k) \quad \forall \quad k \in \mathbb{Z} \quad (2.2b)$$

Der Prozess des Fehlers der Approximation des realen Systems durch die Modellsysteme und das deterministische Störsignal wird mit $\mathbf{n}(k)$ bezeichnet. Er berechnet sich zu

$$\mathbf{n}(k) = \mathbf{y}(k) - \mathbf{x}(k) - \mathbf{x}_*(k) - u(k) \quad (2.3)$$

Bei der Approximation des realen Systems durch die Modellsysteme und das deterministische Störsignal wählt man die Impulsantworten $h_{\kappa}(k)$, $h_{*,\kappa}(k)$ und das Signal $u(k)$ so,

dass das zweite Moment⁴ des Approximationsfehlerprozesses $\mathbf{n}(k)$ innerhalb des Beobachtungszeitraums $0 \leq k < F$ möglichst klein wird. Bei realen Systemen kann man davon ausgehen, dass dieses Moment immer existiert, da der Zufallsprozess $\mathbf{n}(k)$ in der Amplitude begrenzt ist, so dass dessen Verteilungsfunktion unter- bzw. oberhalb bestimmter Werte immer 0 bzw. 1 ist. Das bei der Berechnung des zweiten Moments über $n(k)^2$ gebildete uneigentliche Integral⁵ existiert daher immer. Für jeden Zeitpunkt k erhalten wir eine Minimierungsaufgabe für das zweite Moment des Prozesses $\mathbf{n}(k)$:

$$\begin{aligned} \mathbb{E}\{|\mathbf{n}(k)|^2\} = \mathbb{E}\left\{\left|\mathbf{y}(k) - \sum_{\kappa=k-E}^k \mathbf{v}(\kappa) \cdot h_{\kappa}(k) - \sum_{\kappa=k-E}^k \mathbf{v}(\kappa)^* \cdot h_{*,\kappa}(k) - u(k)\right|^2\right\} \stackrel{!}{=} \text{minimal} \\ \forall \quad 0 \leq k < F \quad \wedge \quad k \in \mathbb{Z}. \end{aligned} \quad (2.4)$$

Um die Lösung dieser Minimierungsaufgabe für einen Zeitpunkt k zu erhalten, leitet man den zu minimierenden Term nach den zu bestimmenden Größen — den Werten der beiden Impulsantworten und des deterministischen Störsignals — partiell ab⁶. Indem man nun alle partiellen Ableitungen zugleich zu null setzt, erhält man ein Gleichungssystem, mit dem man die gesuchten Werte der beiden Impulsantworten und des deterministischen Störsignals bestimmen kann. Beim partiellen Ableiten nach $u(k)$ erhält man die Gleichung

$$u(k) = \mathbb{E}\{\mathbf{y}(k)\} - \sum_{\kappa=k-E}^k \mathbb{E}\{\mathbf{v}(\kappa)\} \cdot h_{\kappa}(k) - \sum_{\kappa=k-E}^k \mathbb{E}\{\mathbf{v}(\kappa)\}^* \cdot h_{*,\kappa}(k). \quad (2.5)$$

Diese Gleichung besagt, dass $u(k)$ gerade so zu wählen ist, dass der verbleibende Approximationsfehlerprozess $\mathbf{n}(k)$ mittelwertfrei ist:

$$\mathbb{E}\{\mathbf{n}(k)\} = 0 \quad \forall \quad k \in \mathbb{Z}. \quad (2.6)$$

Dies zeigt man, indem man nach und nach die Gleichungen (2.3), (2.2) und (2.5) einsetzt, und dann den Erwartungswert der Summe als die Summe der Erwartungswerte der einzelnen Summanden berechnet. Die optimalen Werte $u(k)$ nach Gleichung (2.5) kann man in alle anderen Gleichungen einsetzen. Man erhält dadurch für jeden Zeitpunkt k ein

⁴Also nicht die Varianz, die das zweite *zentrale* Moment ist

⁵Gegebenenfalls ist dieses Integration im Stieltjesschen Sinne durchzuführen.

⁶Der zu minimierende Term ist reell. Diesen Term leitet man zunächst nach allen Realteilen, und dann nach allen Imaginärteilen beider Impulsantworten bzw. des deterministischen Störsignals partiell ab. Man berechnet also ganz konventionell die partiellen Ableitungen einer reellen Funktion mehrerer reeller Variablen. Man erhält so zwei Gleichungen für jeden Wert jeder Impulsantwort und zwei Gleichungen für jeden Wert des deterministischen Störsignals. Die Gleichung, die man bei der partiellen Ableitung nach dem Imaginärteil erhält, multipliziert man anschließend mit j und fasst sie mit der Gleichung, die man bei der partiellen Ableitung nach dem Realteil erhält, zu einer komplexen Gleichung zusammen.

lineares Gleichungssystem zur Bestimmung der Werte der Impulsantworten, das sich in Matrixschreibweise folgendermaßen darstellen lässt:

$$\begin{bmatrix} \vec{h} & \vec{h}_* \end{bmatrix} \cdot \begin{bmatrix} \underline{C}_{\mathbf{v},\mathbf{v}} & \underline{C}_{\mathbf{v},\mathbf{v}^*} \\ \underline{C}_{\mathbf{v}^*,\mathbf{v}} & \underline{C}_{\mathbf{v}^*,\mathbf{v}^*} \end{bmatrix} = \begin{bmatrix} \vec{C}_{\mathbf{y},\mathbf{v}} & \vec{C}_{\mathbf{y},\mathbf{v}^*} \end{bmatrix}. \quad (2.7)$$

Dabei treten die vier $(E+1) \times (E+1)$ Kovarianzmatrizen $\underline{C}_{\mathbf{v},\mathbf{v}}$, $\underline{C}_{\mathbf{v},\mathbf{v}^*}$, $\underline{C}_{\mathbf{v}^*,\mathbf{v}}$ und $\underline{C}_{\mathbf{v}^*,\mathbf{v}^*}$, sowie die zwei $1 \times (E+1)$ Kovarianzvektoren $\vec{C}_{\mathbf{y},\mathbf{v}}$ und $\vec{C}_{\mathbf{y},\mathbf{v}^*}$ und die zwei $1 \times (E+1)$ Zeilenvektoren $\vec{h}$ und $\vec{h}_*$ der Werte der Impulsantworten auf. Wie deren Elemente definiert sind, zeigt die nachfolgende Tabelle:

	Zeile	Spalte	Wert
$\underline{C}_{\mathbf{v},\mathbf{v}}$	i	j	$E\left\{\left(\mathbf{v}(k+1-i) - E\{\mathbf{v}(k+1-i)\}\right) \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)^*\right\}$
$\underline{C}_{\mathbf{v},\mathbf{v}^*}$	i	j	$E\left\{\left(\mathbf{v}(k+1-i) - E\{\mathbf{v}(k+1-i)\}\right) \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)\right\}$
$\underline{C}_{\mathbf{v}^*,\mathbf{v}}$	i	j	$E\left\{\left(\mathbf{v}(k+1-i) - E\{\mathbf{v}(k+1-i)\}\right)^* \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)\right\}$
$\underline{C}_{\mathbf{v}^*,\mathbf{v}^*}$	i	j	$E\left\{\left(\mathbf{v}(k+1-i) - E\{\mathbf{v}(k+1-i)\}\right)^* \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)^*\right\}$
$\vec{C}_{\mathbf{y},\mathbf{v}}$	1	j	$E\left\{\left(\mathbf{y}(k) - E\{\mathbf{y}(k)\}\right) \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)^*\right\}$
$\vec{C}_{\mathbf{y},\mathbf{v}^*}$	1	j	$E\left\{\left(\mathbf{y}(k) - E\{\mathbf{y}(k)\}\right) \cdot \left(\mathbf{v}(k+1-j) - E\{\mathbf{v}(k+1-j)\}\right)\right\}$
$\vec{h}$	1	j	$h_{k+1-j}(k)$
$\vec{h}_*$	1	j	$h_{*,k+1-j}(k)$

Da das lineare Gleichungssystem (2.7) durch partielles Ableiten eines stets positiven Terms entstanden ist, und da die Minimierungsparameter in diesem Term nur linear oder quadratisch auftreten, existiert für jeden einzelnen Zeitpunkt k wenigstens eine Lösung des Gleichungssystems, selbst wenn die Matrix singular ist. Gegebenenfalls kann sich auch ein ein- oder mehrdimensionaler Lösungsraum ergeben. Nur wenn es eine Lösung gibt, die alle Gleichungssysteme für alle Zeitpunkte k zugleich löst, macht es Sinn, zeitinvariante Modellsysteme anzusetzen. Im Fall eines im weiten Sinne zyklstationären Verbundprozesses⁷ aus $\mathbf{v}(k)$ und $\mathbf{y}(k)$ ändert sich sowohl die Matrix des Gleichungssystems als auch der Vektor auf der rechten Seite nicht, wenn man den Zeitpunkt k der Minimierung um ein ganzzahliges Vielfaches der Periode der Zyklstationarität verändert. Es existiert dann immer eine periodisch zeitvariante Lösung für die Impulsantworten der Modellsysteme. Da nur zweite zentrale Momente in die Matrix des Gleichungssystems und den Vektor

⁷Ein stationärer Prozess ist hier der Sonderfall eines zyklstationären Prozesses mit der Periodizität Eins.

auf der rechten Seite eingehen, gibt es selbst dann eine periodisch zeitvariante Lösung für die Impulsantworten der Modellsysteme, wenn sich damit in Gleichung (2.5) eine nicht periodische, zeitabhängige, deterministische Störung $u(k)$ ergibt. Auch falls der Verbundprozess aus $\mathbf{v}(k)$ und $\mathbf{y}(k)$ nicht im weiten Sinn zyklstationär ist, kann es sein, dass es eine periodisch zeitvariante Lösung für die Impulsantworten der Modellsysteme gibt. Dies ist dann jedoch anhand heuristischer Überlegungen zu zeigen, wenn man das RKM zur Vermessung solcher Systeme verwenden will. Im weiteren gehen wir davon aus, dass nur solche Systeme untersucht werden, für die eine periodisch zeitvariante Lösung existiert. Nur dann ist es sinnvoll, diese mit Hilfe des RKM messtechnisch abzuschätzen. Als periodisch zeitvariant wird die Lösung für die Impulsantworten der Modellsysteme dann bezeichnet, wenn eine Verschiebung des Zeitpunktes k , für den die Minimierungsaufgabe zu lösen ist, um eine feste Zeitdifferenz K_H lediglich zu einer zeitlichen Verschiebung der Lösung um dieselbe Zeitdifferenz führt, so dass sich die Lösungen mit K_H periodisch wiederholen. Die Lösungen selbst sind jedoch i. Allg. keine periodischen Folgen. Für die mit dieser Variante des RKM untersuchbaren Systeme muss also

$$h_{\kappa+K_H}(k+K_H) = h_{\kappa}(k) \quad (2.8a)$$

$$\text{und} \quad h_{*,\kappa+K_H}(k+K_H) = h_{*,\kappa}(k) \quad \forall \quad k, \kappa \in \mathbb{Z} \quad (2.8b)$$

gelten. Da wir uns auf kausale Systeme mit zeitlich begrenzter Impulsantwort beschränkt hatten, gilt weiterhin:

$$h_{\kappa}(k) = h_{*,\kappa}(k) = 0 \quad \forall \quad k - \kappa \notin [0; E]. \quad (2.9)$$

Bei solchen Systemen genügt es — wie im Fall des zeitinvarianten Systems — die Erregung im Zeitintervall $-E \leq k < F$ an das zu vermessende System anzulegen, da die Werte der Erregung außerhalb dieses Intervalls nicht in die Minimierung des zweiten Moments des Approximationsfehlers nach Gleichung (2.4) eingehen.

Wenn man bei einem realen System, bei dem der Ausgangsprozess mit dem konjugierten Eingangsprozess in der Art korreliert ist, dass sich bei der eben dargestellten theoretischen Lösung eine von null verschiedene Impulsantwort für das Modellsystem $\mathcal{S}_{*,lin}$ ergibt, lediglich das eine Modellsystem $\mathcal{S}_{lin}$ ansetzt, so erhält man ebenfalls eine Lösung minimaler Approximationsfehlerleistung. Diese Lösung wird in aller Regel aber von der Lösung des vollständigen Systems abweichen und eine höhere Approximationsfehlerleistung aufweisen. Auch wenn man das deterministische Störsignal weglässt, obwohl beim vollständigen Systemmodell $u(k)$ eine von null verschiedene Lösung annehmen würde, wird sich zwar eine Lösung minimaler, aber dennoch größerer Approximationsfehlerleistung ergeben, die von der Lösung bei vollständigem Systemmodell normalerweise abweicht. Da das RKM bestenfalls die theoretische Lösung erwartungstreu abschätzen kann, ist es auch für die

Messung wichtig, das zum realen System passende Systemmodell zu wählen. Wenn man sich nicht sicher ist, ob $u(k)$ oder das Modellsystem $\mathcal{S}_{*,lin}$ weggelassen werden kann, sollte man sich im Zweifel für das vollständige Systemmodell entscheiden. Die Messergebnisse werden dann zeigen, ob eine einfachere Modellierung auch möglich wäre.

Wie beim zeitinvarianten Modellsystem in [1] wird auch hier wieder eine gemäß Gleichung ([1]:2.8) mit M periodische Erregung verwendet. Die Periode M wird dabei als ganzzahliges Vielfaches von K_H gewählt, so dass $M/K_H \in \mathbb{N}$ gilt. Wenn die Erregung im gesamten Zeitintervall $-E \leq k < F$ anliegt, kann man für die Ausgangssignale der periodischen zeitvarianten Modellsysteme

$$\begin{aligned}
 \mathbf{x}(k) &= \sum_{\kappa=-\infty}^{\infty} \mathbf{v}(\kappa) \cdot h_{\kappa}(k) = \sum_{\kappa=0}^{M-1} \sum_{\tilde{\kappa}=-\infty}^{\infty} \mathbf{v}(\kappa + \tilde{\kappa} \cdot M) \cdot h_{\kappa + \tilde{\kappa} \cdot M}(k) = \\
 &= \sum_{\kappa=0}^{M-1} \mathbf{v}(\kappa) \cdot \sum_{\tilde{\kappa}=-\infty}^{\infty} h_{\kappa}(k - \tilde{\kappa} \cdot M) = \sum_{\kappa=0}^{M-1} \mathbf{v}(\kappa) \cdot \tilde{h}_{\kappa}(k) = \\
 &= \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} H(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot \sum_{\kappa=0}^{M-1} \mathbf{v}(\kappa) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot (k-\kappa)} \cdot e^{-j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa} = \\
 &= \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} H(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot \mathbf{V}(\mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}
 \end{aligned} \tag{2.10a}$$

und analog

$$\begin{aligned}
 \mathbf{x}_*(k) &= \sum_{\kappa=-\infty}^{\infty} \mathbf{v}(\kappa)^* \cdot h_{*,\kappa}(k) = \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} H_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot \mathbf{V}(-\mu - \hat{\mu} \cdot \frac{M}{K_H})^* \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \\
 &\quad \forall \quad 0 \leq k < F
 \end{aligned} \tag{2.10b}$$

schreiben. Dabei treten die aus den Impulsantworten durch periodische Fortsetzung gewonnenen periodischen Folgen

$$\begin{aligned}
 \tilde{h}_{\kappa}(k) &= \sum_{\tilde{\kappa}=-\infty}^{\infty} h_{\kappa}(k - \tilde{\kappa} \cdot M) = \\
 &= \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} H(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot (k-\kappa)} \cdot e^{-j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa}
 \end{aligned} \tag{2.11a}$$

und

$$\begin{aligned}
 \tilde{h}_{*,\kappa}(k) &= \sum_{\tilde{\kappa}=-\infty}^{\infty} h_{*,\kappa}(k - \tilde{\kappa} \cdot M) = \\
 &= \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} H_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot (k-\kappa)} \cdot e^{-j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa}
 \end{aligned} \tag{2.11b}$$

sowie die daraus durch diskrete Fouriertransformation bezüglich k und inverse diskrete Fouriertransformation bezüglich κ gewonnenen diskreten Werte der bifrequenten Übertragungsfunktionen

$$H(\mu, \mu_1) = \frac{1}{M} \cdot \sum_{\kappa=0}^{M-1} \sum_{k=0}^{M-1} \tilde{h}_{\kappa}(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu_1 \cdot \kappa} \quad (2.12a)$$

und

$$H_*(\mu, \mu_1) = \frac{1}{M} \cdot \sum_{\kappa=0}^{M-1} \sum_{k=0}^{M-1} \tilde{h}_{*,\kappa}(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu_1 \cdot \kappa} \quad (2.12b)$$

auf. Setzen wir in diese Gleichungen $\mu_1 = \mu + \tilde{\mu} + \hat{\mu} \cdot M/K_H$ ein, so erhalten wir, wenn wir die Periodizität gemäß der Gleichungen (2.8) berücksichtigen, für die Werte der diskreten bifrequenten Übertragungsfunktionen die Aussagen

$$H\left(\mu, \mu + \tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) = \begin{cases} \frac{1}{K_H} \cdot \sum_{k=0}^{M-1} \sum_{\kappa=0}^{K_H-1} \tilde{h}_{\kappa}(k + \kappa) \cdot e^{j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa} \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} & \text{für } \tilde{\mu} = 0 \\ 0 & \text{für } 0 < \tilde{\mu} < \frac{M}{K_H} \end{cases} \quad \forall \quad \mu, \hat{\mu}, \tilde{\mu} \in \mathbb{Z}, \quad (2.13a)$$

und

$$H_*\left(\mu, \mu + \tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) = \begin{cases} \frac{1}{K_H} \cdot \sum_{k=0}^{M-1} \sum_{\kappa=0}^{K_H-1} \tilde{h}_{*,\kappa}(k + \kappa) \cdot e^{j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa} \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} & \text{für } \tilde{\mu} = 0 \\ 0 & \text{für } 0 < \tilde{\mu} < \frac{M}{K_H} \end{cases} \quad \forall \quad \mu, \hat{\mu}, \tilde{\mu} \in \mathbb{Z}, \quad (2.13b)$$

die besagen, dass nur die Werte von null verschieden sein können, deren Argumente μ und μ_1 eine Differenz aufweisen, die ein ganzzahliges Vielfaches von M/K_H ist. Für einen festen Wert von μ werden die Werte der bifrequenten Übertragungsfunktionen mit $\mu_1 = \mu + \hat{\mu} \cdot M/K_H$, die von null verschieden sein können, zu einem Zeilenvektor zusammengefasst:

$$\vec{H}(\mu) = \begin{bmatrix} H(\mu, \mu) \\ H\left(\mu, \mu + \frac{M}{K_H}\right) \\ \vdots \\ H\left(\mu, \mu + (K_H - 1) \cdot \frac{M}{K_H}\right) \\ H_*(\mu, \mu) \\ H_*\left(\mu, \mu + \frac{M}{K_H}\right) \\ \vdots \\ H_*\left(\mu, \mu + (K_H - 1) \cdot \frac{M}{K_H}\right) \end{bmatrix}^T \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (2.14)$$

Die Definition der bifrequenten Übertragungsfunktionen wurde so gewählt, dass für den Fall eines zeitinvarianten Systems mit $K_H=1$ die Hauptdiagonalelemente der bifrequenten Übertragungsfunktionen gerade die Werte $H(\mu \cdot 2\pi/M)$ bzw. $H_*(\mu \cdot 2\pi/M)$ der Übertragungsfunktionen ohne einen weiteren von eins verschiedenen Faktor sind. Es sei wieder darauf hingewiesen, dass ähnlich wie beim Fall eines zeitinvarianten Modellsystems aufgrund des Abtasttheorems durch die für $H(\mu, \mu + \hat{\mu} \cdot M/K_H)$ und $H_*(\mu, \mu + \hat{\mu} \cdot M/K_H)$ approximierten Größen nur die Werte der periodisch fortgesetzten Impulsantworten $\tilde{h}_\kappa(k)$ und $\tilde{h}_{*,\kappa}(k)$ festgelegt sind. Nur wenn heuristische Überlegungen vermuten lassen, dass alle $2 \cdot K_H$ Impulsantworten der beiden realen Systeme auf ein Intervall der Länge M beschränkt sind, können auch die Impulsantworten $h_\kappa(k)$ und $h_{*,\kappa}(k)$ eindeutig aus den $2 \cdot M \cdot K_H$ Optimalwerten der Übertragungsfunktionen bestimmt werden.

In den Gleichungen (2.10) kommt auch noch das diskrete Spektrum der periodischen Erregung, das nach Gleichung ([1]:2.7) definiert ist, vor. Auch einige Zufallsgrößen dieses Spektrums lassen sich zu einen Spaltenvektor

$$\tilde{\mathbf{V}}(\mu) = \begin{bmatrix} \mathbf{V}(\mu) \\ \mathbf{V}(\mu + \frac{M}{K_H}) \\ \vdots \\ \mathbf{V}(\mu + (K_H - 1) \cdot \frac{M}{K_H}) \\ \mathbf{V}(-\mu)^* \\ \mathbf{V}(-\mu - \frac{M}{K_H})^* \\ \vdots \\ \mathbf{V}(-\mu - (K_H - 1) \cdot \frac{M}{K_H})^* \end{bmatrix} \quad \forall \quad \mu = 0 \ (1) \ M-1. \quad (2.15)$$

zusammenfassen. Mit den Vektoren $\tilde{\mathbf{V}}(\mu)$ und $\vec{H}(\mu)$ erhalten wir für den Summenprozess am Ausgang der beiden periodisch zeitvarianten Modellsysteme die kompaktere Schreibweise

$$\mathbf{x}(k) + \mathbf{x}_*(k) = \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad 0 \leq k < F. \quad (2.16)$$

Setzt man $\mathbf{x}(k) + \mathbf{x}_*(k)$ in den Approximationsfehler nach Gleichung (2.3), so erhält man zunächst

$$\mathbf{n}(k) = \mathbf{y}(k) - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} - u(k), \quad (2.17)$$

und damit im Zeitbereich der Messung die modifizierte Minimierungsaufgabe

$$\begin{aligned} E\{|\mathbf{n}(k)|^2\} &= E\left\{\left|\mathbf{y}(k) - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} - u(k)\right|^2\right\} \stackrel{!}{=} \text{minimal} \\ &\forall \quad k = 0 \ (1) \ F-1. \end{aligned} \quad (2.18)$$

Um die Optimallösung für $u(k)$ zu bestimmen, leitet man diese Gleichung für alle Werte von k in der oben beschriebenen Art nach $u(k)$ ab. Die optimale Lösung für die Regressionskoeffizienten $u(k)$ ergibt sich auch hier, wenn der Erwartungswert von $\mathbf{n}(k)$ für alle Zeitpunkte k null wird:

$$u(k) = E\{\mathbf{y}(k)\} - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot E\{\tilde{\mathbf{V}}(\mu)\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad k = 0 \text{ (1) } F-1. \quad (2.19)$$

Dieselbe Lösung erhält man, wenn man in die allgemeinere Lösung (2.5) den im Zeitintervall $-E \leq k < F$ gemäß ([1]:2.8) periodischen Eingangsprozess einsetzt, und die periodische Zeitvarianz des Systems berücksichtigt, indem man nach und nach die Gleichungen (2.8) bis (2.15) verwendet.

Um die Lösungen für die beiden bifrequenten Übertragungsfunktionen zu erhalten, setzen wir zunächst die optimale Lösung für $u(k)$ in den zu minimierenden, reellen Term (2.18) ein. Dann leiten wir diesen Term partiell nach Real- und Imaginärteil der $M \cdot K_H$ Werte der bifrequenten Übertragungsfunktion $H(\tilde{\mu}, \tilde{\mu} + \hat{\mu} \cdot M/K_H)$ ab, uns setzen die Ableitungen mit null gleich. Wir erhalten so insgesamt $2 \cdot M \cdot K_H$ Gleichungen für jeden der F Zeitpunkte k . Jeweils zwei dieser Gleichungen fassen wir danach zu einer komplexen Gleichung zusammen. Wir erhalten so die $F \cdot M \cdot K_H$ komplexen Gleichungen:

$$\begin{aligned} & \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} \left(H(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot E\left\{ \left(\mathbf{V}(\mu + \hat{\mu} \cdot \frac{M}{K_H}) - E\{\mathbf{V}(\mu + \hat{\mu} \cdot \frac{M}{K_H})\} \right) \cdot \right. \right. \\ & \quad \left. \left. \cdot \left(\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) - E\{\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H})\} \right)^* \right\} + \right. \\ & \quad \left. H_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}) \cdot E\left\{ \left(\mathbf{V}(-\mu - \hat{\mu} \cdot \frac{M}{K_H}) - E\{\mathbf{V}(-\mu - \hat{\mu} \cdot \frac{M}{K_H})\} \right) \cdot \right. \right. \\ & \quad \left. \left. \cdot \left(\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) - E\{\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H})\} \right)^* \right\} \right\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\ & = E\left\{ \left(\mathbf{y}(k) - E\{\mathbf{y}(k)\} \right) \cdot \left(\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) - E\{\mathbf{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H})\} \right)^* \right\} \\ & \quad \forall \quad k = 0 \text{ (1) } F-1, \quad \tilde{\mu} = 0 \text{ (1) } M-1 \quad \text{und} \quad \hat{\mu} = 0 \text{ (1) } K_H-1. \quad (2.20a) \end{aligned}$$

Für die partiellen Ableitungen nach Real- und Imaginärteil der $M \cdot K_H$ Werte der anderen bifrequenten Übertragungsfunktion $H_*(\tilde{\mu}, \tilde{\mu} + \hat{\mu} \cdot M/K_H)$ erhalten wir analog die $F \cdot M \cdot K_H$

komplexen Gleichungen:

$$\begin{aligned}
& \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} \left(H\left(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}\right) \cdot \mathbb{E} \left\{ \left(\mathbf{V}\left(\mu + \hat{\mu} \cdot \frac{M}{K_H}\right) - \mathbb{E} \left\{ \mathbf{V}\left(\mu + \hat{\mu} \cdot \frac{M}{K_H}\right) \right\} \right) \cdot \right. \right. \\
& \quad \cdot \left. \left(\mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) - \mathbb{E} \left\{ \mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) \right\} \right) \right\} + \\
& \quad H_*\left(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}\right) \cdot \mathbb{E} \left\{ \left(\mathbf{V}\left(-\mu - \hat{\mu} \cdot \frac{M}{K_H}\right) - \mathbb{E} \left\{ \mathbf{V}\left(-\mu - \hat{\mu} \cdot \frac{M}{K_H}\right) \right\} \right)^* \cdot \right. \\
& \quad \cdot \left. \left(\mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) - \mathbb{E} \left\{ \mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) \right\} \right) \right\} \right) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
& = \mathbb{E} \left\{ \left(\mathbf{y}(k) - \mathbb{E} \left\{ \mathbf{y}(k) \right\} \right) \cdot \left(\mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) - \mathbb{E} \left\{ \mathbf{V}\left(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}\right) \right\} \right) \right\} \\
& \quad \forall \quad k = 0 \text{ (1) } F-1, \quad \tilde{\mu} = 0 \text{ (1) } M-1 \quad \text{und} \quad \hat{\mu} = 0 \text{ (1) } K_H-1. \quad (2.20b)
\end{aligned}$$

Bei den Gleichungen aller partiellen Ableitungen stellen wir fest, dass eine Erhöhung der diskreten Frequenzvariable $\tilde{\mu}$ um ein ganzzahliges Vielfaches n von M/K_H zu derselben Gleichung führt, wie eine Erhöhung der diskreten Variable $\hat{\mu}$ um dieselbe ganze Zahl n . Dies liegt daran, dass in dem gesamten Gleichungssystem nur die Summe $\tilde{\mu} + \hat{\mu} \cdot M/K_H$ vorkommt, und niemals eine der Variablen allein. Wir können daher alle mehrfach aufgeführten Gleichungen eliminieren, indem wir nur die Gleichungen mit $\hat{\mu} = 0$ verwenden. Mit den Vektoren $\vec{H}(\mu)$ nach Gleichung (2.14) und $\vec{V}(\mu)$ nach Gleichung (2.15) ergeben sich die $M \cdot F$ Gleichungssysteme

$$\begin{aligned}
& \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot \mathbb{E} \left\{ \left(\vec{V}(\mu) - \mathbb{E} \left\{ \vec{V}(\mu) \right\} \right) \cdot \left(\vec{V}(\tilde{\mu}) - \mathbb{E} \left\{ \vec{V}(\tilde{\mu}) \right\} \right)^H \right\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
& = \mathbb{E} \left\{ \left(\mathbf{y}(k) - \mathbb{E} \left\{ \mathbf{y}(k) \right\} \right) \cdot \left(\vec{V}(\tilde{\mu}) - \mathbb{E} \left\{ \vec{V}(\tilde{\mu}) \right\} \right)^H \right\} \\
& \quad \forall \quad k = 0 \text{ (1) } F-1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } M-1, \quad (2.21)
\end{aligned}$$

die jeweils auf beiden Seiten des Gleichheitszeichens einen Zeilenvektor der Dimension $1 \times (2 \cdot K_H)$ aufweisen.

Die Lösung für diese Gleichungssysteme ist bei der Wahl einer Fensterfolge, deren Spektrum der Bedingung ([1]:2.27) genügt, identisch mit der Lösung der folgenden M Gleichungssysteme, bei denen auf beiden Seiten ebenfalls ein Zeilenvektor der Dimension $1 \times (2 \cdot K_H)$ steht.

$$\begin{aligned}
& \vec{H}(\mu) \cdot \mathbb{E} \left\{ \left(\vec{V}(\mu) - \mathbb{E} \left\{ \vec{V}(\mu) \right\} \right) \cdot \left(\vec{V}(\mu) - \mathbb{E} \left\{ \vec{V}(\mu) \right\} \right)^H \right\} = \mathbb{E} \left\{ \mathbf{Y}_f(\mu) \cdot \left(\vec{V}(\mu) - \mathbb{E} \left\{ \vec{V}(\mu) \right\} \right)^H \right\} \\
& \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (2.22)
\end{aligned}$$

Im Anhang A.1 wird die Identität der Lösungsräume der Gleichungssysteme (2.21) und (2.22) hergeleitet, wobei dort für den Parameter der Anzahl der Elemente des Zufallsvektors $\tilde{\mathbf{V}}(\mu)$ der Wert $R = 2 \cdot K_H$ einzusetzen ist. Das Gleichungssystem (2.22) hat nun mit $2 \cdot M \cdot K_H$ Gleichungen genauso viele Gleichungen wie Unbekannte. Die darin auftretenden zufälligen Spektralwerte $\mathbf{Y}_f(\mu)$ berechnen sich wieder durch Fensterung und anschließende diskrete Fouriertransformation gemäß Gleichung ([1]:2.25) aus dem Zufallsvektor des Prozesses am Ausgang des realen Systems. Durch die Verwendung einer Fensterfolge, deren Spektrum die in Gleichung ([1]:2.27) angegebene Nullstelleneigenschaft erfüllt, ist es auch hier gelungen ein Gleichungssystem zu erhalten, das in M unabhängige Gleichungssysteme (für jeden Wert von μ) zu je $2 \cdot K_H$ Gleichungen mit je $2 \cdot K_H$ Unbekannten zerfällt. Diese Unbekannten sind diejenigen $2 \cdot K_H$ Werte der beiden bifrequenten Übertragungsfunktion für einen festen Wert μ , die gemäß der Gleichungen (2.13) von null verschieden sein können.

Unter Verwendung der $(2 \cdot K_H) \times (2 \cdot K_H)$ Autokovarianzmatrix

$$\underline{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)} = E \left\{ (\tilde{\mathbf{V}}(\mu) - E\{\tilde{\mathbf{V}}(\mu)\}) \cdot (\tilde{\mathbf{V}}(\mu) - E\{\tilde{\mathbf{V}}(\mu)\})^H \right\} \quad (2.23)$$

des Zufallsvektors $\tilde{\mathbf{V}}(\mu)$ und des $1 \times (2 \cdot K_H)$ Kreuzkovarianzvektors

$$\underline{C}_{\mathbf{Y}_f(\mu), \tilde{\mathbf{V}}(\mu)} = E \left\{ (\mathbf{Y}_f(\mu) - E\{\mathbf{Y}_f(\mu)\}) \cdot (\tilde{\mathbf{V}}(\mu) - E\{\tilde{\mathbf{V}}(\mu)\})^H \right\} \quad (2.24)$$

eines Spektralwertes des Signals am Systemausgang mit einigen der Spektralwerte des Signals am Systemeingang lässt sich das Gleichungssystem (2.22) sehr kompakt als

$$\vec{H}(\mu) \cdot \underline{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)} = \underline{C}_{\mathbf{Y}_f(\mu), \tilde{\mathbf{V}}(\mu)} \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (2.25)$$

schreiben. Im weiteren wollen wir uns auf den Fall beschränken, dass der periodische Prozess am Eingang des realen Systems in der Art gewählt wurde, dass die Autokovarianzmatrix $\underline{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}$ regulär ist. Dann lässt sich diese invertieren, und wir erhalten die Lösung für die $2 \cdot K_H$ Werte der beiden bifrequenten Übertragungsfunktionen für einen festen Wert μ , indem wir beide Seiten des Gleichungssystems (2.25) mit der inversen Matrix von links multiplizieren.

Auch hier wollen wir uns nun wieder der Frage widmen, welche Optimallösungen man erhält, wenn man nicht das zweite Moment des Approximationsfehlerprozesses $\mathbf{n}(k)$ minimiert, sondern stattdessen die M zweiten Momente der Spektralwerte des gefensterten Approximationsfehlerprozesses. Dazu formen wir die nach Gleichung ([1]:2.16) definierten

Spektralwerte des gefensternten Approximationsfehlerprozesses zunächst um:

$$\begin{aligned}
\mathbf{N}_f(\mu) &= \sum_{k=-\infty}^{\infty} \mathbf{n}(k) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
&= \sum_{k=-\infty}^{\infty} (\mathbf{y}(k) - \mathbf{x}(k) - \mathbf{x}_*(k) - u(k)) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
&= \mathbf{Y}_f(\mu) - \frac{1}{M} \cdot \sum_{\check{\mu}=0}^{M-1} \vec{H}(\check{\mu}) \cdot \tilde{\mathbf{V}}(\check{\mu}) \cdot \sum_{k=-\infty}^{\infty} f(k) \cdot e^{j \cdot \frac{2\pi}{M} \cdot (\check{\mu} - \mu) \cdot k} - U_f(\mu) = \\
&= \mathbf{Y}_f(\mu) - \frac{1}{M} \cdot \sum_{\check{\mu}=0}^{M-1} \vec{H}(\check{\mu}) \cdot \tilde{\mathbf{V}}(\check{\mu}) \cdot F((\mu - \check{\mu}) \cdot \frac{2\pi}{M}) - U_f(\mu) = \mathbf{Y}_f(\mu) - \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) - U_f(\mu) \\
&\quad \forall \quad \mu = 0 \text{ (1) } M-1.
\end{aligned} \tag{2.26}$$

Hier wurde nach und nach der Approximationsfehlerprozess nach Gleichung (2.3), der Ausgangsprozess der beiden Modellsysteme nach Gleichung (2.16), das Spektrum des gefensternten Ausgangsprozesses des realen Systems nach Gleichung ([1]:2.25) und das analog definierte Spektrum der gefensternten deterministischen Störung gemäß

$$U_f(\mu) = \sum_{k=-\infty}^{\infty} u(k) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad \mu = 0 \text{ (1) } M-1 \tag{2.27}$$

eingesetzt. Zuletzt wurde berücksichtigt, dass bei Verwendung einer Fensterfolge, deren Spektrum die in Gleichung ([1]:2.27) angegebene Nullstelleneigenschaft erfüllt, in der Summe lediglich ein Summand verbleibt. Die neue Aufgabe der Minimierung der M zweiten Momente der Spektralwerte des gefensternten Approximationsfehlerprozesses lautet nun:

$$\begin{aligned}
\mathbb{E}\{|\mathbf{N}_f(\mu)|^2\} &= \mathbb{E}\left\{\left|\mathbf{Y}_f(\mu) - \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) - U_f(\mu)\right|^2\right\} \stackrel{!}{=} \text{minimal} \\
&\quad \forall \quad \mu = 0 \text{ (1) } M-1.
\end{aligned} \tag{2.28}$$

Im Gegensatz zur ursprünglichen Minimierungsaufgabe (2.18) treten in (2.28) nicht mehr die F Abtastwerte $u(k)$ der deterministischen Störung, sondern nur mehr die M Spektralwerte $U_f(\mu)$ als Parameter der Optimierung auf. Um die optimalen Approximationsparameter $U_f(\mu)$ zu berechnen, leitet man den Term in der Gleichung (2.28) — wieder nach Real- und Imaginärteil von $U_f(\mu)$ getrennt — partiell ab. Die dabei entstehenden Gleichungen lassen sich wieder zu komplexen Gleichungen zusammenfassen. Man erhält so die Optimallösungen

$$U_f(\mu) = \mathbb{E}\{\mathbf{Y}_f(\mu)\} - \vec{H}(\mu) \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\mu)\} \quad \forall \quad \mu = 0 \text{ (1) } M-1. \tag{2.29}$$

Dieselben M Spektralwerte $U_f(\mu)$ erhält man, wenn man die F optimalen Werte $u(k)$ der ursprünglichen Minimierungsaufgabe (2.18) mit einer Fensterfolge fenstert, deren Spektrum der Gleichung ([1]:2.27) genügt, und anschließend fouriertransformiert. Wie man durch Einsetzen der Gleichungen (2.26) und (2.29) verifizieren kann, ergeben sich die optimalen Werte $U_f(\mu)$ in der Art, dass alle M Spektralwerte des gefensterten Approximationsfehlerprozesses mittelwertfrei sind:

$$\begin{aligned}
E\{\mathbf{N}_f(\mu)\} &= E\{\mathbf{Y}_f(\mu) - \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) - U_f(\mu)\} = \\
&= E\left\{\mathbf{Y}_f(\mu) - \vec{H}(\mu) \cdot \tilde{\mathbf{V}}(\mu) - E\{\mathbf{Y}_f(\mu)\} + \vec{H}(\mu) \cdot E\{\tilde{\mathbf{V}}(\mu)\}\right\} = \\
&= E\{\mathbf{Y}_f(\mu)\} - \vec{H}(\mu) \cdot E\{\tilde{\mathbf{V}}(\mu)\} - E\{\mathbf{Y}_f(\mu)\} + \vec{H}(\mu) \cdot E\{\tilde{\mathbf{V}}(\mu)\} = 0 \\
&\quad \forall \quad \mu = 0 \text{ (1) } M-1.
\end{aligned} \tag{2.30}$$

Die Optimallösungen $U_f(\mu)$ setzen wir wieder in die Minimierungsaufgabe (2.28) ein und erhalten so für jeden festen Wert μ eine Minimierungsaufgabe zur Bestimmung der $2 \cdot K_H$ Werte der beiden bifrequenten Übertragungsfunktionen für diesen Wert μ :

$$\begin{aligned}
E\left\{\left|\mathbf{Y}_f(\mu) - E\{\mathbf{Y}_f(\mu)\} - \vec{H}(\mu) \cdot (\tilde{\mathbf{V}}(\mu) - E\{\tilde{\mathbf{V}}(\mu)\})\right|^2\right\} &\stackrel{!}{=} \text{minimal} \\
&\quad \forall \quad \mu = 0 \text{ (1) } M-1.
\end{aligned} \tag{2.31}$$

Indem wir wieder die partiellen Ableitungen berechnen und diese zu null setzten, erhalten wir exakt dieselben M Gleichungssysteme (2.22), und somit dieselben Lösungen für die optimalen Werte der beiden bifrequenten Übertragungsfunktionen, wie bei der ursprünglichen Minimierungsaufgabe (2.18). Dabei spielt es keine Rolle, ob der Approximationsfehlerprozess stationär, zyklstationär oder gar instationär ist, und ob die deterministische Störung konstant, periodisch oder beliebig zeitvariant ist. Wenn man eine andere Fensterfolge verwenden würde, deren Spektrum die in Gleichung ([1]:2.27) angegebene Nullstelleneigenschaft *nicht* erfüllt, würden sich in aller Regel abweichende Gleichungssysteme zur Bestimmung der Werte der beiden bifrequenten Übertragungsfunktionen ergeben, und somit abweichende Optimallösungen, die dann nicht die ursprüngliche Minimierungsaufgabe (2.18) lösen würden. Somit würde auch ein Messverfahren wie das RKM versagen, bei dem die empirischen Werte der beiden bifrequenten Übertragungsfunktionen in der Art bestimmt werden, dass sie die M empirischen zweiten Momente der Spektralwerte des gefensterten Approximationsfehlerprozesses minimieren.

Durch Einsetzen der Gleichungen (2.26), (2.29) und (2.22) lässt sich zeigen, dass die Lösung der theoretischen Regression dem Orthogonalitätsprinzip

$$E\{\mathbf{N}_f(\mu) \cdot \tilde{\mathbf{V}}(\mu)^H\} = \vec{0} \tag{2.32}$$

folgt. Hier ist nun allerdings die Orthogonalität zu allen konjugierten, *nicht* konjugierten und teilweise verschobenen Spektralwerten der Erregung enthalten, die in dem Vektor $\vec{V}(\mu)$ zusammengefasst sind.

2.3 Leistungsdichtespektren und Fensterung

Wir haben gesehen, dass wir nur solche Systeme geeignet modellieren können, für die die Minimierung des Terms in Gleichung (2.4) eine periodisch zeitvariante Lösung für die Werte der beiden Impulsantworten $h_{*,\kappa}(k)$ und $h_{*,\kappa}(k)$ liefert. I. Allg. besagt dies jedoch *nicht*, dass auch die verbleibende Restdispersion, also der Wert $E\{|\mathbf{n}(k)|^2\}$ des zu minimierenden Ausdrucks in Gleichung (2.18), den man durch Einsetzen der optimalen Regressionskoeffizienten gemäß der Gleichung (2.19) und der Lösung des Gleichungssystems (2.25) erhalten, eine Folge ist, die sich abhängig vom Zeitpunkt k der Minimierung periodisch wiederholt. Der Prozess des Approximationsfehlers $\mathbf{n}(k)$ des realen Systems durch die beiden linearen, periodisch zeitvarianten Modellsysteme und die deterministische Störung ist daher i. Allg. instationär. Bei einem instationären Prozess ist es prinzipiell nicht möglich, die zweidimensionale Autokorrelationsfolge $\phi_{\mathbf{n}}(k_1, k_2)$ nach Gleichung (1.1) und die ebenfalls zweidimensionale Kreuzkorrelationsfolge $\psi_{\mathbf{n}}(k_1, k_2)$ nach Gleichung (1.3) mit einer Messung endlicher Dauer zu bestimmen. Man müsste dazu für jedes der unbegrenzt vielen Paare k_1, k_2 ein Ensemble der beiden daran beteiligten Zufallsgrößen $\mathbf{n}(k_1)$ und $\mathbf{n}(k_2)$ zur Verfügung haben, um daraus deren Korrelationen empirisch ermitteln zu können. An einem realen System wird man jedoch immer nur jeweils *eine* konkrete Realisierung für jeden Zeitpunkt zur Verfügung haben. Daher wollen wir uns im weiteren auf die Untersuchung solcher Systeme beschränken, bei denen der Prozess des Approximationsfehlers $\mathbf{n}(k)$ zyklstationär ist. Allenfalls kann die Zyklstationarität des Fehlers durch die in Kapitel 2.1 beschriebene zufällige Auswahl der Intervalle des Zugriffs auf das reale System erreicht werden.

Aufgrund der Zyklstationarität gilt für das zweite Moment des mittelwertfreien Approximationsfehlerprozesses zu den beiden diskreten Zeiten k_1 und k_2 — also für die von zwei Zeitvariablen abhängige und daher zweidimensionale AKF —

$$E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^*\} = E\{\mathbf{n}(k_1 + K_\Phi) \cdot \mathbf{n}(k_2 + K_\Phi)^*\}. \quad (2.33)$$

Dabei ist K_Φ die Periode der Momente des Prozesses $\mathbf{n}(k)$. Es sei hier darauf hingewiesen, dass K_Φ nicht unbedingt mit der Periode K_H der Zeitvarianz der beiden Modellsysteme übereinstimmen muss. Das aus der zweidimensionalen AKF durch zweidimensionale diskrete Fouriertransformation⁸ gewonnene zweidimensionale (bifrequente) LDS lässt sich

⁸Bei der Frequenzvariable Ω_2 wird dabei das Vorzeichen invertiert.

als Überlagerung von zur Gerade $\Omega_2 = \Omega_1$ parallelen Impulslinien im Abstand von $2\pi/K_\Phi$ schreiben.

$$\begin{aligned}
\Phi_{\mathbf{n}}(\Omega_1, \Omega_2) &= \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^*\} \cdot e^{-j \cdot \Omega_1 k_1} \cdot e^{j \cdot \Omega_2 k_2} = \\
&= \sum_{\bar{\mu}=-\infty}^{\infty} \frac{1}{K_\Phi} \cdot \sum_{k=0}^{K_\Phi-1} \dot{\Phi}_{\mathbf{n}}(\Omega_1, k) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} \cdot 2\pi \cdot \delta_0\left(\Omega_2 - \Omega_1 - \bar{\mu} \cdot \frac{2\pi}{K_\Phi}\right) = \\
&= \sum_{\bar{\mu}=-\infty}^{\infty} \dot{\Phi}_{\mathbf{n}}(\Omega_1, \bar{\mu}) \cdot 2\pi \cdot \delta_0\left(\Omega_2 - \Omega_1 - \bar{\mu} \cdot \frac{2\pi}{K_\Phi}\right) \quad (2.34)
\end{aligned}$$

Dabei sind

$$\dot{\Phi}_{\mathbf{n}}(\Omega, k) = \dot{\Phi}_{\mathbf{n}}(\Omega, k + K_\Phi) = \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k + \kappa) \cdot \mathbf{n}(k)^*\} \cdot e^{-j \cdot \Omega \cdot \kappa} \quad (2.35)$$

die von der Zeit k abhängigen und in Ω kontinuierlichen diskreten Fouriertransformierten der Autokorrelationsfolgen zu den Zeitpunkten k , die in Ω mit 2π und in k mit K_Φ periodisch sind. Die diskret invers Fouriertransformierten

$$\dot{\Phi}_{\mathbf{n}}(\Omega, \bar{\mu}) = \dot{\Phi}_{\mathbf{n}}(\Omega, \bar{\mu} + K_\Phi) = \frac{1}{K_\Phi} \cdot \sum_{k=0}^{K_\Phi-1} \dot{\Phi}_{\mathbf{n}}(\Omega, k) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} \quad (2.36)$$

einer Periode in k der Leistungsdichtespektren $\dot{\Phi}_{\mathbf{n}}(\Omega, k)$ bei den Frequenzen $\bar{\mu} \cdot 2\pi/K_\Phi$ sind in $\bar{\mu}$ mit K_Φ periodisch und treten in dem bifrequenten LDS — abgesehen von dem in Frequenzbereich typischerweise auftretenden Vorfaktor 2π — als Stärken der Impulslinien auf, die in der Ω_1, Ω_2 -Ebene in einem Abstand von Vielfachen von $2\pi/K_\Phi$ parallel zur Winkelhalbierenden $\Omega_2 = \Omega_1$ verlaufen. Diese Linien sind in Bild 2.4 für $K_\Phi = 4$ in der Ω_1, Ω_2 -Ebene fettgedruckt dargestellt. Da jeder beliebige Abtastwert — sofern man bei Distributionen überhaupt von Abtastwerten reden mag — des bifrequenten LDS entweder null⁹, oder nicht definiert ist, kann man das bifrequente LDS nicht durch endlich viele Abtastwert beschreiben. Sinnvoller erscheint es schon für jeden der K_Φ Zeitpunkte k endlich viele Abtastwerte der zeitabhängigen Leistungsdichtespektren $\dot{\Phi}_{\mathbf{n}}(\Omega, k)$, oder für jede der K_Φ Impulslinien endlich viele Abtastwerte der vor den Impulslinien auftretenden von $\bar{\mu}$ abhängigen Funktionen $\dot{\Phi}_{\mathbf{n}}(\Omega, \bar{\mu})$ anzugeben.

Da jedoch für die Abtastwerte des zeitabhängigen LDS $\dot{\Phi}_{\mathbf{n}}(\mu \cdot 2\pi/M, k)$ dasselbe gilt, wie für die Abtastwerte des zeitunabhängigen LDS im Fall des stationären Approximationsfehlerprozesses, verwenden wir stattdessen zunächst die zeitabhängigen Werte einer flä-

⁹Im distributionstheoretischen Sinn ist eine Distribution außerhalb eines abgeschlossenen Intervalls null, wenn die Distribution jeder Funktion, die innerhalb des Intervalls null ist, den Wert Null zuordnet.

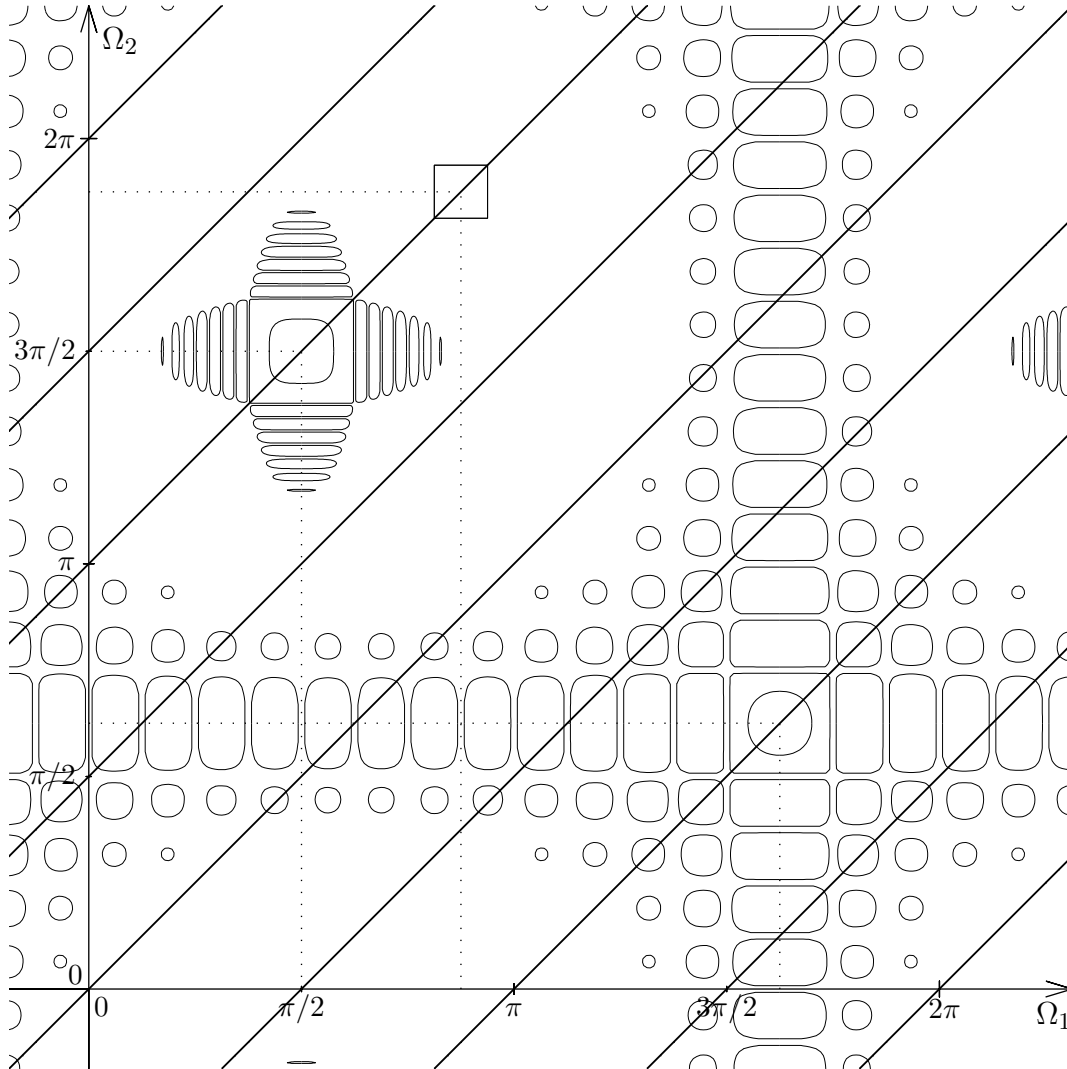


Bild 2.4: Zur näherungsweisen Beschreibung des bifrequenten LDS:

Höhenlinien des Betrags des Produktes $F\left(\mu \cdot \frac{2\pi}{M} - \Omega_1\right) \cdot F\left(\mu \cdot \frac{2\pi}{M} + \tilde{\mu} \cdot \frac{2\pi}{K_\Phi} - \Omega_2\right)^*$ der beiden verschobenen Fensterspektren im Integranden der Gleichung (2.39) für drei verschiedene Fensterfolgen am Beispiel mit $M=16$, $K_\Phi=4$ und

$\mu=7$, $\tilde{\mu}=2$ beim si-Fensters,

$\mu=13$, $\tilde{\mu}=-2$ beim Rechteckfenster mit $F=M$ und

$\mu=4$, $\tilde{\mu}=2$ beim Fensters nach Kapitel [1]:6 mit $F=4 \cdot M$

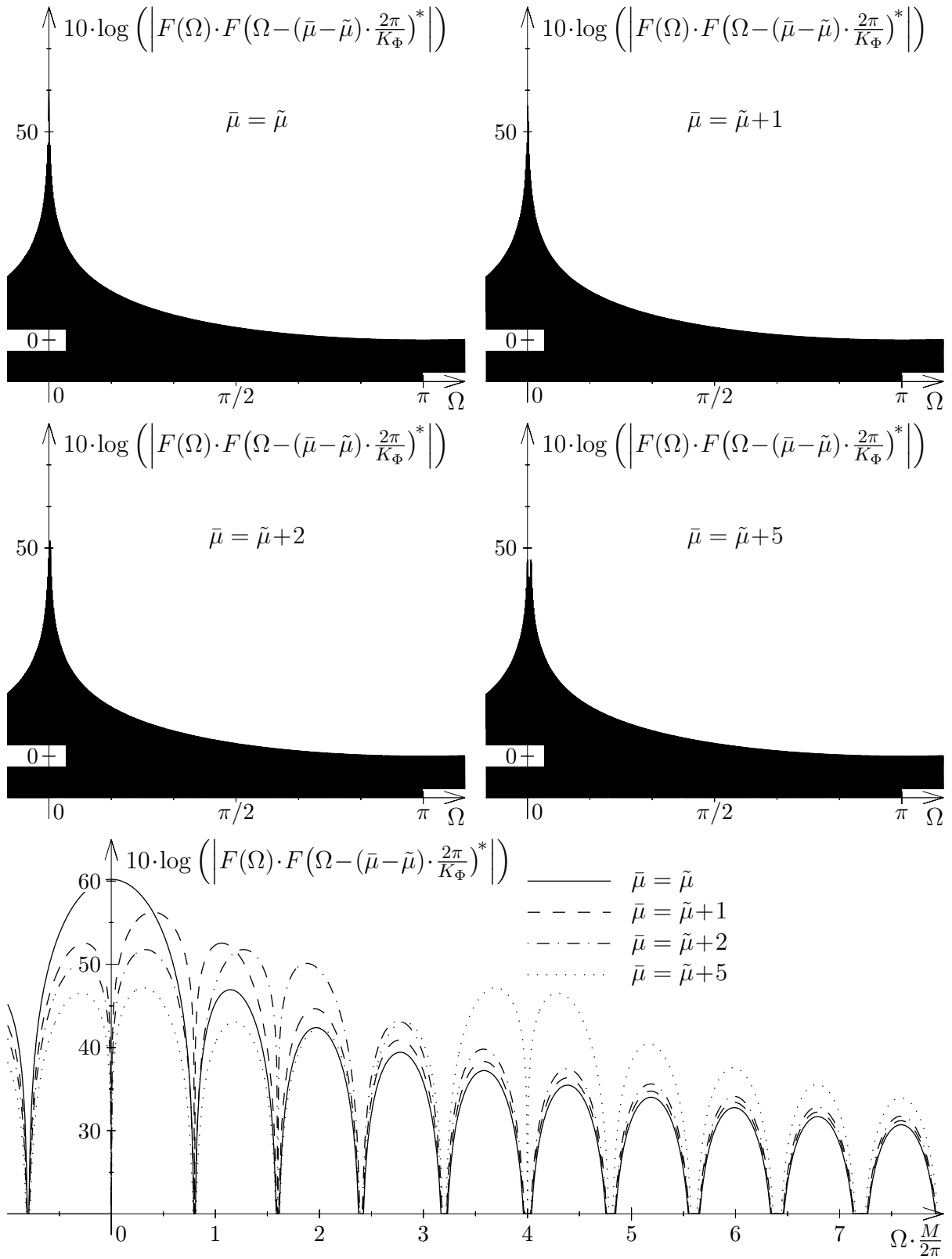


Bild 2.5: Dämpfung der Haupt- und Nebenlinien des bifrequenten LDS am Beispiel des Rechteckfensters mit der Fensterlänge $F = M = 1024$ und mit $M = K_\Phi$.

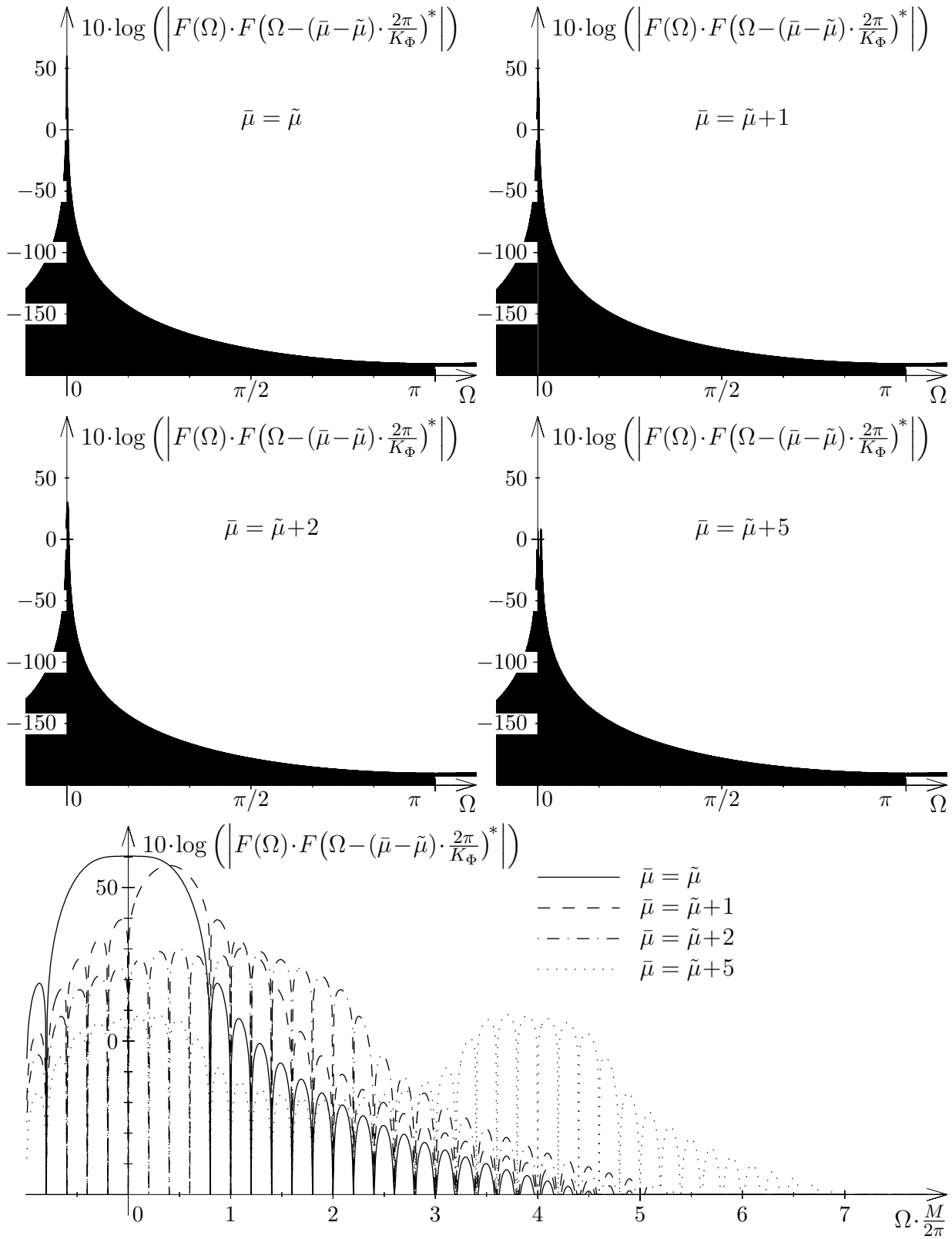


Bild 2.6: Dämpfung der Haupt- und Nebenlinien des bifrequenten LDS am Beispiel des Fensters nach Kapitel [1]:6 mit der Fensterlänge $F = 4 \cdot M$ und mit $M = K_\Phi = 1024$.

chengleichen Stufenapproximation, die analog zu Gleichung ([1]:2.13) als

$$\check{\Phi}_{\mathbf{n}}(\mu, k) = \frac{M}{2\pi} \cdot \int_{\mu \frac{2\pi}{M} - \frac{\pi}{M}}^{\mu \frac{2\pi}{M} + \frac{\pi}{M}} \dot{\Phi}_{\mathbf{n}}(\Omega, k) \cdot d\Omega = \frac{M}{2\pi} \cdot \int_{-\frac{\pi}{M}}^{\frac{\pi}{M}} \dot{\Phi}_{\mathbf{n}}(\mu \cdot \frac{2\pi}{M} - \Omega, k) \cdot d\Omega$$

$$\forall \quad k = 0 \text{ (1) } K_{\Phi} - 1 \quad \text{und} \quad \mu = 0 \text{ (1) } M - 1. \quad (2.37)$$

definiert werden. Um die Analogie zum bifrequenten LDS zu wahren, und um die auf 2π normierten¹⁰ Stärken $\dot{\Phi}_{\mathbf{n}}(\Omega, \bar{\mu})$ der Impulslinien abschätzen zu können, berechnen wir daraus die bezüglich k invers diskret Fouriertransformierten einer Periode von $\bar{\Phi}_{\mathbf{n}}(\mu, k)$:

$$\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) = \frac{1}{K_{\Phi}} \cdot \sum_{k=0}^{K_{\Phi}-1} \check{\Phi}_{\mathbf{n}}(\mu, k) \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k}$$

$$\forall \quad \mu = 0 \text{ (1) } M - 1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } K_{\Phi} - 1. \quad (2.38)$$

Die bereichsweise Integration in Gleichung (2.37) kann man wieder als eine Integration interpretieren, die man für alle $-\pi \leq \Omega < \pi$ durchführt, bei der man aber zuvor aus dem zeitabhängigen LDS mit Hilfe einer verschobenen Rechteckfunktionen den Bereich ausblendet, der den Integrationsgrenzen entspricht. Da sich ein rechteckförmiger Verlauf des Spektrums nicht mit einer endlich langen Zeitfolge realisieren lässt, kann man wie im Fall des stationären Approximationsfehlerprozesses die $M \cdot K_{\Phi}$ Werte $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ nicht als Funktion endlich vieler Werte der zweidimensionalen Autokorrelationsfolge angeben. Daher muss man sich auch hier damit begnügen, die Näherungswerte

$$\begin{aligned} \tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) &= \frac{1}{M} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^* \right\} = \\ &= \frac{1}{(2\pi)^2 \cdot M} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} F(\mu \cdot \frac{2\pi}{M} - \Omega_1) \cdot F(\mu \cdot \frac{2\pi}{M} + \tilde{\mu} \cdot \frac{2\pi}{K_{\Phi}} - \Omega_2)^* \cdot \Phi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot d\Omega_1 \cdot d\Omega_2 = \\ &= \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \frac{1}{2\pi \cdot M} \cdot \int_{-\pi}^{\pi} F(\Omega) \cdot F(\Omega - (\bar{\mu} - \tilde{\mu}) \cdot \frac{2\pi}{K_{\Phi}})^* \cdot \dot{\Phi}_{\mathbf{n}}(\mu \cdot \frac{2\pi}{M} - \Omega, \bar{\mu}) \cdot d\Omega \\ &\forall \quad \mu = 0 \text{ (1) } M - 1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } K_{\Phi} - 1, \end{aligned} \quad (2.39)$$

¹⁰Der Faktor 2π wird hier aus zwei Gründen als Vorfaktor des Dirac-Impulses betrachtet. Zum einen entsteht bei der Fouriertransformation eines periodischen Signals immer ein Impulslinienspektrum, dessen Impulsstärken gegenüber den Fourierreihenkoeffizienten um den Faktor 2π größer sind, und zum anderen ergeben sich so für $K_{\Phi} = 1$ exakt dieselben Werte für die Stufenapproximation wie im Fall eines stationären Prozesses. Alternativ — aber eher unüblich — kann man auch gleich das bifrequente LDS als auf 2π normiert definieren, wie dies beispielsweise in [5] gemacht wurde, und erhält dann unmittelbar die Stufenapproximationswerte als Vorfaktoren der Impulslinien.

die man durch Erwartungswertbildung aus den zufälligen Spektralwerten des gefenster-ten Approximationsfehlerprozesses gewinnt, für $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ angeben zu können. Sinnvollerweise wird man M als ganzzahliges Vielfaches von K_{Φ} wählen, da sich nur in diesem Fall die zufälligen Spektralwerte $\mathbf{N}_f(\mu + \tilde{\mu} \cdot M/K_{\Phi})$ durch eine blockweise Überlagerung und eine anschließende DFT aus den gefenster-ten Zufallsgrößen des Prozesses $\mathbf{n}(k)$ in der in Kapitel [1]:2.2 geschilderten Art berechnen lassen. Wie wir später noch sehen werden, benötigt man zur Berechnung der Messwertkovarianzen auch die Spektralwerte des gefenster-ten Approximationsfehlerprozesses bei den Frequenzen $\mu = -\tilde{\mu} \cdot M/(2 \cdot K_{\Phi})$ mit $\tilde{\mu} \in \mathbb{Z}$. Daher ist M als gerades Vielfaches von K_{Φ} zu wählen, so dass $M/(2 \cdot K_{\Phi}) \in \mathbb{N}$ gelten soll.

Setzt man in die letzte Gleichung einerseits das bifrequente LDS nach Gleichung (2.34) und andererseits eine Fensterfolge ein, deren Spektrum einen rechteckförmigen Betragsfrequenzgang mit einer Amplitude von M und einer Breite von $2\pi/M$ sowie einen beliebigen Phasenfrequenzgang aufweist, so wird durch die beiden Fensterspektren vom bifrequenten LDS nur die Impulslinie mit $\bar{\mu} = \tilde{\mu}$ ausgeblendet, wenn man $M \geq K_{\Phi}$ wählt. Die Integration über Ω_2 liefert im Sinne der Distributionentheorie den Wert der zu integrierenden Funktion bei der Frequenz $\Omega_2 = \Omega_1 + \bar{\mu} \cdot 2\pi/K_{\Phi} = \Omega_1 + \tilde{\mu} \cdot 2\pi/K_{\Phi}$. Bei dem vom Ω_2 abhängigen Fensterspektrum erhält man den Wert des Spektrums bei der Frequenz $\mu \cdot 2\pi/M - \Omega_1$. Dadurch entsteht innerhalb des Integrals über Ω_1 der Faktor $|F(\mu \cdot 2\pi/M - \Omega_1)|^2 = M^2$, der vom Phasenfrequenzgang der Fensterfolge *nicht* abhängt. Ein Vergleich mit der Definition von $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ nach Gleichung (2.38) in Verbindung mit Gleichung (2.37) zeigt, dass man mit einer Fensterfolge mit rechteckförmigem Betragsfrequenzgang gerade diese Werte für $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ erhält. Da der Phasenfrequenzgang des Fensters mit dem rechteckförmigen Betragsfrequenzgang bei dem zyklstationären Prozess für $M \geq K_{\Phi}$ bedeutungslos ist, kann man ihn auch konstant zu null setzen — also eine abgetastete si-Funktion als Fensterfolge verwenden —, und es ergibt sich, wenn man das rechteckige Spektrum des Fensters in den Integrationsgrenzen berücksichtigt, für die diskret Fourier-transformierte der zeitabhängigen Stufenapproximation folgendes:

$$\begin{aligned} \bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) &= \\ &= \frac{1}{M} \cdot \left(\frac{M}{2\pi}\right)^2 \cdot \int_{\mu \frac{2\pi}{M} - \frac{\pi}{M}}^{\mu \frac{2\pi}{M} + \frac{\pi}{M}} \int_{\mu \frac{2\pi}{M} + \tilde{\mu} \frac{2\pi}{K_{\Phi}} - \frac{\pi}{M}}^{\mu \frac{2\pi}{M} + \tilde{\mu} \frac{2\pi}{K_{\Phi}} + \frac{\pi}{M}} \Phi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot d\Omega_2 \cdot d\Omega_1 = \frac{M}{2\pi} \cdot \int_{-\frac{\pi}{M}}^{\frac{\pi}{M}} \Phi_{\mathbf{n}}(\Omega, \tilde{\mu}) \cdot d\Omega \\ \forall \quad \mu &= 0 \ (1) \ M-1, \quad \tilde{\mu} = 0 \ (1) \ K_{\Phi}-1 \quad \text{und} \quad K_{\Phi} \leq M \end{aligned} \quad (2.40)$$

Man erhält also das zweidimensionale Integral über das bifrequente LDS $\Phi_{\mathbf{n}}(\Omega_1, \Omega_2)$ in einem quadratischen Bereich der Ω_1, Ω_2 -Ebene, der um $\mu \cdot 2\pi/M$ in Ω_1 -Richtung und um

$\mu \cdot 2\pi/M + \tilde{\mu} \cdot 2\pi/K_\Phi$ in Ω_2 -Richtung verschoben ist. Das zweidimensionale Integral ist einerseits auf den Flächeninhalt $(2\pi/M)^2$ des Quadrats, und andererseits auf M normiert. Der Bereich, über den integriert wird, ist in Bild 2.4 am Beispiel mit $M=16$, $K_\Phi=4$, $\mu=7$ und $\tilde{\mu}=2$ als Quadrat der Breite $2\pi/M$ eingezeichnet. Der Wert $\bar{\Phi}_n(\mu, \mu + 2 \cdot M/K_\Phi)$ ist das Integral über den Vorfaktor $\dot{\Phi}_n(\Omega, 2)$ des Teils der Impulslinie, der die Diagonale des Quadrats bildet, wobei das Integral noch auf die Seitenlänge des Quadrats normiert ist. Warum es wichtig ist, dass bei dem zweidimensionalen Integral die zusätzliche Normierung auf M , und nicht nur die Normierung auf den Flächeninhalt des Quadrats durchgeführt wird, sei kurz erläutert. Betrachten wir dazu einen zyklstationären Prozess, bei dem die Stärke der Impulslinien für hinreichend große Werte von M — also hinreichend schmale Integrationsbereiche —, als im Bereich der Integration konstant anzusehen ist. Da bei einem zyklstationären Prozess sich die gesamte Energie des LDS auf die zur Winkelhalbierenden der Ω_1, Ω_2 -Ebene parallelen Geraden konzentriert, wird bei einer Veränderung von M der Wert des Integrals vor der Normierung proportional zur Länge der Diagonale des Quadrats, über das integriert wird, — also indirekt proportional zu M — sein. Es ist also *keine* Proportionalität zur Fläche über die integriert wird — also keine Proportionalität zu M^{-2} — vorhanden, wie dies bei einem LDS eines instationären Prozesses der Fall wäre, bei dem das LDS sich auf die gesamte Ω_1, Ω_2 -Ebene verteilt. Nach der Normierung auf die Größe der Fläche, über die integriert wurde, ergibt sich somit eine Proportionalität zu M , was für steigende Werte von M den Impulsliniencharakter des bifrequenten LDS widerspiegelt. Eine Stufenapproximation der Impulslinien, die von der Wahl von M abhängt, ist jedoch nicht das, was wir zur Beschreibung des LDS mittels endlich vieler Werte wünschen. Damit wir das erhalten, was wir eigentlich bestimmen wollen, nämlich Näherungen für die auf 2π normierten Stärken $\dot{\Phi}_n(\Omega, \bar{\mu})$ der Impulslinien, muss man zusätzlich auf M normieren, um so zu erreichen, dass die Werte $\bar{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ für steigende Werte von M gegen diese Werte konvergieren, und nicht divergieren.

Da eine Fensterfolge mit einem rechteckförmigen Betragsquadratspektrum zeitlich unbegrenzt ist, wird man versuchen, eine zeitlich begrenzte Fensterfolge zu verwenden, bei der das Betragsquadrat des Spektrums den rechteckigen Wunschverlauf möglichst gut approximiert. Da jedoch keine zeitlich begrenzte Fensterfolge abrupt vom Durchlassbereich ($|\Omega| \leq \pi/M$), in dem der Betrag des Spektrums nahe bei M liegt, in den Sperrbereich ($2\pi/M \leq |\Omega| \leq \pi$) mit hoher Dämpfung wechselt, empfiehlt es sich M wenigstens doppelt so groß wie K_Φ — wenn möglich natürlich noch größer — zu wählen, so dass außer dem Bereich der einen Impulslinie, der aus dem bifrequenten LDS durch das Betragsquadratspektrum der Fensterfolge herausgeschnitten werden soll, keine weitere Impulslinie bei den beiden verschobenen Spektren der Fensterfolge gleichzeitig durch den Übergangsbereich ($\pi/M \leq |\Omega| \leq 2\pi/M$) der Dämpfung verläuft. So kann eine zusätzliche Verfälschung der Werte $\bar{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ weitgehend vermieden werden. Weil das LDS des

zyklostationären Approximationsfehlerprozesses aus einzelnen parallelen Impulslinien besteht, kommt dann bei Verwendung einer hoch frequenzselektiven Fensterfolge bei der Berechnung von $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ nach Gleichung 2.39 in guter Näherung nur mehr die eine Impulslinie mit $\bar{\mu} = \tilde{\mu}$ zum tragen, so dass nur das Betragsquadrat des Spektrums der Fensterfolge im Integral steht, und daher der Phasenfrequenzgang der Fensterfolge beliebig gewählt werden kann. In Bild 2.4 sind für zwei verschiedene Fensterfolgen die Höhenlinien des Betrags des Produktes der Fensterspektren im zweidimensionalen Integral in Gleichung 2.39 am Beispiel mit $M=16$ und $K_{\Phi}=4$ eingetragen. Das eine Fenster ist das Rechteckfenster mit einer Länge von M und einer Höhe von eins. Hier wurde $\mu=13$, $\tilde{\mu}=-2$ und für die beiden Höhenlinien die Werte $M^2/2$ und $M^2/100$ gewählt. Da die Höhenlinien jeweils geschlossene Kurven bilden, die die Maxima des Betrags des Produktes der Fensterspektren umschließen, erkennt man an der Höhenlinie mit $M^2/2$ die ungefähre Lage des Hauptmaximums und an der Höhenlinie mit $M^2/100$ die ungefähre Lage der Nebenmaxima. Desweiteren erkennt man, dass der Betrag des Produktes der Fensterspektren in weiten Gebieten der Ω_1, Ω_2 -Ebene nicht unter $M^2/100$ absinkt, so dass bei dem Rechteckfenster zu erwarten ist, dass man mit $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ nur eine schlechte Näherung für $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ erhält, wenn einerseits der zu nähernde Wert klein, und andererseits die durch die Ausblendung zu unterdrückenden Anteile groß sind. Das zweite Fenster ist das mit dem Algorithmus nach Kapitel [1]:6 konstruierte Fenster, wobei die Fensterlänge auf $4 \cdot M$ festgelegt wurde. Bei diesem Fenster wurde $\mu=4$, $\tilde{\mu}=2$ und für die beiden Höhenlinien die Werte $M^2/2$ und $M^2/10000$ gewählt. Die zweite Höhenlinie ist hier also um den Faktor 100 kleiner als bei dem Rechteckfenster. Man erkennt dass die Nebenmaxima wesentlich schneller abklingen, und außerdem wesentlich niedriger liegen. Daher ist zu erwarten, dass durch die bessere Unterdrückung der Nebenlinien des LDS, also der Linien, die für den eingestellten Wert von $\tilde{\mu}$ gerade nicht in $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ eingehen sollen, eine deutlich bessere Näherung für $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ erhalten werden kann. Um zu demonstrieren, warum man nicht $M=K_{\Phi}$ wählen sollte, wurde in den Bildern 2.5 und 2.6 bei dieser Einstellung mit $M=K_{\Phi}=1024$ der Betrag des Produktes der Fensterspektren im eindimensionalen Integral in Gleichung (2.39) einerseits für den gewünschten Summanden mit $\bar{\mu} = \tilde{\mu}$ und andererseits für die unerwünschten Summanden mit $\bar{\mu} = \tilde{\mu}+1$, $\bar{\mu} = \tilde{\mu}+2$ und $\bar{\mu} = \tilde{\mu}+5$ als Funktionen über Ω aufgetragen. Wieder wurden dieselben beiden Fensterfolgen, allerdings mit dem Wert $M=1024$ statt $M=16$ verwendet. Da sich die wesentlichen Anteile des Fensterspektrums im Bereich kleiner Frequenzen befinden, ist dieser nochmals vergrößert dargestellt. Bei der Darstellung des gesamten Frequenzbereichs erkennt man, einerseits den deutlich schnelleren und stärkeren Abfall des Spektrums des nach Kapitel [1]:6 berechneten Fensters zu hohen Frequenzen hin, so dass man davon ausgehen kann, dass selbst extrem starke Impulse auf den Nebenlinien gut unterdrückt werden können, die gegenüber der gerade vermessenen Nebenlinie weit genug

entfernt liegen. Andererseits sieht man, dass bei $M=1024$ die Nullstellen der Fensterpektren so nahe beieinander liegen, dass man in der graphischen Darstellung unterhalb einer Linie, die die Nebenmaxima verbindet, nur mehr eine schwarze Fläche erhält. In der Darstellung des Bereichs niedriger Frequenzen ist bei beiden Fensterfolgen zu sehen, dass der Anteil der Nebenlinie mit $\bar{\mu} = \tilde{\mu} + 1$ für die Frequenz $\Omega = \pi/M$ nur mit etwa dem Faktor 0.5 unterdrückt wird. Dies kann auch nicht anders sein, da diese Frequenz bei dem Spektrum der Fensterfolge gerade die Grenze des gewünschten Durchlassbereichs darstellt. Würde man einen schmälere Durchlassbereich wählen, so wäre es praktisch unmöglich, die Bedingung ([1]:2.20) zu erfüllen, nach der die Überlagerung aller um Vielfache von $2\pi/M$ verschobenen Betragsquadrate des Spektrums der Fensterfolge eine Konstante sein soll. Daher sollte man unbedingt vermeiden M und K_Φ gleich zu wählen. Bei dem nach Kapitel [1]:6 berechneten Fenster ergibt sich in Bild 2.6 bereits bei der zweiten Nebenlinie des LDS mit $\bar{\mu} = \tilde{\mu} + 2$ eine ganz brauchbare Unterdrückung von mehr als 30dB im Maximum, und bei $\bar{\mu} = \tilde{\mu} + 5$ ist die Nebenlinie des LDS schon wenigstens um mehr als 50dB abgeschwächt. Wie Bild 2.5 zeigt ist diese Unterdrückung bei dem Rechteckfenster bei weitem nicht so gut, und nimmt zu weiter entfernten Nebenlinien auch nicht so rasch ab.

Im weiteren wird angenommen, dass wir M groß genug gewählt haben und eine hoch frequenzselektive Fensterfolge verwenden, so dass die Erwartungswerte

$$\begin{aligned} & \mathbb{E}\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\hat{\mu})^*\} = \\ &= \frac{1}{(2\pi)^2} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} F\left(\mu \cdot \frac{2\pi}{M} - \Omega_1\right) \cdot F\left(\hat{\mu} \cdot \frac{2\pi}{M} - \Omega_2\right)^* \cdot \Phi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot d\Omega_1 \cdot d\Omega_2 \approx 0 \\ & \quad \forall \quad \hat{\mu} \neq \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \end{aligned} \quad (2.41)$$

für alle Kombinationen von μ und $\hat{\mu}$, die nicht bei der Berechnung der Näherungswerte $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ auftreten, in guter Näherung null sind. Diese Vernachlässigung gilt vor allem dann, wenn diese Erwartungswerte in Summen auftreten, die außerdem noch die Näherungswerte des LDS enthalten, die vom Hauptmaximum des Spektrums der Fensterfolge aus dem LDS herausgeschnitten wurden. Es sei noch erwähnt, dass für die Werte $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ erstens die Symmetrie

$$\begin{aligned} \tilde{\Phi}_{\mathbf{n}}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu\right) &= \frac{1}{M} \cdot \mathbb{E}\left\{\mathbf{N}_f\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot \mathbf{N}_f(\mu)^*\right\} = \\ &= \frac{1}{M} \cdot \mathbb{E}\left\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right)^*\right\}^* = \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right)^*, \end{aligned} \quad (2.42)$$

und zweitens die Ungleichung

$$\left|\tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right)\right|^2 \leq \tilde{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \tilde{\Phi}_{\mathbf{n}}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right). \quad (2.43)$$

gilt. Letzteres zeigt man indem man in Gleichung ([1]:A.14) des Anhangs [1]:A.2 die Substitutionen $\mathbf{X} = \mathbf{N}_f(\mu)$ und $\mathbf{Y} = \mathbf{N}_f(\mu + \tilde{\mu} \cdot M/K_\Phi)^*$ vornimmt.

Die Beschreibung der zweiten Momente des zyklstationären, mittelwertfreien Approximationsfehlerprozesses $\mathbf{n}(k)$ ist nur dann vollständig, wenn man auch noch die Korrelationsfolge $E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)\}$ angibt, die wegen der Zyklstationarität die Periodizität

$$E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)\} = E\{\mathbf{n}(k_1 + K_\Phi) \cdot \mathbf{n}(k_2 + K_\Phi)\}. \quad (2.44)$$

aufweist. Durch zweidimensionale diskrete Fouriertransformation —wieder mit Vorzeicheninvertierung bei der zweiten Frequenzvariable— gewinnen wir daraus die bifrequente kontinuierliche und in beiden Frequenzvariablen mit 2π periodische Funktion

$$\begin{aligned} \Psi_{\mathbf{n}}(\Omega_1, \Omega_2) &= \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} E\{\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)\} \cdot e^{-j \cdot \Omega_1 k_1} \cdot e^{j \cdot \Omega_2 k_2} = \\ &= \sum_{\bar{\mu}=-\infty}^{\infty} \frac{1}{K_\Phi} \cdot \sum_{k=0}^{K_\Phi-1} \dot{\Psi}_{\mathbf{n}}(\Omega_1, k) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} \cdot 2\pi \cdot \delta_0(\Omega_2 - \Omega_1 - \bar{\mu} \cdot \frac{2\pi}{K_\Phi}) = \\ &= \sum_{\bar{\mu}=-\infty}^{\infty} \dot{\Psi}_{\mathbf{n}}(\Omega_1, \bar{\mu}) \cdot 2\pi \cdot \delta_0(\Omega_2 - \Omega_1 - \bar{\mu} \cdot \frac{2\pi}{K_\Phi}), \end{aligned} \quad (2.45)$$

die sich ebenfalls als Überlagerung von zur Gerade $\Omega_2 = \Omega_1$ parallelen Impulslinien im Abstand von $2\pi/K_\Phi$ schreiben lässt. Dabei sind

$$\dot{\Psi}_{\mathbf{n}}(\Omega, k) = \dot{\Psi}_{\mathbf{n}}(\Omega, k + K_\Phi) = \sum_{\kappa=-\infty}^{\infty} E\{\mathbf{n}(k + \kappa) \cdot \mathbf{n}(k)\} \cdot e^{-j \cdot \Omega \cdot \kappa} \quad (2.46)$$

die von der Zeit k abhängigen und in Ω kontinuierlichen diskreten Fouriertransformierten der Kreuzkorrelationsfolgen zu den Zeitpunkten k , die in Ω mit 2π und in k mit K_Φ periodisch sind. Die diskret invers Fouriertransformierten

$$\dot{\Psi}_{\mathbf{n}}(\Omega, \bar{\mu}) = \dot{\Psi}_{\mathbf{n}}(\Omega, \bar{\mu} + K_\Phi) = \frac{1}{K_\Phi} \cdot \sum_{k=0}^{K_\Phi-1} \dot{\Psi}_{\mathbf{n}}(\Omega, k) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} \quad (2.47)$$

einer Periode in k der Kreuzleistungsdichtespektren $\dot{\Psi}_{\mathbf{n}}(\Omega, k)$ bei den Frequenzen $\bar{\mu} \cdot \frac{2\pi}{K_\Phi}$ sind in $\bar{\mu}$ mit K_Φ periodisch und treten in dem bifrequenten KLDS —abgesehen von dem in Frequenzbereich typischerweise auftretenden Vorfaktor 2π — als Stärken der Impulslinien auf, die in der Ω_1, Ω_2 -Ebene in einem Abstand von Vielfachen von $2\pi/K_\Phi$ parallel zur Winkelhalbierenden $\Omega_2 = \Omega_1$ verlaufen. Um eine sinnvolle Aussage über diese kontinuierlichen Funktionen machen zu können, geben wir zunächst wieder die endlich vielen

Werte

$$\begin{aligned} \check{\Psi}_{\mathbf{n}}(\mu, k) &= \frac{M}{2\pi} \cdot \int_{-\frac{\pi}{M}}^{\frac{\pi}{M}} \check{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, k\right) \cdot d\Omega \\ \forall \quad k &= 0 \text{ (1) } K_{\Phi}-1 \quad \text{und} \quad \mu = 0 \text{ (1) } M-1 \end{aligned} \quad (2.48)$$

der flächengleichen Stufenapproximation an. Die daraus durch inverse diskrete Fouriertransformation einer Periode gewonnenen $M \cdot K_{\Phi}$ Werte

$$\begin{aligned} \bar{\Psi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) &= \frac{1}{K_{\Phi}} \cdot \sum_{k=0}^{K_{\Phi}-1} \check{\Psi}_{\mathbf{n}}(\mu, k) \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \\ \forall \quad \mu &= 0 \text{ (1) } M-1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } K_{\Phi}-1 \end{aligned} \quad (2.49)$$

werden durch deren Näherungswerte

$$\begin{aligned} \tilde{\Psi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) &= \frac{1}{M} \cdot \mathbb{E}\left\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right\} = \\ &= \frac{1}{(2\pi)^2 \cdot M} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} F\left(\mu \cdot \frac{2\pi}{M} - \Omega_1\right) \cdot F\left(\Omega_2 - \mu \cdot \frac{2\pi}{M} - \tilde{\mu} \cdot \frac{2\pi}{K_{\Phi}}\right) \cdot \Psi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot d\Omega_1 \cdot d\Omega_2 = \\ &= \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \frac{1}{2\pi \cdot M} \cdot \int_{-\pi}^{\pi} F(\Omega) \cdot F\left((\bar{\mu} - \tilde{\mu}) \cdot \frac{2\pi}{K_{\Phi}} - \Omega\right) \cdot \check{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, \bar{\mu}\right) \cdot d\Omega \\ \forall \quad \mu &= 0 \text{ (1) } M-1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } K_{\Phi}-1, \end{aligned} \quad (2.50)$$

die man durch Erwartungswertbildung aus den zufälligen Spektralwerten des gefensterten Approximationsfehlerprozesses gewinnen kann, ersetzt. Es sei wieder angemerkt, dass für eine reelle Fensterfolge dieselben Fensterspektren im Integral stehen wie in Gleichung (2.39). Auch hier ist festzustellen, dass der Erwartungswert

$$\begin{aligned} &\mathbb{E}\left\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f(-\hat{\mu})\right\} = \\ &= \frac{1}{(2\pi)^2} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} F\left(\mu \cdot \frac{2\pi}{M} - \Omega_1\right) \cdot F\left(\Omega_2 - \hat{\mu} \cdot \frac{2\pi}{M}\right) \cdot \Psi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot d\Omega_1 \cdot d\Omega_2 \approx 0 \\ \forall \quad \hat{\mu} &\neq \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}} \end{aligned} \quad (2.51)$$

für alle Kombinationen von μ und $\hat{\mu}$, die nicht bei der Berechnung von $\tilde{\Psi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)$ auftreten, in guter Näherung null ist, wenn man M groß genug wählt, und eine hoch

frequenzselektive Fensterfolge verwendet. $\tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ weist die Symmetrie

$$\begin{aligned} \tilde{\Psi}_{\mathbf{n}}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, \mu) &= \frac{1}{M} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \mathbf{N}_f(-\mu) \right\} = \\ &= \frac{1}{M} \cdot \mathbb{E} \left\{ \mathbf{N}_f(-\mu) \cdot \mathbf{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right\} = \tilde{\Psi}_{\mathbf{n}}(-\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \end{aligned} \quad (2.52)$$

auf, und erfüllt die Ungleichung

$$\left| \tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right|^2 \leq \tilde{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \tilde{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \quad (2.53)$$

deren Gültigkeit man mit $\mathbf{X} = \mathbf{N}_f(\mu)$ und $\mathbf{Y} = \mathbf{N}_f(-\mu - \tilde{\mu} \cdot M/K_{\Phi})$ in Gleichung ([1]:A.14) des Anhangs [1]:A.2 zeigt.

Wenn man berücksichtigt, dass sich die zweidimensionale Autokorrelationsfolge nach Gleichung (1.1) durch die inverse zweidimensionale Fourierrücktransformation — wieder mit Vorzeicheninvertierung bei Ω_2 — aus dem bifrequenten LDS gemäß

$$\mathbb{E} \{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \} = \frac{1}{(2\pi)^2} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} \Phi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot e^{j \cdot (\Omega_1 k_1 - \Omega_2 k_2)} \cdot d\Omega_1 \cdot d\Omega_2 \quad (2.54)$$

berechnen lässt, kann man die zeitabhängige Varianz des Approximationsfehlers, die man mit $k_1 = k_2 = k$ erhält, als endliche Summe

$$\begin{aligned} \mathbb{E} \{ |\mathbf{n}(k)|^2 \} &= \frac{1}{(2\pi)^2} \cdot \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} \Phi_{\mathbf{n}}(\Omega_1, \Omega_2) \cdot e^{j \cdot (\Omega_1 - \Omega_2) \cdot k} \cdot d\Omega_1 \cdot d\Omega_2 = \\ &= \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \end{aligned} \quad (2.55)$$

angeben. Wenn man die gleiche Summe mit den Näherungswerten $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ bildet, und wenn man M hinreichend groß gewählt hat, so dass man vom bifrequenten LDS die Impulslinien mit $\bar{\mu} \neq \tilde{\mu}$ vernachlässigen kann, erhält man in guter Näherung

$$\begin{aligned} &\frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} = \\ &= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \int_{-\pi}^{\pi} F(\Omega) \cdot F(\Omega - (\bar{\mu} - \tilde{\mu}) \cdot \frac{2\pi}{K_{\Phi}})^* \cdot \tilde{\Phi}_{\mathbf{n}}(\mu \cdot \frac{2\pi}{M} - \Omega, \bar{\mu}) \cdot d\Omega \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \approx \\ &\approx \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \int_{-\pi}^{\pi} |F(\Omega)|^2 \cdot \tilde{\Phi}_{\mathbf{n}}(\mu \cdot \frac{2\pi}{M} - \Omega, \tilde{\mu}) \cdot d\Omega \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} = \\ &= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} |F(\Omega)|^2 \cdot \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \tilde{\Phi}_{\mathbf{n}}(\mu \cdot \frac{2\pi}{M} - \Omega, \tilde{\mu}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \cdot d\Omega = \end{aligned}$$

$$\begin{aligned}
&= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} |F(\Omega)|^2 \cdot \dot{\Phi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, k\right) \cdot d\Omega = \\
&= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} |F(\Omega)|^2 \cdot \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} \cdot e^{-j \cdot \left(\mu \cdot \frac{2\pi}{M} - \Omega\right) \cdot \kappa} \cdot d\Omega = \\
&= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} |F(\mu \cdot \frac{2\pi}{M} - \tilde{\Omega})|^2 \cdot \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} \cdot e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} = \\
&= \frac{1}{2\pi \cdot M^2} \cdot \sum_{\kappa=-\infty}^{\infty} \int_{-\pi}^{\pi} \underbrace{\sum_{\mu=0}^{M-1} |F(\mu \cdot \frac{2\pi}{M} - \tilde{\Omega})|^2}_{= M^2 \text{ nach ([1]:2.20)}} \cdot \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} \cdot e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} = \\
&= \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} \cdot \frac{1}{2\pi} \cdot \int_{-\pi}^{\pi} e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} = \\
&= \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} \cdot \gamma_0(\kappa) = \mathbb{E}\{|\mathbf{n}(k)|^2\}, \tag{2.56}
\end{aligned}$$

wobei vorausgesetzt wurde, dass man eine Fensterfolge verwendet, deren Spektrum der Bedingung ([1]:2.20) genügt.

Anmerkung: Wenn man eine Fensterfolge verwendet, die mit dem in Kapitel [1]:6 vorgestellten Algorithmus berechnet wird, und wenn man die Fensterlänge groß genug wählt, erhält man bereits für $M = 2 \cdot K_{\Phi}$ eine so hohe Sperrdämpfung, dass die eben gemachte Näherung zu Fehlern bei der Berechnung der zeitabhängigen theoretischen Varianz führt, die um Größenordnungen kleiner sind als die zufälligen Abweichungen der empirischen Varianz, die man mit dem RKM selbst für extrem große Mittelungsanzahlen L misst. Es gibt jedoch wenigstens zwei Möglichkeiten, die Näherung bei der Berechnung der theoretischen Varianz zu umgehen. Die eine Möglichkeit besteht darin, den zweidimensionalen Approximationsfehlerprozess $\mathbf{n}(k_1) \cdot \mathbf{n}(k_2)$ zu bilden, und diesen mit einer zweidimensionalen Fensterfolge zu multiplizieren, deren zweidimensionales Spektrum in der Ω_1, Ω_2 -Ebene dort Nulllinien aufweist, wo das zweidimensionale LDS seine Impulslinien hat (mit Ausnahme der Linie $\Omega_1 = \Omega_2$) und die auch noch einige andere Bedingungen erfüllen muss. Eine solche zweidimensionale Fensterfolge kann man z. B. dadurch erhalten, dass man aus der mit dem in Kapitel [1]:6 vorgestellten Algorithmus berechneten eindimensionalen Fensterfolge $f(k)$ die zweidimensionale Fensterfolge $f(k_1) \cdot f(k_2)$ bildet und diese mit der zweidimensionalen Folge faltet, die nur für $0 \leq k_1 = k_2 < K_{\Phi}$ eins ist und sonst null. Die so entstandene zweidimensionale Fensterfolge erfüllt alle Forderungen, die für die Anwendung beim RKM notwendig sind, und ermöglicht ebenfalls eine erwartungstreue Messung der zeitabhängigen Varianz. Da der sich durch die Verwendung des zweidimensional gefensterten Approximationsfehlerprozesses

beim RKM ergebende Mehraufwand nicht unerheblich ist, ist diese Modifikation jedoch nicht zu rechtfertigen. Daher wird auf die ausführliche Darstellung dieser Modifikationen verzichtet. Eine weitere Möglichkeit besteht darin, die Messung mit dem RKM für die Korrelationen aller K_Φ Polyphasenkomponenten [6] des Eingangs- und des Ausgangssignals getrennt durchzuführen, und so die einzelnen Polyphasenkomponenten mit einer entsprechend um den Faktor K_Φ gespreizten Fensterfolge zu fenstern, wie dies in [5] für die Messung von Multiraten Systemen durchgeführt wird. Die dabei auftretenden Prozesse der einzelnen Polyphasenkomponenten sind dann alle stationär, so dass Nebenlinien im bifrequenten LDS der Polyphasenkomponenten dann nicht mehr auftreten.

Analog erhält man, wenn man M hinreichend groß gewählt hat, so dass man bei $\Psi_{\mathbf{n}}(\Omega_1, \Omega_2)$ nach Gleichung (2.25) die Impulslinien mit $\bar{\mu} \neq -\tilde{\mu}$ vernachlässigen kann, für die Summe über alle mit denselben Drehfaktoren wie in Gleichung (2.56) multiplizierten Werte $\tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ eine gute Näherung für den zeitabhängigen Erwartungswert $E\{\mathbf{n}(k)^2\}$.

$$\begin{aligned}
& \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \tilde{\Psi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \sum_{\bar{\mu}=0}^{K_\Phi-1} \int_{-\pi}^{\pi} F(\Omega) \cdot F\left((\bar{\mu} - \tilde{\mu}) \cdot \frac{2\pi}{K_\Phi} - \Omega\right) \cdot \dot{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, \bar{\mu}\right) \cdot d\Omega \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} \approx \\
& \approx \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \int_{-\pi}^{\pi} F(\Omega) \cdot F(-\Omega) \cdot \dot{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, \tilde{\mu}\right) \cdot d\Omega \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} F(\Omega) \cdot F(-\Omega) \cdot \sum_{\tilde{\mu}=0}^{K_\Phi-1} \dot{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, \tilde{\mu}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} \cdot d\Omega = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} F(\Omega) \cdot F(-\Omega) \cdot \dot{\Psi}_{\mathbf{n}}\left(\mu \cdot \frac{2\pi}{M} - \Omega, k\right) \cdot d\Omega = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} |F(\Omega)|^2 \cdot \sum_{\kappa=-\infty}^{\infty} E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\} \cdot e^{-j \cdot \left(\mu \cdot \frac{2\pi}{M} - \Omega\right) \cdot \kappa} \cdot d\Omega = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\mu=0}^{M-1} \int_{-\pi}^{\pi} \left|F\left(\mu \cdot \frac{2\pi}{M} - \tilde{\Omega}\right)\right|^2 \cdot \sum_{\kappa=-\infty}^{\infty} E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\} \cdot e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} = \\
& = \frac{1}{2\pi \cdot M^2} \cdot \sum_{\kappa=-\infty}^{\infty} \int_{-\pi}^{\pi} \underbrace{\sum_{\mu=0}^{M-1} \left|F\left(\mu \cdot \frac{2\pi}{M} - \tilde{\Omega}\right)\right|^2}_{= M^2 \text{ nach ([1]:2.20)}} \cdot E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\} \cdot e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} =
\end{aligned}$$

$$\begin{aligned}
&= \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\} \cdot \frac{1}{2\pi} \cdot \int_{-\pi}^{\pi} e^{-j \cdot \tilde{\Omega} \cdot \kappa} \cdot d\tilde{\Omega} = \\
&= \sum_{\kappa=-\infty}^{\infty} \mathbb{E}\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\} \cdot \gamma_0(\kappa) = \mathbb{E}\{\mathbf{n}(k)^2\}
\end{aligned} \tag{2.57}$$

Dabei wurde vorausgesetzt, dass man eine reelle Fensterfolge verwendet, deren Spektrum der Bedingung ([1]:2.20) genügt.

In diesem Kapitel wurde gezeigt, wie sich ein reales gestörtes System durch das Systemmodell nach Bild 1.1 modellieren lässt. Die sich ergebenden *theoretisch* optimalen Werte der deterministischen Störung und der beiden bifrequenten Übertragungsfunktionen wurden in der Art bestimmt, dass der verbleibende Approximationsfehlerprozess eine minimale Varianz aufweist. Die theoretischen Optimallösungen berechnen sich gemäß der Gleichungen (2.25) und (2.19) bzw. (2.29), falls folgende Voraussetzungen erfüllt sind:

- Das System wird mit einem bereichsweise periodischen Zufallssignal erregt.
- Die Autokovarianzmatrix $\underline{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}$ der Spektralwerte des erregenden Zufallssignals ist für alle M diskreten Frequenzen $\mu \cdot \frac{2\pi}{M}$ regulär.
- Das System lässt sich durch ein lineares periodisch zeitvariantes System approximieren, d. h. die Approximationslösung wiederholt sich periodisch mit dem Zeitpunkt der Approximation.
- Bei der Berechnung der Optimallösungen wird eine Fensterfolge verwendet, deren Spektralwerte bei den diskreten Frequenzen $\mu \cdot \frac{2\pi}{M}$ die Bedingung ([1]:2.27) erfüllen.

Es wurde zur Beschreibung des bifrequenten LDS und des KLDS des Fehlerprozesses in den Gleichungen (2.39) und (2.50) jeweils eine zweidimensionale Spektralfolge endlich vieler Werte angegeben, die sich mit Hilfe einer Fensterung, einer DFT und einer Erwartungswertbildung aus den Ein- und Ausgangsprozessen berechnen lässt. Dabei stellen die folgenden Voraussetzungen sicher, dass die beiden Spektralfolgen in der Lage sind das LDS und das KLDS des Approximationsfehlerprozesses aussagekräftig zu beschreiben.

- Der bei der Approximation verbleibende Fehlerprozess ist zyklstationär.
- Es wird eine Fensterfolge verwendet, deren Spektrum die Bedingung ([1]:2.20) erfüllt, und das eine hohe Sperrdämpfung aufweist.

Sollten diese Voraussetzungen vom realen System nicht a priori erfüllt sein, kann eine zufällige Verschiebung des Zugriffszeitpunktes auf das reale System, wie sie im ersten Unterkapitel beschrieben ist, dafür sorgen, dass die so entstehenden Modellzufallsvektoren die Voraussetzungen erfüllen.

3 Das Rauschklimrmessverfahren mit Fensterung

Nun wollen wir uns der Frage widmen, wie man die mit Hilfe einer Messung Schätzwerte

$$\begin{aligned}
 \hat{H}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_H-1, \\
 \hat{H}_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_H-1, \\
 \hat{u}(k) & \quad \text{mit} \quad k = 0 \ (1) \ F-1, \\
 \hat{U}_f(\mu) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1, \\
 \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1, \\
 \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1
 \end{aligned}$$

für die Optimallösungen

$$\begin{aligned}
 H(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_H-1, \\
 H_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_H-1, \\
 u(k) & \quad \text{mit} \quad k = 0 \ (1) \ F-1 \quad \text{und} \\
 U_f(\mu) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1,
 \end{aligned}$$

sowie für die Näherungen

$$\begin{aligned}
 \tilde{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1, \\
 \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) & \quad \text{mit} \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1
 \end{aligned}$$

der Stufenapproximationen des LDS und des KLDS des Approximationsfehlers gewinnen kann. Wir gehen dazu im wesentlichen so vor, wie dies bei der empirischen Bestimmung der Regressionskoeffizienten üblich ist. Die mit einer Messung für diese Größen gewonnenen Schätzwerte werden im weiteren meist als Messwerte bezeichnet.

3.1 Messung der Übertragungsfunktionen und der deterministischen Störung

Um die Schätzwerte mit Hilfe einer Messung bestimmen zu können, erregen wir das reale System in L Einzelmessungen mit einer Stichprobe vom Umfang L des Zufallsvektors $\vec{v}$ wie dies in Kapitel [1]:3 beschrieben ist. Am Ausgang des Systems messen wir die L Ausgangssignalfolgen $y_\lambda(k)$ mit $\lambda = 1$ (1) L und berechnen daraus durch Fensterung und Fouriertransformation wie in [1] die $M \cdot L$ Spektralwerte $Y_{f,\lambda}(\mu)$.

Nun werden die Schätzwerte für die Regressionskoeffizienten dadurch gewonnen, dass wir anhand dieser Stichproben vom Umfang L die Schätzwerte für die Regressionskoeffizienten in der Art wählen, dass sich der kleinste quadratische Fehler bei der Approximation der Stichprobe des Spektrums des gefensternten Ausgangssignals des realen Systems durch die Stichproben der Spektren der Ausgangssignale der beiden Modellsysteme und die deterministische Modellstörung ergibt. Wir minimieren also nicht die theoretische Varianz des Spektrums $N_f(\mu)$ des gefensternten Approximationsfehlers, sondern das Betragsquadrat der Länge des Vektors, der sich als die Differenz des Stichprobenvektors $\vec{Y}_f(\mu)$ des Spektrums des Ausgangssignals des realen Systems nach Gleichung ([1]:3.12) und der Summe der Stichproben der Spektren der Ausgangssignale der beiden Modellsysteme und des Spektrums der gefensternten deterministischen Störung ergibt. Wir suchen daher eine Ausgleichslösung für das gegenüber Gleichung ([1]:3.3) modifizierte Gleichungssystem

$$\begin{aligned} \frac{1}{M} \cdot \sum_{\tilde{\mu}=0}^{M-1} \sum_{\bar{\mu}=0}^{K_H-1} & \left(\hat{H}(\tilde{\mu}, \tilde{\mu} + \bar{\mu} \cdot \frac{M}{K_H}) \cdot V_\lambda(\tilde{\mu} + \bar{\mu} \cdot \frac{M}{K_H}) + \hat{H}_*(\tilde{\mu}, \tilde{\mu} + \bar{\mu} \cdot \frac{M}{K_H}) \cdot V_\lambda(-\tilde{\mu} - \bar{\mu} \cdot \frac{M}{K_H})^* \right) \cdot \\ & \cdot F((\mu - \tilde{\mu}) \cdot \frac{2\pi}{M}) + \hat{U}_f(\mu) = Y_{f,\lambda}(\mu) \\ \forall \quad \mu = 0 \text{ (1) } M-1 \quad \text{und} \quad \lambda = 1 \text{ (1) } L. \end{aligned} \quad (3.1)$$

Wenn man eine Fensterfunktion verwendet, die der Bedingung ([1]:2.27) genügt, bleibt von der Summe über $\tilde{\mu}$ nur der Summand mit $\tilde{\mu} = \mu$ übrig, und es entstehen in Matrixschreibweise die M Gleichungssysteme

$$\hat{\vec{H}}(\mu) \cdot \tilde{\vec{V}}(\mu) + \hat{U}_f(\mu) \cdot \vec{1} = \vec{Y}_f(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1, \quad (3.2)$$

bei denen auf beiden Seiten jeweils ein $1 \times L$ Zeilenvektor steht. Diese M Gleichungssysteme haben paarweise verschiedene Sätze von Unbekannten. Es tritt jeweils ein Spektralwert $\hat{U}_f(\mu)$ der gefensternten deterministischen Störung, sowie ein Satz von Schätzwerten für die

beiden bifrequenten Übertragungsfunktionen, die zu dem $1 \times 2 \cdot K_H$ Zeilenvektor

$$\hat{\vec{H}}(\mu) = \begin{bmatrix} \hat{H}(\mu, \mu) \\ \hat{H}(\mu, \mu + \frac{M}{K_H}) \\ \vdots \\ \hat{H}(\mu, \mu + (K_H - 1) \cdot \frac{M}{K_H}) \\ \hat{H}_*(\mu, \mu) \\ \hat{H}_*(\mu, \mu + \frac{M}{K_H}) \\ \vdots \\ \hat{H}_*(\mu, \mu + (K_H - 1) \cdot \frac{M}{K_H}) \end{bmatrix}^T \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (3.3)$$

zusammengefasst sind, auf. Der Spektralwert $\hat{U}_f(\mu)$ tritt in allen L Gleichungen (3.1) für eine diskrete Frequenz μ mit dem Koeffizienten Eins auf. Daher tritt in Gleichung (3.2) der Zeilenvektor $\vec{1}$ auf, bei dem alle L Elemente eins sind. Die gesuchten Schätzwerte für die beiden bifrequenten Übertragungsfunktionen werden in den L Gleichungen (3.1) für eine diskrete Frequenz μ mit den unterschiedlichen Spektralwerten der Musterfolgen der Erregung gewichtet. In der $2 \cdot K_H \times L$ Matrix $\tilde{\underline{V}}(\mu)$ sind diese zusammengefasst:

$$\tilde{\underline{V}}(\mu) = \begin{bmatrix} \vec{V}(\mu) \\ \vec{V}(\mu + \frac{M}{K_H}) \\ \vdots \\ \vec{V}(\mu + (K_H - 1) \cdot \frac{M}{K_H}) \\ \vec{V}(-\mu)^* \\ \vec{V}(-\mu - \frac{M}{K_H})^* \\ \vdots \\ \vec{V}(-\mu - (K_H - 1) \cdot \frac{M}{K_H})^* \end{bmatrix} \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (3.4)$$

Dabei sind die Zeilen dieser Matrix die Stichprobenvektoren der Spektralwerte der Erregung, die nach Gleichung ([1]:3.13) definiert sind. Die Matrix selbst ist somit eine konkrete Stichprobe vom Umfang L des nach Gleichung (2.15) definierten Zufallsvektors $\tilde{\vec{V}}(\mu)$.

Die Berechnung der Ausgleichslösung jedes dieser Gleichungssysteme (3.2) führen wir in zwei Schritten durch. Zunächst wird jedes Gleichungssystem jeweils nach dem Anteil aufgelöst, der den gesuchten Messwert des Spektrums der deterministischen Störung enthält, indem man $\hat{\vec{H}}(\mu) \cdot \tilde{\underline{V}}(\mu)$ auf beiden Seiten subtrahiert. Die Ausgleichslösung für die Messwerte $\hat{U}_f(\mu)$ erhalten wir, indem wir zunächst jedes Gleichungssystem mit dem konjugiert transponierten des Vektors von rechts multiplizieren, der mit der gesuchten Größe $\hat{U}_f(\mu)$ von rechts multipliziert wird, und anschließend beide Seiten durch das Quadrat der euklidischen Norm dieses Vektors dividieren. In unserem Fall handelt es sich dabei um den

Einservektor $\vec{1}$, dessen euklidische Norm $\sqrt{L}$ ist.

$$\hat{U}_f(\mu) = \frac{1}{L} \cdot \left(\vec{Y}_f(\mu) - \hat{\vec{H}}(\mu) \cdot \tilde{\underline{V}}(\mu) \right) \cdot \vec{1}^H \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (3.5)$$

Dabei ist

$$\frac{1}{L} \cdot \vec{Y}_f(\mu) \cdot \vec{1}^H = \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (3.6)$$

der empirische Mittelwert der Zufallsgröße $\mathbf{Y}_f(\mu)$. Man kann diesen berechnen, indem man zu einem Akkumulator, den man zu null initialisiert, nach und nach bei jeder der L Einzelmessungen den Wert $Y_{f,\lambda}(\mu)$ addiert, den man durch Fensterung und anschließende DFT des bei der Einzelmessung λ am Systemausgang gemessenen Signals $y_\lambda(k)$ berechnet. Analog — allerdings ohne Fensterung — lässt sich auch der empirische Mittelwert

$$\frac{1}{L} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H = \frac{1}{L} \cdot \sum_{\lambda=1}^L \tilde{V}_\lambda(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (3.7)$$

des Zufallsvektors $\tilde{\underline{V}}(\mu)$ durch Akkumulation der bei den Einzelmessungen verwendeten Eingangsspektren berechnen. Setzt man nun Gleichung (3.5) in die Gleichungssysteme (3.2) ein, so erhält man die Messwerte für die beiden Übertragungsfunktionen:

$$\begin{aligned} \hat{\vec{H}}(\mu) &= \vec{Y}_f(\mu) \cdot \left(\underline{E} - \frac{1}{L} \cdot \vec{1}^H \cdot \vec{1} \right) \cdot \tilde{\underline{V}}(\mu)^H \cdot \left(\tilde{\underline{V}}(\mu) \cdot \left(\underline{E} - \frac{1}{L} \cdot \vec{1}^H \cdot \vec{1} \right) \cdot \tilde{\underline{V}}(\mu)^H \right)^{-1} = \\ &= \frac{1}{L-1} \cdot \vec{Y}_f(\mu) \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} = \hat{\vec{C}}_{\mathbf{Y}_f(\mu), \tilde{\underline{V}}(\mu)} \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \\ &\quad \forall \quad \mu = 0 \text{ (1) } M-1. \end{aligned} \quad (3.8)$$

Die L Zeilenvektoren der Matrix

$$\underline{1}_\perp = \underline{E} - \frac{1}{L} \cdot \vec{1}^H \cdot \vec{1} \quad (3.9)$$

spannen den $L-1$ -dimensionalen Nullraum des Einservektors $\vec{1}$ in der Art auf, dass das Produkt eines beliebigen Vektors $\vec{x}$ mit dieser Matrix den auf den Nullraum projizierten Vektor liefert. Der Anteil des Vektors $\vec{x}$ in Richtung des Einservektors $= \vec{x} \cdot \vec{1}^H \cdot \vec{1} / L$ wird nämlich von dem Vektor selbst abgezogen. Da das Produkt $\vec{1} \cdot \underline{1}_\perp$ den Nullvektor ergibt, ist der Einservektor ein Eigenvektor des Eigenwertes Null. Da jede unitäre Basis — $L-1$ orthonormale Vektoren der Dimension $1 \times L$ — des Nullraums des Einservektors auf sich selbst abgebildet wird, sind die Basisvektoren Eigenvektoren zum $L-1$ -fachen Eigenwert Eins. Die Spur dieser Matrix, die gleich der Summe ihrer Eigenwerte ist, ist daher gleich $L-1$. Diese Matrix $\underline{1}_\perp$ ist abgesehen davon, dass sie hermitesch ist, auch noch idempotent. Sie kann also beliebig oft mit sich selbst multipliziert werden ohne sich dadurch zu ändern.

Das Produkt

$$\begin{aligned}
\hat{\underline{C}}_{\tilde{\underline{V}}(\mu_1), \tilde{\underline{V}}(\mu_2)} &= \frac{1}{L-1} \cdot \tilde{\underline{V}}(\mu_1) \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu_2)^H = \\
&= \frac{1}{L-1} \cdot \tilde{\underline{V}}(\mu_1) \cdot \left(\underline{E} - \frac{1}{L} \cdot \tilde{\underline{1}}^H \cdot \tilde{\underline{1}} \right) \cdot \tilde{\underline{V}}(\mu_2)^H = \\
&= \frac{1}{L-1} \cdot \left(\tilde{\underline{V}}(\mu_1) \cdot \tilde{\underline{V}}(\mu_2)^H - \frac{1}{L} \cdot \tilde{\underline{V}}(\mu_1) \cdot \tilde{\underline{1}}^H \cdot \tilde{\underline{1}} \cdot \tilde{\underline{V}}(\mu_2)^H \right) = \\
&= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L \tilde{\underline{V}}_\lambda(\mu_1) \cdot \tilde{\underline{V}}_\lambda(\mu_2)^H - \frac{1}{L} \cdot \sum_{\lambda=1}^L \tilde{\underline{V}}_\lambda(\mu_1) \cdot \sum_{\lambda=1}^L \tilde{\underline{V}}_\lambda(\mu_2)^H \right) \\
&\quad \forall \quad \mu_1 = 0 \text{ (1) } M-1 \quad \text{und} \quad \mu_2 = 0 \text{ (1) } M-1,
\end{aligned} \tag{3.10}$$

das mit $\mu_1 = \mu_2 = \mu$ in Gleichung (3.8) auftritt, ist die empirische Kovarianzmatrix der beiden Zufallsvektoren $\tilde{\underline{V}}(\mu_1)$ und $\tilde{\underline{V}}(\mu_2)$. Man beachte, dass bei der Berechnung der empirischen Kovarianzmatrix das Konjugieren und Transponieren des zweiten beteiligten Stichprobenzeilenvektors in dieser Definition bereits mit eingeschlossen ist. Dies erfolgt konform mit der Definition der Kovarianzmatrix zweier komplexer Zufallsvektoren, bei der ebenfalls der zweite der daran beteiligten Zufallsvektoren konjugiert wird (siehe Liste der Formelzeichen in [1]). Wie die letzte Umformung zeigt, lässt sich auch die empirische Kovarianzmatrix berechnen, ohne dass dazu die Spektralwerte aller Einzelmessungen abgespeichert werden müssen. Um die erste Summe zu berechnen, addiert man zu einem Akkumulatorfeld, das man zu null initialisiert, nach und nach bei jeder der L Einzelmessungen das dyadische Vektorprodukt $\tilde{\underline{V}}_\lambda(\mu_1) \cdot \tilde{\underline{V}}_\lambda(\mu_2)^H$. Die beiden anderen Summen wurden bereits bei der Berechnung der empirischen Mittelwerte mit Gleichung (3.7) analog berechnet. Nach Gleichung (3.8) benötigt man zur Berechnung der Messwerte der beiden Übertragungsfunktionen auch noch die Kreuzkovarianzvektoren

$$\begin{aligned}
\hat{\underline{C}}_{\underline{Y}_f(\mu_1), \tilde{\underline{V}}(\mu_2)} &= \frac{1}{L-1} \cdot \underline{Y}_f(\mu_1) \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu_2)^H = \\
&= \frac{1}{L-1} \cdot \underline{Y}_f(\mu_1) \cdot \left(\underline{E} - \frac{1}{L} \cdot \tilde{\underline{1}}^H \cdot \tilde{\underline{1}} \right) \cdot \tilde{\underline{V}}(\mu_2)^H = \\
&= \frac{1}{L-1} \cdot \left(\underline{Y}_f(\mu_1) \cdot \tilde{\underline{V}}(\mu_2)^H - \frac{1}{L} \cdot \underline{Y}_f(\mu_1) \cdot \tilde{\underline{1}}^H \cdot \tilde{\underline{1}} \cdot \tilde{\underline{V}}(\mu_2)^H \right) = \\
&= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot \tilde{\underline{V}}_\lambda(\mu_2)^H - \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L \tilde{\underline{V}}_\lambda(\mu_2)^H \right) \\
&\quad \forall \quad \mu_1 = 0 \text{ (1) } M-1 \quad \text{und} \quad \mu_2 = 0 \text{ (1) } M-1,
\end{aligned} \tag{3.11}$$

der Spektralwerte der Signale am Ein- und Ausgang des Systems für $\mu_1 = \mu_2 = \mu$, die sich ebenfalls durch Akkumulation aus den bei den Einzelmessungen λ gemessenen Signalen berechnen lassen.

Da sich die Lösungen (3.8) für die beiden bifrequenten Übertragungsfunktionen nur dann numerisch gut berechnen lassen, wenn die empirischen Kovarianzmatrizen $\hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}$ gut konditioniert sind, muss man zur Messung der beiden Übertragungsfunktionen einen Zufallsvektor $\tilde{\underline{V}}$ verwenden, bei dem die nach Gleichung (2.23) definierten theoretischen Kovarianzmatrizen gut konditioniert sind, so dass auch die empirischen Kovarianzmatrizen nach Gleichung (3.10) mit hoher Wahrscheinlichkeit gut konditioniert sind (siehe Anhang A.2). Auch wenn der Fall einer singulären empirischen Kovarianzmatrix dann so extrem unwahrscheinlich¹ wird, dass er für die praktische Anwendung keinerlei Bedeutung hat, ist es für die theoretische Berechnung der Erwartungswerte und der Varianzen der Messwerte notwendig, die Fälle, bei denen die Stichprobe $\tilde{\underline{V}}(\mu)$ zu einer singulären empirischen Kovarianzmatrix führt, in einer Weise zu handhaben, dass der Erwartungswert und die Varianz des Messwertes existiert. Dazu transformieren wir die empirische Kovarianzmatrix mit Hilfe der unitären Transformationsmatrix $\underline{U}$ kongruent auf ihre Diagonalform:

$$\underline{U} \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)} \cdot \underline{U}^H = \frac{1}{L-1} \cdot \left(\underline{U} \cdot \tilde{\underline{V}}(\mu) \cdot \underline{1}_\perp \right) \cdot \left(\underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \underline{U}^H \right) = \underline{D}. \quad (3.12)$$

Die Diagonalmatrix $\underline{D}$ enthält nur reelle, nichtnegative Diagonalelemente, wovon im singulären Fall ein oder mehrere null sind. Alle Diagonalelemente, die unterhalb einer beliebigen kleinen positiven Schranke liegen, werden auf diese begrenzt. Um wirklich nur die singulären Fälle zu modifizieren, sollte die Schranke wenigstens so klein sein, dass wirklich nur diejenigen Diagonalelemente begrenzt werden, die exakt null sind. Wir erhalten so die Diagonalmatrix $\underline{D}_{\text{Limit}}$. Diese lässt sich dann invertieren. Mit der oben verwendeten unitären kongruenten Transformation erhalten wir daraus eine Matrix $\underline{U}^H \cdot \underline{D}_{\text{Limit}}^{-1} \cdot \underline{U}$, die auch im singulären Fall existiert, und die im nichtsingulären Fall gleich der inversen der empirischen Kovarianzmatrix ist. Mit dieser Matrix berechnen wir die Messwerte $\hat{\underline{H}}(\mu)$ nach Gleichung (3.8).

$$\hat{\underline{H}}(\mu) = \frac{1}{L-1} \cdot \tilde{\underline{Y}}_f(\mu) \cdot \left(\underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \underline{U}^H \right) \cdot \underline{D}_{\text{Limit}}^{-1} \cdot \underline{U} \quad (3.13)$$

Im singulären Fall weist die Matrix $\underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \underline{U}^H$ in den Spalten Nullvektoren auf, die den Diagonalelementen Null entsprechen.

Diese Nullvektoren werden durch die Multiplikation mit den inversen Diagonalelementen der limitierten Diagonalmatrix multipliziert, und bleiben daher Nullvektoren. Der Fehler, der durch die Limitierung der Singulärwerte der empirischen Kovarianzmatrix in den Messwerten der beiden Übertragungsfunktionen entsteht, ist begrenzt, da die Elemente des Stichprobenvektors $\tilde{\underline{Y}}_f(\mu)$ begrenzt sind. Dieser Fehler wird bei der Berechnung

¹Für den Falle einer 1×1 Kovarianzmatrix wird dies im Anhang A.3 diskutiert

des Erwartungswertes und der Varianz des Messwertes mit der Auftrittswahrscheinlichkeit des singulären Falls, der zu diesem fehlerhaften Messwert geführt hat, multipliziert. Wenn die theoretische Kovarianzmatrix gut konditioniert ist, werden die Auftrittswahrscheinlichkeiten aller singulären Fälle so klein, dass deren Beitrag zum Erwartungswert und zur Varianz so gering wird, dass er gegenüber der Berechnung unter Ausschluss aller singulären Fälle vernachlässigt werden kann.

Indem man die Messwerte $\hat{\vec{H}}(\mu)$, die man mit Gleichung (3.8) berechnet, in Gleichung (3.5) einsetzt, erhält man die Messwerte $\hat{U}_f(\mu)$ für das Spektrum der gefensternten deterministischen Störung $u(k)$. Dabei wurde der Anteil des mittleren Spektrums $\underline{\tilde{V}}(\mu) \cdot \vec{1}^H/L$ der Erregung mit den Messwerten der Übertragungsfunktionen multipliziert, und dieses Produkt vom mittleren Spektrum $\vec{Y}_f(\mu) \cdot \vec{1}^H/L$ des gefensternten Ausgangssignals abgezogen. Analog dazu erhält man die Messwerte $\hat{u}(k)$ für die deterministische Störung $u(k)$ dadurch, dass man von dem mittleren Ausgangssignal $\vec{y}(k) \cdot \vec{1}^H/L$ die mit den geschätzten Modellsystemen linear verzerrte mittlere Erregung subtrahiert.

$$\hat{u}(k) = \frac{1}{L} \cdot \sum_{\lambda=1}^L y_{\lambda}(k) - x_{\lambda}(k) - x_{*,\lambda}(k) = \frac{1}{L} \cdot \vec{y}(k) \cdot \vec{1}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{\vec{H}}(\mu) \cdot \underline{\tilde{V}}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad k = 0 \ (1) \ F-1. \quad (3.14)$$

Die F dabei auftretenden $1 \times L$ Zeilenvektoren

$$\vec{y}(k) = [y_1(k), \dots, y_{\lambda}(k), \dots, y_L(k)] \quad \forall \quad k = 0 \ (1) \ F-1 \quad (3.15)$$

setzten sich jeweils aus den L Elementen $y_{\lambda}(k)$ zusammen. Jeder dieser Vektoren ist also eine Stichprobe vom Umfang L der Zufallsgröße $\mathbf{y}(k)$ für einen Zeitpunkt k . Die empirischen Mittelwerte $\vec{y}(k) \cdot \vec{1}^H/L$ lassen sich für alle Zeitpunkte k auch hier wieder berechnen, indem man zu einem Akkumulatorvektor, den man zu null initialisiert, nach und nach bei jeder der L Einzelmessungen die Werte $y_{\lambda}(k)$ des bei der Einzelmessung am Systemausgang gemessenen Signals für alle Zeitpunkte k addiert.

3.2 Erwartungstreue der Messwerte der Übertragungsfunktionen und der deterministischen Störung

Die in [1] angestellte Vorüberlegung zur Berechnung der Erwartungswerte der Messwerte der Übertragungsfunktion ist hier nun zu modifizieren, um das erweiterte Systemmodell mit den beiden periodisch zeitvarianten Modellsystemen und der deterministischen Störung behandeln zu können. Nun ist es Gleichung (2.26) die zeigt, wie sich die M

Zufallsgrößen $\mathbf{N}_f(\mu)$ des Spektrums des gefensterten Approximationsfehlers aus den M Zufallsgrößen $\mathbf{V}(\mu)$ am Eingang und den M Zufallsgrößen $\mathbf{Y}_f(\mu)$ am Ausgang des realen Systems ergeben. Im Gegensatz zu [1] treten nun bei jeder diskreten Frequenz μ mehrere zufällige Spektralwerte der Erregung auf, die jeweils in dem Vektor $\tilde{\mathbf{V}}(\mu)$ nach Gleichung (2.15) zusammengefasst sind. Die weiterhin nach Gleichung ([1]:3.17) definierten M konkreten Stichprobenvektoren $\tilde{\mathbf{N}}_f(\mu)$ vom Umfang L der M Zufallsgrößen $\mathbf{N}_f(\mu)$ erhalten wir, wenn in Gleichung (2.26) statt der Zufallsvektoren $\tilde{\mathbf{V}}(\mu)$ und der Zufallsgrößen $\mathbf{Y}_f(\mu)$ deren konkrete Realisierungen $\tilde{\mathbf{V}}(\mu)$ und $\tilde{\mathbf{Y}}_f(\mu)$ einsetzen.

$$\tilde{\mathbf{N}}_f(\mu) = \tilde{\mathbf{Y}}_f(\mu) - \tilde{\mathbf{H}}(\mu) \cdot \tilde{\mathbf{V}}(\mu) - U_f(\mu) \cdot \tilde{\mathbf{1}} \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (3.16)$$

Hier treten mit $U_f(\mu)$ und mit den Elementen der Vektoren $\tilde{\mathbf{H}}(\mu)$ die unbekannten Optimallösungen der theoretischen Regression auf, die durch die Messung abgeschätzt werden sollen. Die letzte Gleichung lässt sich nach $\tilde{\mathbf{Y}}_f(\mu)$ auflösen und in die Ausgleichslösung (3.8) einsetzen und wir erhalten die Messwerte der Übertragungsfunktionen als Funktion der Stichprobe des Spektrums der Erregung, der Stichprobe des Approximationsfehlerpektrums und der theoretisch optimalen Regressionskoeffizienten:

$$\begin{aligned} \hat{\tilde{\mathbf{H}}}(\mu) &= \frac{1}{L-1} \cdot \left(\tilde{\mathbf{H}}(\mu) \cdot \tilde{\mathbf{V}}(\mu) + U_f(\mu) \cdot \tilde{\mathbf{1}} + \tilde{\mathbf{N}}_f(\mu) \right) \cdot \mathbf{1}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} = \quad (3.17) \\ &= \tilde{\mathbf{H}}(\mu) \cdot \underbrace{\frac{1}{L-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \mathbf{1}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1}}_{=E} + \\ &\quad + \frac{1}{L-1} \cdot U_f(\mu) \cdot \underbrace{\tilde{\mathbf{1}} \cdot \mathbf{1}_{\perp}}_{=\vec{0}} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} + \\ &\quad + \frac{1}{L-1} \cdot \tilde{\mathbf{N}}_f(\mu) \cdot \mathbf{1}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} = \\ &= \tilde{\mathbf{H}}(\mu) + \frac{1}{L-1} \cdot \tilde{\mathbf{N}}_f(\mu) \cdot \mathbf{1}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \quad \forall \quad \mu = 0 \text{ (1) } M-1. \end{aligned}$$

Nun wollen wir die Erwartungswerte der Messwerte der beiden Übertragungsfunktionen bestimmen. Dazu betrachten wir diese Messwerte jeweils als eine konkrete Realisierung — also eine Stichprobe vom Umfang Eins — der Zufallsgrößen $\hat{\tilde{\mathbf{H}}}(\mu)$. Diese Zufallsgrößen erhält man, wenn man statt der konkreten Stichprobenmatrizen $\tilde{\mathbf{V}}(\mu)$ und der konkreten Stichprobenvektoren $\tilde{\mathbf{N}}_f(\mu)$ die mathematischen Stichprobenmatrizen $\tilde{\mathbf{V}}(\mu)$ und die mathematischen Stichprobenvektoren $\tilde{\mathbf{N}}_f(\mu)$, die aufgrund der zufälligen Stichprobenentnahme selbst Zufallsvektoren sind, in die letzte Gleichung einsetzt. Wir bilden also den Erwartungswert über alle möglichen Messungen, die sich jeweils aus L Einzelmessungen zusammensetzen.

Die Erwartungswerte der Messwerte der beiden Übertragungsfunktionen lassen sich nur berechnen, wenn man voraussetzt, dass einerseits alle Stichproben voneinander unabhängig gewonnen wurden, und dass andererseits für jede diskrete Frequenz μ der zufällige Spektralwert $\mathbf{N}_f(\mu)$ von dem Zufallsvektor $\tilde{\mathbf{V}}(\mu)$ bei derselben Frequenz unabhängig ist. Die Unabhängigkeit der Zufallsgrößen, die in dem Zufallsvektor $\tilde{\mathbf{V}}(\mu)$ nach Gleichung (2.15) jeweils zusammengefasst worden sind, muss jedoch *nicht* gefordert werden. Wenn wir nun den Erwartungswert des Produkts einer beliebigen Funktion der Zufallsmatrix $\tilde{\mathbf{V}}(\mu)$ und einer anderen beliebigen Funktion des Zufallsvektors $\tilde{\mathbf{N}}_f(\mu)$ bilden, berechnet sich dieser als das Produkt der beiden Erwartungswerte der einzelnen zufälligen Faktoren. Dies gilt auch für nichtlineare Funktionen, wie z. B. für die Inverse der empirischen Kovarianzmatrix.

Mit Gleichung (3.17) berechnen wir nun die gesuchten Erwartungswerte:

$$\begin{aligned} \mathbb{E}\{\hat{\mathbf{H}}(\mu)\} &= \vec{H}(\mu) + \frac{1}{L-1} \cdot \mathbb{E}\left\{\tilde{\mathbf{N}}_f(\mu) \cdot \mathbf{1}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1}\right\} = \\ &= \vec{H}(\mu) + \frac{1}{L-1} \cdot \underbrace{\mathbb{E}\{\mathbf{N}_f(\mu)\}}_{=0} \cdot \underbrace{\mathbb{E}\left\{\mathbf{1} \cdot \mathbf{1}_{\perp}\right\}}_{=\vec{0}} \cdot \tilde{\mathbf{V}}(\mu)^H \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} = \vec{H}(\mu) \end{aligned} \quad (3.18)$$

Bis hierher hatten wir bei der Berechnung des Erwartungswertes den Fall, dass die empirische Kovarianzmatrix singulär ist, nicht berücksichtigt. Wenn wir diesen Fall in der Art behandeln, indem wir die Singulärwerte der Kovarianzmatrix nach unten limitieren, wie dies im letzten Unterkapitel beschrieben ist, hat dies zwei Auswirkungen auf Erwartungswertberechnung. Zum einen existiert der in der letzten Gleichung angegebene von der Erregung abhängige Erwartungswert nur dann, wenn man die Kovarianzmatrix entsprechend modifiziert, so dass sie sich invertieren lässt. Zum anderen ergibt sich im singulären Fall der Messwertvektor $\hat{\mathbf{H}}(\mu) = \vec{0}$, dessen Erwartungswert ebenfalls der Nullvektor ist. Damit berechnet sich der Erwartungswert zu

$$\begin{aligned} \mathbb{E}\{\hat{\mathbf{H}}(\mu)\} &= \mathbb{E}\left\{\hat{\mathbf{H}}(\mu) \mid \det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)})=0\right\} \cdot P\left(\det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)})=0\right) + \\ &+ \mathbb{E}\left\{\hat{\mathbf{H}}(\mu) \mid \det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}) \neq 0\right\} \cdot P\left(\det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}) \neq 0\right) = \\ &= \vec{0} \cdot P\left(\det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)})=0\right) + \vec{H}(\mu) \cdot P\left(\det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}) \neq 0\right) = \\ &= \vec{H}(\mu) \cdot P\left(\det(\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}) \neq 0\right) \end{aligned} \quad (3.19)$$

In der Praxis wird man zur Erregung des Systems immer Zufallsprozesse verwenden, bei denen die Wahrscheinlichkeit, eine singuläre empirische Kovarianzmatrix zu erhalten, so

extrem klein ist, dass die Wahrscheinlichkeit des gegenteiligen Falles als eins anzusehen ist, und somit sind die Messwerte der beiden bifrequenten Übertragungsfunktionen praktisch erwartungstreu.

Bevor wir mit der Berechnung der Erwartungswerte $\hat{U}_f(\mu)$ der Messwerte des Spektrums der gefensterten deterministischen Störung beginnen, wollen wir zunächst diese nach Gleichung (3.5) definierten Messwerte anders darstellen, indem wir zuerst die Messwerte $\hat{\tilde{H}}(\mu)$ der beiden Übertragungsfunktionen nach Gleichung (3.8) einsetzen. Dann lösen wir Gleichung (3.16) nach den M konkreten Stichprobenvektoren $\tilde{Y}_f(\mu)$ der Spektralwerte des gefensterten Signals am Systemausgang auf und setzen diese ebenfalls ein. Wir erhalten so:

$$\begin{aligned}\hat{U}_f(\mu) &= \frac{1}{L} \cdot \tilde{Y}_f(\mu) \cdot \left(\underline{E} - \frac{\underline{1}_\perp}{L-1} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \right) \cdot \vec{1}^H = \\ &= \frac{1}{L} \cdot \left(\tilde{H}(\mu) \cdot \tilde{\underline{V}}(\mu) + U_f(\mu) \cdot \vec{1} + \tilde{N}_f(\mu) \right) \cdot \left(\underline{E} - \frac{\underline{1}_\perp}{L-1} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \right) \cdot \vec{1}^H = \\ &= U_f(\mu) + \frac{1}{L} \cdot \tilde{N}_f(\mu) \cdot \left(\underline{E} - \frac{\underline{1}_\perp}{L-1} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \right) \cdot \vec{1}^H \\ &\quad \forall \quad \mu = 0 \text{ (1) } M-1.\end{aligned}\tag{3.20}$$

Nun können wir die Erwartungswerte der Messwerte des Spektrums der deterministischen Störung berechnen:

$$\begin{aligned}\mathbb{E}\{\hat{U}_f(\mu)\} &= U_f(\mu) + \frac{1}{L} \cdot \mathbb{E}\left\{ \tilde{N}_f(\mu) \cdot \left(\underline{E} - \frac{\underline{1}_\perp}{L-1} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \right) \cdot \vec{1}^H \right\} = \\ &= U_f(\mu) + \frac{1}{L} \cdot \underbrace{\mathbb{E}\{\tilde{N}_f(\mu)\}}_{=0} \cdot \mathbb{E}\left\{ \vec{1} \cdot \vec{1}^H - \frac{1}{L-1} \cdot \underbrace{\vec{1} \cdot \underline{1}_\perp}_{=\vec{0}} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H \right\} = \\ &= U_f(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1.\end{aligned}\tag{3.21}$$

Auch diese Messwerte sind also erwartungstreu.

Gleichung (2.17) zeigt, wie sich die F Zufallsgrößen $\mathbf{n}(k)$ des gefensterten Approximationsfehlers aus den M Zufallsgrößen $\mathbf{V}(\mu)$ des Spektrums am Eingang des realen Systems und den F Zufallsgrößen $\mathbf{y}(k)$ am Ausgang des realen Systems ergeben. Die F konkreten Stichprobenvektoren $\vec{n}(k)$ vom Umfang L der F Zufallsgrößen $\mathbf{n}(k)$ erhalten wir, wenn in Gleichung (2.17) statt der Zufallsvektoren $\tilde{\underline{V}}(\mu)$ und der Zufallsgrößen $\mathbf{y}(k)$ deren konkrete Realisierungen $\tilde{\underline{V}}(\mu)$ und $\vec{y}(k)$ nach Gleichung (3.15) einsetzen. Nach $\vec{y}(k)$ aufgelöst ergibt sich:

$$\vec{y}(k) = \vec{n}(k) + \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \tilde{H}(\mu) \cdot \tilde{\underline{V}}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} + u(k) \cdot \vec{1}.\tag{3.22}$$

Zur Berechnung der Erwartungswerte $\hat{\mathbf{u}}(k)$ der Messwerte der gefensterten deterministischen Störung ersetzen wir in Gleichung (3.14) den Vektor $\vec{y}(k)$ mit dieser Gleichung, die Messwerte der beiden Übertragungsfunktionen mit Gleichung (3.8) und anschließend den Stichprobenvektor $\vec{Y}_f(\mu)$ mit Gleichung (3.16). So erhalten wir:

$$\begin{aligned}
\hat{u}(k) &= \\
&= \frac{1}{L} \cdot \vec{y}(k) \cdot \vec{1}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \frac{1}{L-1} \cdot \vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
&= \frac{1}{L} \cdot \left(\vec{n}(k) + \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \vec{H}(\mu) \cdot \tilde{\underline{V}}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} + u(k) \cdot \vec{1} \right) \cdot \vec{1}^H - \\
&\quad - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \frac{1}{L-1} \cdot \left(\vec{H}(\mu) \cdot \tilde{\underline{V}}(\mu) + U_f(\mu) \cdot \vec{1} + \vec{N}_f(\mu) \right) \cdot \\
&\quad \cdot \underline{1}_{\perp} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
&= u(k) + \frac{1}{L} \cdot \vec{n}(k) \cdot \vec{1}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \frac{\vec{N}_f(\mu) \cdot \underline{1}_{\perp} \cdot \tilde{\underline{V}}(\mu)^H}{L-1} \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \\
&\quad \forall \quad \mu = 0 \text{ (1) } M-1. \tag{3.23}
\end{aligned}$$

Nun können wir wieder die Erwartungswerte der Messwerte der deterministischen Störung berechnen:

$$\begin{aligned}
\mathbb{E}\{\hat{\mathbf{u}}(k)\} &= u(k) + \underbrace{\mathbb{E}\{\mathbf{n}(k)\}}_{=0} - \\
&- \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \frac{1}{L-1} \cdot \underbrace{\mathbb{E}\{\mathbf{N}_f(\mu)\}}_{=0} \cdot \underbrace{\left\{ \vec{1} \cdot \underline{1}_{\perp} \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \vec{1}^H \right\}}_{=\vec{0}} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\
&= u(k) \quad \forall \quad 0 \leq k < F. \tag{3.24}
\end{aligned}$$

Auch diese Messwerte sind erwartungstreu. Der Leser möge sich selbst davon überzeugen, dass die Auswirkungen der Behandlung des singulären Falls auf die Erwartungswerte der deterministischen Störung und ihrer Spektralwerte in der Praxis bedeutungslos sind.

3.3 Messung der beiden Leistungsdichtespektren

Da man auch hier durch eine Messung keine Stichprobe $\vec{N}_f(\mu)$ des wahren Spektrums des gefensterten Approximationsfehlerprozesses $\mathbf{n}(k)$ gewinnen kann, weil man weder die

wahren Parameter der beiden Modellsysteme, noch das wahre Spektrum der gefenster-ten deterministischen Störung kennt, benötigt man wieder aus den gemessenen Spektren abgeleitete Zufallsgrößen, deren Erwartungswerte gleich den zu messenden Größen sind. Dazu verwenden wir wieder den nach Gleichung ([1]:3.12) definierten gemessenen Stichprobenvektor $\vec{Y}_f(\mu)$, den wir wieder mit einer Matrix $\underline{V}_\perp(\mu)$ wie in Gleichung ([1]:3.25) linear abbilden, um so den Vektor $\hat{\vec{N}}_f(\mu)$ zu erhalten. Mit dem Stichprobenvektor $\vec{N}_f(\mu)$ nach Gleichung (3.16) erhalten wir:

$$\begin{aligned} \hat{\vec{N}}_f(\mu) &= \vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) = \left(\vec{Y}_f(\mu) - \vec{H}(\mu) \cdot \vec{V}(\mu) - U_f(\mu) \cdot \vec{1} \right) \cdot \underline{V}_\perp(\mu) = \vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \\ &\quad \forall \quad \mu = 0 \text{ (1) } M-1. \end{aligned} \quad (3.25)$$

Damit diese Gleichung gilt, muss die Matrix $\underline{V}_\perp(\mu)$ nun außer dem Stichprobenvektor $\vec{V}(\mu)$, den wir in [1] verwendet haben, auch noch alle weiteren Zeilenvektoren der Stichprobenmatrix $\tilde{\underline{V}}(\mu)$, die wir bei der Berechnung der Messwerte der Übertragungsfunktion und der deterministischen Störung verwendet haben, sowie den Einservektor $\vec{1}$ als Eigenvektoren zum Eigenwert Null aufweisen.

$$\tilde{\underline{V}}(\mu) \cdot \underline{V}_\perp(\mu) = \underline{0} \quad \wedge \quad \vec{1} \cdot \underline{V}_\perp(\mu) = \vec{0} \quad (3.26)$$

$\underline{0}$ ist hier eine $(2 \cdot K_H) \times L$ Matrix, deren Elemente ebenso null sind, wie die L Elemente des Zeilenvektors $\vec{0}$. Der Vektor $\hat{\vec{N}}_f(\mu)$ ist dann orthogonal zu den Zeilenvektoren der Stichprobenmatrix $\tilde{\underline{V}}(\mu)$ und zum Einservektor $\vec{1}$ und daher ist er unabhängig von den beiden bifrequenten Übertragungsfunktionen und dem Spektralwert $U_f(\mu)$ der gefenster-ten Störung. Für jede Frequenz μ kann man eine Matrix $\underline{V}_\perp(\mu)$, die die Bedingung (3.26) auf jeden Fall erfüllt, analog zu Gleichung ([1]:3.45), nun aber unter Berücksichtigung der Modellierung des Spektrums der deterministischen Störung, konstruieren. Dazu benö-tigen wir eine Matrix $\check{\underline{V}}(\mu)$, deren Zeilenvektoren voneinander unabhängig sind, so dass sie vollen Rang hat. Die Zeilenvektoren dieser Matrix müssen so gewählt werden, dass alle Zeilenvektoren der Matrix $\tilde{\underline{V}}(\mu)$ in dem Raum liegen, der durch die Zeilenvektoren der Matrix $\check{\underline{V}}(\mu)$ aufgespannt wird, während der Einservektor $\vec{1}$ nicht in diesem Raum liegen darf. Die Elemente der so konstruierten Matrix $\check{\underline{V}}(\mu)$ können zufällig sein, müssen aber von den zufälligen Spektralwerten des gefensterten Approximationsfehlers unabhän-gig ein. Für jede Frequenz μ kann man mit dieser Matrix $\check{\underline{V}}(\mu)$ wie in Gleichung (3.10) eine hermitesche, empirische Kovarianzmatrix

$$\begin{aligned} \hat{\underline{C}}_{\check{\underline{V}}(\mu), \check{\underline{V}}(\mu)} &= \hat{\underline{C}}_{\check{\underline{V}}(\mu), \check{\underline{V}}(\mu)}^H = \frac{1}{L-1} \cdot \check{\underline{V}}(\mu) \cdot \underline{1}_\perp \cdot \check{\underline{V}}(\mu)^H = \\ &= \frac{1}{L-1} \cdot \check{\underline{V}}(\mu) \cdot \left(\underline{E} - \frac{1}{L} \cdot \vec{1}^H \cdot \vec{1} \right) \cdot \check{\underline{V}}(\mu)^H = \\ &= \frac{1}{L-1} \cdot \left(\check{\underline{V}}(\mu) \cdot \check{\underline{V}}(\mu)^H - \frac{1}{L} \cdot \check{\underline{V}}(\mu) \cdot \vec{1}^H \cdot \vec{1} \cdot \check{\underline{V}}(\mu)^H \right) = \end{aligned} \quad (3.27)$$

$$= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L \check{\check{V}}_{\lambda}(\mu) \cdot \check{\check{V}}_{\lambda}(\mu)^H - \frac{1}{L} \cdot \sum_{\lambda=1}^L \check{\check{V}}_{\lambda}(\mu) \cdot \sum_{\lambda=1}^L \check{\check{V}}_{\lambda}(\mu)^H \right)$$

$$\forall \quad \mu = 0 \ (1) \ M-1$$

berechnen, die auf Grund der Konstruktion der Matrix $\check{\check{V}}(\mu)$ regulär ist. Mit dieser empirischen Kovarianzmatrix und der Matrix $\check{\check{V}}(\mu)$ wird dann die idempotente und hermitesche Matrix

$$\underline{V}_{\perp}(\mu) = \underline{V}_{\perp}(\mu)^n = \underline{V}_{\perp}(\mu)^H = \underline{1}_{\perp} - \frac{1}{L-1} \cdot \underline{1}_{\perp} \cdot \check{\check{V}}(\mu)^H \cdot \hat{\underline{C}}_{\check{\check{V}}(\mu), \check{\check{V}}(\mu)}^{-1} \cdot \check{\check{V}}(\mu) \cdot \underline{1}_{\perp}$$

$$\forall \quad \mu = 0 \ (1) \ M-1 \quad \wedge \quad n \in \mathbb{N}, \quad (3.28)$$

die die Bedingung (3.26) erfüllt, für jede Frequenz μ gebildet. Der Rangdefekt der Matrix $\underline{V}_{\perp}(\mu)$ ist um eins größer als die Zeilenanzahl der Matrix $\check{\check{V}}(\mu)$, da sowohl der Einservektor als auch alle Zeilenvektoren der Matrix $\check{\check{V}}(\mu)$ Eigenvektoren zum Eigenwert Null sind. Da alle dazu orthogonalen Vektoren mit der Matrix $\underline{V}_{\perp}(\mu)$ auf sich selbst abgebildet werden, sind sie Eigenvektoren zum Eigenwert Eins. Da die Matrix $\underline{V}_{\perp}(\mu)$ keine anderen Eigenwerte als Null und Eins aufweist, ist die Spur der Matrix gleich dem Rang der Matrix.

Die einfachste Möglichkeit die Matrix $\check{\check{V}}(\mu)$ zu konstruieren besteht darin, mit dem Einservektor $\vec{1}$ zu beginnen und aus der Matrix $\tilde{\underline{V}}(\mu)$ nach und nach Zeilenvektoren hinzuzufügen, die von allen bisher ausgewählten Zeilenvektoren linear unabhängig sind. Dies wiederholt man, bis sich keine weiteren linear unabhängigen Vektoren mehr in der Matrix $\tilde{\underline{V}}(\mu)$ finden lassen. Da die Matrix $\underline{V}_{\perp}(\mu)$ nicht unbedingt den Nullraum der bis dahin ausgewählten Zeilenvektoren der Matrix $\tilde{\underline{V}}(\mu)$ vollständig aufspannen muss, kann man weitere linear unabhängige Zeilenvektoren hinzufügen. Somit kann man, auch wenn im Systemmodell kein Modellsystem, das mit dem konjugierten Eingangssignal erregt wird, vorgesehen worden ist, zur Konstruktion der Matrix $\underline{V}_{\perp}(\mu)$ die Matrix $\check{\check{V}}(\mu)$ verwenden, die sich bei einem vollständigen Systemmodell ergeben würde. Abschließend wird die Zeile der Matrix $\check{\check{V}}(\mu)$, die den Einservektor $\vec{1}$ enthält, wieder entfernt, und mit der so konstruierten Matrix $\check{\check{V}}(\mu)$ wird die Matrix $\underline{V}_{\perp}(\mu)$ mit den Gleichungen (3.27) und (3.28) berechnet. Im weiteren soll die Matrix $\underline{V}_{\perp}(\mu)$ immer in der Art konstruiert sein, dass der Rangdefekt der Matrix $\underline{V}_{\perp}(\mu)$ auf eine von L unabhängige Konstante begrenzt ist, d. h. die Zahl der Zeilenvektoren der Matrix $\check{\check{V}}(\mu)$ darf mit steigender Mittelungsanzahl L nicht über eine von L unabhängige Konstante steigen. Am besten wählt man die Anzahl der Zeilen der Matrix $\check{\check{V}}(\mu)$ konstant.

Prinzipiell muss die Matrix $\underline{V}_{\perp}(\mu)$ auch hier nicht idempotent und hermitesch sein, sofern sie der Bedingung (3.26) genügt, und ihre Elemente von den zufälligen Spektralwerten des

gefensterten Approximationsfehlers unabhängig sind. Im weiteren werde ich mich aber auf die Verwendung dieser idempotenten und hermiteschen Matrizen beschränken, die einen Rangdefekt aufweisen, der von L unabhängig nach oben begrenzt ist. Dadurch ist die Gültigkeit der im weiteren hergeleiteten Ergebnisse sichergestellt. Bei Verwendung anderer Matrizen wäre dies ggf. im Einzelfall zu überprüfen.

Da $\underline{V}_\perp(\mu)$ die Bedingung (3.26) erfüllt, ist nach Gleichung (3.25) der durch die Abbildung mit der Matrix $\underline{V}_\perp(\mu)$ aus dem des Stichprobenvektor $\vec{Y}_f(\mu)$ entstandene und daher bekannte Stichprobenvektor $\hat{\vec{N}}_f(\mu)$ gleich dem Bildvektor des unbekannten Stichprobenvektors $\vec{N}_f(\mu)$, der bei der Abbildung mit derselben Matrix entsteht. Daher verwenden wir in Analogie zu den Messwerten nach Gleichung ([1]:3.34) und ([1]:3.35) im Falle eines stationären Approximationsfehlerprozesses die Messwerte

$$\begin{aligned}\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &= \frac{\hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H)} = \\ &= \frac{\vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H)} = \\ &= \frac{\vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \cdot \vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H)} \\ \forall \quad \mu &= 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1\end{aligned}\tag{3.29a}$$

und

$$\begin{aligned}\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &= \frac{\hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T)} = \\ &= \frac{\vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \cdot \vec{Y}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T)} = \\ &= \frac{\vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \cdot \vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T)} \\ \forall \quad \mu &= 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1.\end{aligned}\tag{3.29b}$$

Diese Messwerte weisen die gleichen Symmetrien auf wie die entsprechenden theoretischen Größen:

$$\hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu) = \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \tag{3.30a}$$

und

$$\hat{\Psi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu) = \hat{\Psi}_n(-\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}). \tag{3.30b}$$

Man wird diese Messwerte *nicht* dadurch berechnen, dass man zunächst die Vektoren $\vec{Y}_f(\mu)$ und die Matrizen $\underline{V}_\perp(\mu)$ bestimmt, und diese dann — wie in den Gleichungen (3.29) angegeben — multipliziert. Dies würde nämlich einen mit der Mittelungsanzahl L quadratisch anwachsenden Speicher erfordern. Im Anhang A.6 wird eine geeignetere Vorgehensweise zur Berechnung der Messwerte $\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$ und $\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$ angegeben, bei der der Speicherbedarf von L unabhängig ist.

Am sinnvollsten erscheint es mir, die Matrizen $\check{\underline{V}}(\mu)$ für die unterschiedlichen Frequenzen μ folgendermaßen zu konstruieren. Zunächst berechnet man sich aus der Periode K_H der Zeitvarianz des Systems und der Periode K_Φ der Zyklostationarität der Störung deren kleinstes gemeinsames Vielfaches

$$K_S = \text{kgv}(K_H, K_\Phi). \quad (3.31)$$

Jede der beiden Frequenzverschiebungen $\tilde{\mu} \cdot M/K_\Phi$ und $\hat{\mu} \cdot M/K_H$ ist dann eine Vielfaches von M/K_S . Danach greifen wir uns für jede Frequenz μ aus den Zufallsgrößen $\mathbf{V}(\mu + \check{\mu} \cdot M/K_S)$ und $\mathbf{V}(-\mu - \check{\mu} \cdot M/K_S)^*$ mehrere Zufallsgrößen heraus, und fassen diese in dem Zufallsvektor $\check{\mathbf{V}}(\mu)$ zusammen. Die analog zu Gleichung (2.23) definierte theoretische Kovarianzmatrix der herausgegriffenen Zufallsgrößen darf dabei keinen Rangdefekt aufweisen, und muss demselben Rang besitzen, wie die theoretische Kovarianzmatrix aller Zufallsgrößen für alle Werte von $\check{\mu}$ aus denen die Zufallsgrößen herausgegriffen wurden. Es sind also bei jeder Frequenz μ so viele Zufallsgrößen herauszugreifen, dass alle linearen Abhängigkeiten innerhalb aller Zufallsgrößen für alle Werte von $\check{\mu}$ eliminiert sind, und trotzdem alle Freiheitsgrade berücksichtigt sind. Die Anzahl der bei der Frequenz μ herausgegriffenen Zufallsgrößen sei mit $K(\mu)$ bezeichnet. Da bei jeder Frequenz maximal $2 \cdot K_S$ Zufallsgrößen herausgegriffen werden können, gilt die Ungleichung

$$K(\mu) \leq 2 \cdot K_S. \quad (3.32)$$

Die Zeilenvektoren der Stichproben der herausgegriffenen Zufallsgrößen fassen wir zu der Matrix $\check{\underline{V}}(\mu)$ zusammen. Auch wenn die theoretische Kovarianzmatrix der bei einer Frequenz μ herausgegriffenen Zufallsgrößen gut konditioniert ist, kann doch der Rang der Matrix $\check{\underline{V}}(\mu)$ kleiner als deren Zeilendimension, und somit die empirische Kovarianzmatrix singulär sein. Verwendet man Zufallsgrößen $\mathbf{V}(\mu + \check{\mu} \cdot M/K_S)$, bei denen alle theoretischen Kovarianzmatrizen der jeweils herausgegriffenen Zufallsgrößen gut konditioniert sind, so wird dieser Fall für hinreichend große Werte von L mit einer Wahrscheinlichkeit auftreten, die so extrem klein ist, dass der Fall einer singulären empirischen Kovarianzmatrix praktisch bedeutungslos ist. Für die theoretischen Herleitungen wird angenommen, dass bei der Matrix $\check{\underline{V}}(\mu)$ ggf. die linear abhängigen Zeilenvektoren durch linear unabhängige Zufallsvektoren ersetzt werden. Die Konstruktion dieser Zufallsvektoren muss dabei in

einer Weise erfolgen, dass deren Elemente — unter Berücksichtigung der Zufälligkeit der Messung und der Matrizen $\check{\underline{V}}(\mu)$ und der Vektoren $\check{\underline{N}}_f(\mu)$ — von den Zufallsgrößentupeln $[\underline{N}_f(\mu), \underline{N}_f(\mu + \check{\mu} \cdot M/K_\Phi)]^T$ und $[\underline{N}_f(\mu), \underline{N}_f(-\mu - \check{\mu} \cdot M/K_\Phi)]^T$ unabhängig sind, um sicherzustellen, dass die im weiteren durchgeführten Berechnungen der Erwartungswerte der Messwerte auch im singulären Fall ihre Gültigkeit behalten.

Mit den so konstruierten Matrizen $\check{\underline{V}}(\mu)$ berechnen wir uns mit Gleichung (3.28) die Matrizen $\underline{V}_\perp(\mu)$. Diese haben dann immer die Spur $L - 1 - K(\mu)$. Dies zeigt man ganz analog zu der Herleitung ab Gleichung ([1]:3.48). Es sei auch angemerkt, dass der Rangdefekt dieser Matrizen gemäß der Ungleichung (3.32) für hinreichend große Werte von L auf eine von L unabhängige obere Schranke begrenzt ist.

Betrachten wir nun zwei Matrizen $\underline{V}_\perp(\mu)$ für zwei Frequenzen μ , die sich um ein ganzzahliges Vielfaches von M/K_S unterscheiden. Die bei beiden Frequenzen herausgegriffenen Zufallsgrößen entstammen dem gleichen Satz der Zufallsgrößen $\underline{V}(\mu + \check{\mu} \cdot M/K_S)$ und $\underline{V}(-\mu - \check{\mu} \cdot M/K_S)^*$. Da bei beiden Frequenzen μ die herausgegriffenen Zufallsgrößen eine reguläre theoretische Kovarianzmatrix größtmöglichen Ranges (= Dimension) aufweisen, lassen sich die herausgegriffenen Zufallsgrößen bei der einen Frequenz als reguläre Linearkombinationen der herausgegriffenen Zufallsgrößen bei der anderen Frequenz darstellen. Damit lässt sich mit Hilfe der Gleichungen (3.27) und (3.28) zeigen, dass bei beiden Frequenzen die nach dem eben beschriebenen Verfahren berechneten idempotenten Matrizen $\underline{V}_\perp(\mu)$ identisch sind. Da bei der Berechnung der Matrizen $\underline{V}_\perp(\mu)$ neben den Zufallsgrößen $\underline{V}(\mu + \check{\mu} \cdot M/K_S)$ auch die konjugierten Zufallsgrößen $\underline{V}(-\mu - \check{\mu} \cdot M/K_S)^*$ bei der negativen Frequenz berücksichtigt worden sind, erfüllen die Matrizen $\underline{V}_\perp(\mu)$ auch die Bedingung ([1]:3.39). Zusammenfassend gilt daher

$$\underline{V}_\perp(\mu) = \underline{V}_\perp(\mu)^H = \underline{V}_\perp(\mu + \bar{\mu} \cdot \frac{M}{K_S}) = \underline{V}_\perp(-\mu - \bar{\mu} \cdot \frac{M}{K_S})^T = \underline{V}_\perp(-\mu)^T = \underline{V}_\perp(-\mu)^* \quad (3.33a)$$

$$\text{und} \quad \text{spur}(\underline{V}_\perp(\mu)) = L - 1 - K(\mu) \quad (3.33b)$$

Es genügt daher, wenn man die Matrizen $\underline{V}_\perp(\mu)$ für $\mu = 0$ (1) $(M - K_S)/(2 \cdot K_S)$ berechnet, da man die Matrizen aller anderen Frequenzen ggf. durch Transponieren daraus erhält.

Ein Teil der Zufallsgrößen aller Werte von $\check{\mu}$ ist bei jeder Frequenz μ an der theoretischen Kovarianzmatrix beteiligt, die bei der Berechnung der theoretischen Werte der Übertragungsfunktionen $H(\mu, \mu + \hat{\mu} \cdot M/K_H)$ und $H_*(\mu, \mu + \hat{\mu} \cdot M/K_H)$ auftritt. Dort wurde gefordert, dass der Teil der Zufallsgrößen $\underline{V}(\mu + \check{\mu} \cdot M/K_S)$, der dort verwendet wird und der den Vektor $\check{\underline{V}}(\mu)$ bildet, so zu wählen ist, dass deren theoretische Kovarianzmatrix bei der Berechnung der theoretischen Werte der Übertragungsfunktionen gut konditioniert und somit regulär ist. Jede der Zufallsgrößen des Zufallsvektors $\check{\underline{V}}(\mu)$ lässt sich daher als Linearkombination der für $\check{\underline{V}}(\mu)$ herausgegriffenen Zufallsgrößen schreiben, da andernfalls

die theoretische Kovarianzmatrix des Zufallsvektors $\check{\check{\mathbf{V}}}(\mu)$ nicht die maximal mögliche Dimension $K(\mu) \times K(\mu)$ hätte, bei der die theoretische Kovarianzmatrix noch regulär ist. Genau dieselben Linearkombinationen, durch die sich die Zufallsgrößen des Zufallsvektors $\check{\check{\mathbf{V}}}(\mu)$ durch die für $\check{\check{\mathbf{V}}}(\mu)$ herausgegriffenen Zufallsgrößen ausdrücken lassen, gelten für die entsprechenden Stichprobenvektoren dieser Zufallsgrößen. Da alle Stichprobenvektoren der für $\check{\check{\mathbf{V}}}(\mu)$ herausgegriffenen Zufallsgrößen Eigenvektoren von $\underline{V}_\perp(\mu)$ zum Eigenwert Null sind, sind auch die Stichprobenvektoren der Zufallsgrößen des Zufallsvektors $\check{\check{\mathbf{V}}}(\mu)$ Eigenvektoren zum gleichen Eigenwert. Die mit der Stichprobenmatrix $\check{\check{\mathbf{V}}}(\mu)$ konstruierten Matrizen $\underline{V}_\perp(\mu)$ erfüllen daher die Bedingung (3.26) wie gefordert.

Wenn man $\underline{V}_\perp(\mu)$ nach dem oben angegebenen Verfahren konstruiert, so dass die Gleichungen (3.33) erfüllt sind, erhält man die einfacheren Messwerte

$$\begin{aligned} \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &= \frac{\hat{N}_f(\mu) \cdot \hat{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot (L - 1 - K(\mu))} = \\ &= \frac{\vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot (L - 1 - K(\mu))} = \frac{\vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{M \cdot (L - 1 - K(\mu))} = \\ &= \frac{1}{M} \cdot \frac{L-1}{L-1-K(\mu)} \cdot \left(\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} - \hat{C}_{\mathbf{Y}_f(\mu), \check{\check{\mathbf{V}}}(\mu)} \cdot \hat{C}_{\check{\check{\mathbf{V}}}(\mu), \check{\check{\mathbf{V}}}(\mu)}^{-1} \cdot \hat{C}_{\check{\check{\mathbf{V}}}(\mu), \mathbf{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} \right) \\ &\quad \forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1 \end{aligned} \quad (3.34a)$$

und

$$\begin{aligned} \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &= \frac{\hat{N}_f(\mu) \cdot \hat{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot (L - 1 - K(\mu))} = \\ &= \frac{\vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{Y}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot (L - 1 - K(\mu))} = \frac{\vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot (L - 1 - K(\mu))} = \\ &= \frac{1}{M} \cdot \frac{L-1}{L-1-K(\mu)} \cdot \left(\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} - \hat{C}_{\mathbf{Y}_f(\mu), \check{\check{\mathbf{V}}}(\mu)} \cdot \hat{C}_{\check{\check{\mathbf{V}}}(\mu), \check{\check{\mathbf{V}}}(\mu)}^{-1} \cdot \hat{C}_{\check{\check{\mathbf{V}}}(\mu), \mathbf{Y}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} \right) \\ &\quad \forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1, \end{aligned} \quad (3.34b)$$

Dabei wurden die empirischen Kovarianzen

$$\begin{aligned} \hat{C}_{\mathbf{Y}_f(\mu_1), \mathbf{Y}_f(\mu_2)} &= \frac{1}{L-1} \cdot \vec{Y}_f(\mu_1) \cdot \underline{1}_\perp \cdot \vec{Y}_f(\mu_2)^H = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot Y_{f,\lambda}(\mu_2)^* - \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_2)^* \right) \\ &\quad \forall \quad \mu_1 = 0 \ (1) \ M-1 \quad \text{und} \quad \mu_2 = 0 \ (1) \ M-1, \end{aligned} \quad (3.35a)$$

und

$$\begin{aligned}\hat{C}_{\mathbf{Y}_f(\mu_1), \mathbf{Y}_f(\mu_2)^*} &= \frac{1}{L-1} \cdot \vec{Y}_f(\mu_1) \cdot \mathbf{1}_{\perp} \cdot \vec{Y}_f(\mu_2)^T = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot Y_{f,\lambda}(\mu_2) - \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_2) \right) \\ &\quad \forall \quad \mu_1 = 0 \ (1) \ M-1 \quad \text{und} \quad \mu_2 = 0 \ (1) \ M-1,\end{aligned}\tag{3.35b}$$

sowie die Kovarianzvektoren

$$\begin{aligned}\hat{\check{C}}_{\check{\mathbf{V}}(\mu_1), \mathbf{Y}_f(\mu_2)} &= \hat{\check{C}}_{\mathbf{Y}_f(\mu_2), \check{\mathbf{V}}(\mu_1)}^H = \frac{1}{L-1} \cdot \check{\mathbf{V}}(\mu_1) \cdot \mathbf{1}_{\perp} \cdot \vec{Y}_f(\mu_2)^H = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L \check{V}_{\lambda}(\mu_1) \cdot Y_{f,\lambda}(\mu_2)^* - \frac{1}{L} \cdot \sum_{\lambda=1}^L \check{V}_{\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_2)^* \right) \\ &\quad \forall \quad \mu_1 = 0 \ (1) \ M-1 \quad \text{und} \quad \mu_2 = 0 \ (1) \ M-1,\end{aligned}\tag{3.35c}$$

und

$$\begin{aligned}\hat{\check{C}}_{\check{\mathbf{V}}(\mu_1), \mathbf{Y}_f(\mu_2)^*} &= \frac{1}{L-1} \cdot \check{\mathbf{V}}(\mu_1) \cdot \mathbf{1}_{\perp} \cdot \vec{Y}_f(\mu_2)^T = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L \check{V}_{\lambda}(\mu_1) \cdot Y_{f,\lambda}(\mu_2) - \frac{1}{L} \cdot \sum_{\lambda=1}^L \check{V}_{\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_2) \right) \\ &\quad \forall \quad \mu_1 = 0 \ (1) \ M-1 \quad \text{und} \quad \mu_2 = 0 \ (1) \ M-1\end{aligned}\tag{3.35d}$$

verwendet.

Dass die mit den Gleichungen (3.34) berechneten Messwerte zusätzlich die Ungleichungen

$$\left| \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right|^2 \leq \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})\tag{3.36a}$$

$$\text{und} \quad \left| \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right|^2 \leq \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}})\tag{3.36b}$$

erfüllen, die den Ungleichungen (2.43) und (2.53) der zu messenden theoretischen Größen entsprechen, ergibt sich einerseits aus der Cauchy-Schwarzschen Ungleichung angewandt auf die Vektorenpaare $\hat{\vec{N}}_f(\mu)$ und $\hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot M/K_{\Phi})$ bzw. $\hat{\vec{N}}_f(\mu)$ und $\hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot M/K_{\Phi})^*$, und andererseits aus der Tatsache, dass bei allen in diesen Ungleichungen auftretenden Messwerten der gleiche Vorfaktor auftritt. Wie wir noch sehen werden, ermöglicht es uns die Erfüllung der letzten beiden Ungleichungen wieder, Konfidenzgebiete für alle Messwerte anzugeben.

Wenn wir die Messwerte wieder als zufällig betrachten, und wenn der Zufallsvektor $\check{\mathbf{V}}(\mu)$, der aus den Zufallsgrößen $\mathbf{V}(\mu + \check{\mu} \cdot M/K_S)$ und $\mathbf{V}(-\mu - \check{\mu} \cdot M/K_S)^*$ des Spektrums der Erregung gebildet ist, die in die Elemente der dann ebenfalls zufälligen Matrizen $\mathbf{V}_\perp(\mu)$ eingehen, unabhängig ist von dem Zufallsgrößenpaar der Spektralwerte $[\mathbf{N}_f(\mu), \mathbf{N}_f(\mu + \check{\mu} \cdot \frac{M}{K_\Phi})]^T$ bzw. $[\mathbf{N}_f(\mu), \mathbf{N}_f(-\mu - \check{\mu} \cdot \frac{M}{K_\Phi})]^T$ der gefensterten Störung des realen Systems, kann man wieder Erwartungswerte, Varianzen und Kovarianzen für die Messwerte berechnen, und erwartungstreue Schätzwerte für die Messwert(ko)varianzen angeben.

Die Erwartungstreue der Messwerte gemäß der Gleichungen (3.29) und somit auch der Messwerte gemäß der Gleichungen (3.34), die einen Spezialfall der erstgenannten darstellen, zeigt man wie im Fall des stationären Approximationsfehlerprozesses in [1] mit den in Anhang A.5 dargestellten Umformungen. Auch hier gehen nur die Hauptdiagonalelemente der Matrizen der bilinearen Formen additiv in die Erwartungswerte ein, da die Nebendiagonalelemente dieser Matrizen bei der Berechnung der bilinearen Form mit Produkten von Stichprobenelementen der Zufallsgrößen $\mathbf{N}_f(\mu)$ und $\mathbf{N}_f(\mu + \check{\mu} \cdot M/K_\Phi)$ bzw. $\mathbf{N}_f(-\mu - \check{\mu} \cdot M/K_\Phi)$ verknüpft werden, die aus unterschiedlichen und unabhängigen Einzelmessungen stammen, und deren Erwartungswerte ($=0$) daher faktorisiert sind. In dem zu bildenden Erwartungswert der nur von den Spektralwerten der Erregung abhängt, lässt sich daher jeweils die Spur der Matrix, auf die der jeweilige Messwert normiert ist, kürzen. Bei der Berechnung des Erwartungswertes des Messwertes $\hat{\Phi}_n(\mu, \mu + \check{\mu} \cdot M/K_\Phi)$ verbleibt der Erwartungswert

$$\frac{1}{M} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\mu + \check{\mu} \cdot \frac{M}{K_\Phi})^* \right\} - \mathbb{E} \{ \mathbf{N}_f(\mu) \} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu + \check{\mu} \cdot \frac{M}{K_\Phi}) \right\}^* \quad (3.37a)$$

während sich bei dem Messwert $\hat{\Psi}_n(\mu, \mu + \check{\mu} \cdot M/K_\Phi)$ der Erwartungswert

$$\frac{1}{M} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(-\mu - \check{\mu} \cdot \frac{M}{K_\Phi}) \right\} - \mathbb{E} \{ \mathbf{N}_f(\mu) \} \cdot \mathbb{E} \left\{ \mathbf{N}_f(-\mu - \check{\mu} \cdot \frac{M}{K_\Phi}) \right\} \quad (3.37b)$$

ergibt. Da die aus dem Approximationsfehlerprozess durch Fensterung und DFT gewonnenen Zufallsgrößen $\mathbf{N}_f(\mu)$ mittelwertfrei sind, stimmen diese Erwartungswerte mit den gewünschten Werten nach Gleichung (2.39) bzw. (2.50) überein.

3.4 Varianzen und Kovarianzen der Messwerte

Mit dem Einheitszeilenvektor

$$\vec{E}_n = [0, \dots, 1, \dots, 0], \quad (3.38)$$

dessen n -tes Element eins ist, während alle anderen Elemente null sind, und der hier $2 \cdot K_H$ Elemente enthält, kann man aus dem Zeilenvektor $\hat{\vec{H}}(\mu)$, der die Messwerte der beiden

bifrequenten Übertragungsfunktionen enthält, den n -ten Messwert herausgreifen, indem man diesen Vektor mit dem transponierten Einheitsvektor multipliziert:

$$\hat{H}_n(\mu) = \hat{\vec{H}}(\mu) \cdot \vec{E}_n^H. \quad (3.39)$$

Wieder kann man den n -ten Messwert $\hat{H}_n(\mu)$ als eine Zufallsgröße betrachten, die aus einer mathematische Stichprobe der Zufallsprozesse am Systemein- und -ausgang entstanden ist. Die Varianz des Messwertes $\hat{H}_n(\mu)$ berechnet sich mit Gleichung (3.17) als Erwartungswert einer quadratischen Form. Die Berechnung erfolgt nach dem im Anhang A.5 beschriebenen Verfahren. Wir erhalten:

$$\begin{aligned} C_{\hat{H}_n(\mu), \hat{H}_n(\mu)} &= E \left\{ \left| \hat{H}_n(\mu) - E \{ \hat{H}_n(\mu) \} \right|^2 \right\} = E \left\{ \left| (\hat{\vec{H}}(\mu) - \vec{H}(\mu)) \cdot \vec{E}_n^H \right|^2 \right\} = \\ &= E \left\{ \frac{1}{(L-1)^2} \cdot \vec{N}_f(\mu) \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \vec{E}_n \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \underline{1}_\perp \cdot \vec{N}_f(\mu)^H \right\} = \\ &= E \left\{ \text{spur} \left(\frac{1}{(L-1)^2} \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \vec{E}_n \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \underline{1}_\perp \right) \right\} \cdot \\ &\quad \cdot \left(E \{ |\vec{N}_f(\mu)|^2 \} - |E \{ \vec{N}_f(\mu) \}|^2 \right) = \\ &= E \left\{ \text{spur} \left(\frac{1}{(L-1)^2} \cdot \vec{E}_n \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \underline{1}_\perp \cdot \underline{1}_\perp \cdot \tilde{\underline{V}}(\mu)^H \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \right) \right\} \cdot \\ &\quad \cdot E \{ |\vec{N}_f(\mu)|^2 \} = \\ &= \frac{1}{(L-1)^2} \cdot E \left\{ \hat{\underline{C}}_n(\mu) \cdot \hat{\underline{C}}_n(\mu)^H \right\} \cdot E \{ |\vec{N}_f(\mu)|^2 \} = \\ &= \frac{M}{L-1} \cdot E \left\{ \vec{E}_n \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \right\} \cdot \tilde{\Phi}_n(\mu, \mu). \end{aligned} \quad (3.40)$$

In einem Zwischenschritt wurde dabei der Zeilenvektor

$$\hat{\underline{C}}_n(\mu) = \vec{E}_n \cdot \hat{\underline{C}}_{\tilde{\underline{V}}(\mu), \tilde{\underline{V}}(\mu)}^{-1} \cdot \tilde{\underline{V}}(\mu) \cdot \underline{1}_\perp \quad (3.41)$$

eingeführt, der durch eine lineare Abbildung des Einheitsvektors $\vec{E}_n$ entsteht, und der später noch von Bedeutung sein wird. Ein Faktor der Messwertvarianz ist der Erwartungswert des n -ten Hauptdiagonalelements der inversen Kovarianzmatrix. Im weiteren setzen wir die oben beschriebene Behandlung des extrem unwahrscheinlichen, aber dennoch möglichen Falls einer singulären empirischen Kovarianzmatrix voraus. In Anhang A.4 wird gezeigt, dass die Hauptdiagonalelemente immer reell und größer als das Inverse des größten Singulärwertes der empirischen Kovarianzmatrix sind und kleiner als das Inverse des kleinsten — nach unten begrenzten — Singulärwertes der empirischen Kovarianzmatrix sind. Somit ist die Existenz des Erwartungswertes des Hauptdiagonalelements der inversen Kovarianzmatrix gesichert. Der von der Erregung unabhängige Vorfaktor $1/(L-1)$

sorgt dafür, dass die Messwertvarianz mit steigender Mittelungsanzahl L gegen null strebt, so dass die Messwerte konsistent sind. Durch die Wahl eines Zufallsvektors $\vec{V}$, bei dem die nach Gleichung (3.10) definierten theoretischen Kovarianzmatrizen gut konditioniert sind, kann erreicht werden, dass die Wahrscheinlichkeit eine schlecht konditionierte empirische Kovarianzmatrix zu erhalten extrem klein wird (siehe Kapitel A.2). Daher wird auch der Erwartungswert des n -ten Hauptdiagonalelements der inversen Kovarianzmatrix — und somit die Messwertvarianz — klein, wenn man einerseits die Varianzen der Zufallsgrößen des Vektors $\vec{V}(\mu)$ möglichst groß wählt, und andererseits darauf achtet, dass die theoretische Kovarianzmatrix dieser Zufallsgrößen gut konditioniert ist.

Für die Kovarianz des n -ten Messwertes $\hat{H}_n(\mu)$ des Messwertvektors $\hat{\mathbf{H}}(\mu)$ der beiden Übertragungsfunktionen ergibt sich analog:

$$\begin{aligned}
C_{\hat{H}_n(\mu), \hat{H}_n(\mu)^*} &= E \left\{ \left(\hat{H}_n(\mu) - E\{\hat{H}_n(\mu)\} \right)^2 \right\} = E \left\{ \left((\hat{\mathbf{H}}(\mu) - \vec{H}(\mu)) \cdot \vec{E}_n^H \right)^2 \right\} = \\
&= E \left\{ \frac{1}{(L-1)^2} \cdot \vec{N}_f(\mu) \cdot \mathbf{1}_\perp \cdot \vec{V}(\mu)^H \cdot \underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \vec{E}_n \cdot (\underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1})^* \cdot \vec{V}(\mu)^* \cdot \mathbf{1}_\perp \cdot \vec{N}_f(\mu)^T \right\} = \\
&= E \left\{ \text{spur} \left(\frac{1}{(L-1)^2} \cdot \mathbf{1}_\perp \cdot \vec{V}(\mu)^H \cdot \underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \vec{E}_n \cdot (\underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1})^* \cdot \vec{V}(\mu)^* \cdot \mathbf{1}_\perp \right) \right\} \cdot \\
&\quad \cdot \left(E\{\mathbf{N}_f(\mu)^2\} - E\{\mathbf{N}_f(\mu)\}^2 \right) = \\
&= E \left\{ \text{spur} \left(\frac{1}{(L-1)^2} \cdot \vec{E}_n \cdot (\underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1})^* \cdot \vec{V}(\mu)^* \cdot \mathbf{1}_\perp \cdot \mathbf{1}_\perp \cdot \vec{V}(\mu)^H \cdot \underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1} \cdot \vec{E}_n^H \right) \right\} \cdot \\
&\quad \cdot E\{\mathbf{N}_f(\mu)^2\} = \\
&= \frac{1}{(L-1)^2} \cdot E \left\{ \hat{\mathbf{C}}_n(\mu) \cdot \hat{\mathbf{C}}_n(\mu)^T \right\}^* \cdot E\{\mathbf{N}_f(\mu)^2\} = \\
&= \frac{1}{L-1} \cdot E \left\{ \vec{E}_n \cdot (\underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1})^* \cdot \underline{\hat{C}}_{\vec{V}(\mu)^*, \vec{V}(\mu)} \cdot \underline{\hat{C}}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1} \cdot \vec{E}_n^H \right\} \cdot E\{\mathbf{N}_f(\mu)^2\}. \quad (3.42)
\end{aligned}$$

Bei der Kovarianz des Messwertes $\hat{H}_n(\mu)$ ist zu beachten, dass hier der Erwartungswert $E\{\mathbf{N}_f(\mu)^2\}$ nicht mehr nur für $\mu=0$ und $\mu=M/2$ nennenswert von null verschieden ist, da er sich nicht mehr nach Gleichung ([1]:3.61) durch einfache Integration des KLDS ergibt, das mit dem auf zwei unterschiedliche Weisen verschobenen Fensterspektrum multipliziert wird. Stattdessen findet hier das entsprechende Doppelintegral nach Gleichung (2.51) mit $\hat{\mu} = -\mu$ Anwendung, bei dem im Integranden ebenfalls das Fensterspektrum mit zwei unterschiedlichen Frequenzverschiebungen auftritt. In diesem Integral steht nun aber statt des eindimensionalen das bifrequente KLDS $\Psi_n(\Omega_1, \Omega_2)$, das nicht nur für $\Omega_1 = \Omega_2$ eine Impulslinie aufweist, wie das bei einem stationären Approximationsfehlerprozess der Fall ist. Deshalb erhalten wir selbst bei Verwendung einer hoch frequenzselektiven Fensterfolge

für den Erwartungswert $E\{\mathbf{N}_f(\mu)^2\}$ nur dann näherungsweise null, wenn μ kein ganzzahliges Vielfaches von $M/2/K_\Phi$ ist. Die Kovarianz des Messwertes $\hat{\mathbf{H}}_n(\mu)$ kann also nur von null verschieden sein, wenn μ ein ganzzahliges Vielfaches von $M/2/K_\Phi$ ist. Die von null verschiedenen Werte der Messwertkovarianz erhalten wir mit Gleichung (2.50):

$$C_{\hat{\mathbf{H}}_n(\mu), \hat{\mathbf{H}}_n(\mu)^*} = \frac{M}{L-1} \cdot E\left\{ \vec{E}_n \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1})^* \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu)^*, \tilde{\mathbf{V}}(\mu)} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \right\} \cdot \hat{\Psi}_n(\mu, -\mu) \\ \forall \quad \mu = \check{\mu} \cdot \frac{M}{2 \cdot K_\Phi} \quad \wedge \quad \check{\mu} \in \mathbb{Z} \quad (3.43)$$

Die wahren Messwert(ko)varianzen lassen sich mit

$$\hat{C}_{\hat{\mathbf{H}}_n(\mu), \hat{\mathbf{H}}_n(\mu)} = \frac{M}{(L-1)^2} \cdot \hat{C}_n(\mu) \cdot \hat{C}_n(\mu)^H \cdot \hat{\Phi}_n(\mu, \mu) = \quad (3.44a) \\ = \frac{M}{L-1} \cdot \vec{E}_n \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \hat{\Phi}_n(\mu, \mu) \\ \forall \quad \mu = 0 \ (1) \ M-1$$

und

$$\hat{C}_{\hat{\mathbf{H}}_n(\mu), \hat{\mathbf{H}}_n(\mu)^*} = \frac{M}{(L-1)^2} \cdot \left(\hat{C}_n(\mu) \cdot \hat{C}_n(\mu)^T \right)^* \cdot \hat{\Psi}_n(\mu, -\mu) = \quad (3.44b) \\ = \frac{M}{L-1} \cdot \vec{E}_n \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1})^* \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu)^*, \tilde{\mathbf{V}}(\mu)} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \vec{E}_n^H \cdot \hat{\Psi}_n(\mu, -\mu) \\ \forall \quad \mu = \check{\mu} \cdot \frac{M}{2 \cdot K_\Phi} \quad \wedge \quad \check{\mu} \in \mathbb{Z}$$

abschätzen. Setzt man hier die Messwerte $\hat{\Phi}_n(\mu, \mu)$ und $\hat{\Psi}_n(\mu, -\mu)$ gemäß der Gleichungen (3.29) in der Form ein, die die Vektoren $\vec{N}_f(\dots)$ enthält, betrachtet man die Schätzwerte wieder als zufällig, und bildet man deren Erwartungswerte mit Anhang A.5, gehen wieder nur die Hauptdiagonalelemente der quadratischen bzw. bilinearen Formen additiv in die Erwartungswerte ein. In dem zu bildenden Erwartungswert, der nur von den Spektralwerten der Erregung abhängt, lässt sich daher jeweils die Spur der Matrix, auf die der jeweilige Messwert normiert ist, kürzen, so dass derselbe vom Spektrum der Erregung abhängige Erwartungswert verbleibt, wie bei den theoretischen Messwert(ko)varianzen. Der bei der Berechnung der Erwartungswerte der Schätzwerte auftretende, von dem Spektrum des gefensterten Approximationsfehlerprozesses abhängige Erwartungswert ist identisch mit dem Erwartungswert, der in den theoretischen Messwert(ko)varianzen zu finden ist. Die Schätzwerte der Messwert(ko)varianzen sind daher erwartungstreu. Wenn man Messwerte $\hat{\Phi}_n(\mu, \mu)$ und $\hat{\Psi}_n(\mu, -\mu)$ verwendet, die mit Matrizen $\underline{V}_\perp(\mu)$ gewonnen wurden, die die Bedingungen (3.33) erfüllen, lässt sich zeigen, dass die konkreten Varianzschätzwerte niemals kleiner als die Beträge der entsprechenden konkreten Kovarianzschätzwerte sind. Dabei berücksichtigt man, dass zum einen die Messwerte $\hat{\Phi}_n(\mu, \mu)$ und $\hat{\Psi}_n(\mu, -\mu)$

die Ungleichung (3.36b) erfüllen (man setzt dort $\mu = \check{\mu} \cdot M/2/K_\Phi$ und $\tilde{\mu} = -\check{\mu}$ ein), und dass zum anderen der Betrag $\left| \hat{\tilde{C}}_n(\mu) \cdot \hat{\tilde{C}}_n(\mu)^T \right|$ eines Skalarprodukts niemals größer ist, als das Produkt der euklidischen Normen der beiden daran beteiligten Vektoren (Cauchy-Schwarzsche Ungleichung), das in diesem Fall gleich dem Skalarprodukt $\hat{\tilde{C}}_n(\mu) \cdot \hat{\tilde{C}}_n(\mu)^H$ ist.

Auch die Messwerte $\hat{\mathbf{U}}_f(\mu)$ betrachten wir nun wieder als Zufallsgrößen, die aus mathematischen Stichproben der Zufallsprozesse am Systemein- und -ausgang entstanden sind. Die Varianzen und die Kovarianzen der Messwerte $\hat{\mathbf{U}}_f(\mu)$ berechnen sich mit Gleichung (3.20) als Erwartungswerte quadratischer bzw. bilinearer Formen, die man mit der in Anhang A.5 beschriebenen Methode umformt. In weiser Voraussicht berechnen wir nicht nur die Varianzen und Kovarianzen der Messwerte, sondern auch gleich alle Kovarianzen zweier Messwerte unterschiedlicher Frequenzen μ_1 und μ_2 als Erwartungswerte bilinearer Formen. Die Messwert(ko)varianzen ergeben sich dann mit $\mu = \mu_1 = \mu_2$. Die Berechnung erfolgt in ähnlicher Weise wie in den Gleichungen [1]:3.30. Wir erhalten:

$$\begin{aligned}
C_{\hat{\mathbf{U}}_f(\mu_1), \hat{\mathbf{U}}_f(\mu_2)} &= \mathbb{E} \left\{ \left(\hat{\mathbf{U}}_f(\mu_1) - \mathbb{E}\{\hat{\mathbf{U}}_f(\mu_1)\} \right) \cdot \left(\hat{\mathbf{U}}_f(\mu_2) - \mathbb{E}\{\hat{\mathbf{U}}_f(\mu_2)\} \right)^H \right\} = \quad (3.45a) \\
&= \mathbb{E} \left\{ \left(\hat{\mathbf{U}}_f(\mu_1) - \mathbf{U}_f(\mu_1) \right) \cdot \left(\hat{\mathbf{U}}_f(\mu_2) - \mathbf{U}_f(\mu_2) \right)^H \right\} = \\
&= \mathbb{E} \left\{ \tilde{\mathbf{N}}_f(\mu_1) \cdot \frac{(L-1) \cdot \underline{\mathbf{E}} - \underline{\mathbf{1}}_\perp \cdot \tilde{\mathbf{V}}(\mu_1)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1)}{L \cdot (L-1)} \cdot \tilde{\mathbf{1}}^H \cdot \right. \\
&\quad \left. \cdot \tilde{\mathbf{1}} \cdot \frac{(L-1) \cdot \underline{\mathbf{E}} - \tilde{\mathbf{V}}(\mu_2)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_2) \cdot \underline{\mathbf{1}}_\perp}{L \cdot (L-1)} \cdot \tilde{\mathbf{N}}_f(\mu_2)^H \right\} = \\
&= \mathbb{E} \left\{ \text{spur} \left(\frac{1}{L^2 \cdot (L-1)^2} \cdot \left((L-1) \cdot \underline{\mathbf{E}} - \underline{\mathbf{1}}_\perp \cdot \tilde{\mathbf{V}}(\mu_1)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1) \right) \cdot \tilde{\mathbf{1}}^H \cdot \right. \right. \\
&\quad \left. \left. \cdot \tilde{\mathbf{1}} \cdot \left((L-1) \cdot \underline{\mathbf{E}} - \tilde{\mathbf{V}}(\mu_2)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_2) \cdot \underline{\mathbf{1}}_\perp \right) \right) \right\} \cdot \\
&\quad \cdot \left(\mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^*\} - \mathbb{E}\{\mathbf{N}_f(\mu_1)\} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_2)^*\} \right) + \mathbb{E}\{\mathbf{N}_f(\mu_1)\} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_2)^*\} = \\
&= \frac{\mathbb{E}\{\hat{\tilde{\mathbf{C}}}_U(\mu_2) \cdot \hat{\tilde{\mathbf{C}}}_U(\mu_1)^H\}}{L^2 \cdot (L-1)^2} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^*\} = \\
&= \left(\frac{1}{L} + \frac{1}{L^2 \cdot (L-1)^2} \cdot \mathbb{E} \left\{ \tilde{\mathbf{1}} \cdot \tilde{\mathbf{V}}(\mu_2)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_2) \cdot \underline{\mathbf{1}}_\perp \cdot \right. \right. \\
&\quad \left. \left. \cdot \underline{\mathbf{1}}_\perp \cdot \tilde{\mathbf{V}}(\mu_1)^H \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1) \cdot \tilde{\mathbf{1}}^H \right\} \right) \cdot \mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^*\}
\end{aligned}$$

und

$$\begin{aligned}
C_{\hat{\mathbf{U}}_f(\mu_1), \hat{\mathbf{U}}_f(\mu_2)^*} &= \mathbb{E} \left\{ \left(\hat{\mathbf{U}}_f(\mu_1) - \mathbb{E}\{\hat{\mathbf{U}}_f(\mu_1)\} \right) \cdot \left(\hat{\mathbf{U}}_f(\mu_2) - \mathbb{E}\{\hat{\mathbf{U}}_f(\mu_2)\} \right) \right\} = \quad (3.45b) \\
&= \mathbb{E} \left\{ \left(\hat{\mathbf{U}}_f(\mu_1) - \mathbf{U}_f(\mu_1) \right) \cdot \left(\hat{\mathbf{U}}_f(\mu_2) - \mathbf{U}_f(\mu_2) \right) \right\} = \\
&= \mathbb{E} \left\{ \tilde{\mathbf{N}}_f(\mu_1) \cdot \frac{(L-1) \cdot \underline{\mathbf{E}} - \underline{\mathbf{1}}_{\perp} \cdot \tilde{\mathbf{V}}(\mu_1)^{\text{H}} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1)}{L \cdot (L-1)} \cdot \vec{\mathbf{1}}^{\text{H}} \cdot \right. \\
&\quad \left. \cdot \vec{\mathbf{1}} \cdot \frac{(L-1) \cdot \underline{\mathbf{E}} - \tilde{\mathbf{V}}(\mu_2)^{\text{T}} \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1})^* \cdot \tilde{\mathbf{V}}(\mu_2)^* \cdot \underline{\mathbf{1}}_{\perp}}{L \cdot (L-1)} \cdot \tilde{\mathbf{N}}_f(\mu_2)^{\text{T}} \right\} = \\
&= \mathbb{E} \left\{ \text{spur} \left(\frac{1}{L^2 \cdot (L-1)^2} \cdot \left((L-1) \cdot \underline{\mathbf{E}} - \underline{\mathbf{1}}_{\perp} \cdot \tilde{\mathbf{V}}(\mu_1)^{\text{H}} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1) \right) \cdot \vec{\mathbf{1}}^{\text{H}} \cdot \right. \right. \\
&\quad \left. \left. \cdot \vec{\mathbf{1}} \cdot \left((L-1) \cdot \underline{\mathbf{E}} - \tilde{\mathbf{V}}(\mu_2)^{\text{T}} \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1})^* \cdot \tilde{\mathbf{V}}(\mu_2)^* \cdot \underline{\mathbf{1}}_{\perp} \right) \right) \right\} \cdot \\
&\quad \cdot \left(\mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)\} - \mathbb{E}\{\mathbf{N}_f(\mu_1)\} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_2)\} \right) + \mathbb{E}\{\mathbf{N}_f(\mu_1)\} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_2)\} = \\
&= \frac{\mathbb{E}\{\hat{\mathbf{C}}_U(\mu_2) \cdot \hat{\mathbf{C}}_U(\mu_1)^{\text{T}}\}^*}{(L-1)^2} \cdot \mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)\} = \\
&= \left(\frac{1}{L} + \frac{1}{L^2 \cdot (L-1)^2} \cdot \mathbb{E} \left\{ \vec{\mathbf{1}} \cdot \tilde{\mathbf{V}}(\mu_2)^{\text{T}} \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_2), \tilde{\mathbf{V}}(\mu_2)}^{-1})^* \cdot \tilde{\mathbf{V}}(\mu_2)^* \cdot \underline{\mathbf{1}}_{\perp} \cdot \right. \right. \\
&\quad \left. \left. \cdot \underline{\mathbf{1}}_{\perp} \cdot \tilde{\mathbf{V}}(\mu_1)^{\text{H}} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu_1), \tilde{\mathbf{V}}(\mu_1)}^{-1} \cdot \tilde{\mathbf{V}}(\mu_1) \cdot \vec{\mathbf{1}}^{\text{H}} \right\} \right) \cdot \mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)\}.
\end{aligned}$$

Dabei wurde jeweils berücksichtigt, dass die optimale theoretische Anpassung des Spektrums der deterministischen Störung nach Gleichung (2.29) dazu führte, dass der Erwartungswert der Zufallsgröße $\mathbf{N}_f(\mu)$ nach Gleichung (2.30) null ist. Wieder wurde in einem Zwischenschritt ein Zeilenvektor

$$\hat{\mathbf{C}}_U(\mu) = \frac{L-1}{L} \cdot \vec{\mathbf{1}} - \frac{\vec{\mathbf{1}} \cdot \tilde{\mathbf{V}}(\mu)^{\text{H}}}{L} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \underline{\mathbf{1}}_{\perp} \quad (3.46)$$

eingeführt, der zur Abschätzung der Beträge der empirischen Varianz und Kovarianz des Messwertes $\hat{\mathbf{U}}_f(\mu)$ gebraucht wird. Wenn man eine hoch frequenzselektive Fensterfolge verwendet, sind die Kovarianzen $\mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^*\}$, die in Gleichung (3.45a) auftreten, gemäß Gleichung (2.41) für die Frequenzen mit $\mu_2 \neq \mu_1 + \tilde{\mu} \cdot M/K_{\Phi}$ gegenüber den Kovarianzen mit $\mu_2 = \mu_1 + \tilde{\mu} \cdot M/K_{\Phi}$ vernachlässigbar klein. Entsprechendes gilt nach Gleichung (2.51) für die Kovarianzen $\mathbb{E}\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)\}$ in Gleichung (3.45b) für die Frequen-

zen mit $\mu_2 \neq -\mu_1 - \tilde{\mu} \cdot M/K_\Phi$. Für die nennenswert von null verschiedenen Messwert(ko)-varianzen erhalten wir:

$$\begin{aligned}
 C_{\hat{U}_f(\mu), \hat{U}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} &= M \cdot \frac{\mathbb{E}\left\{\hat{\mathbf{C}}_U(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\mathbf{C}}_U(\mu)^H\right\}}{(L-1)^2} \cdot \tilde{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) = \\
 &= \frac{M}{L} \cdot \left(1 + \frac{L}{L-1} \cdot \mathbb{E}\left\{\frac{\tilde{\mathbf{1}} \cdot \tilde{\mathbf{V}}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{L} \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1} \cdot \tilde{\mathbf{V}}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \right. \right. \\
 &\quad \left. \left. \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu)}^{-1} \cdot \frac{\tilde{\mathbf{V}}(\mu) \cdot \tilde{\mathbf{1}}^H}{L}\right\}\right) \cdot \tilde{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})
 \end{aligned} \tag{3.47a}$$

und

$$\begin{aligned}
 C_{\hat{U}_f(\mu), \hat{U}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})}^* &= M \cdot \frac{\mathbb{E}\left\{\hat{\mathbf{C}}_U(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\mathbf{C}}_U(\mu)^T\right\}^*}{L^2 \cdot (L-1)^2} \cdot \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) = \\
 &= \frac{M}{L} \cdot \left(1 + \frac{L}{L-1} \cdot \mathbb{E}\left\{\frac{\tilde{\mathbf{1}} \cdot \tilde{\mathbf{V}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{L} \cdot (\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1} \cdot \tilde{\mathbf{V}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}))^* \cdot \right. \right. \\
 &\quad \left. \left. \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu)}^{-1} \cdot \frac{\tilde{\mathbf{V}}(\mu) \cdot \tilde{\mathbf{1}}^H}{L}\right\}\right) \cdot \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})
 \end{aligned} \tag{3.47b}$$

Wie man sieht ist die Messwertkovarianz $C_{\hat{U}_f(\mu), \hat{U}_f(\mu)}^*$ nur dann nennenswert von null verschieden, wenn μ ein ganzzahliges Vielfaches von $\frac{M}{2 \cdot K_\Phi}$ ist, weil es nur dann ein ganzzahliges $\tilde{\mu}$ gibt, bei dem $\mu = -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}$ gilt. Wenn man den Fall ausschließt, dass die empirischen Kovarianzmatrizen des Spektrums der Erregung singulär werden², existieren die von der Erregung abhängigen Erwartungswerte, die in den letzten Gleichungen als Vorfaktoren auftreten, da es sich dabei um zwei Vektoren von empirischen Mittelwerten handelt, die gemeinsam mit einem Produkt empirischer, teils inverser, Kovarianzmatrizen eine bilineare Form bilden. Von diesen Vorfaktoren kann man erwarten, dass sie für steigende Werte von L gegen einen festen endlichen Wert gehen. Die Varianz der Messwerte sinkt dann asymptotisch mit $1/L$, so dass die Messwerte konsistent sind. Die letzten Gleichungen zeigen auch, dass sich die Messwert(ko)varianzen hier aus jeweils zwei Anteilen zusammensetzen. Der eine Anteil ist der Summand mit dem konstanten Vorfaktor M/L , der durch den Approximationsfehlers selbst verursacht wird, während der zweite Anteil, dessen Vorfaktor der von der Erregung abhängige Erwartungswert ist, durch die verrauschte Messung der empirischen Mittelwerte der Erregung, die mit den beiden linearen Modellsystemen gefiltert werden, verursacht wird. Mit $\tilde{\mu}=0$ erhält man in Gleichung (3.47a) die Messwertvarianz. Da in diesem Fall in dem von der Erregung abhängigen Erwartungswert eine positiv semidefinite quadratische Form steht, wird die Messwertvarianz minimal,

²Siehe Messwertvarianz der Übertragungsfunktionen

wenn man einen erregenden Prozess für die Messung verwendet, die keinen zeitabhängigen Mittelwert aufweist. Einen mittelwertfreien erregenden Prozess für die Messung zu verwenden ist natürlich nur dann möglich, wenn man sicher sein kann, dass sich dadurch die für den Betriebszustand typischen, optimalen, theoretischen Regressionskoeffizienten nicht ändern.

Die theoretischen Messwertkovarianzen lassen sich mit

$$\begin{aligned}
\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} &= \frac{M}{(L-1)^2} \cdot \hat{C}_U(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{C}_U(\mu)^H \cdot \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) = \\
&= \frac{M}{L} \cdot \left(1 + \frac{L}{L-1} \cdot \frac{\vec{1} \cdot \tilde{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H}{L} \cdot \hat{C}_{\tilde{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \tilde{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1} \right. \\
&\quad \cdot \hat{C}_{\tilde{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \tilde{V}(\mu)}^{-1} \cdot \hat{C}_{\tilde{V}(\mu), \tilde{V}(\mu)} \cdot \frac{\tilde{V}(\mu) \cdot \vec{1}^H}{L} \left. \right) \cdot \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \\
&\quad \forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1
\end{aligned} \tag{3.48a}$$

und

$$\begin{aligned}
\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} &= \frac{M}{(L-1)^2} \cdot \hat{C}_U(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{C}_U(\mu)^H \cdot \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) = \\
&= \frac{M}{L} \cdot \left(1 + \frac{L}{L-1} \cdot \frac{\vec{1} \cdot \tilde{V}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{L} \cdot (\hat{C}_{\tilde{V}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}), \tilde{V}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1})^* \right. \\
&\quad \cdot \hat{C}_{\tilde{V}(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*, \tilde{V}(\mu)}^{-1} \cdot \hat{C}_{\tilde{V}(\mu), \tilde{V}(\mu)} \cdot \frac{\tilde{V}(\mu) \cdot \vec{1}^H}{L} \left. \right) \cdot \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \\
&\quad \forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1
\end{aligned} \tag{3.48b}$$

abschätzen. Dass diese Schätzwerte erwartungstreu sind, zeigt man genauso wie bei den Schätzwerten der Varianzen und Kovarianzen der Messwerte der Übertragungsfunktion mit der im Anhang A.5 beschriebenen Methode.

Wie gesagt, erhalten wir mit $\mu_1 = \mu_2$ die benötigten Messwert(ko)varianzen. Mit den neuen Frequenzvariablen μ und $\tilde{\mu}$ bedeutet das bei den Messwertvarianzen, dass man $\tilde{\mu} = 0$ einsetzt, und dass man für alle $\mu = 0 \ (1) \ M-1$ von null verschiedene Messwertvarianzen erhalten kann. Bei den Messwertkovarianzen kann man lediglich für $\mu = \tilde{\mu} \cdot M / (2 \cdot K_\Phi)$ mit $\tilde{\mu} = 0 \ (1) \ 2 \cdot K_\Phi - 1$ nennenswert von null verschiedene Werte erhalten, und in der letzten Gleichung ist $\tilde{\mu} = -\tilde{\mu}$ einzusetzen.

Wenn man Messwerte $\hat{\Phi}_n(\mu, \mu)$ und $\hat{\Psi}_n(\mu, -\mu)$ verwendet, die mit Matrizen $\underline{V}_\perp(\mu)$ gewonnen wurden, die den Bedingungen (3.33) genügen, lässt sich analog zu den Schätzwerten der Varianzen und Kovarianzen der Messwerte der Übertragungsfunktion zeigen, dass die

konkreten Varianzschätzwerte niemals kleiner als die Beträge der entsprechenden konkreten Kovarianzschätzwerte sind. Dabei greift man auf die Ungleichung (3.36b) und auf die nach Gleichung (3.46) definierten Vektoren $\hat{\tilde{C}}_U(\dots)$ zurück.

Auch die Messwerte $\hat{\mathbf{u}}(k)$ betrachten wir nun wieder als Zufallsgrößen, die aus mathematischen Stichproben der Zufallsprozesse am Systemein- und -ausgang entstanden sind. Die Varianzen der Messwerte $\hat{\mathbf{u}}(k)$ berechnen sich mit Gleichung (3.23) zu:

$$\begin{aligned} C_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \mathbb{E} \left\{ \left| \hat{\mathbf{u}}(k) - \mathbb{E} \{ \hat{\mathbf{u}}(k) \} \right|^2 \right\} = \mathbb{E} \left\{ \left| \hat{\mathbf{u}}(k) - \mathbf{u}(k) \right|^2 \right\} = \\ &= \mathbb{E} \left\{ \left| \frac{1}{L} \cdot \vec{\mathbf{n}}(k) \cdot \vec{\mathbf{1}}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \frac{\vec{\mathbf{N}}_f(\mu) \cdot \underline{\mathbf{1}}_{\perp} \cdot \tilde{\mathbf{V}}(\mu)^H}{L-1} \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \vec{\mathbf{1}}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \right|^2 \right\}. \end{aligned} \quad (3.49)$$

Als Abkürzung führen wir den Zufallsvektor

$$\hat{\tilde{\mathbf{C}}}_u(\mu) = \frac{\vec{\mathbf{1}} \cdot \tilde{\mathbf{V}}(\mu)^H}{L} \cdot \hat{\tilde{\mathbf{C}}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \tilde{\mathbf{V}}(\mu) \cdot \underline{\mathbf{1}}_{\perp} \quad (3.50)$$

ein und erhalten mit

$$\begin{aligned} C_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \mathbb{E} \left\{ \left| \frac{1}{L} \cdot \vec{\mathbf{n}}(k) \cdot \vec{\mathbf{1}}^H - \frac{1}{M \cdot (L-1)} \cdot \sum_{\mu=0}^{M-1} \vec{\mathbf{N}}_f(\mu) \cdot \hat{\tilde{\mathbf{C}}}_u(\mu)^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \right|^2 \right\} = \\ &= \frac{1}{L^2} \cdot \mathbb{E} \left\{ \vec{\mathbf{n}}(k) \cdot \vec{\mathbf{1}}^H \cdot \vec{\mathbf{1}} \cdot \vec{\mathbf{n}}(k)^H \right\} - \\ &\quad - \frac{1}{M \cdot L \cdot (L-1)} \cdot \sum_{\mu=0}^{M-1} \mathbb{E} \left\{ \vec{\mathbf{n}}(k) \cdot \vec{\mathbf{1}}^H \cdot \hat{\tilde{\mathbf{C}}}_u(\mu) \cdot \vec{\mathbf{N}}_f(\mu)^H \right\} \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} - \\ &\quad - \frac{1}{M \cdot L \cdot (L-1)} \cdot \sum_{\mu=0}^{M-1} \mathbb{E} \left\{ \vec{\mathbf{N}}_f(\mu) \cdot \hat{\tilde{\mathbf{C}}}_u(\mu)^H \cdot \vec{\mathbf{1}} \cdot \vec{\mathbf{n}}(k)^H \right\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} + \\ &\quad + \frac{1}{M^2 \cdot (L-1)^2} \cdot \sum_{\mu_1=0}^{M-1} \sum_{\mu_2=0}^{M-1} \mathbb{E} \left\{ \vec{\mathbf{N}}_f(\mu_1) \cdot \hat{\tilde{\mathbf{C}}}_u(\mu_1)^H \cdot \hat{\tilde{\mathbf{C}}}_u(\mu_2) \cdot \vec{\mathbf{N}}_f(\mu_2)^H \right\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot (\mu_1 - \mu_2) \cdot k} \end{aligned} \quad (3.51)$$

eine Summe von Erwartungswerten von bilinearen Formen. Die Erwartungswerte der bilinearen Formen lassen sich mit Anhang A.5 berechnen. Dabei erkennt man dass die beiden Einfachsummen in der letzten Gleichung keinen Beitrag liefern. Als nächstes berücksichtigen wir, dass die Zufallsgrößen $\vec{\mathbf{n}}(k)$ und $\vec{\mathbf{N}}_f(\mu)$ mittelwertfrei sind, und dass die Kovarianzen $\mathbb{E} \{ \vec{\mathbf{N}}_f(\mu_1) \cdot \vec{\mathbf{N}}_f(\mu_2)^* \}$ nach Gleichung (2.41) für die Frequenzen mit $\mu_2 \neq \mu_1 + \tilde{\mu} \cdot M / K_{\Phi}$ gegenüber den Kovarianzen mit $\mu_2 = \mu_1 + \tilde{\mu} \cdot M / K_{\Phi}$ vernachlässigbar klein sind, wenn wir

eine hoch frequenzselektive Fensterfolge verwenden. Damit ergibt sich:

$$\begin{aligned}
C_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \frac{1}{L} \cdot \mathbb{E}\{|\mathbf{n}(k)|^2\} + \\
&+ \frac{1}{M^2 \cdot (L-1)^2} \cdot \sum_{\mu_1=0}^{M-1} \sum_{\mu_2=0}^{M-1} \mathbb{E}\left\{\hat{\mathbf{C}}_u(\mu_2) \cdot \hat{\mathbf{C}}_u(\mu_1)^H\right\} \cdot \mathbb{E}\left\{\mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^*\right\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot (\mu_1 - \mu_2) \cdot k} \approx \\
&\approx \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} + \\
&+ \frac{1}{M \cdot (L-1)^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \mathbb{E}\left\{\hat{\mathbf{C}}_u\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot \hat{\mathbf{C}}_u(\mu)^H\right\} \cdot \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M \cdot (L-1)^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \left(\frac{L-1}{L} \cdot \vec{1} \cdot \vec{1}^H \cdot \frac{L-1}{L} + \mathbb{E}\left\{\hat{\mathbf{C}}_u\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot \hat{\mathbf{C}}_u(\mu)^H\right\} \right) \cdot \\
&\quad \cdot \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M \cdot (L-1)^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \mathbb{E}\left\{\left(\frac{L-1}{L} \cdot \vec{1} - \hat{\mathbf{C}}_u\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right)\right) \cdot \left(\vec{1}^H \cdot \frac{L-1}{L} - \hat{\mathbf{C}}_u(\mu)^H\right)\right\} \cdot \\
&\quad \cdot \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M \cdot (L-1)^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \mathbb{E}\left\{\hat{\mathbf{C}}_U\left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot \hat{\mathbf{C}}_U(\mu)^H\right\} \cdot \tilde{\Phi}_{\mathbf{n}}\left(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}\right) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} C_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} \\
&\quad \forall \quad k = 0 \text{ (1) } F-1,
\end{aligned}$$

wobei nach und nach die Gleichungen (2.55), (3.46) und (3.47a) Verwendung fanden.

Ganz analog erhalten wir aus den Messwertkovarianzen nach Gleichung (3.47b) für die Kovarianzen der Messwerte $\hat{\mathbf{u}}(k)$ die Näherungen

$$\begin{aligned}
C_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)^*} &\approx \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} C_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} \\
&\quad \forall \quad k = 0 \text{ (1) } F-1.
\end{aligned}$$

Diese Messwert(ko)varianzen schätzen wir dadurch ab, dass wir die Werte der wahren Kovarianzen der Messwerte $\hat{U}_f(\mu)$ durch deren Schätzwerte ersetzen:

$$\hat{C}_{\hat{u}(k), \hat{u}(k)} = \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})} \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \quad (3.53a)$$

$$\hat{C}_{\hat{u}(k), \hat{u}(k)^*} = \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \hat{C}_{\hat{U}_f(\mu), \hat{U}_f(-\mu-\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^*} \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \quad (3.53b)$$

$$\forall \quad k = 0 \text{ (1) } F-1$$

Da sich diese Schätzwerte als eine endliche Überlagerung erwartungstreuer und konsistenter Schätzwerte berechnen, schätzen sie die Näherung der Messwert(ko)varianzen der deterministischen Störung ebenfalls erwartungstreu und konsistent ab. Wenn wir zur Berechnung der Schätzwerte $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu+\tilde{\mu} \cdot M/K_{\Phi})}$ und $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(-\mu-\tilde{\mu} \cdot M/K_{\Phi})^*}$ die Messwerte $\hat{\Phi}_{\mathbf{n}}(\dots, \dots)$ und $\hat{\Psi}_{\mathbf{n}}(\dots, -\dots)$ verwenden, die mit Matrizen $\underline{V}_{\perp}(\dots)$ gewonnen wurden, die die Bedingungen (3.33) erfüllen, lässt sich auch für die Schätzwerte der Messwert(ko)varianzen zeigen, dass die konkreten Varianzschätzwerte niemals kleiner als die Beträge der entsprechenden konkreten Kovarianzschätzwerte sind. Im Unterkapitel A.9.2 der Anhangs wird diese Aussage bewiesen.

Nun bestimmen wir die Varianzen und Kovarianzen der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot M/K_{\Phi})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot M/K_{\Phi})$ indem wir mit den Gleichungen (3.29) die vier zweiten zentralen Momente

$$\begin{aligned} C_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})} &= \mathbb{E} \left\{ \left| \hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}) - \tilde{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right|^2 \right\} = \\ &= \mathbb{E} \left\{ \left| \vec{N}_f(\mu) \cdot \frac{\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{M \cdot \text{spur}(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H)} \cdot \vec{N}_f(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H \right|^2 \right\} - \\ &\quad - \frac{1}{M^2} \cdot \left| \mathbb{E} \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^* \right\} \right|^2, \end{aligned} \quad (3.54a)$$

$$\begin{aligned} C_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^*} &= \mathbb{E} \left\{ \left(\hat{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}) - \tilde{\Phi}_{\mathbf{n}}(\mu, \mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \right)^2 \right\} = \\ &= \mathbb{E} \left\{ \left(\vec{N}_f(\mu) \cdot \frac{\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{M \cdot \text{spur}(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H)} \cdot \vec{N}_f(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H \right)^2 \right\} - \\ &\quad - \frac{1}{M^2} \cdot \mathbb{E} \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^* \right\}^2, \end{aligned} \quad (3.54b)$$

$$\begin{aligned}
C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} &= E \left\{ \left| \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) - \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \right|^2 \right\} = \\
&= E \left\{ \left| \vec{N}_f(\mu) \cdot \frac{\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T)} \cdot \vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \right|^2 \right\} - \\
&\quad - \frac{1}{M^2} \cdot \left| E \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \right\} \right|^2
\end{aligned} \tag{3.54c}$$

und

$$\begin{aligned}
C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} &= E \left\{ \left(\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) - \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \right)^2 \right\} = \\
&= E \left\{ \left(\vec{N}_f(\mu) \cdot \frac{\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T}{M \cdot \text{spur}(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T)} \cdot \vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \right)^2 \right\} - \\
&\quad - \frac{1}{M^2} \cdot E \left\{ \mathbf{N}_f(\mu) \cdot \mathbf{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \right\}^2
\end{aligned} \tag{3.54d}$$

$$\forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi-1$$

berechnen. Dazu unterscheiden wir zwei Fälle.

Im ersten Fall ist μ ein ganzzahliges Vielfaches von $M/(2 \cdot K_\Phi)$.

$$\mu = \hat{\mu} \cdot \frac{M}{2 \cdot K_\Phi} \quad \text{mit} \quad \hat{\mu} \in \mathbb{Z} \tag{3.55a}$$

In diesem Fall ist nicht nur die diskrete Frequenz $\mu + \tilde{\mu} \cdot M/K_\Phi$ um ein ganzzahliges Vielfaches von M/K_Φ gegenüber μ verschoben. Auch die negativen diskreten Frequenzen $-\mu$ und $-\mu - \tilde{\mu} \cdot M/K_\Phi$ liegen in demselben Frequenzraster:

$$-\mu = -\hat{\mu} \cdot \frac{M}{2 \cdot K_\Phi} = \hat{\mu} \cdot \frac{M}{2 \cdot K_\Phi} - \hat{\mu} \cdot \frac{M}{K_\Phi} = \mu - \hat{\mu} \cdot \frac{M}{K_\Phi} \tag{3.55b}$$

$$-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi} = -\hat{\mu} \cdot \frac{M}{2 \cdot K_\Phi} - \tilde{\mu} \cdot \frac{M}{K_\Phi} = \hat{\mu} \cdot \frac{M}{2 \cdot K_\Phi} + \check{\mu} \cdot \frac{M}{K_\Phi} = \mu + \check{\mu} \cdot \frac{M}{K_\Phi} \tag{3.55c}$$

$$\text{mit} \quad \check{\mu} = -\hat{\mu} - \tilde{\mu} \in \mathbb{Z}.$$

Daher können nach den Gleichungen (2.41) und (2.51) in diesem Fall alle möglichen Kovarianzen zweier Zufallsgrößen, die man sich aus den vier Zufallsgrößen $\mathbf{N}_f(\mu)$, $\mathbf{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$, $\mathbf{N}_f(-\mu)$, $\mathbf{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$ und den dazu konjugierten Zufallsgrößen herausgreift, von null verschieden sein. Ebenso finden sich Schätzwerte für das M -fache all dieser Kovarianzen unter den in den Gleichungen (3.29) angegebenen Messwerten.

Nun können wir jeweils mit der ersten Zeile des Vektorgleichungssystems (A.41) des Anhangs A.8 die vier gesuchten theoretischen Messwertvarianzen und -kovarianzen berechnen, indem wir in den Gleichungen (A.40) die in Tabelle 3.1 aufgelisteten Substitutionen vornehmen, und die Ergebnisse auf M normieren. Wir erhalten:

Berechnung von	$\vec{N}_1$	$\vec{N}_2$	$\vec{N}_3$	$\vec{N}_4$
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$\vec{N}_f(\mu)$	$\vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$\vec{N}_f(\mu)^*$	$\vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$\vec{N}_f(\mu)$	$\vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$\vec{N}_f(\mu)$	$\vec{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$\vec{N}_f(\mu)$	$\vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$\vec{N}_f(\mu)^*$	$\vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$\vec{N}_f(\mu)$	$\vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$\vec{N}_f(\mu)$	$\vec{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$
Berechnung von	N_1	N_2	N_3	N_4
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$N_f(\mu)$	$N_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$N_f(\mu)^*$	$N_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$N_f(\mu)$	$N_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$N_f(\mu)$	$N_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$N_f(\mu)$	$N_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$N_f(\mu)^*$	$N_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$N_f(\mu)$	$N_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$N_f(\mu)$	$N_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$
Berechnung von	$\underline{V}_1$	$\underline{V}_2$	$\underline{V}_3$	$\underline{V}_4$
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$\underline{V}_\perp(\mu)^*$	$\underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$
$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$\underline{V}_\perp(\mu)^*$	$\underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$
$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*}$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$	$\underline{V}_\perp(\mu)$	$\underline{V}_\perp(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$

Tabelle 3.1: Substitutionen in den Gleichungen (A.40) bzw. (A.42). Die Variablen c_1 , c_2 und c_3 berechnen sich gemäß des unteren Teils der Tabelle A.1 mit den hier angegebenen Substitutionen für $\underline{V}_1$ bis $\underline{V}_4$.

$$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = \quad (3.56a)$$

$$= E\{\mathbf{[h]} \cdot \mathbf{[d]}^{-2}\} \cdot |\tilde{\Psi}_n(\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 + E\{\mathbf{[d]}^{-1}\} \cdot \tilde{\Phi}_n(\mu, \mu) \cdot \tilde{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$$

$$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = \quad (3.56b)$$

$$= E\{\mathbf{[i]} \cdot \mathbf{[d]}^{-2}\} \cdot \tilde{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 + E\{\mathbf{[j]} \cdot \mathbf{[d]}^{-2}\} \cdot \tilde{\Psi}_n(\mu, -\mu) \cdot \tilde{\Psi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^*$$

$$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = \quad (3.56c)$$

$$= E\{\mathbf{[k]} \cdot \mathbf{[g]}^{-2}\} \cdot |\tilde{\Phi}_n(\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 + E\{\mathbf{[g]}^{-1}\} \cdot \tilde{\Phi}_n(\mu, \mu) \cdot \tilde{\Phi}_n(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$$

$$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = \quad (3.56d)$$

$$= E\{\mathbf{[l]} \cdot \mathbf{[g]}^{-2}\} \cdot \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 + E\{\mathbf{[m]} \cdot \mathbf{[g]}^{-2}\} \cdot \tilde{\Psi}_n(\mu, -\mu) \cdot \tilde{\Psi}_n(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$$

$$\forall \quad \mu = 0 \left(\frac{M}{2 \cdot K_\Phi} \right) M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_\Phi - 1.$$

Die bei der Berechnung der Messwert(ko)varianzen auftretenden Terme $\mathbf{[d]}$, $\mathbf{[g]}$, $\mathbf{[h]}$, $\mathbf{[i]}$, $\mathbf{[j]}$, $\mathbf{[k]}$, $\mathbf{[l]}$ und $\mathbf{[m]}$ lassen sich mit Hilfe der Tabelle³ 3.2 aus den bei den unterschiedlichen Frequenzen verwendeten Matrizen $\mathbf{V}_\perp(\mu)$ berechnen. In Anhang A.7 wird gezeigt, dass alle in Tabelle 3.2 aufgelisteten Terme asymptotisch mit der Mittelungsanzahl L proportional steigen. Alle Erwartungswerte in den Gleichungen (3.56), und somit auch die Messwertvarianzen, fallen daher asymptotisch indirekt proportional mit L . Somit sind die Messwerte $\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ und $\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ konsistent. Falls man die nach dem oben beschriebenen Verfahren konstruierten Matrizen verwendet, die die Gleichungen (3.33) erfüllen, so kann man mit den Gleichungen (3.55) die Gleichheit der Matrizen

$$\begin{aligned} \mathbf{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^* &= \mathbf{V}_\perp(\mu)^* = \mathbf{V}_\perp(-\mu + \hat{\mu} \cdot \frac{M}{K_\Phi})^* = \mathbf{V}_\perp(-\mu + \hat{\mu} \cdot \frac{M}{K_\Phi})^T = \\ &= \mathbf{V}_\perp(\mu - \hat{\mu} \cdot \frac{M}{K_\Phi}) = \mathbf{V}_\perp(\mu) = \mathbf{V}_\perp(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}). \end{aligned} \quad (3.57)$$

zeigen. In diesem Fall sind alle in Tabelle 3.2 auftretenden Matrizen gleich der Matrix $\mathbf{V}_\perp(\mu)$, alle Matrixprodukte lassen sich aufgrund der Idempotenz als eine Matrix schrei-

³In der Tabelle 3.2 sind die aus den konkreten Stichproben berechneten nicht zufälligen, und daher nicht fettgedruckten Werte eingetragen, die bei der Berechnung der konkreten Schätzwerte der Messwert(ko)varianzen benötigt werden. Bei der hier vorliegenden Substitution sind jedoch die aus den mathematischen Stichproben abgeleiteten zufälligen Werte einzusetzen. Da sich diese in derselben Art berechnen, wie die konkreten Werte aus den konkreten Stichproben, wurde darauf verzichtet, dieselbe Tabelle noch einmal in Fettdruck abzudrucken.

$\text{a)} = \text{spur}\left(\underline{V}_{\perp}(\mu)\right)$ $\text{b)} = \text{spur}\left(\underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{c)} = \text{spur}\left(\underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$
$\text{d)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{e)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^*\right)$ $\text{f)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{g)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^*\right)$
$\text{h)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^* \cdot \underline{V}_{\perp}(\mu)^* \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{i)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{j)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^* \cdot \underline{V}_{\perp}(\mu)^*\right)$ $\text{k)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}(\mu)^* \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^*\right)$ $\text{l)} = \text{spur}\left(\underline{V}_{\perp}(\mu)^* \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)\right)$ $\text{m)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^* \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}(\mu)^*\right)$
$\text{n)} = \text{spur}\left(\underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu)^*\right)$ $\text{o)} = \text{spur}\left(\underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^*\right)$ $\text{p)} = \text{spur}\left(\underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{V}_{\perp}\left(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^*\right)$

Tabelle 3.2: Substitutionen in den Gleichungen (3.56), (3.60) und der Tabelle (3.3). Im Anhang A.6 wird gezeigt, wie sich alle in dieser Tabelle aufgelisteten Matrixspuren aus den bereits bei der Berechnung der Messwerte des LDS und KLDS verwendeten empirischen Kovarianzmatrizen berechnen lassen, wenn man die Matrizen $\underline{V}_{\perp}(\mu)$ mit Hilfe der Gleichungen (3.27) und (3.28) berechnet.

ben, und alle Matrixspuren sind nicht zufällig und gleich dem Wert $L-1-K(\mu)$. Die in den Gleichungen (3.56) vor den Produkten, Quadraten und Betragsquadraten der theoretischen Werte des LDS bzw. KLDS als Vorfaktoren auftretenden Erwartungswerte sind alle gleich $(L-1-K(\mu))^{-1}$ und somit von der Erregung unabhängig.

Im Anhang A.8 ist ebenfalls angegeben, wie man die Varianzen und Kovarianzen der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ abschätzen kann. Dazu setzt man die in Tabelle 3.1 angegebenen Substitutionen in die Gleichungen (A.40) ein. Die Zufallsgrößen, die jeweils die ersten Elemente der Zufallsvektoren (A.47) sind, schätzen die theoretischen Kovarianzen der dort auftretenden empirischen Kovarianzschätzwerte erwartungstreu ab. Diese empirischen Kovarianzschätzwerte sind mit den angegebenen Substitutionen gerade die M -fachen LDS bzw. KLDS-Messwerte, und somit sind die auf M normierten konkreten Realisierungen der jeweils ersten Zufallsgrößen der Zufallsvektoren (A.47) die gesuchten konkreten Schätzwerte für die Messwert(ko)varianzen.

$$\hat{C}_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} = [\boxed{\text{A}}, \boxed{\text{B}}, \boxed{\text{C}}] \cdot \begin{bmatrix} |\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\ |\hat{\Psi}_{\mathbf{n}}(\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\ \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \end{bmatrix} \quad (3.58a)$$

$$\hat{C}_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^*} = [\boxed{\text{D}}, \boxed{\text{E}}] \cdot \begin{bmatrix} \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^2 \\ \hat{\Psi}_{\mathbf{n}}(\mu, -\mu) \cdot \hat{\Psi}_{\mathbf{n}}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^* \end{bmatrix} \quad (3.58b)$$

$$\hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} = [\boxed{\text{F}}, \boxed{\text{G}}, \boxed{\text{H}}] \cdot \begin{bmatrix} |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\ |\hat{\Phi}_{\mathbf{n}}(\mu, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\ \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \end{bmatrix} \quad (3.58c)$$

$$\hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^*} = [\boxed{\text{I}}, \boxed{\text{J}}] \cdot \begin{bmatrix} \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^2 \\ \hat{\Psi}_{\mathbf{n}}(\mu, -\mu) \cdot \hat{\Psi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \end{bmatrix} \quad (3.58d)$$

$$\forall \quad \mu = 0 \left(\frac{M}{2 \cdot K_{\Phi}} \right) M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ K_{\Phi}-1$$

Die hier auftretenden Terme $\boxed{\text{A}}$ bis $\boxed{\text{J}}$ sind im oberen Teil der Tabelle 3.3 zusammengestellt. Die dabei verwendeten Abkürzungen $\boxed{\text{a}}$ bis $\boxed{\text{p}}$ lassen sich mit Tabelle 3.2 aus

$\underline{V}_1(\mu)$	idempotent und hermitesch	erfüllt Gleichungen (3.33)
$\underline{A}$	$= \frac{2 \cdot \underline{d} \cdot \underline{e} \cdot \underline{h} - \underline{a} \cdot \underline{b} \cdot \underline{h}^2 - \underline{d}^2 \cdot \underline{e}^2}{\underline{a} \cdot \underline{b} \cdot \underline{d}^2 \cdot \underline{e}^2 + 2 \cdot \underline{d} \cdot \underline{e} \cdot \underline{h} - 2 \cdot \underline{d}^2 \cdot \underline{e}^2 - \underline{a} \cdot \underline{b} \cdot \underline{h}^2}$	$= \frac{-2}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{B}$	$= \frac{\underline{a} \cdot \underline{b} \cdot \underline{e}^2 \cdot \underline{h} - \underline{d} \cdot \underline{e}^3}{\underline{a} \cdot \underline{b} \cdot \underline{d}^2 \cdot \underline{e}^2 + 2 \cdot \underline{d} \cdot \underline{e} \cdot \underline{h} - 2 \cdot \underline{d}^2 \cdot \underline{e}^2 - \underline{a} \cdot \underline{b} \cdot \underline{h}^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{C}$	$= \frac{\underline{a} \cdot \underline{b} \cdot \underline{d} \cdot \underline{e}^2 - \underline{a} \cdot \underline{b} \cdot \underline{e} \cdot \underline{h}}{\underline{a} \cdot \underline{b} \cdot \underline{d}^2 \cdot \underline{e}^2 + 2 \cdot \underline{d} \cdot \underline{e} \cdot \underline{h} - 2 \cdot \underline{d}^2 \cdot \underline{e}^2 - \underline{a} \cdot \underline{b} \cdot \underline{h}^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{D}$	$= \frac{\underline{i} \cdot \underline{n} \cdot \underline{o} - 2 \cdot \underline{j} ^2}{\underline{d}^2 \cdot \underline{n} \cdot \underline{o} + \underline{i} \cdot \underline{n} \cdot \underline{o} - 2 \cdot \underline{j} ^2}$	$= \frac{L-3-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{E}$	$= \frac{\underline{j} \cdot \underline{n} \cdot \underline{o}}{\underline{d}^2 \cdot \underline{n} \cdot \underline{o} + \underline{i} \cdot \underline{n} \cdot \underline{o} - 2 \cdot \underline{j} ^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{F}$	$= \frac{2 \cdot \underline{g} \cdot \underline{f} \cdot \underline{k} - \underline{a} \cdot \underline{c} \cdot \underline{k}^2 - \underline{g}^2 \cdot \underline{f}^2}{\underline{a} \cdot \underline{c} \cdot \underline{g}^2 \cdot \underline{f}^2 + 2 \cdot \underline{g} \cdot \underline{f} \cdot \underline{k} - 2 \cdot \underline{g}^2 \cdot \underline{f}^2 - \underline{a} \cdot \underline{c} \cdot \underline{k}^2}$	$= \frac{-2}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{G}$	$= \frac{\underline{a} \cdot \underline{c} \cdot \underline{f}^2 \cdot \underline{k} - \underline{g} \cdot \underline{f}^3}{\underline{a} \cdot \underline{c} \cdot \underline{g}^2 \cdot \underline{f}^2 + 2 \cdot \underline{g} \cdot \underline{f} \cdot \underline{k} - 2 \cdot \underline{g}^2 \cdot \underline{f}^2 - \underline{a} \cdot \underline{c} \cdot \underline{k}^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{H}$	$= \frac{\underline{a} \cdot \underline{c} \cdot \underline{g} \cdot \underline{f}^2 - \underline{a} \cdot \underline{c} \cdot \underline{f} \cdot \underline{k}}{\underline{a} \cdot \underline{c} \cdot \underline{g}^2 \cdot \underline{f}^2 + 2 \cdot \underline{g} \cdot \underline{f} \cdot \underline{k} - 2 \cdot \underline{g}^2 \cdot \underline{f}^2 - \underline{a} \cdot \underline{c} \cdot \underline{k}^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{I}$	$= \frac{\underline{l} \cdot \underline{n} \cdot \underline{p} - 2 \cdot \underline{m} ^2}{\underline{g}^2 \cdot \underline{n} \cdot \underline{p} + \underline{l} \cdot \underline{n} \cdot \underline{p} - 2 \cdot \underline{m} ^2}$	$= \frac{L-3-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{J}$	$= \frac{\underline{m} \cdot \underline{n} \cdot \underline{p}}{\underline{g}^2 \cdot \underline{n} \cdot \underline{p} + \underline{l} \cdot \underline{n} \cdot \underline{p} - 2 \cdot \underline{m} ^2}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L+1-K(\mu))}$
$\underline{K}$	$= \frac{1}{1 - \underline{a} \cdot \underline{b}}$	$= \frac{-1}{(L-2-K(\mu)) \cdot (L-K(\mu))}$
$\underline{L}$	$= \frac{\underline{a} \cdot \underline{b}}{\underline{a} \cdot \underline{b} \cdot \underline{d} - \underline{d}}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L-K(\mu))}$
$\underline{M}$	$= \frac{\underline{i}}{\underline{d}^2 + \underline{i}}$	$= \frac{1}{L-K(\mu)}$
$\underline{N}$	$= \frac{1}{1 - \underline{a} \cdot \underline{c}}$	$= \frac{-1}{(L-2-K(\mu)) \cdot (L-K(\mu))}$
$\underline{O}$	$= \frac{\underline{a} \cdot \underline{c}}{\underline{a} \cdot \underline{c} \cdot \underline{g} - \underline{g}}$	$= \frac{L-1-K(\mu)}{(L-2-K(\mu)) \cdot (L-K(\mu))}$
$\underline{P}$	$= \frac{\underline{l}}{\underline{g}^2 + \underline{l}}$	$= \frac{1}{L-K(\mu)}$

Tabelle 3.3: Substitutionen in den Gleichungen (3.58) und (3.61). $\underline{a}$ – $\underline{p}$ siehe Tabelle 3.2.

den bei den unterschiedlichen Frequenzen verwendeten Matrizen $\underline{V}_\perp(\mu)$ berechnen. Falls man die nach dem oben beschriebenen Verfahren konstruierten Matrizen verwendet, die die Gleichungen (3.33) erfüllen, sind die in den Termen $\underline{\mathbb{A}}$ bis $\underline{\mathbb{J}}$ auftretenden Matrizen alle gleich. Es ergeben sich dann für die Abkürzungen $\underline{\mathbb{a}}$ bis $\underline{\mathbb{p}}$ immer dieselben Matrixspuren $L-1-K(\mu)$ und damit für die Terme $\underline{\mathbb{A}}$ bis $\underline{\mathbb{J}}$ die in der rechten Spalte der Tabelle 3.3 angegebenen Quotienten. Von den mit diesen Quotienten berechneten Schätzwerten für die Messwert(ko)varianzen kann man zeigen, dass der konkrete Schätzwert der Messwertvarianz jeweils größer oder gleich dem Betrag des konkreten Schätzwertes der Messwertkovarianz ist. Dazu setzt man in die vier Schätzwerte der Messwert(ko)varianzen in den Gleichungen (3.58) die Messwerte gemäß der Gleichungen (3.34) jeweils in der Form ein, die die Vektoren $\hat{\vec{N}}_f(\dots)$ enthält. Mit der für $L \geq 3+K(\mu)$ stets positiven Konstante $\alpha = M^2 \cdot (L-1-K(\mu))^2 \cdot (L-2-K(\mu)) \cdot (L+1-K(\mu))$ ist dann zu zeigen, dass

$$\begin{aligned}
 & \alpha \cdot \hat{C}_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = \tag{3.59a} \\
 & = -2 \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(\mu)^H + \\
 & \quad + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{\vec{N}}_f(\mu)^H + \\
 & \quad + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^H \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \geq \\
 & \geq \left| (L-3-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{\vec{N}}_f(\mu)^T + \right. \\
 & \quad \left. + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^T \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{\vec{N}}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \right| = \\
 & = \alpha \cdot \left| \hat{C}_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} \right|
 \end{aligned}$$

und

$$\begin{aligned}
 & \alpha \cdot \hat{C}_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = \tag{3.59b} \\
 & = -2 \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{\vec{N}}_f(\mu)^H + \\
 & \quad + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(\mu)^H + \\
 & \quad + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^H \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^H \geq \\
 & \geq \left| (L-3-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(\mu)^T + \right. \\
 & \quad \left. + (L-1-K(\mu)) \cdot \hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^T \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{\vec{N}}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})^T \right| = \\
 & = \alpha \cdot \left| \hat{C}_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} \right|
 \end{aligned}$$

$$\forall \quad \mu = 0 \left(\frac{M}{2 \cdot K_\Phi} \right) M-1 \quad \text{und} \quad \tilde{\mu} = 0 \quad (1) \quad K_\Phi - 1$$

gilt. Im Anhang A.9.1 wird gezeigt, dass diese Ungleichungen erfüllt sind. Dazu ist in Ungleichung (A.54) des Anhangs $a = L - 1 - K(\mu)$, $\vec{X} = \hat{N}_f(\mu)$ und $\vec{Y} = \hat{N}_f(\mu + \tilde{\mu} \cdot M/K_\Phi)^*$ bzw. $\vec{Y} = \hat{N}_f(-\mu - \tilde{\mu} \cdot M/K_\Phi)$ einzusetzen. Die Gleichung im Anhang gilt nur für $a \geq 2$, was hier bedeutet, dass die Mittelungsanzahl L mindestens $3 + K(\mu)$ sein muss.

Nun wollen wir den Fall behandeln, dass die negativen diskreten Frequenzen $-\mu$ und $-\mu - \tilde{\mu} \cdot M/K_\Phi$ *nicht* um ein ganzzahliges Vielfaches von M/K_Φ gegenüber μ verschoben sind, und somit die Bedingung (3.55a) nicht erfüllt ist. In diesem Fall sind nach den Gleichungen (2.41) und (2.51) viele der möglichen Kovarianzen zweier Zufallsgrößen, die man sich aus den Zufallsgrößen $\mathbf{N}_f(\mu)$, $\mathbf{N}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$, $\mathbf{N}_f(-\mu)$, $\mathbf{N}_f(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi})$ und den dazu konjugierten Zufallsgrößen herausgreift, bei Verwendung einer hoch frequenzselektiven Fensterfolge in guter Näherung null. Für diese Kovarianzen haben wir daher auch keine Messwerte berechnet. Daher verwenden wir nun die in Tabelle 3.1 aufgelisteten Substitutionen in den Gleichungen (A.42) — und nicht in den Gleichungen (A.40) wie im Fall (3.55a) —, um so eine reduzierte Vektorgleichung (A.41) zur Berechnung der vier gesuchten theoretischen Messwert(ko)varianzen und -kovarianzen zu erhalten. Jeweils die auf M normierte erste Zeile dieser reduzierten Vektorgleichung liefert uns die theoretischen Messwert(ko)varianzen:

$$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = E\{\mathbf{d}^{-1}\} \cdot \tilde{\Phi}_n(\mu, \mu) \cdot \tilde{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \quad (3.60a)$$

$$C_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = E\{\mathbf{i} \cdot \mathbf{d}^{-2}\} \cdot \tilde{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 \quad (3.60b)$$

$$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = E\{\mathbf{g}^{-1}\} \cdot \tilde{\Phi}_n(\mu, \mu) \cdot \tilde{\Phi}_n(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \quad (3.60c)$$

$$C_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = E\{\mathbf{l} \cdot \mathbf{g}^{-2}\} \cdot \tilde{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 \quad (3.60d)$$

$$\forall \quad \tilde{\mu} = 0 \text{ (1) } K_\Phi - 1 \quad \text{und} \quad \mu = 1 \text{ (1) } M - 1 \quad \text{ohne} \quad \mu = 0 \text{ (} \frac{M}{2 \cdot K_\Phi} \text{) } M - 1.$$

Die bei der Berechnung der Messwert(ko)varianzen auftretenden Terme $\mathbf{d}$, $\mathbf{g}$, $\mathbf{i}$ und $\mathbf{l}$ sind als Spuren der bei den unterschiedlichen Frequenzen verwendeten Matrizen $\mathbf{V}_\perp(\mu)$ zu berechnen, wie dies in Tabelle 3.2 angegeben ist. Alle Erwartungswerte in den Gleichungen (3.60), und somit auch die Messwertvarianzen, fallen asymptotisch indirekt proportional mit L (siehe Anhang A.7). Somit sind die Messwerte $\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ und $\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot M/K_\Phi)$ konsistent. Falls man die nach dem oben beschriebenen Verfahren konstruierten Matrizen $\mathbf{V}_\perp(\mu)$ verwendet, die die Gleichungen (3.33) erfüllen, sind die alle Erwartungswerte, die als Faktoren vor den Produkten, Quadraten und Betragsquadraten der theoretischen Werte des LDS bzw. KLDS auftreten, auch hier gleich $(L - 1 - K(\mu))^{-1}$ und somit von der Erregung unabhängig. Jeweils eine auf M normierte konkrete Realisierung des ersten Elementes des erwartungstreuen Zufallsvektors (A.47) liefert uns die

gesuchten konkreten Schätzwerte für die Messwert(ko)varianzen:

$$\hat{C}_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = [\mathbb{K}, \mathbb{L}] \cdot \begin{bmatrix} |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \\ \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) \end{bmatrix} \quad (3.61a)$$

$$\hat{C}_{\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = [\mathbb{M}] \cdot \hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 \quad (3.61b)$$

$$\hat{C}_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})} = [\mathbb{N}, \mathbb{O}] \cdot \begin{bmatrix} |\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \\ \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(-\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}, -\mu - \tilde{\mu} \cdot \frac{M}{K_\Phi}) \end{bmatrix} \quad (3.61c)$$

$$\hat{C}_{\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^*} = [\mathbb{P}] \cdot \hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})^2 \quad (3.61d)$$

$$\forall \quad \tilde{\mu} = 0 \quad (1) \quad K_\Phi - 1 \quad \text{und} \quad \mu = 1 \quad (1) \quad M - 1 \quad \text{ohne} \quad \mu = 0 \quad (\frac{M}{2 \cdot K_\Phi}) \quad M - 1.$$

Die bei der Berechnung der Schätzwerte der Messwert(ko)varianzen auftretenden Terme $\mathbb{K}$ bis $\mathbb{P}$ entnimmt man dem unteren Teil der Tabelle 3.3. Dabei lassen sich die dort auftretenden Abkürzungen $\mathbb{a}$, $\mathbb{b}$, $\mathbb{c}$, $\mathbb{d}$, $\mathbb{g}$, $\mathbb{i}$ und $\mathbb{l}$ mit Hilfe der Tabelle 3.2 aus den bei den unterschiedlichen Frequenzen verwendeten Matrizen $V_\perp(\mu)$ berechnen. Falls man die nach dem oben beschriebenen Verfahren konstruierten Matrizen verwendet, sind die in den Termen $\mathbb{K}$ bis $\mathbb{P}$ auftretenden Matrixspuren gemäß der Gleichungen (3.33b) alle gleich $L - 1 - K(\mu)$. Es ergeben sich dann für diese Terme die in der rechten Spalte der Tabelle 3.3 angegebenen Quotienten. Bei den mit diesen Quotienten berechneten Schätzwerten kann man zeigen, dass jeweils der konkrete Schätzwert der Messwertvarianz größer oder gleich dem Betrag des konkreten Schätzwertes der Messwertkovarianz ist, falls die Mittelungsanzahl L mindestens $3 + K(\mu)$ ist. Wir beginnen mit den Ungleichungen (3.36) und formen diese nach und nach um, bis wir die gewünschten Ungleichungen für die Schätzwerte der Messwert(ko)varianzen erhalten:

$$\begin{aligned} \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &\geq |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \\ (L - 1 - K(\mu)) \cdot \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) &\geq (L - 1 - K(\mu)) \cdot |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \\ (L - 1 - K(\mu)) \cdot \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) - |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 &\geq \\ &\geq (L - 1 - K(\mu)) \cdot |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 - |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \\ (L - 1 - K(\mu)) \cdot \hat{\Phi}_n(\mu, \mu) \cdot \hat{\Phi}_n(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}) - |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 &\geq \\ &\geq (L - 2 - K(\mu)) \cdot |\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})|^2 \end{aligned}$$

$$\begin{aligned}
& \frac{[-1, L-1-K(\mu)]}{(L-2-K(\mu)) \cdot (L-K(\mu))} \cdot \left[\frac{|\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2}{\hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} \right] \geq \\
& \geq \frac{L-2-K(\mu)}{(L-2-K(\mu)) \cdot (L-K(\mu))} \cdot |\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& \hat{C}_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} \geq |\hat{C}_{\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^*| \quad (3.62a)
\end{aligned}$$

und

$$\begin{aligned}
& \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \geq |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& (L-1-K(\mu)) \cdot \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \geq (L-1-K(\mu)) \cdot |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& (L-1-K(\mu)) \cdot \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) - |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \geq \\
& \geq (L-1-K(\mu)) \cdot |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 - |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& (L-1-K(\mu)) \cdot \hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) - |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \geq \\
& \geq (L-2-K(\mu)) \cdot |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& \frac{[-1, L-1-K(\mu)]}{(L-2-K(\mu)) \cdot (L-K(\mu))} \cdot \left[\frac{|\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2}{\hat{\Phi}_{\mathbf{n}}(\mu, \mu) \cdot \hat{\Phi}_{\mathbf{n}}(-\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}}, -\mu - \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} \right] \geq \\
& \geq \frac{L-2-K(\mu)}{(L-2-K(\mu)) \cdot (L-K(\mu))} \cdot |\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})|^2 \\
& \hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})} \geq |\hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^*| \quad (3.62b)
\end{aligned}$$

$$\forall \quad \tilde{\mu} = 0 \ (1) \ K_{\Phi} - 1 \quad \text{und} \quad \mu = 1 \ (1) \ M - 1 \quad \text{ohne} \quad \mu = 0 \ (\frac{M}{2 \cdot K_{\Phi}}) \ M - 1$$

Damit können wir für alle Messwerte Schätzwerte für deren Varianzen und Kovarianzen angeben. Auch hier war lediglich bei den Messwerten des LDS und des KLDS eine Verbundnormalverteilung angenommen worden, um die dort auftretenden vierten Momente auf deren bereits gemessenen zweiten Momente zurückführen zu können. Alle anderen Messwert(ko)varianzen konnten ohne die Kenntnis ihrer Verbundverteilung berechnet werden. Die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu_1, \mu_2)$ des bifrequenten LDS mit gleichen Frequenzen $\mu_1 = \mu_2$ sind die einzigen Messwerte die immer reell sind. Deren Kovarianzschätzwerte sind immer gleich ihren Varianzschätzwerten. Bei diesen Messwerten kann man, wenn man annimmt, dass sie normalverteilt sind, wieder mit Hilfe der komplementären Fehlerfunktion Konfidenzintervalle nach Gleichung ([1]:3.73) abschätzen, indem man in Gleichung ([1]:3.72)

die Schätzwerte $\hat{C}_{\hat{\Phi}_n(\mu, \mu), \hat{\Phi}_n(\mu, \mu)}$ der Messwertvarianzen einsetzt. Alle anderen Messwerte sind echt komplex. Wenn man von deren Real- und Imaginärteilen annimmt, dass sie verbundnormalverteilt sind, ergeben sich wieder Konfidenzellipsen, deren Halbachsen man mit Gleichung ([1]:3.80) abschätzen kann, wenn man dort die entsprechenden Messwertvarianz- und -kovarianzschätzwerte einsetzt.

4 Weitere Messwerte

4.1 Messwerte der zeitvarianten Impulsantworten

Aus den Abtastwerten $H(\mu, \mu + \hat{\mu} \cdot M/K_H)$ der sich bei der Lösung der Regression theoretisch ergebenden bifrequenten Übertragungsfunktion kann man durch eine DFT der Länge K_H bezüglich $\hat{\mu}$ und eine inverse DFT der Länge M bezüglich μ die mit M periodisch fortgesetzte, zeitvariante Impulsantwort $\tilde{h}_\kappa(k + \kappa)$ berechnen, die die zeitvariante Impulsantwort $h_\kappa(k + \kappa)$ nur dann vollständig beschreibt, wenn diese zeitlich auf das Intervall $k - \kappa \in [0, M - 1]$ begrenzt ist. Der Fall einer zeitinvarianten Impulsantwort $h_\kappa(k + \kappa) = h(k) \quad \forall k \in [0, E]$ ergibt sich im weiteren mit $K_H = 1$ und wird daher nicht separat betrachtet. Wenn wir auch das von dem konjugierten Eingangssignal erregte Modellsystem verwenden, berechnet sich die mit M periodisch fortgesetzte, zeitvariante Impulsantwort $\tilde{h}_{*,\kappa}(k + \kappa)$ analog aus den Abtastwerten der theoretischen Übertragungsfunktion $H_*(\mu, \mu + \hat{\mu} \cdot M/K_H)$. Aus den Messwerten $\hat{H}(\mu, \mu + \hat{\mu} \cdot M/K_H)$ und ggf. $\hat{H}_*(\mu, \mu + \hat{\mu} \cdot M/K_H)$ lassen sich in derselben Art Messwerte $\hat{\tilde{h}}_\kappa(k + \kappa)$ bzw. $\hat{\tilde{h}}_{*,\kappa}(k + \kappa)$ für die mit M periodisch fortgesetzten, zeitvarianten Impulsantworten berechnen.

$$\hat{\tilde{h}}_\kappa(k + \kappa) = \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} \hat{H}\left(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}\right) \cdot e^{-j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad (4.1a)$$

$$\hat{\tilde{h}}_{*,\kappa}(k + \kappa) = \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \sum_{\hat{\mu}=0}^{K_H-1} \hat{H}_*\left(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H}\right) \cdot e^{-j \cdot \frac{2\pi}{K_H} \cdot \hat{\mu} \cdot \kappa} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad (4.1b)$$

$$\forall \quad k + \kappa = 0 \text{ (1) } M - 1 \quad \text{und} \quad \kappa = 0 \text{ (1) } K_H - 1.$$

Diese Messwerte sind erwartungstreu, da einerseits die Messwerte der Übertragungsfunktionen erwartungstreu sind, und andererseits die zweidimensionale DFT eine Linearkombination mit konstanten Koeffizienten ist.

Die theoretischen Messwertvarianzen und -kovarianzen erhält man wieder, indem man jeweils den Erwartungswert des Betragsquadrats und des Quadrats der Messwertabweichung berechnet. Die Messwertabweichung lässt sich als zweidimensionale DFT der in Gleichung (3.18) angegebenen Messwertfehler der Übertragungsfunktionen berechnen. Wenn man eine hoch frequenzselektive Fensterfolge verwendet, kann man in den so berechneten

Vierfachsummen der theoretischen Messwert(ko)varianzen wieder diejenigen Summanden vernachlässigen, die eine der in den Gleichungen (2.41) und (2.51) genannten Kovarianzen $E\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f(\hat{\mu})^*\}$ oder $E\{\mathbf{N}_f(\mu) \cdot \mathbf{N}_f(-\hat{\mu})\}$ der Spektralwerte des Approximationsfehlerprozesses als Faktor enthalten. Bei den verbleibenden Summanden kann man die Spektralwertkovarianzen des Approximationsfehlers durch die M -fachen theoretischen Werte $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ ersetzen. Neben den Spektralwertkovarianzen des Approximationsfehlers treten in den Summanden auch noch die Erwartungswerte der Elemente des Produkts dreier Matrizen als Faktoren auf. Diese Matrizen sind zwei inverse empirische Kovarianzmatrizen, die von links und rechts an eine weitere Kovarianzmatrix heranmultipliziert werden. Die Elemente dieser Matrizen hängen nur von den Spektralwerten der Erregung ab. Verzichtet man bei diesen Faktoren auf die Erwartungswertbildung, indem man stattdessen wieder die konkreten Realisierungen dieser Faktoren verwendet, und schätzt man noch die theoretischen Werte des bifrequenten LDS und KLDS durch ihre Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ ab, so erhält man die Schätzwerte

$$\begin{aligned} \hat{C}_{\hat{\mathbf{h}}_{\kappa}(k+\kappa), \hat{\mathbf{h}}_{\kappa}(k+\kappa)} &= \frac{1}{(L-1) \cdot M} \cdot \\ &\cdot \sum_{\mu=0}^{\frac{M}{K_H}-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \vec{w}_{K_H, \kappa} \cdot \hat{C}_{\tilde{\mathbf{V}}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \tilde{\mathbf{V}}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^{-1} \cdot \hat{C}_{\tilde{\mathbf{V}}(\mu+\tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \hat{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \vec{w}_{K_H, \kappa}^H \cdot \\ &\cdot \sum_{\hat{\mu}=0}^{K_H-1} \hat{\Phi}_{\mathbf{n}}(\mu + \hat{\mu} \cdot \frac{M}{K_H}, \mu + \hat{\mu} \cdot \frac{M}{K_H} + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \\ &\quad \forall \quad k + \kappa = 0 \text{ (1) } M-1 \quad \text{und} \quad \kappa = 0 \text{ (1) } K_H-1 \end{aligned} \quad (4.2a)$$

für die Messwertvarianzen und die Schätzwerte

$$\begin{aligned} \hat{C}_{\hat{\mathbf{h}}_{\kappa}(k+\kappa), \hat{\mathbf{h}}_{\kappa}(k+\kappa)^*} &= \frac{1}{(L-1) \cdot M} \cdot \\ &\cdot \sum_{\mu=0}^{\frac{M}{K_H}-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \vec{w}_{K_H, \kappa}^* \cdot (\hat{C}_{\tilde{\mathbf{V}}(-\mu-\tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \tilde{\mathbf{V}}(-\mu-\tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^{-1})^* \cdot \hat{C}_{\tilde{\mathbf{V}}(-\mu-\tilde{\mu} \cdot \frac{M}{K_{\Phi}})^*, \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \hat{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}^{-1} \cdot \vec{w}_{K_H, \kappa}^H \cdot \\ &\cdot \sum_{\hat{\mu}=0}^{K_H-1} \hat{\Psi}_{\mathbf{n}}(\mu + \hat{\mu} \cdot \frac{M}{K_H}, \mu + \hat{\mu} \cdot \frac{M}{K_H} + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \\ &\quad \forall \quad k + \kappa = 0 \text{ (1) } M-1 \quad \text{und} \quad \kappa = 0 \text{ (1) } K_H-1. \end{aligned} \quad (4.2b)$$

für die Messwertkovarianzen der Impulsantwort des von $v(k)$ erregten Systems. Die Erwartungstreue dieser Schätzwerte lässt sich für die einzelnen Summanden — und somit für die gesamte Summe — analog zum Fall eines stationären Approximationsfehlers zeigen.

Der Vektor

$$\vec{w}_{K_H, \kappa} = \left[1, e^{j \cdot \frac{2\pi}{K_H} \cdot \kappa}, \dots, e^{-j \cdot \frac{2\pi}{K_H} \cdot (K_H-1) \cdot \kappa}, \vec{0} \right], \quad (4.3a)$$

enthält vorne die Drehfaktoren der inversen DFT der Länge K_H , sowie am Ende den Nullzeilenvektor $\vec{0}$ mit K_H Elementen. Damit werden aus dem zwischen den Vektoren liegenden Matrixprodukt der linke, obere Block herausgeschnitten. Die Messwert(ko)-varianzen der Impulsantwort des von $v(k)^*$ erregten Systems erhält man ebenfalls mit den Gleichungen (4.2) mit denselben Kovarianzmatrizen, allerdings ist hier der Vektor der DFT-Drehfaktoren *am Anfang* um den Nullzeilenvektor $\vec{0}$ mit K_H Elementen zu verlängern

$$\vec{w}_{K_H, \kappa} = \left[\vec{0}, 1, e^{j \cdot \frac{2\pi}{K_H} \cdot \kappa}, \dots, e^{-j \cdot \frac{2\pi}{K_H} \cdot (K_H-1) \cdot \kappa} \right], \quad (4.3b)$$

wodurch aus dem dazwischenliegenden Matrixprodukt der rechte, untere Block herausgeschnitten wird.

Es ist bei Verwendung von Messwerten $\hat{\Phi}_{\mathbf{n}}(\dots, \dots)$ und $\hat{\Psi}_{\mathbf{n}}(\dots, \dots)$, die mit Matrizen $V_{\perp}(\mu)$ gewonnen wurden, die die Bedingungen (3.33) erfüllen, zu vermuten, dass auch bei diesen Messwerten die Schätzwerte der Kovarianzen betraglich nicht größer werden als die Schätzwerte der Varianzen. Ein Beweis dafür wurden nicht erbracht. Er dürfte ähnlich ablaufen, wie bei den Messwert(ko)varianzen der deterministischen Störung in Kapitel A.9.2 des Anhangs.

4.2 Messwerte der Auto- und Kreuzkorrelationsfolgen

Unter der Voraussetzung, dass sowohl die Autokovarianzfolge $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\}$ als auch die Kreuzkovarianzfolge $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\}$ auf das Intervall $\kappa \in (-M/2; M/2)$ beschränkt ist, und dass man eine reelle Fensterfolge verwendet, bei der die zweidimensionale Fensterautokorrelationsfolge

$$d_k(\kappa) = \frac{K_{\Phi}}{M} \cdot \sum_{\tilde{k}=-\infty}^{\infty} f(k + \tilde{k} \cdot K_{\Phi})^* \cdot f(k + \kappa + \tilde{k} \cdot K_{\Phi}) \quad (4.4)$$

für $k = 0 \text{ (1) } K_{\Phi} - 1$ und $\kappa = 1 - M/2 \text{ (1) } M/2 - 1$ keine Nullstellen aufweist, kann man aus den Messwerten $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ erwartungstreue Schätzwerte für die beiden Kovarianzfolgen berechnen, indem man die Messwerte zunächst einer DFT der Länge K_{Φ} bezüglich $\tilde{\mu}$ und dann einer inversen DFT der Länge M bezüglich μ unterwirft. Anschließend wird das Ergebnis der zweidimensionalen DFT durch die Fensterautokorrelationsfolge $d_k(\kappa)$ dividiert.

$$\begin{aligned}
\hat{\phi}_{\mathbf{n}}(k+\kappa, k) &= \text{Schätzwert für } E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\} = \\
&\frac{1}{M \cdot d_k(\kappa)} \cdot \sum_{\mu=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_{\Phi}-1} \hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi}) \cdot e^{-j \cdot \frac{2\pi}{K_{\Phi}} \cdot \tilde{\mu} \cdot k} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot \kappa} \\
\forall \quad \kappa &= 1 - \frac{M}{2} \text{ (1) } \frac{M}{2} - 1 \quad \text{und} \quad k = 0 \text{ (1) } K_H - 1. \quad (4.5)
\end{aligned}$$

Die Berechnung der Messwerte $\hat{\psi}_{\mathbf{n}}(k+\kappa, k)$ für $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\}$ erfolgt analog. Es ist in der letzten Gleichung lediglich $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ durch $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ zu ersetzen. Die Erwartungstreue dieser Messwerte ergibt sich u. a. aus der Tatsache, dass die Multiplikation der Kovarianzfolgen $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\}$ und $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\}$ mit der Fensterautokorrelationsfolge $d_k(\kappa)$ gerade der in den Gleichungen (2.39) und (2.50) dargestellten zweidimensionalen Faltung der beiden Spektren der Fensterfolge mit dem zweidimensionalen LDS bzw. KLDS entspricht. Man erhält also die theoretischen Kovarianzfolgen durch dieselbe zweidimensionale DFT wie die gemessenen Folgen, indem man in die letzte Gleichung $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ bzw. $\tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ für die entsprechenden, erwartungstreuen Messwerte einsetzt.

Auch für die Messwerte der beiden Korrelationsfolgen kann man Schätzwerte für deren Varianz und Kovarianz aus den bisher berechneten Messwerten und Matrixspuren der Tabelle 3.2 gewinnen, wenn man wieder davon ausgeht, dass die Spektralwerte $\mathbf{N}_f(\mu)$ unterschiedlicher Frequenzen jeweils verbundnormalverteilte Quadrupel bilden. Auf eine ausführliche Darstellung der Herleitung dieser Schätzwerte wird verzichtet. Es soll nun lediglich die prinzipielle Vorgehensweise für deren Berechnung am Beispiel der Varianzen der Messwerte der Autokovarianzfolge $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\}$ verbal erläutert werden. Die Messwertvarianzen berechnen sich prinzipiell als die Differenz der Erwartungswerte der Betragsquadrate der Messwerte und der Betragsquadrate der Erwartungswerte der Messwerte. Die Erwartungswerte der Messwerte stimmen mit den theoretischen Werten überein. Wenn man die Messwerte nach Gleichung (3.29) in die Messwerte der Autokovarianzfolge nach Gleichung (4.5) einsetzt, und von diesen anschließend das Betragsquadrat bildet, erhält man eine Vierfachsumme, bei der jeder Summand das Produkt zweier bilinearer Formen ist. Der Erwartungswert dieser Vierfachsumme ist die Summe der Erwartungswerte der Summanden. Um die Erwartungswerte der Summanden berechnen zu können, bedarf es wieder einer Fallunterscheidung bezüglich der Indizes, die als Frequenzen in den Stichprobenvektoren des gefensterten Approximationsfehlerprozesses auftreten. Man muss nun vier Fälle unterscheiden. Im ersten Fall liegen alle vier diskreten Frequenzen im Raster $0 \left(\frac{M}{K_{\Phi}}\right) M-1$, oder alle vier im Raster $\frac{M}{2 \cdot K_{\Phi}} \left(\frac{M}{K_{\Phi}}\right) M-1$, und es können somit alle möglichen Kovarianzen zweier zufälliger Fehlerspektralwerte der

vier an dem Produkt der bilinearen Formen beteiligten Spektralwerte von null verschieden sein. Man kann bei diesen Summanden deren Erwartungswerte, sowie Schätzwerte für diese, mit Hilfe der Gleichungen (A.41) und (A.47) des Anhangs A.8 erhalten, indem man dort die Matrizen und Vektoren nach Gleichung (A.40) einsetzt, wobei man diese wiederum durch die Matrizen $\underline{\mathbf{V}}_{\perp}(\mu)$ und Vektoren $\vec{\mathbf{N}}_f(\mu)$ sowie die Zufallsgrößen $\mathbf{N}_f(\mu)$ in derselben Art wie bei der Berechnung der Messwertvarianzen des LDS mit Tabelle 3.1 substituiert. Man erhält so für den Erwartungswert eines Summanden jeweils die Summe dreier Produkte theoretischer Spektralwertkovarianzen. Im zweiten Fall liegen alle vier diskreten Frequenzen in einem Frequenzraster, das gegenüber dem Frequenzraster des ersten Falls um *kein* ganzzahliges Vielfaches von $M/2/K_{\Phi}$ verschoben ist. Dann sind einige der möglichen theoretischen Kovarianzen bei Verwendung einer hoch frequenzselektiven Fensterfolge in guter Näherung null, und in den Gleichungen (A.41) und (A.47) sind die Matrizen und Vektoren nach Gleichung (A.42) einzusetzen, so dass sich reduzierte Gleichungssysteme (A.39) ergeben. In dritten Fall liegen die beiden zufälligen Stichprobenvektoren der zufälligen Spektralwerte des Approximationsfehlerprozesses bei der ersten bilinearen Form in einem Frequenzraster des zweiten Falls, und die beiden Stichprobenvektoren der zweiten bilinearen Form in einem Frequenzraster, das gegenüber dem anderen Frequenzraster gerade die negativen Frequenzen enthält. Die Substitution der Matrizen und Vektoren gemäß der Gleichungen (A.43) in den Gleichungen (A.41) und (A.47) liefert dann für den Erwartungswert des Summanden die Summe zweier Produkte theoretischer Spektralwertkovarianzen. Im vierten Fall, wenn die beiden Stichprobenvektoren der ersten bilinearen Form ebenso wie die beiden Stichprobenvektoren der zweiten bilinearen Form in einem Frequenzraster mit dem Frequenzabstand M/K_{Φ} liegen, wobei die beiden Frequenzraster aber keine gemeinsamen Frequenzen enthalten, können nur sehr wenige der möglichen Spektralwertkovarianzen nennenswert von null verschieden sein. Es sind dann die Matrizen und Vektoren nach Gleichung (A.44) in die Gleichungen (A.41) und (A.47) einzusetzen. Summanden, die keinem der vier Fälle zuzuordnen sind, kommen in der Vierfachsumme nicht vor, da bei der DFT der Länge K_{Φ} bei der Berechnung der Messwerte der Autokorrelationsfolge nur solche bilineare Formen auftreten, bei denen die beiden Stichprobenvektoren in einem Frequenzraster mit dem Frequenzabstand M/K_{Φ} liegen. Bei allen Summanden der Vierfachsumme treten nur solche Messwerte für das LDS und KLDS auf, die bereits gemessen worden sind, und auch nur die Matrixspuren nach Tabelle 3.2, die man zur Abschätzung der Messwert(ko)varianzen sowieso schon berechnet hat. Ob die so berechneten Schätzwerte für die Kovarianzen der Messwerte der Korrelationsfolgen $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\}$ und $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\}$ betraglich niemals größer sind als die entsprechenden Varianzschätzwerte, wurde nicht untersucht.

5 Spektralschätzung mit zeitabhängigem ersten Moment

Wir wollen nun die Spektralschätzung des Zufallsprozesses $\mathbf{y}(k)$ als einen Sonderfall des RKM interpretieren. Diesen erhalten wir, wenn wir bei einem System ohne Eingang nur die spektralen Eigenschaften des Ausgangsprozesses vermessen wollen. Im System nach Bild 1.1 setzen wir folglich die Erregung zu null, und die Optimierungsparameter der Werte der beiden Impulsantworten treten nicht mehr in den zu minimierenden Termen nach Gleichung (2.4) auf. Unser Modellsystem besteht dann nur mehr aus der deterministischen Störung $u(k)$ und der zufälligen Störung $\mathbf{n}(k)$. Wir beschränken uns in diesem Kapitel auf den Fall, dass sich bei der Minimierung ein stationärer Approximationsfehlerprozess ergibt, für den wir wie in [1] dessen eindimensionales LDS und dessen ebenfalls eindimensionales KLDS durch die je M Werte $\tilde{\Phi}_{\mathbf{n}}(\mu)$ und $\tilde{\Psi}_{\mathbf{n}}(\mu)$ beschreiben wollen.

5.1 Spektralschätzung komplexer Prozesse

Die Minimierungsaufgabe lautet nun:

$$\mathbb{E}\{|\mathbf{n}(k)|^2\} = \mathbb{E}\{|\mathbf{y}(k) - u(k)|^2\} \stackrel{!}{=} \text{minimal} \quad \forall \quad k = 0 \text{ (1) } F-1. \quad (5.1)$$

Wenn wir diesen Ausdruck für jeden Zeitpunkt k jeweils nach dem Real- und Imaginärteil von $u(k)$ partiell ableiten, und die beiden sich dann ergebenden reellen Gleichungen wieder zu einer komplexen Gleichung zusammenfassen, erhalten wir mit

$$u(k) = \mathbb{E}\{\mathbf{y}(k)\} \quad (5.2)$$

das zeitabhängige erste Moment von $\mathbf{y}(k)$ als die F optimalen Regressionskoeffizienten. Die Werte $\tilde{\Phi}_{\mathbf{n}}(\mu)$ sind nur mehr von der Wahl der Werte der deterministischen Störung abhängig.

$$\begin{aligned} \tilde{\Phi}_{\mathbf{n}}(\mu) &= \frac{1}{M} \cdot \mathbb{E}\{|\mathbf{N}_f(\mu)|^2\} = \frac{1}{M} \cdot \mathbb{E}\left\{\left|\sum_{k=-\infty}^{\infty} \mathbf{n}(k) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}\right|^2\right\} = \\ &= \frac{1}{M} \cdot \mathbb{E}\left\{\left|\sum_{k=-\infty}^{\infty} (\mathbf{y}(k) - u(k)) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}\right|^2\right\} = \frac{1}{M} \cdot \mathbb{E}\{|\mathbf{Y}_f(\mu) - U_f(\mu)|^2\} \\ &\quad \forall \quad \mu = 0 \text{ (1) } M-1 \end{aligned} \quad (5.3)$$

Die Minimierung dieser Werte legt die M Optimierungsparameter $U_f(\mu)$ fest. Dieselben Werte $U_f(\mu)$ erhält man auch durch Fensterung und anschließende DFT aus den F optimalen Werten von $u(k)$, die das zweite Moment des Approximationsfehlers $\mathbf{n}(k)$ minimieren. Die Wahl der optimalen Parameter führt auch hier dazu, dass die ersten Momente $E\{\mathbf{n}(k)\}$ und $E\{\mathbf{N}_f(\mu)\}$ null werden.

Bei der Spektralschätzung muss die Nullstellenbedingung ([1]:2.27) für das Spektrum der Fensterfolge nicht mehr erfüllt werden. Es empfiehlt sich jedoch auch hier eine Fensterfolge zu verwenden, die der Bedingung ([1]:2.20) genügt, um auch im Spektralbereich die richtige Varianz als Mittelwert aller M Werte $\tilde{\Phi}_{\mathbf{n}}(\mu)$ zu erhalten. Da diese Forderung von der im Kapitel [1]:6 vorgestellten Fensterfolge ebenso erfüllt wird, wie auch die bei der Spektralschätzung gewünschte hohe Frequenzselektivität, ist die Verwendung der damit berechneten Fensterfolge immer dann zu empfehlen, wenn das zu messende LDS über der Frequenz stark schwankt.

Der Messwert $\hat{u}(k)$ der gefensterten deterministischen Störung zum Zeitpunkt k ist die Ausgleichslösung der Approximation der Stichprobe vom Umfang L des Ausgangssignals des realen Systems zum Zeitpunkt k durch den bezüglich der Einzelmessungen λ konstanten Wert $\hat{u}(k)$. Es ergibt sich der empirische Mittelwert

$$\hat{u}(k) = \frac{1}{L} \cdot \sum_{\lambda=1}^L y_{\lambda}(k) = \frac{1}{L} \cdot \vec{y}(k) \cdot \vec{1}^H \quad \forall \quad k = 0 \text{ (1) } F-1 \quad (5.4)$$

als Ausgleichslösung. Die Messwerte $\hat{U}_f(\mu)$ für das Spektrum der gefensterten deterministischen Störung erhalten wir wieder dadurch, dass wir den quadratischen Fehler bei der Approximation der Stichprobe des Spektrums des gefensterten Ausgangssignals des realen Systems durch die Werte der empirischen Regressionskoeffizienten $\hat{U}_f(\mu)$ minimieren. Als Ausgleichslösung erhalten wir die empirischen Mittelwerte der Spektralwerte des gefensterten Systemausgangsprozesses:

$$\hat{U}_f(\mu) = \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu) = \frac{1}{L} \cdot \vec{Y}_f(\mu) \cdot \vec{1}^H \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (5.5)$$

Die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ und $\hat{\Psi}_{\mathbf{n}}(\mu)$ erhalten wir durch Projektion des Stichprobenvektors $\vec{Y}_f(\mu)$ in einen Unterraum, der zum Vektor $\vec{1}$ orthogonal ist. Sinnvollerweise erfolgt diese Projektion mit Hilfe der Matrix $\underline{1}_{\perp}$ nach Gleichung (3.9), deren Spur $L-1$ ist. Als erwartungstreue Messwerte erhalten wir die empirischen Varianzen

$$\begin{aligned} \hat{\Phi}_{\mathbf{n}}(\mu) &= \frac{\hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^H}{M \cdot (L-1)} = \frac{\vec{N}_f(\mu) \cdot \underline{1}_{\perp} \cdot \vec{N}_f(\mu)^H}{M \cdot (L-1)} = \frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(\mu)^H}{M \cdot (L-1)} = \frac{1}{M} \cdot \hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu)} \\ &\quad \forall \quad \mu = 0 \text{ (1) } M-1 \end{aligned} \quad (5.6a)$$

und Kovarianzen

$$\hat{\Psi}_{\mathbf{n}}(\mu) = \frac{\hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(-\mu)^T}{M \cdot (L-1)} = \frac{\vec{N}_f(\mu) \cdot \underline{1}_{\perp} \cdot \vec{N}_f(-\mu)^T}{M \cdot (L-1)} = \frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(-\mu)^T}{M \cdot (L-1)} = \frac{1}{M} \cdot \hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(-\mu)^*}$$

$$\forall \quad \mu = 0 \ (1) \ M-1 \quad (5.6b)$$

der Spektralwerte des gefensterten Prozesses $\mathbf{y}(k)$, die die Ungleichung ([1]:3.38) immer erfüllen, und die ohne Zwischenspeicherung der Spektralwerte aller Einzelmessungen berechnet werden können.

Die Varianz und Kovarianz der Messwerte $\hat{U}_f(\mu)$ ergibt sich zu

$$C_{\hat{U}_f(\mu), \hat{U}_f(\mu)} = E\left\{|\hat{U}_f(\mu) - E\{\hat{U}_f(\mu)\}|^2\right\} = E\left\{|\hat{U}_f(\mu) - U_f(\mu)|^2\right\} = E\left\{\left|\vec{N}_f(\mu) \cdot \frac{\vec{1}^H}{L}\right|^2\right\} =$$

$$= E\left\{\text{spur}\left(\frac{\vec{1}^H \cdot \vec{1}}{L^2}\right)\right\} \cdot \left(E\{|\mathbf{N}_f(\mu)|^2\} - |E\{\mathbf{N}_f(\mu)\}|^2\right) - |E\{\mathbf{N}_f(\mu)\}|^2 =$$

$$= \frac{1}{L} \cdot E\{|\mathbf{N}_f(\mu)|^2\} = \frac{M}{L} \cdot \tilde{\Phi}_{\mathbf{n}}(\mu) \quad (5.7a)$$

und

$$C_{\hat{U}_f(\mu), \hat{U}_f(\mu)^*} = E\left\{(\hat{U}_f(\mu) - E\{\hat{U}_f(\mu)\})^2\right\} = E\left\{(\hat{U}_f(\mu) - U_f(\mu))^2\right\} = E\left\{\left(\vec{N}_f(\mu) \cdot \frac{\vec{1}^H}{L}\right)^2\right\} =$$

$$= E\left\{\text{spur}\left(\frac{\vec{1}^H \cdot \vec{1}}{L^2}\right)\right\} \cdot \left(E\{\mathbf{N}_f(\mu)^2\} - E\{\mathbf{N}_f(\mu)\}^2\right) - E\{\mathbf{N}_f(\mu)\}^2 =$$

$$= \frac{1}{L} \cdot E\{\mathbf{N}_f(\mu)^2\} \begin{cases} = \frac{M}{L} \cdot \tilde{\Psi}_{\mathbf{n}}(\mu) & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \approx 0 & \text{sonst,} \end{cases} \quad (5.7b)$$

und kann mit

$$\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)} = \frac{M}{L} \cdot \hat{\Phi}_{\mathbf{n}}(\mu) \quad \forall \quad \mu = 0 \ (1) \ M-1 \quad (5.8a)$$

und

$$\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)^*} = \frac{M}{L} \cdot \hat{\Psi}_{\mathbf{n}}(\mu) \quad \text{für} \quad \mu \in \{0; \frac{M}{2}\} \quad (5.8b)$$

erwartungstreu abgeschätzt werden, wobei die Kovarianzschätzwerte betraglich niemals größer als die Varianzschätzwerte sind. Daher kann man mit diesen Varianz- und Kovarianzschätzwerten die Schätzwerte der Halbachsen der Konfidenzellipsen nach Gleichung ([1]:3.80) für ein gewünschtes Konfidenzniveau angeben, wobei man dort die eben berechneten Schätzwerte für die Messwert(ko)varianzen statt der Messwert(ko)varianz-schätzwerte $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ und $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)^*}$ einsetzt.

Bei der Berechnung der Varianzen und Kovarianzen der Messwerte $\hat{u}(k)$ der gefensterten deterministischen Störung ist nun keine Näherung mehr nötig. Die Messwertabweichung

$$\begin{aligned} \hat{u}(k) - u(k) &= \frac{1}{L} \cdot \sum_{\lambda=1}^L y_{\lambda}(k) - u(k) = \frac{1}{L} \cdot \sum_{\lambda=1}^L (n_{\lambda}(k) + u(k)) - u(k) = \frac{1}{L} \cdot \sum_{\lambda=1}^L n_{\lambda}(k) \\ \forall \quad k &= 0 \text{ (1) } F-1 \end{aligned} \quad (5.9)$$

enthält nun in Gegensatz zu Gleichung (3.23) keine Terme mehr, die von den verrauschten Messwerten der beiden Übertragungsfunktionen abhängen. Daher ergibt sich bei einem stationären Approximationsfehlerprozess bei Verwendung einer Fensterfolge, die der Bedingung ([1]:2.20) genügt, die zeitunabhängige Messwertvarianz

$$C_{\hat{u}(k), \hat{u}(k)} = E\left\{|\hat{u}(k) - u(k)|^2\right\} = \frac{1}{L} \cdot E\{|\mathbf{n}(k)|^2\} = \frac{1}{L \cdot M} \cdot \sum_{\mu=0}^{M-1} \tilde{\Phi}_{\mathbf{n}}(\mu) \quad (5.10a)$$

und die ebenfalls zeitunabhängige Messwertkovarianz

$$C_{\hat{u}(k), \hat{u}(k)^*} = E\left\{(\hat{u}(k) - u(k))^2\right\} = \frac{1}{L} \cdot E\{\mathbf{n}(k)^2\} = \frac{1}{L \cdot M} \cdot \sum_{\mu=0}^{M-1} \tilde{\Psi}_{\mathbf{n}}(\mu). \quad (5.10b)$$

Als Summe erwartungstreuer Messwerte sind die Schätzwerte

$$\hat{C}_{\hat{u}(k), \hat{u}(k)} = \frac{1}{L \cdot M} \cdot \sum_{\mu=0}^{M-1} \hat{\Phi}_{\mathbf{n}}(\mu) \quad (5.11a)$$

und

$$\hat{C}_{\hat{u}(k), \hat{u}(k)^*} = \frac{1}{L \cdot M} \cdot \sum_{\mu=0}^{M-1} \hat{\Psi}_{\mathbf{n}}(\mu) \quad (5.11b)$$

ebenfalls erwartungstreu. Dass der Schätzwert der Kovarianz nie betragslich größer ist als der Schätzwert der Varianz, zeigt man indem man die Summe bei der Berechnung der Schätzwerte in Gleichung (5.11a) durch die Summe der arithmetischen Mittel der Messwerte $\hat{\Phi}_{\mathbf{n}}(-\mu)$ und $\hat{\Phi}_{\mathbf{n}}(\mu)$ ersetzt. Diese Summe ist immer größer als die Summe der entsprechenden geometrischen Mittel, deren Summanden nach Gleichung ([1]:3.38) größer als die Beträge von $\hat{\Psi}_{\mathbf{n}}(\mu)$ sind. Da der Betrag einer Summe kleiner oder gleich der Summe der Beträge der Summanden ist, ist der Betrag des Kovarianzschätzwertes nie größer als der Varianzschätzwert. Die Schätzwerte der zeitunabhängigen Halbachsen der Konfidenzellipse erhält man mit dem gewünschten Konfidenzniveau und den eben berechneten Messwert(ko)varianzschätzwerten, indem man diese statt der Messwert(ko)-varianzschätzwerte $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ und $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)^*}$ in Gleichung ([1]:3.80) einsetzt.

Bei der Berechnung der Varianzen und Kovarianzen der Messwerte $\hat{\Phi}_n(\mu)$ und $\hat{\Psi}_n(\mu)$ ist nun ein Rangdefekt von eins einzusetzen, weil nun die Matrix $\underline{1}_\perp$ statt der Matrix $\underline{V}_\perp(\mu)$, zur Berechnung der Messwerte verwendet wird. Da bei der Berechnung der Messwert(ko)varianzen in den Gleichungen ([1]:3.69) und ([1]:3.70) ein Rangdefekt von Zwei eingesetzt wurde, ist dort lediglich L durch $L+1$ zu substituieren. Wir erhalten daher die theoretischen Messwert(ko)varianzen

$$\begin{aligned} C_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} &= \frac{1}{L-1} \cdot |\tilde{\Psi}_n(\mu)|^2 + \frac{1}{L-1} \cdot \tilde{\Phi}_n(\mu)^2, \\ C_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)} &= \frac{2}{L-1} \cdot \tilde{\Phi}_n(\mu)^2 \\ \text{und } C_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*} &= \frac{2}{L-1} \cdot \tilde{\Psi}_n(\mu)^2 \\ \text{für } \mu &\in \left\{0; \frac{M}{2}\right\} \end{aligned} \quad (5.12a)$$

bzw.

$$\begin{aligned} C_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} &= \frac{1}{L-1} \cdot \tilde{\Phi}_n(\mu)^2, \\ C_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)} &= \frac{1}{L-1} \cdot \tilde{\Phi}_n(-\mu) \cdot \tilde{\Phi}_n(\mu) \\ \text{und } C_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*} &= \frac{1}{L-1} \cdot \tilde{\Psi}_n(\mu)^2 \\ \text{für } \mu &= 1 \ (1) \ M-1 \quad \wedge \quad \mu \neq \frac{M}{2}, \end{aligned} \quad (5.12b)$$

die wir mit

$$\begin{aligned} \hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} &= \frac{L-3}{(L+1) \cdot (L-2)} \cdot \hat{\Phi}_n(\mu)^2 + \frac{L-1}{(L+1) \cdot (L-2)} \cdot |\hat{\Psi}_n(\mu)|^2, \\ \hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)} &= \frac{2 \cdot (L-1)}{(L+1) \cdot (L-2)} \cdot \hat{\Phi}_n(\mu)^2 - \frac{2}{(L+1) \cdot (L-2)} \cdot |\hat{\Psi}_n(\mu)|^2 \\ \text{und } \hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*} &= \frac{2}{L+1} \cdot \hat{\Psi}_n(\mu)^2 \\ \text{für } \mu &\in \left\{0; \frac{M}{2}\right\} \end{aligned} \quad (5.13a)$$

bzw.

$$\begin{aligned} \hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} &= \frac{1}{L} \cdot \hat{\Phi}_n(\mu)^2, \\ \hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)} &= \frac{L-1}{L \cdot (L-2)} \cdot \hat{\Phi}_n(-\mu) \cdot \hat{\Phi}_n(\mu) - \frac{1}{L \cdot (L-2)} \cdot |\hat{\Psi}_n(\mu)|^2 \\ \text{und } \hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*} &= \frac{1}{L} \cdot \hat{\Psi}_n(\mu)^2 \\ \text{für } \mu &= 1 \ (1) \ M-1 \quad \wedge \quad \mu \neq \frac{M}{2} \end{aligned} \quad (5.13b)$$

erwartungstreu abschätzen. Auch hier sind die Schätzwerte der Kovarianz $\hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu), \hat{\Psi}_{\mathbf{n}}(\mu)}^*$ nie betragsmäßig größer als die Schätzwerte der Varianz $\hat{C}_{\hat{\Psi}_{\mathbf{n}}(\mu), \hat{\Psi}_{\mathbf{n}}(\mu)}$. Für die Konfidenzintervalle ([1]:3.73) der reellen Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ setzt man die Schätzwerte der halben Intervallbreite nach Gleichung ([1]:3.72) mit den eben berechneten Messwertvarianzschätzwerten ein. Die Schätzwerte der Halbachsen der Konfidenzellipsen der Messwerte $\hat{\Psi}_{\mathbf{n}}(\mu)$ erhält man mit dem gewünschten Konfidenzniveau und den eben berechneten Messwert(ko)varianzschätzwerten, indem man diese statt der Messwert(ko)varianzschätzwerte $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ und $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}^*$ in die Gleichungen ([1]:3.80) einsetzt.

5.2 Spektralschätzung reeller Prozesse

Die Spektralschätzung reeller Prozesse läuft im wesentlichen genau wie eben beschrieben ab. Auf einige Besonderheiten, die sich bei reellen Prozessen ergeben, soll nun hier eingegangen werden.

Bei einem reellen Zufallsprozess am Systemausgang erhalten wir als Optimallösung der Minimierung (5.1) eine reelle deterministische Modellstörung, die sich wie im Fall eines komplexwertigen Systems nach Gleichung (5.2) als das erste Moment $E\{\mathbf{y}(k)\}$ des zu messenden Prozesses ergibt. Somit erhalten wir einen reellen mittelwertfreien Modellzufallsvektor für den Approximationsfehlerprozess $\mathbf{n}(k) = \mathbf{y}(k) - u(k)$. Da bei einem reellen Modellzufallsvektor die Autokorrelationsfolge $E\{\mathbf{n}(k) \cdot \mathbf{n}(k+\kappa)\}$ und die Korrelationsfolge $E\{\mathbf{n}(k) \cdot \mathbf{n}(k+\kappa)\}$ identisch, reell und geradesymmetrisch sind, ist das KLDS $\Psi_{\mathbf{n}}(\Omega)$ reell und geradesymmetrisch und identisch mit dem LDS $\Phi_{\mathbf{n}}(\Omega)$. Auch die Näherungen

$$\begin{aligned}
 \tilde{\Phi}_{\mathbf{n}}(\mu) &= \frac{1}{M} \cdot E\{|\mathbf{N}_f(\mu)|^2\} = \frac{1}{M} \cdot E\{|\mathbf{N}_f(-\mu)|^2\} = \tilde{\Phi}_{\mathbf{n}}(-\mu) = \\
 &= \frac{1}{M} \cdot E\{\mathbf{N}_f(-\mu) \cdot \mathbf{N}_f(\mu)\} = \tilde{\Psi}_{\mathbf{n}}(\mu) = \frac{1}{M} \cdot E\left\{\left|\sum_{k=-\infty}^{\infty} \mathbf{n}(k) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}\right|^2\right\} = \\
 &= \frac{1}{M} \cdot E\left\{\left|\sum_{k=-\infty}^{\infty} (\mathbf{y}(k) - u(k)) \cdot f(k) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}\right|^2\right\} = \frac{1}{M} \cdot E\{|\mathbf{Y}_f(\mu) - U_f(\mu)|^2\} \\
 &\quad \forall \quad \mu = 0 \text{ (1) } \frac{M}{2}
 \end{aligned} \tag{5.14}$$

der entsprechenden Stufenapproximationen $\bar{\Phi}_{\mathbf{n}}(\mu)$ sind bei Verwendung einer reellen Fensterfolge identisch, reell und geradesymmetrisch, und nur von der Wahl der M Optimierungsparameter $U_f(\mu)$ abhängig. Diese ergeben sich mit

$$U_f(\mu) = U_f(-\mu)^* = E\{\mathbf{Y}_f(\mu)\} \quad \forall \quad \mu = 0 \text{ (1) } \frac{M}{2} \tag{5.15}$$

als die Mittelwerte der Spektralwerte des gemessenen und gefensterten Zufallsvektors am Systemausgang, und sind somit dieselben Werte, die man auch durch Fensterung und anschließende DFT aus den F optimalen Werten von $u(k)$, die das zweite Moment des Approximationsfehlers $\mathbf{n}(k)$ minimieren, erhält. Daher sind auch die Spektralwerte des gefensterten Approximationsfehlers mittelwertfrei.

Auch bei der Spektralschätzung reeller Prozesse muss die Nullstellenbedingung ([1]:2.27) für das Spektrum der Fensterfolge nicht mehr erfüllt werden. Es empfiehlt sich jedoch auch hier eine reelle Fensterfolge zu verwenden, die der Bedingung ([1]:2.20) genügt, um auch im Spektralbereich die richtige Varianz als Mittelwert aller M Werte $\tilde{\Phi}_{\mathbf{n}}(\mu)$ zu erhalten. Da diese Forderung von der im Kapitel 6 vorgestellten Fensterfolge ebenso erfüllt wird, wie auch die bei der Spektralschätzung gewünschte hohe Frequenzselektivität erzielt wird, ist die Verwendung der damit berechneten Fensterfolge immer dann zu empfehlen, wenn das zu messende LDS über der Frequenz stark schwankt.

Die Messwerte $\hat{u}(k)$ berechnen sich wie bei der Spektralschätzung eines komplexen Prozesses nach Gleichung (5.4) als die empirischen Mittelwerte des Approximationsfehlerprozesses, sind jedoch hier immer reell. Die Messwerte für das Spektrum der gefensterten deterministischen Störung nach Gleichung (5.5) sind dementsprechend die empirischen Mittelwerte des Spektrums des gefensterten Prozesses $\mathbf{y}(k)$. Sie sind symmetrisch ($\hat{U}_f(\mu) = \hat{U}_f(-\mu)^*$) und minimieren den quadratischen Fehler der Ausgleichslösung der Gleichungssysteme

$$\hat{U}_f(\mu) \cdot \vec{1} = \vec{Y}_f(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (5.16)$$

Die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ erhalten wir aus dem Vektor $\hat{\vec{N}}_f(\mu)$, der durch Projektion des Stichprobenvektors $\vec{Y}_f(\mu)$ in einen Unterraum, der zum Vektor $\vec{1}$ orthogonal ist, entsteht. Sinnvollerweise erfolgt die Projektion mit der idempotenten Matrix $\mathbf{1}_{\perp}$ nach Gleichung (3.9), deren Spur mit $L-1$ maximal ist. Wir erhalten damit die erwartungstreuen Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu) = \hat{\Phi}_{\mathbf{n}}(-\mu) = \hat{\Psi}_{\mathbf{n}}(\mu)$ gemäß Gleichung (5.6), die ohne Zwischenspeicherung der Spektralwerte aller Einzelmessungen berechnet werden können.

Die Varianz der Messwerte $\hat{U}_f(\mu)$ ergibt sich wie bei einem komplexwertigen System gemäß Gleichung (5.7a) und kann mit (5.8a) abgeschätzt werden. Die Messwerte $\hat{U}_f(\mu)$ sind für die beiden diskreten Frequenzen $\mu=0$ und $\mu=M/2$ immer reell. Daher ist für diese beiden Frequenzen die Messwertkovarianz gleich der Messwertvarianz, was sich auch in Gleichung (5.7b) mit $\hat{\Phi}_{\mathbf{n}}(\mu) = \hat{\Psi}_{\mathbf{n}}(\mu)$ ergibt. Man gibt daher für diese beiden Messwerte Konfidenzintervalle nach Gleichung ([1]:3.73) an, deren halbe Intervallbreiten sich nach Gleichung ([1]:3.72) mit dem eben angegebenen Schätzwerten der Messwertvarianzen abschätzen lassen. Bei allen anderen Frequenzen ist die Messwertkovarianz wegen der geforderten Stationarität des Prozesses $\mathbf{n}(k)$ in guter Näherung null. Die Radian

der sich ergebenden Konfidenzkreise der komplexen Messwerte schätzt man mit dem gewünschten Konfidenzniveau und den Schätzwerten der Messwertvarianzen ab, indem man in Gleichung ([1]:3.80c) die Messwertvarianzen $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ durch die eben berechneten Messwertvarianzen $\hat{C}_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu)}$ ersetzt.

Bei der Berechnung der Varianzen und Kovarianzen der Messwerte $\hat{\mathbf{u}}(k)$ der gefensterten deterministischen Störung ist auch beim reellwertigen System keine Näherung mehr erforderlich. Die Messwertabweichung ist nun reell, berechnet aber sich wie im Fall einer komplexen deterministischen Störung nach Gleichung (5.9). Bei einem stationären Approximationsfehlerprozess und bei Verwendung einer Fensterfolge, die der Bedingung ([1]:2.20) genügt, ergibt sich dieselbe zeitunabhängige Messwertvarianz nach Gleichung (5.10a) wie im komplexen Fall. Als Summe erwartungstreuer Messwerte ist der Schätzwert der zeitunabhängigen Messwertvarianz, die sich wieder nach Gleichung (5.11a) berechnet, ebenfalls erwartungstreu. Somit kann man hier ein zeitunabhängiges Konfidenzintervall nach Gleichung ([1]:3.73) abschätzen, wobei dort die halbe Intervallbreite nach Gleichung ([1]:3.72) mit dem gewünschten Konfidenzniveau und mit der nach Gleichung (5.11a) berechneten Messwertvarianz statt der Messwertvarianz $\hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)}$ abgeschätzt wird.

Bei der Berechnung der Varianzen der Messwerte $\hat{\Phi}_n(\mu)$ ist nun in den Gleichungen (5.12) die Gleichheit $\tilde{\Phi}_n(\mu) = \tilde{\Psi}_n(\mu)$ einzusetzen:

$$C_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} = \begin{cases} \frac{2}{L-1} \cdot \tilde{\Phi}_n(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \frac{1}{L-1} \cdot \tilde{\Phi}_n(\mu)^2 & \text{für } \mu = 1 \text{ (1) } \frac{M-1}{2}. \end{cases} \quad (5.17)$$

Da auch für die Messwerte $\hat{\Phi}_n(\mu) = \hat{\Psi}_n(\mu)$ gilt, schätzt man die Messwertvarianz mit den Gleichungen (5.13) als

$$\hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)} = \begin{cases} \frac{2}{L+1} \cdot \hat{\Phi}_n(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \frac{1}{L} \cdot \hat{\Phi}_n(\mu)^2 & \text{für } \mu = 1 \text{ (1) } \frac{M-1}{2} \end{cases} \quad (5.18)$$

ab. Damit lassen sich die Schätzwerte der Konfidenzintervalle nach Gleichung ([1]:3.72) und ([1]:3.73) angeben.

6 Reellwertige Systeme

In diesem Kapitel wird darauf eingegangen, wie das in Kapitel [1]:4 vorgestellte Messverfahren zur Messung reellwertiger realer Systeme benutzt werden kann, wenn auch die deterministische Störung $u(k)$ modelliert werden soll. Außerdem wird kurz auf drei weitere Varianten des RKM zur Messung reellwertiger Systeme eingegangen. Die in diesem Kapitel gemachten Untersuchungen beziehen sich auf den Fall *eines* zeitinvarianten Modellsystems ($K_H = 1$), bei dem das Modellsystem $\mathcal{S}_{*,lin}$ in Bild 1.1 weggelassen wird. Bei einem reellwertigen System macht es nämlich keinen Sinn, ein System zu modellieren, das von der konjugierten Erregung gespeist wird. Von dem Approximationsfehlerprozess wird angenommen, dass er stationär ist ($K_\Phi = 1$). Die Modifikationen, die sich bei einem periodisch zeitvarianten Modellsystem oder einem zyklstationären Approximationsfehlerprozess ergeben würden, möge sich der Leser anhand der in Kapitel 2 und 3 und der in diesem Kapitel angegebenen Überlegungen selbst herleiten.

6.1 Erste Variante des RKM zur Messung reellwertiger Systeme

Die Art der Erregung des Systems und der Fensterung und Fouriertransformation des Ausgangsprozesses bleibt gegenüber Kapitel [1]:4 unverändert. Auch die in den Gleichungen ([1]:4.1) und ([1]:4.2) genannten Symmetrien der Spektren bleiben erhalten. Der Approximationsfehler, dessen zweites Moment zur Systemapproximation minimiert wird, ist nun jedoch nach Gleichung (2.17) definiert, für die sich im Fall eines zeitinvarianten Modellsystems

$$\mathbf{n}(k) = \mathbf{y}(k) - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} H\left(\mu \cdot \frac{2\pi}{M}\right) \cdot \mathbf{V}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} - u(k), \quad (6.1)$$

ergibt. Der zu minimierende Term ([1]:2.10) ist ebenfalls um das Modell der deterministischen Störung zu erweitern:

$$E\{|\mathbf{n}(k)|^2\} = E\left\{\left|\mathbf{y}(k) - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} H\left(\mu \cdot \frac{2\pi}{M}\right) \cdot \mathbf{V}(\mu) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} - u(k)\right|^2\right\} \stackrel{!}{=} \text{minimal}. \quad (6.2)$$

Wenn man voraussetzt, dass die Varianzen aller M Zufallsgrößen $\mathbf{V}(\mu)$ von null verschieden sind, erhält man mit

$$H\left(\mu \cdot \frac{2\pi}{M}\right) = \frac{\mathbb{E}\left\{\left(\mathbf{V}(\mu) - \mathbb{E}\{\mathbf{V}(\mu)\}\right)^* \cdot \mathbf{Y}_f(\mu)\right\}}{\mathbb{E}\left\{\left|\mathbf{V}(\mu) - \mathbb{E}\{\mathbf{V}(\mu)\}\right|^2\right\}} \quad \forall \quad \mu = 0 \ (1) \ M-1 \quad (6.3)$$

die Lösung für die theoretischen Werte der Übertragungsfunktion, die hier ebenfalls die Symmetrie ([1]:4.3) aufweist. Damit zeigen auch die Optimalwerte des Spektrums der gefensterten deterministischen Störung nach Gleichung (2.29), für die sich hier

$$U_f(\mu) = \mathbb{E}\{\mathbf{Y}_f(\mu)\} - \mathbb{E}\{\mathbf{V}(\mu)\} \cdot H\left(\mu \cdot \frac{2\pi}{M}\right) \quad \forall \quad \mu = 0 \ (1) \ M-1 \quad (6.4)$$

ergibt, dieselbe Symmetrie und die deterministische Störung nach Gleichung (2.19), für die sich hier

$$u(k) = \mathbb{E}\{\mathbf{y}(k)\} - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} H\left(\mu \cdot \frac{2\pi}{M}\right) \cdot \mathbb{E}\{\mathbf{V}(\mu)\} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad k = 0 \ (1) \ F-1. \quad (6.5)$$

ergibt, ist reell. Somit ist auch der Approximationsfehlerprozess

$$\mathbf{n}(k) = \mathbf{y}(k) - \mathbf{x}(k) - u(k) \quad (6.6)$$

reell. Bei einem reellwertigen System liefert also die optimale Approximierung immer ein reellwertiges Modellsystem, dem sich ausgangsseitig eine reelle Störung überlagert, deren erstes Moment $u(k)$ ebenfalls reell ist. Zur vollständigen Beschreibung der zweiten zentralen Momente des Approximationsfehlerprozesses genügt daher die Angabe der reellen Autokorrelationsfolge, aus der man im Fall eines stationären Approximationsfehlerprozesses durch diskrete Fouriertransformation das gradesymmetrische reelle LDS $\Phi_{\mathbf{n}}(\Omega)$ gewinnt. Das bei einem komplexwertigen System noch anzugebende KLDS $\Psi_{\mathbf{n}}(\Omega)$ ist hier identisch mit dem LDS. Auf die Messung der Stufenapproximation $\bar{\Psi}_{\mathbf{n}}(\mu)$ oder deren Näherung $\tilde{\Psi}_{\mathbf{n}}(\mu)$ kann daher bei einem reellwertigen System verzichtet werden. Auch bei einem reellwertigen realen System liefert die Minimierung des zweiten Moments des Approximationsfehlers eine Lösung, bei der die Orthogonalität des Spektrums der Erregung und des Spektrums des gefensterten Approximationsfehlers gegeben ist, so dass auch hier Gleichung ([1]:2.30) erfüllt ist.

Um Schätzwerte für die optimalen Regressionskoeffizienten durch eine Messung zu erhalten, erregt man das System mit dem in Kapitel [1]:3 beschriebenen Verfahren, also mit L reellen Testsignalsequenzen, die bereichsweise mit M periodisch fortgesetzt sind. Diese dürfen hier jedoch mit einem Zufallsvektor $\vec{\mathbf{v}}$ erzeugt werden, der ein zeitabhängiges erstes Moment aufweist. Was man misst, ist nur der Realteil der L Stichprobenelemente des Ausgangssignals des realen Systems. Der Imaginärteil der L Systemausgangssignale

wird bei der Berechnung der Messwerte und ihrer Varianzen und Kovarianzen zu null gesetzt. Da jedes Stichprobenelement (= Signalabschnitt einer Einzelmessung) sowohl am Systemein- als auch am -ausgang reell ist, weisen beide Stichprobenvektoren $\hat{\vec{V}}(\mu)$ und $\hat{\vec{Y}}_f(\mu)$ die in Gleichung ([1]:4.4) genannte Symmetrie auf. Die M Gleichungssysteme (3.2) vereinfachen sich hier zu

$$\hat{H}(\mu) \cdot \vec{V}(\mu) + \hat{U}_f(\mu) \cdot \vec{1} = \vec{Y}_f(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (6.7)$$

Wir erhalten die M Ausgleichslösungen

$$\hat{H}(\mu) = \hat{C}_{\vec{Y}_f(\mu), \vec{V}(\mu)} \cdot \hat{C}_{\vec{V}(\mu), \vec{V}(\mu)}^{-1} \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (6.8)$$

und

$$\hat{U}_f(\mu) = \frac{1}{L} \cdot \left(\vec{Y}_f(\mu) - \hat{H}(\mu) \cdot \vec{V}(\mu) \right) \cdot \vec{1}^H \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (6.9)$$

mit den empirischen Varianzen

$$\begin{aligned} \hat{C}_{\vec{V}(\mu_1), \vec{V}(\mu_2)} &= \frac{1}{L-1} \cdot \vec{V}(\mu_1) \cdot \underline{1}_\perp \cdot \vec{V}(\mu_2)^H = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L V_\lambda(\mu_1) \cdot V_\lambda(\mu_2)^* - \frac{1}{L} \cdot \sum_{\lambda=1}^L V_\lambda(\mu_1) \cdot \sum_{\lambda=1}^L V_\lambda(\mu_2)^* \right) \\ &\quad \forall \quad \mu_1 = 0 \text{ (1) } M-1 \quad \text{und} \quad \mu_2 = 0 \text{ (1) } M-1, \end{aligned} \quad (6.10)$$

und den empirischen Kovarianzen

$$\begin{aligned} \hat{C}_{\vec{Y}_f(\mu_1), \vec{V}(\mu_2)} &= \frac{1}{L-1} \cdot \vec{Y}_f(\mu_1) \cdot \underline{1}_\perp \cdot \vec{V}(\mu_2)^H = \\ &= \frac{1}{L-1} \cdot \left(\sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot V_\lambda(\mu_2)^* - \frac{1}{L} \cdot \sum_{\lambda=1}^L Y_{f,\lambda}(\mu_1) \cdot \sum_{\lambda=1}^L V_\lambda(\mu_2)^* \right) \\ &\quad \forall \quad \mu_1 = 0 \text{ (1) } M-1 \quad \text{und} \quad \mu_2 = 0 \text{ (1) } M-1 \end{aligned} \quad (6.11)$$

Diese Ausgleichslösungen besitzen dann ebenfalls die Symmetrieeigenschaften

$$\hat{H}(\mu) = \hat{H}(-\mu)^* \quad \text{bzw.} \quad \hat{U}_f(\mu) = \hat{U}_f(-\mu)^*. \quad (6.12)$$

Daher genügt es bei geradem M — wovon wir im weiteren ausgehen — die jeweils $M/2 + 1$ Werte der Lösung für $\mu = 0 \text{ (1) } M/2$ zu berechnen. Mit den Ausgleichslösungen der Übertragungsfunktion erhält man die reellen Messwerte der deterministischen Störung:

$$\hat{u}(k) = \frac{1}{L} \cdot \vec{y}(k) \cdot \vec{1}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{H}(\mu) \cdot \vec{V}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \quad \forall \quad k = 0 \text{ (1) } F-1. \quad (6.13)$$

Da wir bisher keine Modifikationen im Messverfahren vorgenommen haben, braucht die Erwartungstreue der Messwerte $\hat{H}(\mu)$, $\hat{U}_f(\mu)$ und $\hat{u}(k)$ nicht gesondert gezeigt zu werden.

Die beim Einsetzen der Messwerte in die Gleichungssysteme (6.7) verbleibenden Fehlervektoren $\hat{\vec{N}}_f(\mu)$ der Ausgleichslösung, die sich wieder nach Gleichung (3.25) berechnen, sind ebenfalls in derselben Art symmetrisch wie alle anderen Stichprobenvektoren der Spektren der Signale am reellwertigen System.

Auch bei einem reellwertigen System verwenden wir für die Abschätzung der Stufenapproximation des LDS die immer reellen und erwartungstreuen Messwerte $\hat{\Phi}_n(\mu)$ nach Gleichung (3.29a). Dabei verwenden wir weiterhin die Matrix $\underline{V}_\perp(\mu)$, die hier nun die beiden Vektoren $\vec{1}$ und $\vec{V}(\mu)$ als Eigenvektoren zum Eigenwert Null haben muss:

$$\vec{V}(\mu) \cdot \underline{V}_\perp(\mu) = \vec{0} \quad \wedge \quad \vec{1} \cdot \underline{V}_\perp(\mu) = \vec{0}. \quad (6.14)$$

Bei der Konstruktion dieser Matrix nach Gleichung (3.28) ergibt sich nun jedoch ein wesentlicher Unterschied. Da bei einem reellwertigen System die beiden Zeilenvektoren $\vec{V}(\mu)$ und $\vec{V}(-\mu)^*$ der immer gleich sind, enthält die Matrix $\check{\underline{V}}(\mu)$ nur einen Zeilenvektor:

$$\check{\underline{V}}(\mu) = \vec{V}(\mu) \quad \forall \quad \mu = 0 \ (1) \ M-1 \quad (6.15)$$

Die weitere Berechnung der Matrix $\underline{V}_\perp(\mu)$ kann dann wieder mit Hilfe der Gleichungen (3.27) und (3.28) erfolgen und ergibt:

$$\begin{aligned} \underline{V}_\perp(\mu) &= \underline{V}_\perp(\mu)^n = \underline{V}_\perp(\mu)^H = \underline{1}_\perp - \frac{1}{L-1} \cdot \underline{1}_\perp \cdot \vec{V}(\mu)^H \cdot \hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}^{-1} \cdot \vec{V}(\mu) \cdot \underline{1}_\perp \\ &\forall \quad \mu = 0 \ (1) \ \frac{M}{2} \quad \wedge \quad n \in \mathbb{N}. \end{aligned} \quad (6.16)$$

Diese Matrix erfüllt die Gleichungen (3.33) und hat die Spur $L-2$. Indem wir die modifiziert konstruierte Matrix $\underline{V}_\perp(\mu)$ in die Gleichung (3.29a) einsetzen, erhalten wir mit $K(\mu)=1$ die in Gleichung (3.34a) angegebenen Messwerte:

$$\begin{aligned} \hat{\Phi}_n(\mu) &= \frac{\hat{\vec{N}}_f(\mu) \cdot \hat{\vec{N}}_f(\mu)^H}{M \cdot (L-2)} = \\ &= \frac{\vec{Y}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{Y}_f(\mu)^H}{M \cdot (L-2)} = \frac{\vec{N}_f(\mu) \cdot \underline{V}_\perp(\mu) \cdot \vec{N}_f(\mu)^H}{M \cdot (L-2)} = \\ &= \frac{1}{M} \cdot \frac{L-1}{L-2} \cdot \left(\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu)} - \hat{C}_{\mathbf{Y}_f(\mu), \mathbf{V}(\mu)} \cdot \hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}^{-1} \cdot \hat{C}_{\mathbf{Y}_f(\mu), \mathbf{V}(\mu)}^* \right) = \\ &= \frac{1}{M} \cdot \frac{L-1}{L-2} \cdot \left(\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu)} - |\hat{H}(\mu)|^2 \cdot \hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)} \right) \\ &\forall \quad \mu = 0 \ (1) \ M-1. \end{aligned} \quad (6.17)$$

Dabei wurden zuletzt die empirischen Varianzen nach Gleichung (3.35a) mit $\mu_1 = \mu_2$ verwendet.

Die Varianz der Messwerte der Übertragungsfunktion berechnet sich mit Gleichung (3.40). Der Vektor $\hat{\vec{H}}(\mu)$ enthält in unserem Fall nur ein Element, nämlich den Messwert $\hat{H}(\mu)$.

Auch die Kovarianzmatrix $\hat{\mathbf{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}$ ist hier nun eine skalare Größe, nämlich die nach Gleichung (6.10) berechnete Varianz $\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$ des Spektralwertes der Erregung. Der Einheitsvektor $\vec{E}_n$, dessen n -tes Element eins ist, entartet hier zu dem skalaren Wert 1. Damit ergibt sich die Messwertvarianz

$$C_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)} = \frac{M}{L-1} \cdot \mathbb{E}\{\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}^{-1}\} \cdot \tilde{\Phi}_n(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (6.18)$$

Die Messwerte der Übertragungsfunktion sind für die beiden diskreten Frequenzen $\mu=0$ und $\mu=M/2$ immer reell. Für diese beiden Messwerte ist die Messwertkovarianz nach Gleichung (3.43) gleich der eben berechneten Messwertvarianz. Für alle anderen Frequenzen ist die Messwertkovarianz in guter Näherung null. Als erwartungstreue Schätzwerte für die Messwertvarianz erhalten wir mit Gleichung (3.44a):

$$\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)} = \frac{M}{L-1} \cdot \hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}^{-1} \cdot \hat{\Phi}_n(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (6.19)$$

Gleichung (3.44b) liefert für die beiden Frequenzen $\mu=0$ und $\mu=M/2$ exakt dieselben Schätzwerte für die Messwertkovarianz. Für alle anderen Frequenzen wird die Messwertkovarianz mit null abgeschätzt. Damit wird für die beiden Frequenzen $\mu=0$ und $\mu=M/2$ die Länge der kürzeren Halbachse der Konfidenzellipse zu null. Es ist daher für diese beiden diskreten Frequenzen sinnvoll, statt der Konfidenzellipsen, Konfidenzintervalle analog zu den Konfidenzintervallen der LDS-Messwerte nach Gleichung ([1]:3.73) anzugeben. Die halbe Intervallbreite wird mit Gleichung ([1]:3.72) abgeschätzt, wobei hier die Schätzwerte der Varianzen von $\hat{\Phi}_n(\mu)$ durch die Schätzwerte der Varianzen von $\hat{\mathbf{H}}(\mu)$ zu ersetzen sind. Bei den komplexen Messwerten der Übertragungsfunktion aller anderen Frequenzen wird die Messwertkovarianz bei Verwendung einer hoch frequenzselektiven Fensterfolge gegenüber der Messwertvarianz wieder vernachlässigbar klein, so dass man auch beim reellwertigen System Konfidenzkreise erhält, deren Radius man mit Gleichung ([1]:3.80c) aus der Messwertvarianz und dem gewünschten Konfidenzniveau abschätzt.

Für die Messwerte $\hat{\mathbf{U}}_f(\mu)$ gilt dasselbe. Auch hier sind die Messwerte für $\mu=0$ und $\mu=M/2$ reell, so dass man Konfidenzintervalle angibt, deren halbe Intervallbreiten sich nach Gleichung ([1]:3.72) mit dem Schätzwerten der Messwertvarianzen nach Gleichung (3.48a) abschätzen lassen. In unseren Fall ergibt sich die Messwertvarianz mit Gleichung (3.47a)

zu:

$$C_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu)} = \frac{M}{L-1} \cdot \mathbb{E}\left\{\frac{\vec{\mathbf{V}}(\mu) \cdot \vec{\mathbf{V}}(\mu)^H}{L \cdot \hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}}\right\} \cdot \tilde{\Phi}_n(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (6.20)$$

Die Messwertvarianz wird minimal, wenn der Erwartungswert des Quotienten aus dem empirischen, nichtzentralen zweiten Moment $\vec{\mathbf{V}}(\mu) \cdot \vec{\mathbf{V}}(\mu)^H / L$ und dem empirischen, zentralen zweiten Moment $\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$ minimal wird. Daher sollte man, wenn möglich, als Erregung eine Zufallsprozess wählen, dessen Spektralwerte mittelwertfrei sind. Gleichung (3.48a) liefert die erwartungstreuen Schätzwerte

$$\hat{C}_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu)} = \frac{M}{L-1} \cdot \frac{\vec{V}(\mu) \cdot \vec{V}(\mu)^H}{L \cdot \hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}} \cdot \hat{\Phi}_{\mathbf{n}}(\mu) \quad \forall \quad \mu = 0 \text{ (1) } M-1. \quad (6.21)$$

Die Radien der Konfidenzkreise der komplexen Messwerte $\hat{\mathbf{U}}_f(\mu)$ aller anderen Frequenzen schätzt man mit dem gewünschten Konfidenzniveau und den Schätzwerten der Messwertvarianzen ab, indem man in Gleichung ([1]:3.80c) die Messwertvarianzen $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ durch die in Gleichung (6.21) berechneten Werte ersetzt.

Die Messwerte $\hat{\mathbf{u}}(k)$ sind für alle Zeitpunkte $k = 0 \text{ (1) } F-1$ reell. Mit Gleichung (3.52) ergibt sich die zeitunabhängige Näherung

$$C_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} \approx \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} C_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu)} \quad \forall \quad k = 0 \text{ (1) } F-1 \quad (6.22)$$

für die Varianz der Messwerte $\hat{\mathbf{u}}(k)$, die mit Gleichung (3.53a) aus den Schätzwerten der Varianzen der Messwerte $\hat{\mathbf{U}}_f(\mu)$ zu

$$\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} = \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \hat{C}_{\hat{\mathbf{U}}_f(\mu), \hat{\mathbf{U}}_f(\mu)} \quad \forall \quad k = 0 \text{ (1) } F-1, \quad (6.23)$$

abgeschätzt werden kann. Somit kann man hier ein zeitunabhängiges Konfidenzintervall nach Gleichung ([1]:3.73) abschätzen, wobei dort die halbe Intervallbreite nach Gleichung ([1]:3.72) mit der Messwertvarianz nach Gleichung (6.23) eingesetzt wird.

Die Varianzen der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ lassen sich mit den Gleichungen (3.56a) und (3.60a) berechnen. Da wir Matrizen $\mathbf{V}_{\perp}(\mu)$ verwenden, die die Gleichungen (3.33) erfüllen, sind die in den Gleichungen (3.56a) und (3.60a) vor den Produkten und Betragsquadraten der theoretischen Werte des LDS bzw. KLDS als Vorfaktoren auftretenden Erwartungswerte alle gleich $(L-1-K(\mu))^{-1} = (L-2)^{-1}$. Unter Berücksichtigung der für reelle Prozesse immer gültigen Beziehung $\tilde{\Phi}_{\mathbf{n}}(\mu) = \tilde{\Phi}_{\mathbf{n}}(-\mu) = \tilde{\Psi}_{\mathbf{n}}(\mu)$ erhalten wir

$$C_{\hat{\Phi}_{\mathbf{n}}(\mu), \hat{\Phi}_{\mathbf{n}}(\mu)} = \begin{cases} \frac{2}{L-2} \cdot \tilde{\Phi}_{\mathbf{n}}(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \frac{1}{L-2} \cdot \tilde{\Phi}_{\mathbf{n}}(\mu)^2 & \text{für } \mu = 1 \text{ (1) } \frac{M-1}{2}. \end{cases} \quad (6.24)$$

Wir verwenden die in den Gleichungen (3.58a) und (3.61a) angegebenen erwartungstreuen Schätzwerte, für die sich in unserem Fall

$$\hat{C}_{\hat{\Phi}_{\mathbf{n}}(\mu), \hat{\Phi}_{\mathbf{n}}(\mu)} = \begin{cases} \frac{2}{L} \cdot \hat{\Phi}_{\mathbf{n}}(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \frac{1}{L-1} \cdot \hat{\Phi}_{\mathbf{n}}(\mu)^2 & \text{für } \mu = 1 \text{ (1) } \frac{M-1}{2} \end{cases} \quad (6.25)$$

ergibt. Mit deren Hilfe schätzt man die halbe Breite der Konfidenzintervalle nach Gleichung ([1]:3.73) mit Gleichung ([1]:3.72) ab.

6.2 Weitere Varianten zur Messung reellwertiger Systeme

Es wurden drei weitere Varianten des RKM zur Messung reellwertiger Systeme untersucht. Diese sollen nun kurz vorgestellt werden, ohne dabei alle Messwerte, deren Erwartungswerte, Varianzen und Kovarianzen explizit anzugeben.

Bei allen drei Varianten werden die am realen reellwertigen System bei aufeinanderfolgenden Einzelmessungen auftretenden Ein- und Ausgangssignale jeweils als Real- und Imaginärteil zweier komplexer Signalsequenzen zusammengefasst, und jeweils als eine konkrete Realisierung zweier entsprechender komplexer Zufallsvektoren am Ein- und Ausgang des Systems angesehen. Mit diesen komplexen Zufallsvektoren kann nun die Approximation mit einem komplexwertigen Modellsystem, mit einer komplexen deterministischen Störung und mit einem komplexen Approximationsfehler vorgenommen werden, wie dies in Kapitel 2.2 geschehen ist. Dabei muss die Erregung des Systems und die gegebenenfalls zufällige Auswahl der Zeitintervalle der Einzelmessungen so gewählt werden, dass sich Optimallösungen der theoretischen komplexen Regression ergeben, die sich in die Optimallösungen der theoretischen reellen Regression der in Unterkapitel 6.1 geschilderten Systemapproximation mit reellen Zufallsvektoren überführen lassen. Diese Optimallösungen sind schließlich die Größen, die man zu messen wünscht. Wenn schon die theoretischen Optimallösungen der komplexen Regression abweichen, wird man nicht erwarten können, dass die Erwartungswerte der Messergebnisse — also der empirisch bestimmten optimalen Regressionskoeffizienten — die richtigen Optimallösungen der reellen Regression sein werden. Desweiteren muss sich ein Zufallsvektor für den komplexen Approximationsfehlerprozess ergeben, dessen Real- und Imaginärteil unabhängig sind, wobei diese beiden Anteile dieselbe Verbundverteilung besitzen müssen, wie der Zufallsvektor des reellen Approximationsfehlerprozesses bei der Systemapproximation in Unterkapitel 6.1. Bei der reellen Regression waren die Stichprobenelemente des Zufallsvektors des reellen Approximationsfehlerprozesses zweier aufeinanderfolgender Einzelmessungen unabhängig. Da die beiden Stichprobenelemente zweier aufeinanderfolgender Einzelmessungen nun als Real- und Imaginärteil eines komplexen Stichprobenelements betrachtet werden, müssen die Real- und Imaginärteile des komplexen Zufallsvektors unabhängig sein und die gleiche Verbundverteilung besitzen wie der reelle Zufallsvektor, wenn man erwarten will, dass die aus den Stichproben berechneten Messwerte der stochastischen Eigenschaften des Approximationsfehlerprozesses die gewünschten Erwartungswerte besitzen. Nur wenn sich eine Erregung finden lässt, die diese Forderungen erfüllt, kann man bei den drei folgenden Varianten des RKM erwarten, dass die Messergebnisse zu einer adäquaten Beschreibung des realen Systems im typischen Betriebszustand geeignet sind. Bei den ersten beiden Varian-

ten kann dies gegebenenfalls unter Beachtung der zufälligen Auswahl der Zeitintervalle der Messung erreicht werden, indem man den Real- und Imaginärteil der Erregung als eine konkrete Stichprobe aus einem Zufallsvektor mit unabhängigen Real- und Imaginärteilstichproben gleicher Verbundverteilung gewinnt. Man gewinnt also die Testsignalsequenzen auf dieselbe Weise, wie bei dem in Unterkapitel 6.1 dargestellten Messverfahren, interpretiert aber die Signalsequenzen aufeinanderfolgender Einzelmessungen nicht mehr als zwei vollständige Signale sondern als Real- und Imaginärteil eines einzigen Signals. Bei der dritten Variante ist dies nicht möglich, da man hier — wie wir noch sehen werden — eine spezielle komplexe Erregung, die aus abhängigen Real- und Imaginärteilstichproben gewonnen wird, verwendet.

Im weiteren werden die theoretisch optimalen Regressionskoeffizienten der Systemapproximation mit dem reellwertigen Modellsystem nach Unterkapitel 6.1 mit dem Index reell gekennzeichnet, während die bisher verwendeten Formelzeichen die entsprechenden Größen der Systemapproximation mit den komplexen Zufallsvektoren bezeichnen. Bei allen Varianten ergibt sich unter den eben genannten Voraussetzungen ein theoretisches komplexes Modellsystem, bei dem die Anteile der Übertragungsfunktion, die die Real- und Imaginärteile der komplexen Zufallsvektoren des Ein- und Ausgangs kreuzweise verknüpfen, null sind, weil derartige Verknüpfungen bei einer Systemapproximation mit reellen Zufallsvektoren nicht vorkommen können, weil die aufeinanderfolgenden Einzelmessungen unabhängig sind. Die Verknüpfung der Realteile der komplexen Zufallsvektoren des Ein- und Ausgangs ist identisch mit der Verknüpfung der Imaginärteile der komplexen Zufallsvektoren des Ein- und Ausgangs. Die wiederum ist identisch mit der Verknüpfung $H_{\text{reell}}(\Omega)$ der reellen Zufallsvektoren des Ein- und Ausgangs bei der Systemapproximation mit reellen Signalen. Daher erhält man bei allen drei weiteren RKM-Varianten zur Messung reellwertiger Systeme ein komplexwertiges Modellsystem mit

$$H(\Omega) = H_{\text{reell}}(\Omega) = H_{\text{reell}}(-\Omega)^* = H(-\Omega)^*, \quad (6.26)$$

wenn die oben genannten Voraussetzungen eingehalten werden. Da die Regressionskoeffizienten $u_{\text{reell}}(k)$ der deterministischen Störung bei der Regression mit den reellen Zufallsvektoren bei dem komplexen Systemmodell im Real- und Imaginärteil des komplexen Zufallsvektors des Systemausgangs identisch vorhanden sind, gilt für die Optimallösungen der Regressionskoeffizienten $u(k)$ der Regression mit den komplexen Modellzufallsvektoren

$$\begin{aligned} u(k) &= u_{\text{reell}}(k) + j \cdot u_{\text{reell}}(k) = (1+j) \cdot u_{\text{reell}}(k) = \\ &= (1+j) \cdot u_{\text{reell}}(k)^* = j \cdot (1-j) \cdot u_{\text{reell}}(k)^* = j \cdot ((1+j) \cdot u_{\text{reell}}(k))^* = j \cdot u(k)^*. \end{aligned} \quad (6.27)$$

Für die Regressionskoeffizienten $U_f(\mu)$ des Spektrums der gefensterten komplexen deterministischen Störung gilt entsprechend

$$\begin{aligned}
U_f(\mu) &= U_{f,\text{reell}}(\mu) + j \cdot U_{f,\text{reell}}(\mu) = (1+j) \cdot U_{f,\text{reell}}(\mu) = \\
&= (1+j) \cdot U_{f,\text{reell}}(-\mu)^* = j \cdot ((1+j) \cdot U_{f,\text{reell}}(-\mu))^* = j \cdot U_f(-\mu)^*. \quad (6.28)
\end{aligned}$$

Wenn die Voraussetzungen für die Anwendbarkeit der drei RKM-Varianten gegeben sind, gilt für die zweiten Momente des Spektrums des gefensterten komplexen Approximationsfehlers

$$\tilde{\Phi}_{\mathbf{n}}(\mu) = \tilde{\Phi}_{\mathbf{n}}(-\mu) = 2 \cdot \tilde{\Phi}_{\mathbf{n},\text{reell}}(\mu) = 2 \cdot \tilde{\Psi}_{\mathbf{n},\text{reell}}(\mu) \quad (6.29a)$$

und $\tilde{\Psi}_{\mathbf{n}}(\mu) = 0,$ (6.29b)

wobei $\tilde{\Phi}_{\mathbf{n},\text{reell}}(\mu) = \tilde{\Psi}_{\mathbf{n},\text{reell}}(\mu)$ die entsprechenden zweiten Momente des Spektrums des gefensterten reellen Approximationsfehlers sind, der sich bei der Approximation des realen reellwertigen Systems mit den reellen Zufallsvektoren ergibt.

Bei der ersten der drei weiteren RKM-Varianten wird die Messung nun genauso durchgeführt, wie wenn ein komplexwertiges System vorliegen würde. Man berechnet also zunächst die Messwerte, die im Kapitel 3 angegeben worden sind. Da die konkreten Realisierungen der Signale der Einzelmessungen nun zu komplexen Signalsequenzen zusammengefasst werden, weisen deren Spektren nicht mehr die Symmetrien auf, die den Spektren reeller Signale eigen sind. Somit weisen auch die Messwerte nicht die Symmetrien der entsprechenden theoretischen Größen auf. Lediglich die Erwartungswerte der Messwerte entsprechen diesen Symmetrien, da die Erwartungstreue der Messwerte weiterhin gegeben ist. Man berechnet sich daher aus den bisher gewonnenen Messwerten des komplexen Systemmodells Messwerte für die eigentlich gesuchten Regressionskoeffizienten des reellen Systemmodells, die die gewünschten Symmetrien aufweisen. Man berechnet sich daher die Messwerte

$$\hat{H}_{\text{reell}}(\mu) = \frac{\hat{H}(\mu) + \hat{H}(-\mu)^*}{2}, \quad (6.30a)$$

$$\hat{U}_{f,\text{reell}}(\mu) = \frac{1}{2} \cdot \left(\frac{\hat{U}_f(\mu)}{1+j} + \frac{j \cdot \hat{U}_f(-\mu)^*}{1+j} \right) = \frac{1}{4} \cdot \left((1-j) \cdot \hat{U}_f(\mu) + (1+j) \cdot \hat{U}_f(-\mu)^* \right), \quad (6.30b)$$

$$\hat{u}_{\text{reell}}(k) = \frac{1}{2} \cdot \left(\frac{\hat{u}(k)}{1+j} + \frac{j \cdot \hat{u}(k)^*}{1+j} \right) = \frac{1}{4} \cdot \left((1-j) \cdot \hat{u}(k) + (1+j) \cdot \hat{u}(k)^* \right) \quad (6.30c)$$

und

$$\hat{\Phi}_{\mathbf{n},\text{reell}}(\mu) = \frac{\hat{\Phi}_{\mathbf{n}}(\mu) + \hat{\Phi}_{\mathbf{n}}(-\mu)}{4}. \quad (6.30d)$$

Wenn man von diesen Messwerten nun die Varianzen berechnet, wird man feststellen, dass man bestenfalls dieselben Messwertvarianzen erhält, die man bei der in Unterkapitel 6.1 dargestellten Version des RKM zur Messung reellwertiger Systeme erhält, wenn man dort insgesamt eine Einzelmessung am realen System weniger durchführt. Diese minimale Messwertvarianz erhält man, wenn die Spektralwerte der komplexen Erregung bei positiven und negativen Frequenzen gleiche Varianz aufweisen, und unabhängig sind. Bei

dieser Variante des RKM mit der in den letzten Gleichungen dargestellten abschließenden Messwertmittelung ist für je zwei Einzelmessungen am realen System sowohl eine DFT des komplexen Eingangssignals als auch eine DFT des komplexen Ausgangssignals zu berechnen. In Kapitel 8 wird gezeigt, dass man auch bei der in Unterkapitel 6.1 vorgestellten RKM-Variante mit derselben Anzahl an diskreten Fouriertransformationen auskommt. Da die letztgenannte Variante jedoch bei minimal besserer Messwertvarianz einen geringeren Speicherbedarf aufweist, und deren Anwendbarkeit auch nicht solch starken Restriktionen hinsichtlich der nicht immer gegebenen Unabhängigkeit der Einzelmessungen unterliegt, ist diese vorzuziehen.

Bei einer weiteren RKM Variante nützt man die Symmetrien, die in den Regressionskoeffizienten des komplexen Systemmodells vorhanden sein müssen dadurch aus, dass man diese schon bei der Aufstellung der Gleichungen, für die die Ausgleichslösungen berechnet werden sollen, berücksichtigt. Mit den komplexen Stichprobenvektoren

$$\vec{V}(\mu) = \left[V_1(\mu) + j \cdot V_2(\mu), \dots, V_{2\tilde{\lambda}-1}(\mu) + j \cdot V_{2\tilde{\lambda}}(\mu), \dots, V_{L/2-1}(\mu) + j \cdot V_{L/2}(\mu) \right] \quad (6.31a)$$

und

$$\vec{Y}_f(\mu) = \left[Y_{f,1}(\mu) + j \cdot Y_{f,2}(\mu), \dots, Y_{f,2\tilde{\lambda}-1}(\mu) + j \cdot Y_{f,2\tilde{\lambda}}(\mu), \dots, Y_{f,L/2-1}(\mu) + j \cdot Y_{f,L/2}(\mu) \right] \quad (6.31b)$$

$$\forall \quad \mu = 0 \text{ (1) } M-1$$

die jeweils aus $L/2$ Elementen bestehen, erhält man zunächst für alle M Frequenzen die Gleichungen (6.7), die jetzt aber alle die Dimension $1 \times (L/2)$ haben, und die M Unbekannten $\hat{H}(\mu)$ sowie die M Unbekannten $\hat{U}_f(\mu)$ des komplexen Modellsystems enthalten. Indem man hier μ durch $-\mu$ ersetzt und die Gleichungen konjugiert, erhält man mit

$$\hat{H}(-\mu)^* \cdot \vec{V}(-\mu)^* + \hat{U}_f(-\mu)^* \cdot \vec{1} = \vec{Y}_f(-\mu)^* \quad (6.32)$$

ebenfalls wieder M redundante Gleichungen der Dimension $1 \times (L/2)$ mit denselben Unbekannten. Nun fordert man, dass die empirischen Regressionskoeffizienten dieselben Symmetrien aufweisen müssen, wie die theoretische Regressionskoeffizienten nach Gleichung (6.27) und (6.29). Man setzt daher in die vorigen Gleichungen und die Gleichungen (6.7) die Terme

$$\hat{H}(\mu) = \hat{H}(-\mu)^* = \hat{H}_{\text{reell}}(\mu) \quad (6.33a)$$

$$\text{und} \quad \hat{U}_f(\mu) = j \cdot \hat{U}_f(-\mu)^* = (1+j) \cdot \hat{U}_{f,\text{reell}}(\mu) \quad (6.33b)$$

ein, und fasst dann alle Gleichungen mit denselben Unbekannten zusammen. Man erhält so die $M/2+1$ Gleichungen:

$$\hat{H}_{\text{reell}}(\mu) \cdot \left[\vec{V}(\mu), \vec{V}(-\mu)^* \right] + \hat{U}_{f,\text{reell}}(\mu) \cdot \left[(1+j) \cdot \vec{1}, (1-j) \cdot \vec{1} \right] = \left[\vec{Y}_f(\mu), \vec{Y}_f(-\mu)^* \right] \quad (6.34)$$

$$\forall \quad \mu = 0 \text{ (1) } \frac{M}{2}.$$

Die $M/2-1$ Gleichungen für $\mu = \frac{M}{2}+1$ (1) $M-1$ lassen sich durch Konjugieren und Vertauschen der Reihenfolge der Elemente Zeilenvektoren in die angegebenen Gleichungen überführen und sind daher redundant. Mit diesen Gleichungen kann man nun die gesuchten Messwerte als Ausgleichslösungen bestimmen. In den F Gleichungen (6.13) für die F Unbekannten $\hat{u}(k)$ der deterministischen Störung des komplexen Modellsystems ist die Mittelungsanzahl $L/2$ statt L einzusetzen. Außerdem ist dort für die komplexen empirischen Mittelwerte

$$\frac{2 \cdot \vec{V}(\mu) \cdot \vec{1}^H}{L} = \frac{2}{L} \cdot \sum_{\tilde{\lambda}=1}^{L/2} \left(V_{2 \cdot \tilde{\lambda}-1}(\mu) + j \cdot V_{2 \cdot \tilde{\lambda}}(\mu) \right) \quad \forall \quad \mu = 0 \text{ (1) } M-1 \quad (6.35a)$$

und

$$\frac{2}{L} \cdot \vec{y}(k) \cdot \vec{1}^H = \frac{2}{L} \cdot \sum_{\tilde{\lambda}=1}^{L/2} \left(y_{2 \cdot \tilde{\lambda}-1}(k) + j \cdot y_{2 \cdot \tilde{\lambda}}(k) \right) \quad \forall \quad k = 0 \text{ (1) } F-1 \quad (6.35b)$$

einzusetzen. Anschließend werden die Gleichungen (6.13) konjugiert und mit j multipliziert. Der Laufindex μ der Summe wird noch durch den negativen Laufindex $-\mu$ substituiert, und man erhält unter Ausnutzung der M Periodizität der Summanden die F Gleichungen

$$\begin{aligned} j \cdot \hat{u}(k)^* &= j \cdot \frac{2}{L} \cdot \vec{y}(k)^* \cdot \vec{1}^T - \frac{2}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{H}(\mu)^* \cdot j \cdot \vec{V}(\mu)^* \cdot \vec{1}^T \cdot e^{-j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\ &= j \cdot \frac{2}{L} \cdot \vec{y}(k)^* \cdot \vec{1}^T - \frac{2}{M \cdot L} \cdot \sum_{\mu=1-M}^0 \hat{H}(-\mu)^* \cdot j \cdot \vec{V}(-\mu)^* \cdot \vec{1}^T \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\ &= j \cdot \frac{2}{L} \cdot \vec{y}(k)^* \cdot \vec{1}^T - \frac{2}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{H}(-\mu)^* \cdot j \cdot \vec{V}(-\mu)^* \cdot \vec{1}^T \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \\ &\quad \forall \quad k = 0 \text{ (1) } F-1. \end{aligned} \quad (6.36)$$

Wieder fordert man, dass auch die empirischen Regressionskoeffizienten der deterministischen Störung des komplexen Systemmodells dieselben Symmetrien aufweisen, wie die theoretischen Regressionskoeffizienten nach Gleichung (6.28). Man setzt daher in die letzten Gleichungen und die Gleichungen (6.13) die Terme

$$\hat{u}(k) = j \cdot \hat{u}(k)^* = (1+j) \cdot \hat{u}_{\text{reell}}(k) \quad (6.37)$$

für die zu bestimmenden Werte der deterministischen Störung und Gleichung (6.33a) für die empirischen Werte der Übertragungsfunktion ein, und fasst bei den Gleichungen (6.13) und (6.36) alle Gleichungen mit denselben Unbekannten $\hat{u}_{\text{reell}}(k)$ zusammen. Man erhält

die F Gleichungen der Dimension 1×2

$$\begin{aligned} (1+j) \cdot \hat{u}_{\text{reell}}(k) \cdot \vec{1} &= \\ &= \frac{2}{L} \cdot \left[\vec{y}(k) \cdot \vec{1}^H, j \cdot \vec{y}(k)^* \cdot \vec{1}^T \right] - \frac{2}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{H}_{\text{reell}}(\mu) \cdot \left[\vec{V}(\mu) \cdot \vec{1}^H, j \cdot \vec{V}(-\mu)^* \cdot \vec{1}^T \right] \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \end{aligned} \quad (6.38)$$

deren Ausgleichslösung sich zu

$$\begin{aligned} \hat{u}_{\text{reell}}(k) &= \\ &= \frac{\left[\vec{y}(k) \cdot \vec{1}^H, j \cdot \vec{y}(k)^* \cdot \vec{1}^T \right] \cdot \vec{1}^H}{L \cdot (1+j)} - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \hat{H}_{\text{reell}}(\mu) \cdot \frac{\left[\vec{V}(\mu) \cdot \vec{1}^H, j \cdot \vec{V}(-\mu)^* \cdot \vec{1}^T \right] \cdot \vec{1}^H}{L \cdot (1+j)} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\ &= \frac{1}{L} \cdot \sum_{\tilde{\lambda}=1}^{L/2} (y_{2 \cdot \tilde{\lambda}-1}(k) + y_{2 \cdot \tilde{\lambda}}(k)) - \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \hat{H}_{\text{reell}}(\mu) \cdot \frac{1}{L} \cdot \sum_{\tilde{\lambda}=1}^{L/2} (V_{2 \cdot \tilde{\lambda}-1}(\mu) + V_{2 \cdot \tilde{\lambda}}(\mu)) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} = \\ &= \frac{1}{L} \cdot \vec{y}_{\text{reell}}(k) \cdot \vec{1}^H - \frac{1}{M \cdot L} \cdot \sum_{\mu=0}^{M-1} \hat{H}_{\text{reell}}(\mu) \cdot \vec{V}_{\text{reell}}(\mu) \cdot \vec{1}^H \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k} \end{aligned} \quad (6.39)$$

berechnet. Bis auf den Unterschied, dass die M Messwerte der Übertragungsfunktion nun die Ausgleichslösungen anderer Gleichungssysteme sind, berechnen sich die Messwerte der deterministische Störung also identisch wie bei der Version des RKM, die in Unterkapitel 6.1 vorgestellt wurde. Die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ des Betragsquadrats des Spektrums des gefensterten komplexen Approximationsfehlers erhalten wir, indem wir das Betragsquadrat der euklidischen Länge des Differenzvektors, der sich beim Einsetzen der Ausgleichslösung jeweils als Differenz der rechten und der linken Seite in den Gleichungen (6.34) ergibt, durch die um 2 reduzierte Spaltendimension des Gleichungssystems dividieren. Die Spaltendimension des Gleichungssystems ist die Anzahl L der Einzelmessungen am realen System. Diese ist um 2 zu reduzieren, da die linke Seite jeder Gleichung jeweils einen zweidimensionalen Unterraum des L -dimensionalen Raums aufspannt.

$$\begin{aligned} \hat{\Phi}_{\mathbf{n}}(\mu) &= \frac{1}{L-2} \cdot \\ &\cdot \left(\left[\vec{Y}_f(\mu), \vec{Y}_f(-\mu)^* \right] - \hat{H}_{\text{reell}}(\mu) \cdot \left[\vec{V}(\mu), \vec{V}(-\mu)^* \right] - \hat{U}_{f,\text{reell}}(\mu) \cdot \left[(1+j) \cdot \vec{1}, (1-j) \cdot \vec{1} \right] \right) \cdot \\ &\cdot \left(\left[\vec{Y}_f(\mu), \vec{Y}_f(-\mu)^* \right] - \hat{H}_{\text{reell}}(\mu) \cdot \left[\vec{V}(\mu), \vec{V}(-\mu)^* \right] - \hat{U}_{f,\text{reell}}(\mu) \cdot \left[(1+j) \cdot \vec{1}, (1-j) \cdot \vec{1} \right] \right)^H \end{aligned} \quad (6.40)$$

Da die Differenzvektoren, die sich mit $-\mu$ statt mit μ ergeben, lediglich konjugiert und permutiert sind, sind diese gleich lang, und die Symmetrie, die die theoretischen Werte nach Gleichung (6.29) aufweisen, ist auch bei diesen Messwerten vorhanden. Desweiteren

besagt die Gleichung (6.29), dass die zweiten Momente des Spektrums des gefensternten reellen Approximationsfehlers halb so groß sind wie die zweiten Momente des Spektrums des gefensternten komplexen Approximationsfehlers. Daher sind die Messwerte

$$\hat{\Phi}_{\mathbf{n},\text{reell}}(\mu) = \frac{1}{2} \cdot \hat{\Phi}_{\mathbf{n}}(\mu) \quad (6.41)$$

erwartungstreue Schätzwerte für die gesuchten Größen $\tilde{\Phi}_{\mathbf{n},\text{reell}}(\mu)$, wenn die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ die Größen $\tilde{\Phi}_{\mathbf{n}}(\mu)$ erwartungstreu abschätzen.

Diese Variante des RKM zur Messung reellwertiger Systeme besteht nun also darin, mit Hilfe der komplexen und nicht symmetrischen Stichprobenvektoren gemäß der Gleichungen (6.31) die Messwerte $\hat{H}_{\text{reell}}(\mu)$ und $\hat{U}_{f,\text{reell}}(\mu)$ als Ausgleichslösungen der Gleichungssysteme (6.34) zu bestimmen. Mit diesen Ausgleichslösungen werden einerseits mit Gleichung (6.39) die Messwerte $\hat{u}_{\text{reell}}(k)$ als Ausgleichslösungen der Gleichungssysteme (6.38), mit den empirischen Mittelwerten gemäß der Gleichungen (6.35), und andererseits die Messwerte $\hat{\Phi}_{\mathbf{n},\text{reell}}(\mu)$ nach Gleichung (6.40) und (6.41) berechnet. Man kann nun zeigen, dass sich die Gleichungssysteme (6.34) durch Multiplikation mit einer unitären Matrix von rechts und durch Division durch $\sqrt{2}$ auf die Gleichungssysteme (6.7) überführen lassen, die bei der im Unterkapitel 6.1 vorgestellte RKM-Variante auftreten. Daher liefern diese beiden RKM-Varianten wenigstens theoretisch — also abgesehen von Rechengenauigkeiten — dieselben Messwerte. Daher sind auch die Erwartungswerte, Varianzen und Kovarianzen der Messwerte beider RKM-Varianten gleich. Da das eben beschriebene Verfahren aber sowohl hinsichtlich des Speicherbedarfs als auch hinsichtlich der Anzahl der Gleitkommaoperationen (flops) aufwendiger ist, als das Verfahren nach Unterkapitel 6.1, ist letztgenanntes vorzuziehen.

Bei den bisher vorgestellten drei RKM-Varianten ein reellwertiges System zu messen, kann man immer zur Erregung in zwei aufeinanderfolgenden Einzelmessungen Signalsequenzen verwenden, die konkrete Realisierungen unabhängiger reeller Zufallsvektoren sind. Andererseits kann man bei fast allen realen Systemen in zwei aufeinanderfolgenden Einzelmessungen auch den Real- und Imaginärteil einer komplexen Signalsequenz zur Erregung verwenden, die eine konkrete Realisierung eines komplexen Zufallsvektors ist, der abhängige Real- und Imaginärteilvektoren aufweist, ohne dass dabei die Gefahr besteht, dass dadurch die Messergebnisse verfälscht werden. Nur bei solchen Systemen lässt sich die nun folgende und letzte dem Autor bekannte RKM-Variante einsetzen. Bei dieser Variante werden zur Erregung sogenannte Halbbandsignale verwendet. Bei diesen Signalen sind die Zufallsgrößen der Spektralwerte der Erregung für negative Frequenzen null.

$$\mathbf{V}(\mu) = 0 \quad \forall \quad \mu = \frac{M}{2} + 1 \text{ (1) } M-1. \quad (6.42)$$

Die Messwerte der Übertragungsfunktion werden mit Gleichung (6.7) nur für die positiven Frequenzen mit $\mu = 1 \text{ (1) } M/2-1$ berechnet. Die Messwerte der restlichen Frequenzen

$\mu = M/2 + 1$ (1) $M - 1$ ergeben sich durch konjugierte Spiegelung. Bei diesen Frequenzen kann man die Messwerte $\hat{U}_{f,\text{reell}}(\mu)$ und $\hat{\Phi}_{\mathbf{n},\text{reell}}(\mu)$ über eine reine Spektralschätzung bestimmen, da dort kein erregungsabhängiger Signalanteil am Ausgang der linearen Modellsystems vorhanden ist. Bei den beiden Frequenzen $\mu=0$ und $\mu=M/2$ muss man die vollständige Messung der Variante nach Unterkapitel 6.1 verwenden, wenn man an allen Messwerten interessiert ist. Mit diesen Messwerten kann man dann auch die Messwerte der deterministischen Störung nach Gleichung (6.13) berechnen. Wenn man sich die Messwertvarianzen dieses Verfahrens berechnet, so stellt man fest, dass man bei gleicher Anzahl von Einzelmessungen fast die doppelte Messwertvarianz erhält. Bei gleicher Messgenauigkeit hat man daher also fast doppelt so viele Einzelmessung durchzuführen. Daher ist diese Variante für die Messung der Übertragungsfunktion eines reellen Systems im gesamten Frequenzbereich $\mu = 1$ (1) $M - 1$ nicht empfehlenswert.

Das zuletzt angeführte Verfahren ist das in [4] vorgestellte Halbbandmessverfahren, wobei hier lediglich die deterministische Störung ergänzt ist. Dieses Messverfahren stellt dort den Spezialfall eines Messverfahren dar, mit dessen Hilfe man mit relativ geringem Aufwand die Übertragungsfunktion des Modellsystems innerhalb eines schmalen Frequenzbandes mit hoher Frequenzauflösung messen kann. Es wird dann allgemein mit schmalbandigen Teilbandsignalen erregt, und bei der Messung der Ausgangssignale eine geeignete Unterabtastung vorgenommen, wodurch bei gleichem Frequenzabstand der Messwerte der Übertragungsfunktion eine immense Reduktion des Rechenaufwands erzielt wird. Leider können bei dieser Schmalband-RKM-Variante die Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ des Spektrums des Approximationsfehlers wegen der Unterabtastung keine Aussage über das LDS einer im realen System vorhandenen Störung liefern, sondern lediglich zur Abschätzung der Messwertvarianz benutzt werden. Da aber die Hauptvorteile der Fensterung gerade bei der Messung des LDS liegen, und daher der Einsatz einer Fensterfolge bei der schmalbandigen RKM-Variante nur in den seltensten Fällen sinnvoll erscheint, und da die Halbband-Variante bei gleicher Messwertvarianz einen höheren Rechenaufwand aufweist, als die vorher beschriebenen Varianten, überlasse ich es dem Leser sich herzuleiten, wie sich dieses vor allem in der schmalbandigen Variante empfehlenswerte Messverfahren um die Fensterung und die Modellierung der deterministischen Störung erweitern lässt.

7 Ablaufübersicht für eine Variante des RKM

Es zeigt sich, dass alle am Anfang des Kapitels 3 aufgelisteten Messwerte und die dazugehörigen Konfidenzgebiete aus Werten berechnet werden können, die man auf einfache Weise durch die Akkumulation von einfachen Produkten über alle L Einzelmessungen erhält. Wenn man eine Fensterfolge verwendet, deren Spektrum die nach Gleichung ([1]:2.27) geforderten äquidistanten Nullstellen am Einheitskreis aufweist, ist es auch beim RKM mit Fensterung nicht notwendig, alle Erregungen bzw. Systemantworten aller Einzelmessungen abzuspeichern, um am Ende der Messung mit rechenintensiven Matrixoperationen die Messwerte und Konfidenzgebiete zu erhalten. Hier möchte ich nun auflisten, in welchen Schritten die Messung mit dem RKM mit Fensterung in der Praxis durchgeführt werden kann, und wieviel Speicher dazu benötigt wird. Dabei möchte ich mich auf den Fall eines komplexwertigen realen Systems beschränken, das sich, wie in Bild 1.1 gezeigt, durch ein zeitinvariantes ($K_H=1$) Modellsystem $\mathcal{S}_{lin}$, beschrieben durch die Übertragungsfunktion $H(\mu)$, und ein zweites zeitinvariantes Modellsystem $\mathcal{S}_{*,lin}$, das von dem konjugierten Eingangssignal erregt wird und durch die Übertragungsfunktion $H_*(\mu)$ beschrieben wird, modellieren lässt. Da es sich hier um eindimensionale Übertragungsfunktionen handelt, wird auf die in den Kapiteln 2 und 3 verwendete Doppelindizierung verzichtet. Hier enthält der nach Gleichung (2.15) definierte Zufallsvektor $\vec{V}(\mu)$ lediglich die zwei Zufallsgrößen $\mathbf{V}(\mu)$ und $\mathbf{V}(-\mu)^*$. Die komplexe, deterministische und zeitabhängige Störung $u(k)$ wird modelliert. Von dem Approximationsfehlerprozess wird angenommen, dass er stationär ($K_\Phi=1$) ist, so dass die Messung eines eindimensionalen LDS $\tilde{\Phi}_n(\mu)$ und eines ebenfalls eindimensionalen KLDS $\tilde{\Psi}_n(\mu)$ ausreicht. Auch hier wird auf die Doppelindizierung verzichtet. Der Zufallsspaltenvektor $\check{\vec{V}}(\mu)$, der bei der Berechnung der Messwerte für das LDS und das KLDS verwendet wird ist gleich dem Zufallsspaltenvektor $\tilde{\vec{V}}(\mu)$. Desweiteren sei M gerade, und somit die Näherungen, bei denen die Ausblendeigenschaft des Spektrums der Fensterfolge eingegangen ist, bei Verwendung einer hoch frequenzselektiven Fensterfolge sehr gut erfüllt sind.

In der folgenden Auflistung sind die Punkte, die mit $\emptyset$ gekennzeichnet sind, Kommentare und beschreiben meist die Berechnung von Werten, die sich aus den davor berechneten Werten unmittelbar (z. B. aufgrund einer Symmetrieeigenschaft) ergeben. Sie brauchen

in einem Programm *nicht* explizit berechnet zu werden, da diese Werte bei den weiteren Berechnungen ggf. durch die bereits berechneten Werte ersetzt werden. Die Gleichheitszeichen „=“ sind meist als Wertzuweisungen, und nicht als Gleichungen im mathematischen Sinne zu verstehen, wie dies in den meisten Programmiersprachen üblich ist. Wie bisher ist bei Multiplikationen ggf. die Matrixmultiplikation gemeint, und der Exponent $(\dots)^{-1}$ bedeutet bei Matrizen eine Invertierung. Unter einem Speicherplatz, wird im weiteren der Speicherbedarf verstanden, der bei dem verwendeten Rechner für die Speicherung einer reellen Gleitkommazahl benötigt wird.

1. Festlegen der Anzahl der zu messenden Frequenzen M .
2. Festlegen der Einschwingtakte E des zu messenden Systems. Wenn man die Einschwingzeit nicht auf Grund heuristischer Überlegungen exakt vorhersehen kann (z. B. Messung eines FIR-Systems), sollte man diese Zeit so groß wählen, dass auch im schlimmsten Fall damit zu rechnen ist, dass die transienten Vorgänge innerhalb dieser Zeit soweit abgeklungen sind, dass diese bei der Messung bedeutungslos werden. Im Zweifelsfall sollte man die Messung mit einer veränderten Einschwingzeit wiederholen, um festzustellen, ob sich die Messergebnisse dadurch verändern.
3. Wahl und Berechnung einer Fensterfolge der Länge F . In den Kapiteln [1]:6 und 10.1 werden Algorithmen vorgestellt, mit denen man Fensterfolgen berechnen kann, die alle Bedingungen erfüllen, die an die Fensterfolge gestellt wurden. Bei diesen Fensterfolgen ist die Fensterlänge $F = N \cdot M$ ein ganzzahliges Vielfaches N der Anzahl M der Frequenzpunkte, für die die Messwerte gewonnen werden sollen. Die Berechnung des Fensters erfolgt am besten mit dem in Kapitel [1]:6 vorgestellten Algorithmus, der die Variante des Fensters berechnet, das mir am geeignetsten für die Verwendung beim RKM erscheint. Bei diesem Fenster ist es so, dass mit steigendem N die Frequenzselektivität und die Potenz des Abfalls des Betrags des Spektrums der Fensterfolge für zunehmende Frequenzen ansteigt.

$$\mathbf{f}(k) = \mathbf{fenster}(N, M) \quad \text{für} \quad k = 0 \text{ (1) } F-1$$

Da die Fensterfolge reell ist, kann sie auf einem Speicher $\mathbf{f}$ mit F Speicherplätzen abgelegt werden.

- ∅) Es kann auch jede andere Fensterfolge beliebiger Länge verwendet werden. Wird jedoch die Bedingungen ([1]:2.27) nicht erfüllt, so ist die Messung mit einem anderen Algorithmus durchzuführen, der die Pseudoinverse einer wesentlich größeren Matrix eines Gleichungssystems, das dem Gleichungssystem ([1]:3.4) ähnelt, berechnet. Wird die Bedingung ([1]:2.20) nicht erfüllt, oder verwendet man eine Fensterfolge mit geringer Frequenzselektivität, so werden die Messwerte unnötig stark verrauscht

sein, die Messwerte $\hat{\Phi}_n(\mu)$ und $\hat{\Psi}_n(\mu)$ werden die gewünschten Größen $\bar{\Phi}_n(\mu)$ und $\bar{\Psi}_n(\mu)$ nur schlecht annähern, und die Schätzwerte der Varianzen und Kovarianzen, die eine Näherung enthalten, die nur für hoch frequenzselektive Fensterfolgen gültig ist, werden nicht mehr erwartungstreu sein.

4. Bereitstellen von 8 Akkumulatorfeldern. Alle Akkumulatorfelder werden zu null initialisiert:

$$\begin{aligned}
 \text{Akku_y}(k) &= 0 & \text{für } k &= 0 \text{ (1) } F-1 \\
 \text{Akku_V}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M-1 \\
 \text{Akku_VQ}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M-1 \\
 \text{Akku_VV}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M/2 \\
 \text{Akku_YQ}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M-1 \\
 \text{Akku_YY}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M/2 \\
 \text{Akku_YV}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M-1 \\
 \text{Akku_VY}(\mu) &= 0 & \text{für } \mu &= 0 \text{ (1) } M-1
 \end{aligned}$$

Das Akkumulatorfeld

Akku_y	benötigt	$2 \cdot F$	Speicherplätze,
Akku_V	benötigt	$2 \cdot M$	Speicherplätze,
Akku_VQ	benötigt	M	Speicherplätze,
Akku_VV	benötigt	$M+2$	Speicherplätze,
Akku_YQ	benötigt	M	Speicherplätze,
Akku_YY	benötigt	$M+2$	Speicherplätze,
Akku_YV	benötigt	$2 \cdot M$	Speicherplätze und
Akku_VY	benötigt	$2 \cdot M$	Speicherplätze.

5. Zufällige Auswahl der L Zeitintervalle (Index λ) in denen das reale System erregt werden soll. Die zufällige Auswahl kann dabei in der Art erfolgen, wie dies im Kapitel 2.1 beschrieben worden ist. In vielen Fällen kann man die zufällige Auswahl auch in der Art realisieren, dass man zwischen zwei Einzelmessungen eine Pause einer zufälligen Länge einschiebt. In diesem Fall braucht L nicht zu Beginn der Messung festgelegt werden. Oft (z. B. bei Rechnersimulationen) kann auf die zufällige Auswahl der Zeitintervalle ganz verzichtet werden. Bei Rechnersimulationen kann man die zufällige Lage der Zeitintervalle meist durch eine zufällige Zeitverschiebung der simulierten Störung realisieren.
6. Initialisierung des Zählers der Einzelmessungen.

$$\lambda = 0$$

7. Inkrementieren des Zählers der Einzelmessungen.

$$\lambda = \lambda + 1$$

8. Erzeugung des Testsignals für die Einzelmessung. Mit Hilfe eines geeigneten Zufallsgenerators wird eine Periode der Länge M des periodischen Zufallssignals für $k = 0 (1) M-1$ erzeugt. Dieses wird auf einem Speicher $\mathbf{v}$ mit $2 \cdot M$ Speicherplätzen geschrieben. Die DFT des Testsignals $\mathbf{v}$ liefert die Spektralwerte $V_\lambda(\mu)$, die auf dem Speicher $\mathbf{V}$ mit $2 \cdot M$ Speicherplätzen abgespeichert werden, wobei die Spektralwerte, die sich evtl. noch von der vorigen Einzelmessung auf diesem Speicher befinden, überschrieben werden.

$$\mathbf{V}(\mu) = \text{fft}_M(\mathbf{v}(k)) \quad \text{für} \quad \mu = 0 (1) M-1$$

Wahlweise kann auch das Spektrum mit Hilfe eines Zufallsgenerators gewonnen werden, und das Testsignal durch eine inverse DFT erzeugt werden.

$$\mathbf{v}(k) = \text{ifft}_M(\mathbf{V}(\mu)) \quad \text{für} \quad k = 0 (1) M-1$$

Das Testsignal $\mathbf{v}$ wird in dem für diese Einzelmessung zu Beginn ausgewählten Zeitintervall auf den Eingang des zu messenden Systems gegeben indem es in beide Richtungen periodisch fortgesetzt wird. Dieses Zeitintervall wurde mit $k \in [-E; F-1]$ bezeichnet, wobei k die auf dieses Intervall bezogene Zeit ist. Es ist darauf zu achten, dass die gewünschten stochastischen Eigenschaften des Testsignals auch in den Umgebungen der Zeitpunkte erfüllt sind, an denen die einzelnen Perioden $\mathbf{v}$ des Signals aneinandergesetzt wurden. Dies kann vor allen bei der Generierung des Signals im Zeitbereich Probleme bereiten.

9. Messung der F Abtastwerte $y_\lambda(k)$ des Systemausgangssignals der Einzelmessung λ im Zeitintervall $k \in [0; F-1]$. Dieser Signalausschnitt wird auf dem Speicher $\mathbf{y}$ mit $2 \cdot F$ Speicherplätzen abgelegt, wobei die Signalwerte, die sich evtl. noch von der vorigen Einzelmessung auf diesem Speicher befinden, überschrieben werden.

$$\mathbf{y}(k) = y_\lambda(k) \quad \text{für} \quad k = 0 (1) F-1$$

10. Fensterung des Systemausgangssignals und anschließende DFT. Dies liefert nach Gleichung ([1]:2.25) das Spektrum $Y_{f,\lambda}(\mu)$. Dieses Spektrum wird auf dem Speicher $\mathbf{Y}$ mit $2 \cdot M$ Speicherplätzen abgelegt, wobei die Spektralwerte, die sich evtl. noch von der vorigen Einzelmessung auf diesem Speicher befinden, überschrieben werden.

$$y_{f,\lambda}(k) = \sum_{\kappa=0}^N \mathbf{f}(k+\kappa \cdot M) \cdot \mathbf{y}(k+\kappa \cdot M) \quad \text{für} \quad k = 0 (1) M-1$$

$$\mathbf{Y}(\mu) = \text{fft}_M(y_{f,\lambda}(k)) \quad \text{für} \quad \mu = 0 (1) M-1$$

11. Akkumulieren des Spektrums der Erregung. Die Werte $V_\lambda(\mu)$ werden zum Inhalt des Akkumulatorfeldes **Akku_V** addiert.

$$\mathbf{Akku_V}(\mu) = \mathbf{Akku_V}(\mu) + \mathbf{V}(\mu) \quad \text{für} \quad \mu = 0 (1) M-1$$

12. Die Betragsquadrate $|V_\lambda(\mu)|^2$ werden zum Inhalt des Akkumulatorfeldes **Akku_VQ** addiert.

$$\text{Akku_VQ}(\mu) = \text{Akku_VQ}(\mu) + |V(\mu)|^2 \quad \text{für } \mu = 0 \text{ (1) } M-1$$

13. Die Produkte $V_\lambda(\mu) \cdot V_\lambda(-\mu)$ werden zum Inhalt des Akkumulatorfeldes **Akku_VV** addiert.

$$\text{Akku_VV}(\mu) = \text{Akku_VV}(\mu) + V(\mu) \cdot V(M-\mu) \quad \text{für } \mu = 1 \text{ (1) } \frac{M}{2}$$

$$\text{Akku_VV}(0) = \text{Akku_VV}(0) + V(0)^2$$

$$\emptyset) \quad \text{Akku_VV}(M-\mu) = \text{Akku_VV}(\mu) \quad \text{sonst}$$

14. Akkumulieren des Systemausgangssignals. Die Werte $y_\lambda(k)$ werden zum Inhalt des Akkumulatorfeldes **Akku_y** addiert.

$$\text{Akku_y}(k) = \text{Akku_y}(k) + y(k) \quad \text{für } k = 0 \text{ (1) } F-1$$

15. Die Betragsquadrate $|Y_{f,\lambda}(\mu)|^2$ werden zum Inhalt des Akkumulatorfeldes **Akku_YQ** addiert.

$$\text{Akku_YQ}(\mu) = \text{Akku_YQ}(\mu) + |Y(\mu)|^2 \quad \text{für } \mu = 0 \text{ (1) } M-1$$

16. Die Produkte $Y_{f,\lambda}(\mu) \cdot Y_{f,\lambda}(-\mu)$ werden zum Inhalt des Akkumulatorfeldes **Akku_YY** addiert.

$$\text{Akku_YY}(\mu) = \text{Akku_YY}(\mu) + Y(\mu) \cdot Y(M-\mu) \quad \text{für } \mu = 1 \text{ (1) } \frac{M}{2}$$

$$\text{Akku_YY}(0) = \text{Akku_YY}(0) + Y(0)^2$$

$$\emptyset) \quad \text{Akku_YY}(M-\mu) = \text{Akku_YY}(\mu) \quad \text{sonst}$$

17. Die Produkte $Y_{f,\lambda}(\mu) \cdot V_\lambda(\mu)^*$ werden zum Inhalt des Akkumulatorfeldes **Akku_YV** addiert.

$$\text{Akku_YV}(\mu) = \text{Akku_YV}(\mu) + Y(\mu) \cdot V(\mu)^* \quad \text{für } \mu = 0 \text{ (1) } M-1$$

18. Die Produkte $V_\lambda(-\mu) \cdot Y_{f,\lambda}(\mu)$ werden zum Inhalt des Akkumulatorfeldes **Akku_VY** addiert.

$$\text{Akku_VY}(\mu) = \text{Akku_VY}(\mu) + V(M-\mu) \cdot Y(\mu) \quad \text{für } \mu = 1 \text{ (1) } M-1$$

$$\text{Akku_VY}(0) = \text{Akku_VY}(0) + V(0) \cdot Y(0)$$

19. Weitere Einzelmessungen durchführen, indem man die Punkte 7 bis 18 wiederholt. Entweder man führt eine konstante Anzahl von Einzelmessungen durch, oder man entscheidet anhand eines Kriteriums, das man aus den Messwerten gewinnt, und dessen Berechnung in den folgenden Punkten beschrieben wird, ob weitere Einzelmessungen durchzuführen sind. Besonders geeignet erscheint es mir, nach einer Mindestanzahl von Einzelmessungen, die Singulärwerte der im Punkt 23 berechneten empirischen 2×2 Kovarianzmatrizen zu berechnen, und so lange weitere Einzelmessungen durchzuführen, bis die Singulärwerte bei allen diskreten Frequenzen μ innerhalb eines Toleranzbereiches um die theoretischen Werte der Singulärwerte des Spektrums der Erregung liegen. Man kann jedoch an jeder Stelle der folgenden Messwertberechnung zu Punkt 7 zurückspringen und weitere Einzelmessungen durchführen, wenn man z. B. mit der Qualität der bisher erzielten Messergebnisse nicht zufrieden ist.
20. Wenn die Mittelungsanzahl L nicht zu Beginn der Messung festgelegt worden ist, wird L auf die tatsächlich durchgeführte Anzahl der Einzelmessungen gesetzt.

$$L = \lambda$$

Auch für die Berechnung eines Abbruchkriteriums muss immer die Anzahl der bis dahin tatsächlich durchgeführten Einzelmessungen verwendet werden.

21. Berechnung der empirischen Varianzen $\hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$ des Spektrums der Erregung. Es handelt sich dabei um die Hauptdiagonalelemente der Kovarianzmatrizen nach Gleichung (3.10) bzw. (3.27). Das $L \cdot (L-1)$ -fache der empirischen Varianzen wird auf einem Speicher `C_VQ` mit M Speicherplätzen abgelegt.

$$\text{C_VQ}(\mu) = L \cdot \text{Akku_VQ}(\mu) - |\text{Akku_V}(\mu)|^2 \quad \text{für } \mu = 0 \text{ (1) } M-1$$

22. Berechnung der empirischen Kovarianzen $\hat{C}_{\mathbf{V}(\mu), \mathbf{V}(-\mu)^*}$ des Spektrums der Erregung bei positiver Frequenz und des konjugierten Spektrums bei negativer Frequenz. Es handelt sich dabei um die Nebendiagonalelemente der Kovarianzmatrizen nach Gleichung (3.10) bzw. (3.27). Das $L \cdot (L-1)$ -fache der empirischen Kovarianzen wird auf einem Speicher `C_VV` mit $M+2$ Speicherplätzen abgelegt.

$$\text{C_VV}(\mu) = L \cdot \text{Akku_VV}(\mu) - \text{Akku_V}(\mu) \cdot \text{Akku_V}(M-\mu) \quad \text{für } \mu = 1 \text{ (1) } \frac{M}{2}$$

$$\text{C_VV}(0) = L \cdot \text{Akku_VV}(0) - \text{Akku_V}(0)^2$$

$$\emptyset) \quad \text{C_VV}(M-\mu) = \text{C_VV}(\mu) \quad \text{sonst}$$

23. Berechnung der Singulärwerte $s(\mu)$ der empirischen Kovarianzmatrizen $\hat{\underline{C}}_{\check{\mathbf{V}}(\mu), \check{\mathbf{V}}(\mu)}$ nach Gleichung (3.10) bzw. (3.27). Die Singulärwerte sind die immer nichtnegativen, reellen Lösungen der folgenden quadratischen Gleichungen

$$\left(\mathbf{C_VQ}(\mu) - L \cdot (L-1) \cdot s(\mu) \right) \cdot \left(\mathbf{C_VQ}(M-\mu) - L \cdot (L-1) \cdot s(\mu) \right) = |\mathbf{C_VV}(\mu)|^2$$

für $\mu = 1 \ (1) \ \frac{M}{2}$

$$s(0) = \frac{\mathbf{C_VQ}(0) \pm |\mathbf{C_VV}(0)|}{L \cdot (L-1)}$$

$$\emptyset) \quad s(M-\mu) = s(\mu) \quad \text{sonst}$$

Anhand der Singulärwerte kann entschieden werden, ob weitere Einzelmessungen durchgeführt werden sollen, indem man zum Punkt 7 zurückspringt.

24. Berechnung der inversen empirischen Kovarianzmatrizen $\hat{\underline{C}}_{\check{\mathbf{V}}(\mu), \check{\mathbf{V}}(\mu)}^{-1} = \hat{\underline{C}}_{\check{\mathbf{V}}(\mu), \check{\mathbf{V}}(\mu)}^{-1}$. Für die inversen Kovarianzmatrizen mit $\mu = 1 \ (1) \ M/2-1$ werden je vier Speicherplätze benötigt, da diese 2×2 Matrizen hermitesch sind und reelle Hauptdiagonalelemente aufweisen. Bei den beiden Frequenzen $\mu=0$ und $\mu = M/2$ sind die Diagonalelemente der inversen Kovarianzmatrizen gleich, so dass man insgesamt also mindestens $2 \cdot M + 2$ Speicherplätze benötigt. Hier werden die Bindungen der Elemente der inversen Kovarianzmatrizen jedoch nicht ausgenutzt. Es werden die durch $L \cdot (L-1)$ dividierten, inversen, empirischen Kovarianzmatrizen auf den Feldern $\mathbf{K_VV}$ mit $4 \cdot M + 8$ Speicherplätzen abgelegt.

$$\mathbf{K_VV}(\mu) = \begin{bmatrix} \mathbf{K_VV}(\mu, 1, 1) & \mathbf{K_VV}(\mu, 1, 2) \\ \mathbf{K_VV}(\mu, 2, 1) & \mathbf{K_VV}(\mu, 2, 2) \end{bmatrix} = \begin{bmatrix} \mathbf{C_VQ}(\mu) & \mathbf{C_VV}(\mu) \\ \mathbf{C_VV}(\mu)^* & \mathbf{C_VQ}(M-\mu) \end{bmatrix}^{-1}$$

für $\mu = 1 \ (1) \ \frac{M}{2}$

$$\mathbf{K_VV}(0) = \begin{bmatrix} \mathbf{K_VV}(0, 1, 1) & \mathbf{K_VV}(0, 1, 2) \\ \mathbf{K_VV}(0, 2, 1) & \mathbf{K_VV}(0, 2, 2) \end{bmatrix} = \begin{bmatrix} \mathbf{C_VQ}(0) & \mathbf{C_VV}(0) \\ \mathbf{C_VV}(0)^* & \mathbf{C_VQ}(0) \end{bmatrix}^{-1}$$

$$\emptyset) \quad \mathbf{K_VV}(M-\mu) = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \cdot \mathbf{K_VV}(\mu)^T \cdot \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad \text{sonst}$$

25. Berechnung des empirischen Mittelwertes des Spektrums des gefensterten Signals am Ausgang. Das L -fache dieser Werte wird auf einem Speicher $\mathbf{Y_mittel}$ mit

$2 \cdot M$ Speicherplätzen abgelegt.

$$\bar{y}_f(k) = \sum_{\kappa=0}^N f(k + \kappa \cdot M) \cdot \text{Akku}_y(k + \kappa \cdot M) \quad \text{für } k = 0 \ (1) \ M-1$$

$$\text{Y_mittel}(\mu) = \text{fft}_M(\bar{y}_f(k)) \quad \text{für } \mu = 0 \ (1) \ M-1$$

26. Berechnung der empirischen Varianzen $\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu)}$ des Spektrums des gefensterten Systemausgangssignals mit Gleichung (3.35a). Das $L \cdot (L-1)$ -fache der empirischen Varianzen wird auf einem Speicher C_YQ mit M Speicherplätzen abgelegt.

$$\text{C_YQ}(\mu) = L \cdot \text{Akku_YQ}(\mu) - |\text{Y_mittel}(\mu)|^2 \quad \text{für } \mu = 0 \ (1) \ M-1$$

27. Berechnung der empirischen Kovarianzen $\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(-\mu)^*}$ des Spektrums des gefensterten Systemausgangssignals bei positiver und des konjugierten Spektrums bei negativer Frequenz mit Gleichung (3.35b). Das $L \cdot (L-1)$ -fache der empirischen Kovarianzen wird auf einem Speicher C_YY mit $M+2$ Speicherplätzen abgelegt.

$$\text{C_YY}(\mu) = L \cdot \text{Akku_YY}(\mu) - \text{Y_mittel}(\mu) \cdot \text{Y_mittel}(M-\mu) \quad \text{für } \mu = 1 \ (1) \ \frac{M}{2}$$

$$\text{C_YY}(0) = L \cdot \text{Akku_YY}(0) - \text{Y_mittel}(0)^2$$

$$\emptyset) \quad \text{C_YY}(M-\mu) = \text{C_YY}(\mu) \quad \text{sonst}$$

28. Berechnung der empirischen Kreuzkovarianzen $\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{V}(\mu)}$ des Spektrums des gefensterten Systemausgangssignals und der Erregung. Es handelt sich dabei um die ersten Elemente der Kovarianzvektoren nach Gleichung (3.11). Das $L \cdot (L-1)$ -fache der empirischen Kreuzkovarianzen wird auf einem Speicher C_YV mit $2 \cdot M$ Speicherplätzen abgelegt.

$$\text{C_YV}(\mu) = L \cdot \text{Akku_YV}(\mu) - \text{Y_mittel}(\mu) \cdot \text{Akku_V}(\mu)^* \quad \text{für } \mu = 0 \ (1) \ M-1$$

29. Berechnung der empirischen Kreuzkovarianzen $\hat{C}_{\mathbf{V}(-\mu), \mathbf{Y}_f(\mu)^*}$ des konjugierten Spektrums des gefensterten Systemausgangssignals bei positiver und des Spektrums der Erregung bei negativer Frequenz. Es handelt sich dabei um die zweiten Elemente der Kovarianzvektoren nach Gleichung (3.11). Das $L \cdot (L-1)$ -fache der empirischen Kreuzkovarianzen wird auf einem Speicher C_VY mit $2 \cdot M$ Speicherplätzen abgelegt.

$$\text{C_VY}(\mu) = L \cdot \text{Akku_VY}(\mu) - \text{Akku_V}(M-\mu) \cdot \text{Y_mittel}(\mu) \quad \text{für } \mu = 1 \ (1) \ M-1$$

$$\text{C_VY}(0) = L \cdot \text{Akku_VY}(0) - \text{Akku_V}(0) \cdot \text{Y_mittel}(0)$$

30. Berechnung der Messwerte $\hat{H}(\mu)$ und $\hat{H}_*(\mu)$ der beiden Übertragungsfunktionen mit Gleichung (3.8). Die Werte der Übertragungsfunktionen werden auf zwei Speichern H und HS mit je $2 \cdot M$ Speicherplätzen abgelegt.

$$\begin{bmatrix} H(\mu) & HS(\mu) \end{bmatrix} = \begin{bmatrix} C_{YV}(\mu) & C_{VY}(\mu) \end{bmatrix} \cdot K_{VV}(\mu) \quad \text{für } \mu = 0 \ (1) \ M-1$$

31. Berechnung des empirischen Mittelwertes des Spektrums der Summe der Signale am Ausgang der beiden linearen Modellsysteme. Das L -fache dieser Werte wird auf einem Speicher X_mittel mit $2 \cdot M$ Speicherplätzen abgelegt.

$$X_mittel(\mu) = \begin{bmatrix} H(\mu) & HS(\mu) \end{bmatrix} \cdot \begin{bmatrix} Akku_V(\mu) \\ Akku_V(M-\mu)^* \end{bmatrix} \quad \text{für } \mu = 1 \ (1) \ M-1$$

$$X_mittel(0) = \begin{bmatrix} H(0) & HS(0) \end{bmatrix} \cdot \begin{bmatrix} Akku_V(0) \\ Akku_V(0)^* \end{bmatrix}$$

32. Berechnung der Messwerte $\hat{U}_f(\mu)$ des Spektrums der gefensterten Störung gemäß Gleichung (3.5). Diese Werte werden für die weiteren Berechnungen nicht mehr benötigt.

$$U(\mu) = \frac{1}{L} \cdot (Y_mittel(\mu) - X_mittel(\mu)) \quad \text{für } \mu = 0 \ (1) \ M-1$$

33. Berechnung des empirischen Mittelwertes des Signals am Ausgang des linearen Modellsystems. Das L -fache dieser Werte wird auf einem Speicher x_mittel mit $2 \cdot M$ Speicherplätzen abgelegt.

$$x_mittel(k) = \text{ifft}_M(X_mittel(\mu)) \quad \text{für } \mu = 0 \ (1) \ M-1 \quad \wedge \quad k = 0 \ (1) \ M-1$$

34. Berechnung der Messwerte $\hat{u}(k)$ der deterministischen Störung mit Hilfe der Gleichung (3.14). Diese Werte werden für die weiteren Berechnungen nicht mehr benötigt.

$$u(k) = \frac{1}{L} \cdot (Akku_y(k) - x_mittel(k - \kappa \cdot M))$$

$$\text{für } k = 0 \ (1) \ F-1 \quad \wedge \quad 0 \leq k - \kappa \cdot M < M$$

35. Berechnung der Messwerte $\hat{\Phi}_n(\mu)$ nach Gleichung (3.34a) zur Beschreibung des LDS. Diese Werte werden auf einem Speicher Φ mit M Speicherplätzen abgelegt.

$$\Phi(\mu) = \frac{C_{YQ}(\mu) - [C_{YV}(\mu), C_{VY}(\mu)] \cdot K_{VV}(\mu) \cdot [C_{YV}(\mu), C_{VY}(\mu)]^H}{M \cdot L \cdot (L-3)}$$

$$\text{für } \mu = 0 \ (1) \ M-1$$

36. Berechnung der Messwerte $\hat{\Psi}_{\mathbf{n}}(\mu)$ nach Gleichung (3.34b) zur Beschreibung des KLDS. Diese Werte werden auf einem Speicher **Psi** mit $M+2$ Speicherplätzen abgelegt.

$$\mathbf{Psi}(\mu) = \frac{\mathbf{C_YY}(\mu) - [\mathbf{C_YV}(\mu), \mathbf{C_VY}(\mu)] \cdot \mathbf{K_VV}(\mu) \cdot [\mathbf{C_VY}(M-\mu), \mathbf{C_YV}(M-\mu)]^T}{M \cdot L \cdot (L-3)}$$

$$\text{für } \mu = 1 \text{ (1) } \frac{M}{2}$$

$$\mathbf{Psi}(0) = \frac{\mathbf{C_YY}(0) - [\mathbf{C_YV}(0), \mathbf{C_VY}(0)] \cdot \mathbf{K_VV}(0) \cdot [\mathbf{C_VY}(0), \mathbf{C_YV}(0)]^T}{M \cdot L \cdot (L-3)}$$

$$\emptyset) \quad \mathbf{Psi}(M-\mu) = \mathbf{Psi}(\mu) \quad \text{sonst}$$

37. Berechnung der Varianzen $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)}$ der Messwerte der Übertragungsfunktion $H(\mu)$ gemäß Gleichung (3.44a). Diese Werte werden auf einem Speicher **C_HQ** mit M Speicherplätzen abgelegt.

$$\mathbf{C_HQ}(\mu) = M \cdot L \cdot \mathbf{K_VV}(\mu, 1, 1) \cdot \mathbf{Phi}(\mu) \quad \text{für } \mu = 0 \text{ (1) } M-1$$

38. Berechnung der Varianzen $\hat{C}_{\hat{\mathbf{H}}_*(\mu), \hat{\mathbf{H}}_*(\mu)}$ der Messwerte der Übertragungsfunktion $H_*(\mu)$ gemäß Gleichung (3.44a). Diese Werte werden auf einem Speicher **C_HSQ** mit M Speicherplätzen abgelegt.

$$\mathbf{C_HSQ}(\mu) = M \cdot L \cdot \mathbf{K_VV}(\mu, 2, 2) \cdot \mathbf{Phi}(\mu) \quad \text{für } \mu = 0 \text{ (1) } M-1$$

39. Berechnung der Kovarianzen $\hat{C}_{\hat{\mathbf{H}}(\mu), \hat{\mathbf{H}}(\mu)^*}$ der Messwerte der Übertragungsfunktion $H(\mu)$ gemäß Gleichung (3.44b). Bei Verwendung einer hoch frequenzselektiven Fensterfolge sind nur die Werte für $\mu=0$ und $\mu = M/2$ nennenswert von null verschieden. In Gleichung (3.44b) ergibt das Produkt aus der Matrix in der Mitte und der rechten, inversen Matrix die Permutationsmatrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$. Die Einheitsvektoren links und rechts greifen ein Element der verbleibenden inversen, konjugierten Matrix heraus. Die Kovarianzen werden auf einem Speicher **C_HH** mit 4 Speicherplätzen abgelegt.

$$\mathbf{C_HH}(\mu) = M \cdot L \cdot \mathbf{K_VV}(\mu, 2, 1)^* \cdot \mathbf{Psi}(\mu) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

40. Berechnung der Kovarianzen $\hat{C}_{\hat{\mathbf{H}}_*(\mu), \hat{\mathbf{H}}_*(\mu)^*}$ der Messwerte der Übertragungsfunktion $H_*(\mu)$ gemäß Gleichung (3.44b). Bei Verwendung einer hoch frequenzselektiven Fensterfolge sind nur die Werte für $\mu=0$ und $\mu = M/2$ nennenswert von null verschieden. In Gleichung (3.44b) ergibt das Produkt aus der Matrix in der Mitte und der rechten, inversen Matrix die Permutationsmatrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$. Die Einheitsvektoren

links und rechts greifen ein Element der verbleibenden inversen, konjugierten Matrix heraus. Die Kovarianzen werden auf einem Speicher **C_HSHS** mit 4 Speicherplätzen abgelegt.

$$\mathbf{C_HSHS}(\mu) = M \cdot L \cdot \mathbf{K_VV}(\mu, 1, 2)^* \cdot \mathbf{Psi}(\mu) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

41. Berechnung der Varianzen $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)}$ der Messwerte des Spektrums der gefenster-ten Störung gemäß Gleichung (3.48a). In dieser Gleichung ergibt das Produkt aus der Matrix in der Mitte und der rechten, inversen Matrix die Einheitsmatrix. Die Messwertvarianzen werden auf einem Speicher **C_UQ** mit M Speicherplätzen abge-legt.

$$\mathbf{C_UQ}(\mu) = \frac{M}{L} \cdot \left(1 + \begin{bmatrix} \mathbf{Akku_V}(\mu) \\ \mathbf{Akku_V}(M-\mu)^* \end{bmatrix}^H \cdot \mathbf{K_VV}(\mu) \cdot \begin{bmatrix} \mathbf{Akku_V}(\mu) \\ \mathbf{Akku_V}(M-\mu)^* \end{bmatrix} \right) \cdot \mathbf{Phi}(\mu)$$

für $\mu = 1 \ (1) \ M-1$

$$\mathbf{C_UQ}(0) = \frac{M}{L} \cdot \left(1 + \begin{bmatrix} \mathbf{Akku_V}(0) \\ \mathbf{Akku_V}(0)^* \end{bmatrix}^H \cdot \mathbf{K_VV}(0) \cdot \begin{bmatrix} \mathbf{Akku_V}(0) \\ \mathbf{Akku_V}(0)^* \end{bmatrix} \right) \cdot \mathbf{Phi}(0)$$

42. Berechnung der Kovarianzen $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)^*}$ der Messwerte des Spektrums der gefens-terten Störung gemäß Gleichung (3.48b). Bei Verwendung einer hoch frequenzse-lektiven Fensterfolge sind nur die Kovarianzen für $\mu=0$ und $\mu = M/2$ nennenswert von null verschieden. Da jedoch für die Berechnung der Kovarianzen der determi-nistischen Störung im Zeitbereich alle Kovarianzen $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(-\mu)^*}$ benötigt werden, und diese für die beiden angegebenen Frequenzen mit den obengenannten überein-stimmen, werden gleich alle berechnet. In Gleichung (3.48b) ergibt das Produkt aus der Matrix in der Mitte und der rechten, inversen Matrix die Permutationsmatrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$. Die Kovarianzen werden auf einem Speicher **C_UU** mit $M+2$ Speicherplätzen abgelegt.

$$\mathbf{C_UU}(\mu) = \frac{M}{L} \cdot \left(1 + \begin{bmatrix} \mathbf{Akku_V}(\mu) \\ \mathbf{Akku_V}(M-\mu)^* \end{bmatrix}^H \cdot \mathbf{K_VV}(\mu) \cdot \begin{bmatrix} \mathbf{Akku_V}(\mu) \\ \mathbf{Akku_V}(M-\mu)^* \end{bmatrix} \right) \cdot \mathbf{Psi}(\mu)$$

für $\mu = 1 \ (1) \ \frac{M}{2}$

$$\mathbf{C_UU}(0) = \frac{M}{L} \cdot \left(1 + \begin{bmatrix} \mathbf{Akku_V}(0) \\ \mathbf{Akku_V}(0)^* \end{bmatrix}^H \cdot \mathbf{K_VV}(0) \cdot \begin{bmatrix} \mathbf{Akku_V}(0) \\ \mathbf{Akku_V}(0)^* \end{bmatrix} \right) \cdot \mathbf{Psi}(0)$$

$$\emptyset) \quad \mathbf{C_UU}(M-\mu) = \mathbf{C_UU}(\mu) \quad \text{sonst}$$

43. Berechnung der Varianz $\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)}$ der Messwerte der Störung mit Hilfe der Gleichung (3.53a). Hier handelt es sich um den reellen Wert der zeitunabhängigen Näherung. Er wird auf einem Speicherplatz $\mathbf{C_uQ}$ abgelegt.

$$\mathbf{C_uQ} = \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \mathbf{C_uQ}(\mu)$$

44. Berechnung der Kovarianz $\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)^*}$ der Messwerte der Störung mit Hilfe der Gleichung (3.53b). Hier handelt es sich um den komplexen Wert der zeitunabhängigen Näherung. Er wird auf zwei Speicherplätzen $\mathbf{C_uu}$ abgelegt.

$$\mathbf{C_uu} = \frac{1}{M^2} \cdot \sum_{\mu=0}^{M-1} \mathbf{C_uu}(\mu)$$

45. Berechnung der Varianzen $\hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)}$ der LDS-Messwerte nach Gleichung (3.58a) und (3.61a). Diese Werte werden auf einem Speicher $\mathbf{C_PhiQ}$ mit M Speicherplätzen abgelegt.

$$\mathbf{C_PhiQ}(\mu) = \frac{L-5}{(L-1) \cdot (L-4)} \cdot \mathbf{Phi}(\mu)^2 + \frac{L-3}{(L-1) \cdot (L-4)} \cdot |\mathbf{Psi}(\mu)|^2$$

für $\mu \in \{0; \frac{M}{2}\}$

$$\mathbf{C_PhiQ}(\mu) = \frac{1}{L-2} \cdot \mathbf{Phi}(\mu)^2 \quad \text{für } \mu = 1 \text{ (1) } M-1 \quad \wedge \quad \mu \neq \frac{M}{2}$$

46. Berechnung der Varianzen $\hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)}$ der KLDS-Messwerte nach Gleichung (3.58c) und (3.61c). Diese Werte werden auf einem Speicher $\mathbf{C_PsiQ}$ mit $M/2+1$ Speicherplätzen abgelegt.

$$\mathbf{C_PsiQ}(\mu) = \frac{2 \cdot (L-3)}{(L-1) \cdot (L-4)} \cdot \mathbf{Phi}(\mu)^2 - \frac{2}{(L-1) \cdot (L-4)} \cdot |\mathbf{Psi}(\mu)|^2$$

für $\mu \in \{0; \frac{M}{2}\}$

$$\mathbf{C_PsiQ}(\mu) = \frac{(L-3)}{(L-2) \cdot (L-4)} \cdot \mathbf{Phi}(\mu) \cdot \mathbf{Phi}(M-\mu) - \frac{1}{(L-2) \cdot (L-4)} \cdot |\mathbf{Psi}(\mu)|^2$$

für $\mu = 1 \text{ (1) } \frac{M}{2}-1$

$$\emptyset) \quad \mathbf{C_PsiQ}(M-\mu) = \mathbf{C_PsiQ}(\mu) \quad \text{sonst}$$

47. Berechnung der Kovarianzen $\hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*}$ der KLDS-Messwerte nach Gleichung (3.58d) und (3.61d). Diese Werte werden auf einem Speicher $\mathbf{C_PsiPsi}$ mit $M+2$ Speicherplätzen abgelegt.

$$\begin{aligned} \mathbf{C_PsiPsi}(\mu) &= \frac{2}{L-1} \cdot \mathbf{Psi}(\mu)^2 & \text{für } \mu \in \left\{0; \frac{M}{2}\right\} \\ \mathbf{C_PsiPsi}(\mu) &= \frac{1}{L-2} \cdot \mathbf{Psi}(\mu)^2 & \text{für } \mu = 1 \text{ (1) } \frac{M}{2} - 1 \end{aligned}$$

$$\emptyset) \quad \mathbf{C_PsiPsi}(M-\mu) = \mathbf{C_PsiPsi}(\mu) \quad \text{sonst}$$

48. Festlegen des Parameters α des Konfidenzniveaus $1-\alpha$.

49. Konfidenzgebiete der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$. Bei allen Frequenzen außer bei $\mu=0$ und $\mu = M/2$ ergeben sich Kreise um die Messwerte $\hat{H}(\mu)$. Der Radius dieser Kreise ist gleich dem Betrag der Halbachsen der Konfidenzellipsen gemäß der Gleichungen ([1]:3.80):

$$\begin{aligned} \mathbf{A_1_H}(\mu) &= \sqrt{-\ln(\alpha) \cdot \mathbf{C_HQ}(\mu)} \\ &\text{für } \mu = 1 \text{ (1) } M-1 \quad \wedge \quad \mu \neq \frac{M}{2} \\ \mathbf{A_2_H}(\mu) &= j \cdot \mathbf{A_1_H}(\mu) \end{aligned}$$

Bei den Frequenzen $\mu=0$ und $\mu = M/2$ ergeben sich Konfidenzellipsen um die Messwerte $\hat{H}(\mu)$. Die Halbachsen der Ellipsen sind:

$$\begin{aligned} \mathbf{A_1_H}(\mu) &= \sqrt{-\ln(\alpha) \cdot (\mathbf{C_HQ}(\mu) + |\mathbf{C_HH}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{\mathbf{C_HH}(\mu)\}} \\ \mathbf{A_2_H}(\mu) &= j \cdot \sqrt{-\ln(\alpha) \cdot (\mathbf{C_HQ}(\mu) - |\mathbf{C_HH}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{\mathbf{C_HH}(\mu)\}} \\ &\text{für } \mu \in \left\{0; \frac{M}{2}\right\} \end{aligned}$$

50. Konfidenzgebiete der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$. Bei allen Frequenzen außer bei $\mu=0$ und $\mu = M/2$ ergeben sich Kreise um die Messwerte $\hat{H}_*(\mu)$. Der Radius dieser Kreise ist gleich dem Betrag der Halbachsen der Konfidenzellipsen gemäß der Gleichungen ([1]:3.80):

$$\begin{aligned} \mathbf{A_1_HS}(\mu) &= \sqrt{-\ln(\alpha) \cdot \mathbf{C_HSQ}(\mu)} \\ &\text{für } \mu = 1 \text{ (1) } M-1 \quad \wedge \quad \mu \neq \frac{M}{2} \\ \mathbf{A_2_HS}(\mu) &= j \cdot \mathbf{A_1_HS}(\mu) \end{aligned}$$

Bei den Frequenzen $\mu=0$ und $\mu = M/2$ ergeben sich Konfidenzellipsen um die Messwerte $\hat{H}_*(\mu)$. Die Halbachsen der Ellipsen sind:

$$\begin{aligned} A_{1_HS}(\mu) &= \sqrt{-\ln(\alpha) \cdot (C_{HSQ}(\mu) + |C_{HSHS}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{HSHS}(\mu)\}} \\ A_{2_HS}(\mu) &= j \cdot \sqrt{-\ln(\alpha) \cdot (C_{HSQ}(\mu) - |C_{HSHS}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{HSHS}(\mu)\}} \\ &\text{für } \mu \in \{0; \frac{M}{2}\} \end{aligned}$$

51. Konfidenzgebiete der Messwerte des Spektrums der gefensterten Störung. Bei allen Frequenzen außer bei $\mu=0$ und $\mu = M/2$ ergeben sich Kreise um die Messwerte $\hat{U}_f(\mu)$. Der Radius dieser Kreise ist gleich dem Betrag der Halbachsen der Konfidenzellipsen, die sich analog zu den Gleichungen ([1]:3.80) berechnen:

$$\begin{aligned} A_{1_U}(\mu) &= \sqrt{-\ln(\alpha) \cdot C_{UQ}(\mu)} \\ &\text{für } \mu = 1 \text{ (1) } M-1 \quad \wedge \quad \mu \neq \frac{M}{2} \\ A_{2_U}(\mu) &= j \cdot A_{1_U}(\mu) \end{aligned}$$

Bei den Frequenzen $\mu=0$ und $\mu = M/2$ ergeben sich Konfidenzellipsen um die Messwerte $\hat{U}_f(\mu)$. Die Halbachsen der Ellipsen sind:

$$\begin{aligned} A_{1_U}(\mu) &= \sqrt{-\ln(\alpha) \cdot (C_{UQ}(\mu) + |C_{UU}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{UU}(\mu)\}} \\ A_{2_U}(\mu) &= j \cdot \sqrt{-\ln(\alpha) \cdot (C_{UQ}(\mu) - |C_{UU}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{UU}(\mu)\}} \\ &\text{für } \mu \in \{0; \frac{M}{2}\} \end{aligned}$$

52. Konfidenzgebiete der Messwerte der deterministischen Störung. Die Konfidenzgebiete werden für alle Zeitpunkte gleich abgeschätzt. Die zeitunabhängigen Halbachsen der Ellipsen berechnen sich analog zu den Gleichungen ([1]:3.80) und sind:

$$\begin{aligned} A_{1_u} &= \sqrt{-\ln(\alpha) \cdot (C_{uQ} + |C_{uu}|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{uu}\}} \\ A_{2_u} &= j \cdot \sqrt{-\ln(\alpha) \cdot (C_{uQ} - |C_{uu}|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{uu}\}} \end{aligned}$$

53. Konfidenzintervalle der LDS-Messwerte. Da diese Messwerte immer reell sind, ergeben sich Konfidenzintervalle $\left[\hat{\Phi}_n(\mu) - \hat{A}_\Phi(\mu); \hat{\Phi}_n(\mu) + \hat{A}_\Phi(\mu) \right]$, die symmetrisch zum Messwert liegen. Die halbe Intervallbreite berechnet sich mit Gleichung ([1]:3.72) zu:

$$A_{\Phi}(\mu) = \sqrt{2} \cdot \sqrt{C_{\Phi}Q(\mu)} \cdot \text{erfc}^{-1}(\alpha) \quad \text{für } \mu = 0 \text{ (1) } M-1$$

54. Konfidenzgebiete der KLDS-Messwerte. Bei allen Frequenzen ergeben sich Konfidenzellipsen um die Messwerte $\hat{\Psi}_{\mathbf{n}}(\mu)$. Die Halbachsen der Ellipsen, die sich analog zu den Gleichungen ([1]:3.80) berechnen, sind:

$$\begin{aligned}
 A_{1_Psi}(\mu) &= \sqrt{-\ln(\alpha) \cdot (C_{PsiQ}(\mu) + |C_{PsiPsi}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{PsiPsi}(\mu)\}} \\
 A_{2_Psi}(\mu) &= j \cdot \sqrt{-\ln(\alpha) \cdot (C_{PsiQ}(\mu) - |C_{PsiPsi}(\mu)|)} \cdot e^{\frac{j}{2} \cdot \angle\{C_{PsiPsi}(\mu)\}} \\
 &\quad \text{für } \mu = 0 \text{ (1) } \frac{M}{2} \\
 \emptyset) \quad &\begin{aligned} A_{1_Psi}(\mu) &= A_{1_Psi}(M-\mu) \\ A_{2_Psi}(\mu) &= A_{2_Psi}(M-\mu) \end{aligned} \quad \text{sonst}
 \end{aligned}$$

Unter den Voraussetzungen, die im Laufe der Herleitung des RKM mit Fensterung für komplexwertige Systeme gemacht worden sind, sind die so gewonnenen Messwerte $\hat{\mathbf{H}}(\mu)$, $\hat{\mathbf{U}}_f(\mu)$, $\hat{\mathbf{u}}(k)$, $\hat{\Phi}_{\mathbf{n}}(\mu)$ und $\hat{\Psi}_{\mathbf{n}}(\mu)$ erwartungstreu und konsistent. Die Qualität der Messung kann mit Hilfe der angegebenen Konfidenzgebiete und Intervalle abgeschätzt werden.

8 Aufwandsabschätzung für andere Varianten des RKM

In diesem Kapitel soll nun untersucht werden, wie der im letzten Kapitel geschilderte Ablauf zur Berechnung der Messwerte zu modifizieren ist, wenn man das periodisch zeitvariante Modellsystem mit $K_H > 1$ bei Vorliegen einer zyklstationären Störung mit $K_\Phi > 1$ verwendet. Im Anschluss daran werden die Varianten mit den reduzierten Systemmodellen nach [1] ebenso untersucht, wie auch die Varianten zur Messung an reellwertigen Systemen nach Kapitel 6 und zur Spektralschätzung nach Kapitel 5.

Wird ein periodisch zeitvariantes Modellsystem angesetzt, und ist der Approximationsfehlerprozess zyklstationär, so benötigt man zusätzliche Akkumulatorfelder um die Kovarianzen der Spektralwerte bei der Frequenz μ mit den Spektralwerten bei den um Vielfache von M/K_S verschobenen Frequenzen empirisch ermitteln zu können. Man muss nun in Punkt 4 die Akkumulatorfelder **Akku_VQ**, **Akku_VV**, **Akku_YQ**, **Akku_YY**, **Akku_YV** und **Akku_VY** mit K_S -facher Größe bereitstellen, um alle Produkte

$$\begin{aligned} &V_\lambda(\mu) \cdot V_\lambda\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^*, \\ &V_\lambda(\mu) \cdot V_\lambda\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right), \\ &Y_{f,\lambda}(\mu) \cdot Y_{f,\lambda}\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^*, \\ &Y_{f,\lambda}(\mu) \cdot Y_{f,\lambda}\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right), \\ &Y_{f,\lambda}(\mu) \cdot V_\lambda\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^* \text{ und} \\ &Y_{f,\lambda}(\mu) \cdot V_\lambda\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right) \end{aligned}$$

aller Einzelmessungen in der Messschleife der Punkte 7 bis 18 akkumulieren zu können. In den Punkten 21, 22 und 26 bis 29 wird man dann K_S mal so viele entsprechende $L \cdot (L-1)$ -fache Kovarianzen **C_VQ**, **C_VV**, **C_YQ**, **C_YY**, **C_YV** und **C_VY** erhalten, indem man von den L -fachen Akkumulatoren **Akku_VQ**, **Akku_VV**, **Akku_YQ**, **Akku_YY**, **Akku_YV** und **Akku_VY** die Produkte

$$\begin{aligned} &\text{Akku}_V(\mu) \cdot \text{Akku}_V\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^*, \\ &\text{Akku}_V(\mu) \cdot \text{Akku}_V\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right), \\ &Y_{\text{mittel}}(\mu) \cdot Y_{\text{mittel}}\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^*, \\ &Y_{\text{mittel}}(\mu) \cdot Y_{\text{mittel}}\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right), \\ &Y_{\text{mittel}}(\mu) \cdot \text{Akku}_V\left(\mu + \hat{\mu} \cdot \frac{M}{K_S}\right)^* \text{ und} \\ &Y_{\text{mittel}}(\mu) \cdot \text{Akku}_V\left(-\mu - \hat{\mu} \cdot \frac{M}{K_S}\right) \text{ subtrahiert.} \end{aligned}$$

Da K_S ein ganzzahliges Vielfaches von K_H ist, benötigt man zur Berechnung der Messwerte der beiden bifrequenten Übertragungsfunktionen mit Hilfe der Gleichung (3.8) keine weiteren $L \cdot (L-1)$ -fachen Kovarianzen. Alle dort in dem Kovarianzvektor $\hat{\underline{C}}_{\mathbf{Y}_f(\mu), \tilde{\mathbf{V}}(\mu)}$ und der Kovarianzmatrix $\hat{\underline{C}}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}$ vorkommenden Kovarianzen sind, abgesehen von dem Faktor $L \cdot (L-1)$, der sich sowieso herauskürzt, in den Feldern $\mathbf{C_VQ}$, $\mathbf{C_VV}$, $\mathbf{C_YV}$ und $\mathbf{C_VY}$ enthalten.

Es sei darauf hingewiesen, dass die Zufallsvektoren $\tilde{\mathbf{V}}(\mu)$ nach Gleichung (2.15) bei einer Erhöhung von μ um ein ganzzahliges Vielfaches von M/K_H lediglich permutiert werden. Die dabei auftretende Permutationsmatrix lässt sich mit Hilfe der Permutationsmatrix

$$\underline{Q} = \begin{bmatrix} 0 & 1 & 0 & \cdots & \cdots & 0 \\ \vdots & 0 & 1 & 0 & \cdots & 0 \\ \vdots & & \ddots & \ddots & \ddots & \vdots \\ \vdots & & & 0 & 1 & 0 \\ 0 & \cdots & \cdots & \cdots & 0 & 1 \\ 1 & 0 & \cdots & \cdots & \cdots & 0 \end{bmatrix} \quad (8.1)$$

gemäß

$$\underline{P} = \begin{bmatrix} \underline{Q} & \underline{0} \\ \underline{0} & \underline{Q} \end{bmatrix} \quad (8.2)$$

konstruieren. Für die Zufallsvektoren gilt dann:

$$\tilde{\mathbf{V}}(\mu) = \tilde{\mathbf{V}}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) = \underline{P}^{\hat{\mu}} \cdot \tilde{\mathbf{V}}(\tilde{\mu}). \quad (8.3)$$

Daher sind die Kovarianzmatrizen aller Gleichungssysteme (2.25) mit Werten von μ , die sich um ein ganzzahliges Vielfaches von M/K_H unterscheiden, ebenfalls permutierte Versionen der Kovarianzmatrix die man mit $0 \leq \mu < M/K_H$ erhält.

$$\begin{aligned} \underline{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)} &= \mathbb{E} \left\{ \left(\tilde{\mathbf{V}}(\mu) - \mathbb{E} \{ \tilde{\mathbf{V}}(\mu) \} \right) \cdot \left(\tilde{\mathbf{V}}(\mu) - \mathbb{E} \{ \tilde{\mathbf{V}}(\mu) \} \right)^H \right\} = \\ &= \mathbb{E} \left\{ \left(\tilde{\mathbf{V}}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) \} \right) \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) \} \right)^H \right\} = \\ &= \mathbb{E} \left\{ \left(\underline{P}^{\hat{\mu}} \cdot \tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \underline{P}^{\hat{\mu}} \cdot \tilde{\mathbf{V}}(\tilde{\mu}) \} \right) \cdot \left(\underline{P}^{\hat{\mu}} \cdot \tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \underline{P}^{\hat{\mu}} \cdot \tilde{\mathbf{V}}(\tilde{\mu}) \} \right)^H \right\} = \\ &= \mathbb{E} \left\{ \underline{P}^{\hat{\mu}} \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu}) \} \right) \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu}) \} \right)^H \cdot (\underline{P}^{\hat{\mu}})^H \right\} = \\ &= \underline{P}^{\hat{\mu}} \cdot \mathbb{E} \left\{ \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu}) \} \right) \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E} \{ \tilde{\mathbf{V}}(\tilde{\mu}) \} \right)^H \right\} \cdot (\underline{P}^{\hat{\mu}})^H = \\ &= \underline{P}^{\hat{\mu}} \cdot \underline{C}_{\tilde{\mathbf{V}}(\tilde{\mu}), \tilde{\mathbf{V}}(\tilde{\mu})} \cdot (\underline{P}^{\hat{\mu}})^H \end{aligned} \quad (8.4)$$

Dieselbe Permutationssymmetrie gilt auch für die Stichprobenmatrizen

$$\tilde{V}(\mu) = \tilde{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}) = \underline{P}^{\hat{\mu}} \cdot \tilde{V}(\tilde{\mu}) \quad (8.5)$$

und die empirischen Kovarianzmatrizen

$$\hat{C}_{\tilde{V}(\mu), \tilde{V}(\mu)} = \hat{C}_{\tilde{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H}), \tilde{V}(\tilde{\mu} + \hat{\mu} \cdot \frac{M}{K_H})} = \underline{P}^{\hat{\mu}} \cdot \hat{C}_{\tilde{V}(\tilde{\mu}), \tilde{V}(\tilde{\mu})} \cdot (\underline{P}^{\hat{\mu}})^H \quad (8.6)$$

in den Lösungen (3.8) für die beiden bifrequenten Übertragungsfunktionen mit Werten von μ , die sich um ein ganzzahliges Vielfaches von M/K_H unterscheiden. Eine weitere Symmetrie ergibt sich für negative Frequenzen. Mit der Antidiagonalmatrix

$$\underline{\overline{D}} = \begin{bmatrix} 0 & 0 & \cdots & \cdots & 0 & 0 & 1 \\ 0 & 0 & \cdots & \cdots & 0 & 1 & 0 \\ \vdots & \vdots & & \ddots & 1 & 0 & 0 \\ \vdots & \vdots & \ddots & \ddots & \ddots & \vdots & \vdots \\ 0 & 0 & 1 & \ddots & & \vdots & \vdots \\ 0 & 1 & 0 & \cdots & \cdots & 0 & 0 \\ 1 & 0 & 0 & \cdots & \cdots & 0 & 0 \end{bmatrix} \quad (8.7)$$

gilt für die Zufallsvektoren der Erregung und deren Stichprobenmatrizen

$$\tilde{\mathbf{V}}(-\mu)^* = \underline{\overline{D}} \cdot \underline{P} \cdot \tilde{\mathbf{V}}(\mu) \quad \text{und} \quad \tilde{V}(-\mu)^* = \underline{\overline{D}} \cdot \underline{P} \cdot \tilde{V}(\mu), \quad (8.8)$$

und somit ergeben sich die Permutationssymmetrien

$$\begin{aligned} C_{\tilde{V}(-\mu), \tilde{V}(-\mu)}^* &= \underline{\overline{D}} \cdot \underline{P} \cdot C_{\tilde{V}(\mu), \tilde{V}(\mu)} \cdot \underline{P}^H \cdot \underline{\overline{D}}^H \quad \text{und} \\ \hat{C}_{\tilde{V}(-\mu), \tilde{V}(-\mu)}^* &= \underline{\overline{D}} \cdot \underline{P} \cdot \hat{C}_{\tilde{V}(\mu), \tilde{V}(\mu)} \cdot \underline{P}^H \cdot \underline{\overline{D}}^H \end{aligned} \quad (8.9)$$

für die theoretischen und die empirischen Kovarianzmatrizen. Man braucht zur Berechnung der beiden bifrequenten Übertragungsfunktionen letztendlich nur etwa $M/(2 \cdot K_H)$ empirische Kovarianzmatrizen, deren Dimension hier jeweils $(2 \cdot K_H) \times (2 \cdot K_H)$ ist, zu invertieren.

Zur Berechnung der Messwerte des LDS und des KLDS gemäß der Gleichungen (3.34) benötigt man u. a. auch die Matrizen $\underline{V}_{\perp}(\mu)$ nach Gleichung (3.28). Wenn man diese Matrizen nach dem in Kapitel 3 beschriebenen Verfahren konstruiert, so dass die Gleichungen (3.33) erfüllt sind, ist es bei der Berechnung der Messwerte erforderlich, die in

den Matrizen $V_{\perp}(\mu)$ nach Gleichung (3.28) auftretenden empirischen Kovarianzmatrizen $\hat{C}_{\check{V}(\mu), \check{V}(\mu)}$ nach Gleichung (3.27) zu bestimmen, die alle empirischen Kovarianzen eines Satzes von linear unabhängigen Zufallgrößen untereinander enthalten, die aus den Zufallgrößen $V(\mu + \check{\mu} \cdot \frac{M}{K_S})$ und $V(-\mu - \check{\mu} \cdot \frac{M}{K_S})^*$ für alle $\check{\mu} = 0 \dots K_S - 1$ entnommen wurden. Nimmt man an, dass all diese Zufallgrößen linear unabhängig sind, so ist $K(\mu) = 2 \cdot K_S$ und es sind für jede der Frequenzen $\mu = 0 \dots (M - K_S)/(2 \cdot K_S)$ genau $4 \cdot K_S^2$ empirische Kovarianzen des Spektrums der Erregung zu berechnen. Die Kovarianzmatrizen für die restlichen Frequenzen erhält man dann durch geeignete Permutation aus den Kovarianzmatrizen $\hat{C}_{\check{V}(\mu), \check{V}(\mu)}$ bei den angegebenen Frequenzen. Wie man das $L \cdot (L - 1)$ -fache der in den Kovarianzmatrizen auftretenden empirischen Kovarianzen mit Hilfe der Akkumulatorfelder `Akku_VQ`, `Akku_VV` und `Akku_V` berechnet, wurde oben bereits angegeben. Aus den $L \cdot (L - 1)$ -fachen Kovarianzen `C_VQ` und `C_VV` sucht man sich jeweils die für die Frequenz μ in der Kovarianzmatrix enthaltenen empirischen Kovarianzen heraus.

K_S ist als das kleinste gemeinsame Vielfache der beiden Perioden K_H und K_{Φ} definiert. Somit kann die Zahl der Elemente der empirischen Kovarianzmatrix groß werden, wenn beide Perioden teilerfremd sind. Da die empirische Kovarianzmatrix invertiert werden muss, wird dann nicht nur der Speicherbedarf, sondern auch der Rechenaufwand relativ groß. Bei vielen periodisch zeitvarianten Systemen ist die Periode der überlagerten zufälligen Störung mit der Periode des zeitvarianten Systems zyklstationär. In diesem Fall hält sich der Speicher- und Rechenaufwand noch in vertretbaren Grenzen.

Zur Berechnung der Messwerte $\hat{\Phi}_n(\mu, \mu + \check{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\hat{\Psi}_n(\mu, \mu + \check{\mu} \cdot \frac{M}{K_{\Phi}})$ benötigt man für jede Frequenz μ alle empirischen Kreuzkovarianzen der Spektralwerte des gefensterten Ausgangssignals $\vec{Y}_f(\mu + \check{\mu} \cdot \frac{M}{K_{\Phi}})$ und der Zufallgrößen $V(\mu + \check{\mu} \cdot \frac{M}{K_S})$ und $V(-\mu - \check{\mu} \cdot \frac{M}{K_S})^*$ des Spektrums der Erregung. Wie man das $L \cdot (L - 1)$ -fache dieser empirischen Kovarianzen `C_YV` und `C_VY` mit Hilfe der Akkumulatorfelder `Akku_YV`, `Akku_VY`, `Akku_V` und `Y_mittel` berechnet, ist oben angegeben. Auch hier sucht man sich die für die Messwerte der Frequenz μ jeweils benötigten Kreuzkovarianzen heraus. Desweiteren werden für jede Frequenz μ noch Schätzwerte für alle Kovarianzen der Zufallsgrößen $\vec{Y}_f(\mu + \check{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\vec{Y}_f(-\mu - \check{\mu} \cdot \frac{M}{K_{\Phi}})$ untereinander benötigt, die man analog für jede Frequenz μ mit Hilfe der Akkumulatorfelder `Akku_YQ`, `Akku_YY` und `Y_mittel` berechnet, und deren $L \cdot (L - 1)$ -faches man mit etwa K_{Φ} -fachen Speichbedarf auf den Feldern `C_YQ` und `C_YY` abspeichert.

Unabhängig davon, welches Modellsystem man verwendet (zeitinvariant oder periodisch zeitvariant und evtl. das von dem konjugierten Eingangssignal erregte Modellsystem), sind die für die Berechnung der Übertragungsfunktionen und der deterministischen Störung sowie ihres Spektrums benötigten Kovarianzen bereits in den Kovarianzen, die man zur Berechnung des LDS und KLDS braucht, enthalten.

Es sei noch darauf hingewiesen, dass die Matrixprodukte der bilinearen Formen, die als Vorfaktoren vor den Messwerten $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ bei den Messwertkovarianzen der Übertragungsfunktionen gemäß der Gleichungen (3.44) stehen, sich wegen der obengenannten Permutationseigenschaften ebenfalls durch eine Permutation ineinander überführen lassen, wenn man μ um ein Vielfaches von M/K_H erhöht. Bei den Messwertkovarianzen des Spektrums der deterministischen Störung gemäß der Gleichungen (3.48) bewirkt eine Verschiebung von μ um ein Vielfaches von M/K_H gar keine Veränderung der bilinearen Formen vor den Messwerten des LDS und KLDS. Im Fall, dass K_H ein ganzzahliges Vielfaches von K_{Φ} ist, ergeben sich somit Vorfaktoren, die von $\tilde{\mu}$ unabhängig sind. Bei den Messwertvarianzen ist nach der Permutation auch noch die empirischen Kovarianzmatrix in der Mitte des Matrixprodukts in der Gleichung (3.48a) gleich der empirischen Kovarianzmatrix $\hat{C}_{\tilde{\mathbf{V}}(\mu), \tilde{\mathbf{V}}(\mu)}$, so dass diese sich mit einer der beiden inversen Kovarianzmatrizen kompensiert. Wenn man das von der konjugierten Erregung gespeiste Modellsystem mit der bifrequenten Übertragungsfunktion $H_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$ verwendet, erhält man nach einer weiteren Permutation bei den Messwertkovarianzen nach Gleichung (3.48b) dieselben bilinearen Formen wie bei den Messwertvarianzen nach Gleichung (3.48a), und sowohl bei den Messwertvarianzen, als auch bei den Messwertkovarianzen, kann die eine Inverse wieder gekürzt werden, so dass $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ mit demselben Vorfaktor in die Messwertvarianzen eingeht, wie $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ in die Messwertkovarianzen.

Nun soll angedeutet werden, welche Vereinfachungen an dem im Kapitel 7 geschilderten Ablauf zur Berechnung der Messwerte vorgenommen werden können, wenn man die reduzierten Systemmodelle nach Kapitel [1]:3 bis [1]:4 und Kapitel 5 und 6 verwendet. Wie in Kapitel 7 beschränken wir uns auch hier auf den Fall zeitinvarianter Modellsysteme $\mathcal{S}_{lin}$ und $\mathcal{S}_{*,lin}$ mit $K_H = 1$ und nehmen an, dass ein stationärer Approximationsfehlerprozess mit $K_{\Phi} = 1$ vorliegt.

Wenn man davon ausgehen kann, dass ein mittelwertfreier Approximationsfehler vorliegt, kann man auf die Modellierung der deterministischen Störung verzichten, wie dies in [1] geschehen ist. Ein erstes Moment im Ausgangssignal kann unter dieser Annahme nur durch ein durch das lineare Modellsystem verzerrtes erstes Moment der Erregung verursacht sein. Da das erste Moment der Erregung in derselben Art verzerrt wird, wie die in den Stichprobenelementen enthaltenen zufälligen Abweichungen von diesem ersten Moment, braucht man den Mittelwert des Ausgangssignals nun nicht getrennt behandeln. Die Akkumulatorfelder `Akku_y` und `Akku_V`, die bisher zur Bestimmung der empirischen Mittelwerte am Systemein- und -ausgang verwendet wurden, werden daher nicht mehr benötigt, so dass sich ein gegenüber dem vollständigen Systemmodell reduzierter Rechenaufwand ergibt. Zur Berechnung der Messwert und der Schätzwerte der Messwert(ko)varianzen werden

bei dem reduzierten Systemmodell ohne deterministische Störung die Akkumulatorfelder

Akku_VQ	mit	M	Speicherplätzen,
Akku_VV	mit	$M+2$	Speicherplätzen,
Akku_YQ	mit	M	Speicherplätzen,
Akku_YY	mit	$M+2$	Speicherplätzen,
Akku_YV	mit	$2 \cdot M$	Speicherplätzen und
Akku_VY	mit	$2 \cdot M$	Speicherplätzen

benötigt. Das L -fache dieser Akkumulatoren tritt in allen Berechnungen nun an die Stelle der $L \cdot (L-1)$ -fachen Kovarianzen und Kovarianzmatrizen. Die Punkte 11, 14, 21, 22, 25 – 29, 31 – 34, 41 – 44, 51 und 52 entfallen ersatzlos. Bei der Berechnung der Vorfaktoren der Messwerte $\hat{\Phi}_n(\mu)$ und $\hat{\Psi}_n(\mu)$ in den Punkten 35 und 36 ist jeweils ein um 1 erhöhter Matrixrang zu berücksichtigen, indem man dort im Nenner den Faktor $(L-3)$ durch $(L-2)$ ersetzt. Auch bei der Berechnung der Varianzen und Kovarianzen dieser Messwerte in den Punkten 45 bis 47 ist der erhöhte Matrixrang zu berücksichtigen, indem man dort jeweils L durch $L+1$ ersetzt.

In [1] wurde nur das eine zeitinvariante Modellsystem $\mathcal{S}_{lin}$ angesetzt. Auch bei dieser Variante benötigt man alle eben genannten Akkumulatorfelder. Die Berechnung der Übertragungsfunktion in Punkt 30 erfolgt hier jedoch nach Gleichung ([1]:3.14):

$$H(\mu) = \text{Akku_YV}(\mu) / \text{Akku_VQ}(\mu).$$

In Punkt 37 ergibt sich mit Gleichung ([1]:3.59) die Messwertvarianz

$$C_{\text{HQ}}(\mu) = M / \text{Akku_VQ}(\mu) \cdot \text{Phi}(\mu),$$

und mit Gleichung ([1]:3.64) in Punkt 39 die Messwertkovarianz

$$C_{\text{HH}}(\mu) = M \cdot \text{Akku_VQ}(\mu)^* / \text{Akku_VQ}(\mu)^2 \cdot \text{Psi}(\mu).$$

Abgesehen davon, dass die Punkte 38, 40 und 50 nun entfallen, bleibt die Berechnung der weiteren Messwerte und ihrer Varianzen und Kovarianzen unverändert, wobei die Erhöhung des Matrixranges, wie eben beschrieben, zu berücksichtigen ist.

Bei dem Spektralschätzverfahren komplexer Prozesse in Kapitel 5, bei dem das erste und das zweite zentrale Moment getrennt berechnet werden, wird kein lineares Modellsystem angesetzt. Zur Berechnung der Messwerte und der Schätzwerte der Messwert(ko)varianzen werden dann nur die Akkumulatorfelder

Akku_y	mit	$2 \cdot F$	Speicherplätzen,
Akku_YQ	mit	M	Speicherplätzen und
Akku_YY	mit	$M+2$	Speicherplätzen

benötigt. Die anderen Akkumulatorfelder werden nicht benötigt, da diese zur Akkumulation von Produkten benötigt wurden, die Faktoren enthielten, die von der Erregung abhängig waren. Im Fall der reinen Spektralschätzung gibt es entweder keine Erregung, oder sie wird gefiltert dem zu vermessenden Zufallsprozess zugeschlagen. Das erste Moment ist die deterministische Störung, die sich nach Gleichung (5.4) in Punkt 34 nun ohne $\mathbf{x_mittel}(k - \kappa \cdot M)$ zu

$$\mathbf{u}(k) = \mathbf{Akku_y}(k)/L$$

berechnet. In Punkt 32 wird deren Spektrum analog mit Gleichung (5.5) als

$$U(\mu) = \mathbf{Y_mittel}(\mu)/L$$

berechnet. Die Varianzen und Kovarianzen des Spektrums der deterministischen Störung werden in den Punkten 41 und 42 nun mit den Gleichungen (5.8) wesentlich einfacher als

$$\mathbf{C_UQ}(\mu) = \frac{M}{L} \cdot \mathbf{Phi}(\mu) \quad \text{und} \quad \mathbf{C_UU}(\mu) = \frac{M}{L} \cdot \mathbf{Psi}(\mu)$$

berechnet. Die Varianzen und Kovarianzen der deterministischen Störung selbst werden unverändert in den Punkten 43 und 44 bestimmt. Die Berechnung der Messwerte $\hat{\Phi}_n(\mu)$ und $\hat{\Psi}_n(\mu)$ in den Punkten 35 und 36 erfolgt hier gemäß der Gleichungen (5.6):

$$\mathbf{Phi}(\mu) = \frac{\mathbf{C_YQ}(\mu)}{M \cdot L \cdot (L-1)} \quad \text{und} \quad \mathbf{Psi}(\mu) = \frac{\mathbf{C_YY}(\mu)}{M \cdot L \cdot (L-1)}.$$

Bei der Berechnung der Varianzen und Kovarianzen dieser Messwerte in den Punkten 45 bis 47 ist nun ein um 2 erhöhter Matrixrang zu berücksichtigen. Indem man dort jeweils L durch $L+2$ ersetzt erhält man die in den Gleichungen (5.13) angegebenen Werte.

Wird bei der Spektralschätzung eines komplexen Prozesses $\mathbf{y}(k)$ das erste Moment nicht separat gemessen, so wird auch der Akkumulator $\mathbf{Akku_y}$ nicht benötigt, da dieser ja für die Bestimmung des empirischen Mittelwertes vorgesehen ist. Man ersetzt die Punkte 35 und 36 durch die in den Gleichungen ([1]:5.1) und ([1]:5.2) angegebenen Terme und erhält:

$$\mathbf{Phi}(\mu) = \frac{\mathbf{Akku_YQ}(\mu)}{M \cdot L} \quad \text{und} \quad \mathbf{Psi}(\mu) = \frac{\mathbf{Akku_YY}(\mu)}{M \cdot L}.$$

Da hier nun ein um 3 erhöhter Matrixrang zu berücksichtigen ist, ist bei der Berechnung der Varianzen und Kovarianzen dieser Messwerte in den Punkten 45 bis 47 jeweils L durch $L+3$ zu ersetzen, um die in den Gleichungen ([1]:5.4) angegebenen Werte zu erhalten.

Bei Verwendung eines Mehrtonsignals ergibt sich der Vorteil, dass die Summe über alle Betragsquadratspektren des Eingangssignals aller Einzelmessungen, die sonst mit Hilfe des Akkumulators $\mathbf{Akku_VQ}$ berechnet werden muss, entfallen kann, da diese dann das L -fache des bei allen Einzelmessungen konstanten Betragsquadrats des erregenden Spektrums ist.

Bei Verwendung des Chirpsignals ergibt sich hinsichtlich des Rechenaufwands nur dann ein weiterer wesentlicher Vorteil, wenn man lediglich $\hat{H}(\mu)$, $\hat{H}_*(\mu)$, $\hat{U}_f(\mu)$ und $\hat{u}(k)$ messen will. Dann werden zum Einen statt des Akkumulators **Akku_V** mit $2 \cdot M$ Speicherplätzen nur zwei Speicherplätze für die Akkumulation des Real- und Imaginärteils des zufälligen Drehfaktors benötigt. Zum Anderen, wird man auf den Akkumulatoren **Akku_YV** und **Akku_VY** nicht die in Kapitel 7 angegebenen Spektralwertprodukte aufsummieren, sondern stattdessen jeweils die ebenfalls M Werte $y_{f,\lambda}(k)$ des gefensterten und blockweise überlagerten Zeitsignals der Einzelmessungen, die man mit dem zufälligen Drehfaktor $e^{-j \cdot \hat{\varphi}_\lambda}$ bzw. $e^{j \cdot \hat{\varphi}_\lambda}$ dieser Einzelmessung multipliziert. Die beiden für die weitere Berechnung der Messwerte notwendigen DFTs werden dann nur einmal am Ende der Messschleife (Punkt 7 bis 18 in Kapitel 7) durchgeführt. Das Ergebnis der einen DFT wird man anschließend mit dem Spektrum des konstanten, nicht zufälligen Spektralanteils des Chirpsignals multiplizieren, und das Ergebnis der anderen DFT mit dem konjugierten Spektrum. Wenn man jedoch das LDS, das KLDS oder die Messwertvarianzen und -kovarianzen bestimmen will, benötigt man auch die Akkumulatorfelder **Akku_YQ** und **Akku_YY**. Da zu deren Berechnung bei jeder Einzelmessung die DFT notwendig ist, die eben vermieden werden sollte, bringt in diesem Fall die Verwendung des Chirpsignals — abgesehen von dem guten Crest-Faktor — keine Vorteile.

Die im weiteren dargestellten Modifikationen des Messablaufs nach Kapitel 7 ergeben sich bei der Messung eines reellwertigen Systems, bzw. bei der Spektralschätzung reeller Prozesse, wobei auch für die reduzierten Modellsysteme kurz auf die Besonderheiten, die sich bei der Berechnung der Messwerte reeller Systeme und Signale ergeben, eingegangen werden soll.

Zur Berechnung der Messwert werden bei einem vollständigen reellwertigen Modellsystem die Akkumulatorfelder

Akku_y	mit	F	Speicherplätzen,
Akku_V	mit	M	Speicherplätzen,
Akku_VQ	mit	$M/2 + 1$	Speicherplätzen,
Akku_YQ	mit	$M/2 + 1$	Speicherplätzen und
Akku_YV	mit	M	Speicherplätzen

benötigt. Dabei wurde berücksichtigt, dass die Akkumulatorfelder **Akku_VV**, **Akku_YY** und **Akku_VY** nicht benötigt werden. Diese würden nämlich denselben Inhalt wie die Akkumulatorfelder **Akku_VQ**, **Akku_YQ** und **Akku_YV** enthalten. Bei der Berechnung des Speicherbedarfs der Akkumulatorfelder wurde die Symmetrie und die evtl. vorhandene Reellwertigkeit der Spektralwerte und Zeitsignale mit einkalkuliert. Im Messablaufs nach Kapitel 7 brauchen die Punkte 13, 16, 18, 22 – 24, 27, 29, 36, 38 – 40, 42, 44, 46, 47, 50 und 54 nicht berechnet zu werden. Die Werte der Übertragungsfunktion berechnen sich in Punkt

30 nun mit Gleichung (6.8) zu:

$$H(\mu) = C_{YV}(\mu)/C_{VQ}(\mu).$$

Für den L -fachen Mittelwert des Spektrums am Ausgang des linearen Modellsystems ist in Punkt 31

$$X_{\text{mittel}}(\mu) = H(\mu) \cdot \text{Akku}_V(\mu)$$

einzusetzen. Bei der Berechnung der Messwerte $\hat{\Phi}_n(\mu)$ ist nun der Matrixrang $L-2$ zu verwenden, und somit erhält man mit Gleichung (6.17) in Punkt 35 die Messwerte:

$$\text{Phi}(\mu) = \frac{C_{YQ}(\mu) - |C_{YV}(\mu)|^2 / C_{VQ}(\mu)}{M \cdot L \cdot (L-2)}.$$

In Punkt 37 ergibt sich nach Gleichung (6.19) für die Varianz der Messwerte der Übertragungsfunktion:

$$C_{HQ}(\mu) = M \cdot L / C_{VQ}(\mu) \cdot \text{Phi}(\mu),$$

und mit Gleichung (6.21) wird in Punkt 41 die Varianz der Messwerte des Spektrums der deterministischen Störung zu:

$$C_{UQ}(\mu) = M \cdot \text{Akku}_V(\mu) / C_{VQ}(\mu) \cdot \text{Phi}(\mu).$$

Die Varianz der Messwerte des LDS wird mit Gleichung (6.25) in Punkt 45 zu

$$\begin{aligned} C_{\text{Phi}Q}(\mu) &= \frac{2}{L} \cdot \text{Phi}(\mu)^2 \quad \text{für } \mu \in \{0; \frac{M}{2}\} \\ \text{und } C_{\text{Phi}Q}(\mu) &= \frac{1}{L-1} \cdot \text{Phi}(\mu)^2 \quad \text{sonst} \end{aligned}$$

abgeschätzt. Da die Werte der Übertragungsfunktion für die beiden Frequenzen $\mu = 0$ und $\mu = \frac{M}{2}$ immer reell sind, ergeben sich hier keine Konfidenzellipsen sondern Konfidenzintervalle deren Breite man analog zu Gleichung ([1]:3.72) in Punkt 49 als

$$A_H(\mu) = \sqrt{2} \cdot \sqrt{C_{HQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

angibt. Für dieselben Frequenzen sind auch die Spektralwerte der deterministischen Störung immer reell. Daher werden in Punkt 51 die Breiten der Konfidenzintervalle mit

$$A_U(\mu) = \sqrt{2} \cdot \sqrt{C_{UQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

abgeschätzt. Auch für die reelle deterministischen Störung werden in Punkt 52 die zeitunabhängigen Konfidenzintervalle

$$A_u(\mu) = \sqrt{2} \cdot \sqrt{C_{uQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha)$$

angegeben. Die Konfidenzintervalle der LDS-Messwerte in Punkt 53 bleiben unverändert. Alle Punkte der Messablaufsliste, die zur Berechnung solcher Werte dienen, die durch eine einfache Symmetrie auf bereits berechnete Werte zurückführbar sind, kann man überspringen. Sollten diese dann nicht berechneten Werte in weiteren Kalkulationen benötigt werden, so sind sie durch die entsprechenden symmetrischen Werte zu ersetzen.

Bei komplexwertigen Systemen ist bei jeder Einzelmessung jeweils eine DFT für das gefensterte Systemein- und -ausgangssignal durchzuführen. Da diese Signale nun reell sind, kann man die beiden DFTs zweier aufeinanderfolgender Einzelmessungen jeweils zu einer DFT zusammenfassen. Für das Systemausgangssignal sei dies kurz erläutert. Aus den beiden reellen gefensterten Signalen zweier aufeinanderfolgender Einzelmessungen wird das komplexe Signal

$$\tilde{y}_{f,\lambda}(k) = y_{f,\lambda-1}(k) + j \cdot y_{f,\lambda}(k) \quad \forall \quad \lambda = 2(2)L \quad (8.10)$$

gebildet und einer DFT unterworfen, so dass man $\tilde{Y}_{f,\lambda}(\mu)$ erhält. Mit Hilfe der Symmetrieeigenschaften der Spektren reeller Signale kann man die Spektren der einzelnen Anteile nach der DFT des komplexen Signals wieder trennen.

$$\begin{aligned} Y_{f,\lambda-1}(\mu) &= \frac{\tilde{Y}_{f,\lambda}(\mu) + \tilde{Y}_{f,\lambda}(-\mu)^*}{2} \\ Y_{f,\lambda}(\mu) &= \frac{\tilde{Y}_{f,\lambda}(\mu) - \tilde{Y}_{f,\lambda}(-\mu)^*}{2 \cdot j} \\ \forall \quad \lambda &= 2(2)L \end{aligned} \quad (8.11)$$

Will man sich die Division durch 2 ersparen, so verwendet man einfach die halbierte Fensterfolge. Die Division durch j ist keine echte komplexe Division. Sie macht lediglich aus dem Imaginärteil den Realteil und aus dem Realteil den negativen Imaginärteil. Insgesamt kann man sich so die Hälfte aller bei der DFT auftretenden komplexen Multiplikationen sparen.

Wenn man bei einem reellwertigen System die deterministische Störung nicht modelliert, werden zur Berechnung der Messwert die Akkumulatorfelder **Akku_y** und **Akku_V** nicht benötigt, die für die Bestimmung der empirischen Mittelwerte der Spektren der Signale am Systemein- und -ausgang vorgesehen waren. Da wir jetzt annehmen, dass das erste Moment des Approximationsfehlers null ist, gehen wir davon aus, dass ein erstes Moment im Ausgangssignal durch ein erstes Moment in der Erregung des linearen Modellsystems verursacht wird. Da das erste Moment der Erregung durch das lineare Modellsystem genauso verzerrt wird wie die zufälligen Anteile der Erregung, braucht man den Mittelwert des Ausgangssignals nicht separat behandeln. Zur Berechnung der Messwert und der Schätzwerte der Messwert(ko)varianzen werden bei dem reduzierten Systemmodell ohne

deterministische Störung nur die Akkumulatorfelder

Akku_VQ	mit	$M/2 + 1$	Speicherplätzen,
Akku_YQ	mit	$M/2 + 1$	Speicherplätzen und
Akku_YV	mit	M	Speicherplätzen

benötigt. Das L -fache dieser Akkumulatoren tritt in allen Berechnungen nun an die Stelle der $L \cdot (L-1)$ -fachen Kovarianzen und Kovarianzmatrizen. Zusätzlich entfallen hier noch die Punkte 11, 14, 21, 25, 26, 28, 31 – 34, 41, 43, 51 und 52 ersatzlos. Bei der Berechnung der Vorfaktoren der Messwerte $\hat{\Phi}_n(\mu)$ in Punkt 35 ist gegenüber dem Fall, dass bei einem reellen System die deterministische Störung modelliert wird, jeweils ein um 1 erhöhter Matrixrang zu berücksichtigen, indem man dort im Nenner den Faktor $(L-2)$ durch $(L-1)$ ersetzt. So erhält man die in Gleichung ([1]:4.8) angegebene Messwerte

$$\text{Phi}(\mu) = \frac{\text{Akku_YQ}(\mu) - |\text{Akku_YV}(\mu)|^2 / \text{Akku_VQ}(\mu)}{M \cdot (L-1)}.$$

Auch bei der Berechnung der Varianzen dieser Messwerte in Punkt 45 ist der erhöhte Matrixrang zu berücksichtigen, indem man dort jeweils L durch $L+1$ ersetzt. Die führt zu den in den Gleichungen ([1]:4.10) genannten Werten

$$\begin{aligned} \text{C_PhiQ}(\mu) &= \frac{2}{L+1} \cdot \text{Phi}(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \text{und } \text{C_PhiQ}(\mu) &= \frac{1}{L} \cdot \text{Phi}(\mu)^2 & \text{sonst.} \end{aligned}$$

Da die Werte der Übertragungsfunktion für die beiden Frequenzen $\mu = 0$ und $\mu = \frac{M}{2}$ immer reell sind, ergeben sich hier keine Konfidenzellipsen sondern Konfidenzintervalle deren Breite man analog zu Gleichung ([1]:3.72) in Punkt 49 als

$$\text{A_H}(\mu) = \sqrt{2} \cdot \sqrt{\text{C_HQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

angibt. Die Konfidenzintervalle der LDS-Messwerte in Punkt 53 bleiben unverändert. Auch hier kann man jeweils die DFT der Signale zweier Einzelmessungen gemäß Gleichung (8.10) zu einer DFT zusammenfassen, und mit Gleichung (8.11) die Spektren der Signale der Einzelmessungen wieder trennen.

Zur Berechnung der Messwerte werden bei der Spektralschätzung reeller Prozesse mit getrennter Messung des ersten und des zweiten zentralen Moments nur die zwei Akkumulatorfelder

Akku_y	mit	F	Speicherplätzen und
Akku_YQ	mit	$M/2 + 1$	Speicherplätzen

benötigt, die man wie in den Punkten 14 und 15 angegeben akkumuliert. Weiterhin unverändert sind noch die Punkte 25 und 26 zu berechnenden. Punkt 32 liefert mit Gleichung

(5.5) das Spektrum der deterministischen Störung als

$$U(\mu) = Y_{\text{mittel}}(\mu)/L.$$

Die Varianzen des Spektrums der deterministischen Störung werden in Punkt 41 nun mit Gleichung (5.8a) wesentlich einfacher als

$$C_{UQ}(\mu) = \frac{M}{L} \cdot \text{Phi}(\mu)$$

berechnet. Die zeitunabhängige Varianz der deterministischen Störung selbst wird unverändert in Punkt 43 bestimmt. Die Berechnung der Messwerte $\hat{\Phi}_n(\mu)$ in Punkt 35 erfolgt hier gemäß der Gleichung (5.6a):

$$\text{Phi}(\mu) = \frac{C_{YQ}(\mu)}{M \cdot L \cdot (L-1)}.$$

Bei der Berechnung der Varianzen dieser Messwerte in Punkt 45 führt der geänderte Matrixrang zu den in Gleichung (5.18) angegebenen Werten

$$\begin{aligned} C_{\text{Phi}Q}(\mu) &= \frac{2}{L+1} \cdot \text{Phi}(\mu)^2 & \text{für } \mu \in \{0; \frac{M}{2}\} \\ \text{und } C_{\text{Phi}Q}(\mu) &= \frac{1}{L} \cdot \text{Phi}(\mu)^2 & \text{sonst.} \end{aligned}$$

Für die beiden Frequenzen $\mu = 0$ und $\mu = \frac{M}{2}$ sind die Spektralwerte der deterministischen Störung immer reell. Daher werden in Punkt 51 die Breiten der Konfidenzintervalle mit

$$A_U(\mu) = \sqrt{2} \cdot \sqrt{C_{UQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha) \quad \text{für } \mu \in \{0; \frac{M}{2}\}$$

abgeschätzt. Auch für die reelle deterministischen Störung werden in Punkt 52 die zeitunabhängigen Konfidenzintervalle

$$A_u(\mu) = \sqrt{2} \cdot \sqrt{C_{uQ}(\mu)} \cdot \text{erfc}^{-1}(\alpha)$$

angegeben. Die Konfidenzintervalle der LDS-Messwerte in Punkt 53 bleiben unverändert. Auch hier kann man jeweils die DFT der Signale zweier Einzelmessungen gemäß Gleichung (8.10) zu einer DFT zusammenfassen, und mit Gleichung (8.11) die Spektren der Signale der Einzelmessungen wieder trennen.

Verzichtet man bei der Spektralschätzung reeller Prozesse auf die Messung des ersten Momentes des Prozesses $\mathbf{y}(k)$, weil man von diesem weiß, dass er mittelwertfrei ist, so braucht man nur das eine Akkumulatorfeld **Akku_YQ**, da das Akkumulatorfeld **Akku_y** für die Bestimmung des empirischen Mittelwertes vorgesehen ist, der nun mit null abgeschätzt wird. Die Berechnung der Messwerte $\hat{\Phi}_n(\mu)$ in Punkt 35 erfolgt hier gemäß der

Gleichung ([1]:5.1):

$$\text{Phi}(\mu) = \frac{\text{Akku_YQ}(\mu)}{M \cdot L}.$$

Bei der Berechnung der Varianzen dieser Messwerte in Punkt 45 führt der geänderte Matrixrang zu den in Gleichung ([1]:5.7) angegebenen Werten

$$\begin{aligned} \text{C_PhiQ}(\mu) &= \frac{2}{L+2} \cdot \text{Phi}(\mu)^2 && \text{für } \mu \in \left\{0; \frac{M}{2}\right\} \\ \text{und } \text{C_PhiQ}(\mu) &= \frac{1}{L+1} \cdot \text{Phi}(\mu)^2 && \text{sonst.} \end{aligned}$$

Die Konfidenzintervalle der LDS-Messwerte in Punkt 53 bleiben unverändert. Auch hier kann man jeweils die DFT der Signale zweier Einzelmessungen gemäß Gleichung (8.10) zu einer DFT zusammenfassen, und mit Gleichung (8.11) die Spektren der Signale der Einzelmessungen wieder trennen.

9 Beispiele für RKM Messergebnisse

9.1 Varianz und Kovarianz der Messwerte der deterministischen Störung

Bei der Berechnung der Messwert(ko)varianzen der deterministischen Störung trat in Gleichung (3.52) eine Doppelsumme der zwei-dimensionalen DFT über ein Produkt von Erwartungswerten auf. Von dieser Doppelsumme wurden bei der Berechnung der Messwert(ko)varianzen nur die Summanden berücksichtigt, bei denen μ_2 um ein ganzzahliges Vielfaches von M/K_Φ gegenüber μ_1 verschoben ist. Es wurde festgestellt, dass die anderen Summanden aufgrund der Ausblendeigenschaft des Betragsquadrats des Spektrums des Fensters vernachlässigt werden können. Da das Betragsquadrat des Spektrums eines endlich langen Fensters keinen ideal rechteckförmigen Verlauf haben kann, ist zwischen dem Durchlassbereich und dem Sperrbereich des Fensterspektrums immer ein Übergangsbereich vorhanden, so dass man im Integral nach Gleichung (2.41) auch für $\hat{\mu} = \mu + \tilde{\mu} \cdot M/K_\Phi \pm 1$ noch nennenswert von null verschiedene Terme erhalten wird. Wenn man für das Spektrum der Erregung einen nichtzentralen ($E\{\mathbf{V}(\mu)\} \neq 0$) Zufallsvektor $\vec{\mathbf{V}}$ verwendet, dessen Kovarianzen $C_{\mathbf{V}(\mu_2), \mathbf{V}(\mu_1)}$ bzw. $C_{\mathbf{V}(\mu_2)^*, \mathbf{V}(\mu_1)}$ für $\mu_2 = \mu_1 + \tilde{\mu} \cdot M/K_\Phi \pm 1$ von null verschieden sind, werden die in erster Näherung vernachlässigten Summanden bewirken, dass die Messwert(ko)varianzen von $\hat{\mathbf{u}}(k)$ nicht mehr mit K_Φ periodisch sind, sondern Anteile mit den benachbarten Kreisfrequenzen $\Omega = \tilde{\mu} \cdot \frac{2\pi}{K_\Phi} \pm \frac{2\pi}{M}$ aufweisen. Dies soll nun am Beispiel des zeitinvarianten Systems ($K_H=1$) mit der Übertragungsfunktion $H(\Omega)=1$ demonstriert werden.

Zur Erregung verwenden wir die mit $M=64$ periodisch fortgesetzten Impulssequenzen $\mathbf{v}(k)$, die für $k = 1$ (1) $M-1$ immer null sind, und nur für $k = 0$ einen normalverteilten komplexen Zufallswert mit dem Mittelwert Eins, der Varianz Eins und dem komplexen Korrelationskoeffizienten $E\{\mathbf{v}(0)^2\}/E\{|\mathbf{v}(0)|^2\} = 0,7+0,3 \cdot j$ enthalten. Bei dieser Art der Erregung wird das Spektrum $V_\lambda(\mu)$ bei jeder Einzelmessung von μ unabhängig. Daher sind auch die Terme, die in den Gleichungen (3.47) vor dem LDS bzw. KLDS der gefensterten Störung stehen, und somit in die Gleichungen (3.52) eingehen, von den beiden beteiligten Frequenzen unabhängig und von null verschieden. Eine deterministische Störung wird bei der Messung nicht eingespeist, so dass $u(k) = 0$ gilt. Dem periodisch fortgesetzten erregenden Zufallssignal wird eine stationäre mittelwertfreie gaußverteilte kom-

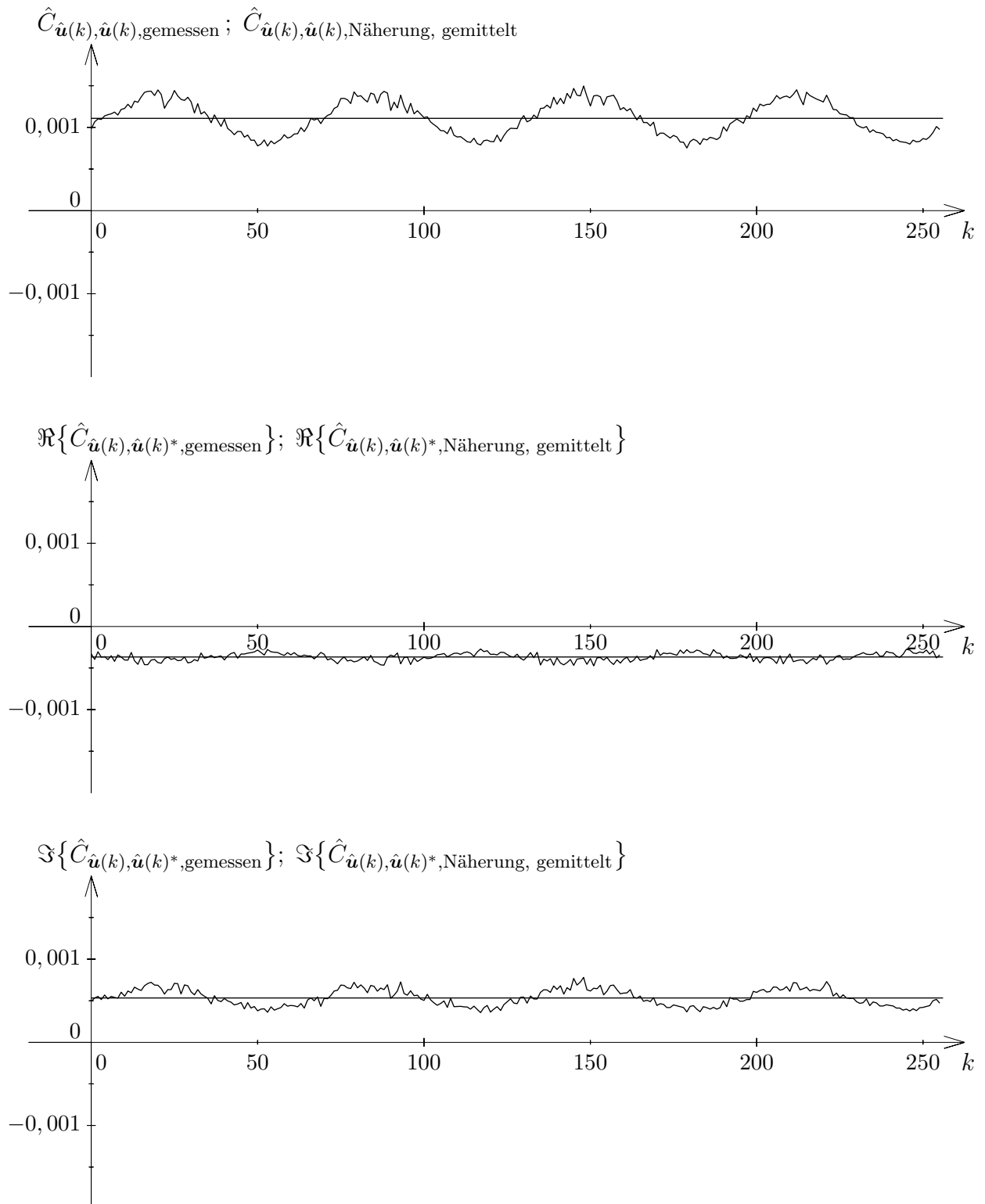


Bild 9.1: Messwert(ko)varianzen von $\hat{\mathbf{u}}(k)$.

Messung mit: $M = 64$, $E = 0$, $L = 20$ und Fenster nach Kapitel [1]:6 mit $N = 4$.

Mittelung über 1000 komplette RKM-Messungen.

plexe Störung ($K_\Phi = 1$) mit der Varianz 0,01 und dem komplexen Korrelationskoeffizienten $E\{\mathbf{n}(k)^2\}/E\{|\mathbf{n}(k)|^2\} = -0,5 + 0,5 \cdot j$ überlagert. Die Rauschwerte für unterschiedliche Werte von k sind unkorreliert. Das LDS dieser Störung ist konstant 0,01 und das KLDS ist konstant $-0,005 + 0,005 \cdot j$. Nun wurde eine Messung mit $L = 20$ Einzelmessungen mit einem Fenster nach Kapitel [1]:6 mit $N = 4$ durchgeführt, und diese 999 mal wiederholt. Die Messwert(ko)varianzen wurden nun aus den 1000 Schätzwerten $\hat{u}(k)$ für jeden Zeitpunkt $k = 0(1)F - 1$ empirisch bestimmt, und sind in Bild 9.1 über der diskreten Zeit k aufgetragen. Desweiteren sind die nach den Gleichungen (3.53) berechneten und über alle 1000 kompletten RKM-Messungen gemittelten zeitunabhängigen Näherungen der Messwert(ko)varianzen eingetragen. Deutlich erkennt man, dass neben den zeitlich konstanten Gleichanteilen, die durch die zeitunabhängigen Näherungen offensichtlich recht gut abgeschätzt werden, auch noch Anteile mit der Kreisfrequenz $\Omega = \pm 2\pi/64$ vorhanden sind, während die höherfrequenten Anteile offensichtlich rauschartiger Natur sind. Es sei noch darauf hingewiesen, dass man die systematischen Fehler, die man durch die zeitunabhängige Näherung der Messwert(ko)varianzen macht, bei Verwendung eines sinnvolleren erregenden Zufallsvektors $\vec{V}$ (mittelwertfreie und unkorrelierte Spektralwerte) vermeiden kann.

9.2 Messung eines komplexen, periodisch zeitvarianten Systems mit zyklstationärer Störung

Bei diesem Beispiel handelt es sich um die Simulation eines Stereodecoders, bei dem ein Stereo-Multiplex-Signal, wie es im UKW-Rundfunk üblich ist, mit 76 kHz abgetastet und A/D gewandelt wird, und anschließend mit einem digitalen FIR-Filter in der Art interpoliert wird, dass einerseits die beiden Signale des linken und rechten Kanals mit einer Abtastrate von jeweils 38 kHz getrennt vorliegen, und andererseits der Pilotton bei 19 kHz unterdrückt wird. Das Signal des linken Kanals können wir als den Realteil eines komplexen Signals auffassen, dessen Imaginärteil das Signal des rechten Kanals ist. Als Erregung $v_\lambda(k)$ wird bei jeder Einzelmessung ein Ausschnitt der Länge $M = 512$ eines mittelwertfreien, normalverteilten Zufallssignals verwendet. Die Varianz der komplexen Erregung wird zeitunabhängig auf $C_{\mathbf{v}(k), \mathbf{v}(k)} = 0,072$ festgelegt. Die Abtastwerte der Erregung wurden für unterschiedliche Zeitpunkte unkorreliert gewählt, während für den Real- und Imaginärteil für gleiche Zeitpunkte ein komplexer Korrelationskoeffizient von $C_{\mathbf{v}(k), \mathbf{v}(k)^*} / C_{\mathbf{v}(k), \mathbf{v}(k)} = j \cdot 0,7$ eingestellt wurde. Dies entspricht bei den reellen Signalen des rechten und des linken Stereokanals einem reellen Korrelationskoeffizienten von 0,7, wobei die Signale beider Kanäle dieselbe Varianz aufweisen.

Aus diesen M Abtastwerten wird nun das analoge Stereo-Multiplex-Signal für die Abtastzeitpunkte des AD-Wandlers im Stereodecoder berechnet. Dabei wird angenommen, dass die Werte $v_\lambda(k)$ die Abtastwerte im Abstand $1/38$ kHz einer Periode eines periodischen bandbegrenzten analogen komplexen Signals seien. In beiden Stereokanälen wird zunächst ein analoges Tiefpassfilter elften Grades mit Tschebyscheff-Verhalten im Durchlassbereich bei der Berechnung des Stereo-Multiplex-Signals simuliert. Wenn mit $v_{TP,L}(t)$ und $v_{TP,R}(t)$ die beiden reellen, analogen Signale an den Ausgängen der beiden Tiefpassfilter bezeichnet sind, erhält man das Stereo-Multiplex-Signal gemäß

$$\begin{aligned} v_{Stereo}(t) = & (v_{TP,L}(t) + v_{TP,R}(t)) + \\ & + (v_{TP,L}(t) - v_{TP,R}(t)) \cdot \sin(2\pi \cdot 38 \text{ kHz}) + \\ & + 0,1 \cdot \sin(2\pi \cdot 19 \text{ kHz}). \end{aligned} \quad (9.1)$$

Die Pilottonamplitude 0,1 und die Streuung der Erregung wurden dabei in einem Verhältnis gewählt, wie es im UKW-Rundfunk üblich ist. Aufgrund der Periodizität der Erregung lässt sich das Stereo-Multiplex-Signal für jeden beliebigen Zeitpunkt im Rahmen der Rechengenauigkeit als stationäre Lösung exakt berechnen, obwohl die analogen Tiefpässe keine zeitlich begrenzte Impulsantwort besitzen. Der mit 76 kHz getaktete AD-Wandler im Stereodecoder wird dadurch simuliert, dass man sowohl für die Zeitpunkte $t = (k+1/4+\Delta k_{sync})/38 \text{ kHz}$ als auch für die Zeitpunkte $t = (k+3/4+\Delta k_{sync})/38 \text{ kHz}$ mit $k \in [-E; F-1]$ die Abtastwerte des Stereo-Multiplex-Signals berechnet. Für jeden Zeitpunkt k erhält man also zwei Abtastwerte. Dabei berücksichtigt der Summand Δk_{sync} eine evtl. vorhandene Fehlsynchronisation des AD-Wandlers im Stereodecoder. In der Simulation wurde der Wert $\Delta k_{sync} = 0,01$ eingestellt. Um einen Linearitätsfehler des AD-Wandlers zu simulieren wurde vor der gleichmäßigen und symmetrischen 16-Bit-Quantisierung, die auch im Beispiel des Kapitels [1]:7.7 — dort aber mit 12-Bit — verwendet wird, die nichtlineare Verzerrung

$$0,01 \cdot v_{Stereo}(t)^3 - 0,002 \cdot v_{Stereo}(t)^2 + v_{Stereo}(t) + 0,002 \quad (9.2)$$

des Stereo-Multiplex-Signals vorgenommen. Jeder zweite Abtastwert des Ausgangssignals des AD-Wandlers wird mit einem linearphasigen FIR-Filter interpoliert, und liefert uns den Realteil (= linker Stereokanal) des Ausgangssignals $y_\lambda(k)$ des simulierten Systems. Die dazwischenliegenden Abtastwerte des Ausgangssignals des AD-Wandlers werden ebenfalls mit einem linearphasigen FIR-Filter interpoliert und bilden den Imaginärteil (= rechter Stereokanal) von $y_\lambda(k)$. Dabei wird das FIR-Filter verwendet, dessen Koeffizienten gerade die zeitlich gespiegelten Koeffizienten des Filters des Realteilstpfades sind. Als Koeffizienten der beiden FIR-Filter wurden die Abtastwerte der kontinuierlichen Fens-terautokorrelationsfunktion $d_\infty(t)$ verwendet, die in Kapitel 10.7 eingeführt wird. Die

Fourierreihenkoeffizienten des Filters kann man mit dem Programm im Unterkapitel 11.4 mit den Parametern $N=34$, $A_0=3$ und $A_1=15$ ohne eine weitere frei wählbare Nullstelle s_0 berechnen. Die Abtastwerte der Fensterautokorrelationsfunktion wurden mit dem Programm in Unterkapitel 11.7 für die normierten Zeitpunkte $t = (k+1/4)/38$ im Intervall $[-1; 1]$ berechnet. Die Fouriertransformierte der beiden FIR-Filter hat bei der halben Abtastfrequenz 19 kHz eine Nullstelle, die für eine hohe Dämpfung des Pilottons sorgt. Desweiteren ist der Betrag der Fouriertransformierten der FIR-Filter für die Frequenzen 0 (1) 15 kHz — diese liegen innerhalb der UKW-Bandbreite — bei geeigneter Normierung eins.

Damit ist das simulierte System beschrieben, und es wird nun kurz auf die Parametereinstellung der RKM-Messung eingegangen. Da das Stereo-Multiplex-Signal im AD-Wandler des Stereodecoders mit der vierfachen Frequenz des Pilottons abgetastet wird, und jeweils zwei reelle Abtastwerte einen komplexen Wert des Ausgangssignals $y_\lambda(k)$ liefern, führt eine Erhöhung von k um zwei zu einer unveränderten Aussteuerung der leicht nichtlinearen Kennlinie des AD-Wandler. Daher kann angenommen werden, dass sich der Stereodecoder durch ein periodisch zeitvariantes System mit der Periode $K_H=2$ gut modellieren lässt. Vom Approximationsfehler kann erwartet werden, dass er mit der Periode $K_\Phi=2$ zyklstationär ist. Da zu vermuten ist, dass der simulierte Synchronisationsfehler bei der Abtastung des Stereo-Multiplex-Signals zu einem u. U. unsymmetrischen Übersprechen der beiden Stereo-Kanäle führt, wurde auch das von der konjugierten Erregung $v(k)^*$ gespeiste Modellsystem $\mathcal{S}_{*,lin}$ angesetzt. Bei der Berechnung der Messwerte für das LDS und KLDS wurde die Variante gemäß der Gleichungen (3.34) gewählt. Da bei der Simulation des Systems bereits die Abtastwerte des periodischen Stereo-Multiplex-Signals berechnet werden, braucht die Einschwingzeit der Tschebyscheff-Tiefpässe bei der Messung nicht berücksichtigt zu werden. Es genügt daher wenn man die Einschwingzeit E größer gleich der Länge 76 der FIR-Interpolationsfilter wählt. Bei der Messung wurde $E=2048$ eingestellt. Als Fensterfolge beim RKM wurde die in [1]:6 vorgestellte Fensterfolge mit $N=4$ verwendet. Es wurde über $L=50000$ Einzelmessungen gemittelt. Die Länge der DFT betrug $M=512$, so dass alle Spektren im Raster $2\pi/M$ gemessen wurden. Das Konfidenzniveau wurde auf $1-\alpha = 90\%$ festgelegt.

Bild 9.2 zeigt die gemessene deterministische Störung. Im unteren Teilbild ist einerseits der Betrag des gemessenen Spektrums — verrauschte Kurve — und andererseits die gemessene Messwertstreuung — glatte Kurve — halblogarithmisch dargestellt. Außer bei der Frequenz $\Omega=0$ liegen alle Messwerte in der Größenordnung der Messwertstreuung, und sind daher als null anzusehen, oder zumindest wohl kleiner als -90 dB. Da sowohl die deterministische Störung als auch der zufällige Approximationsfehler die beiden FIR-Interpolationsfilter durchlaufen, ist eine Absenkung des Spektrums und der Messwert-

streuung im Bereich um die Frequenz $\Omega = \pi \hat{=} 19 \text{ kHz}$ zu beobachten. Bei $\Omega = 0$ ergibt sich der Wert 2,9121 dB. Somit ist die deterministische Störung praktisch zeitlich konstant. Dies erkennt man auch in den beiden oberen Teilbildern. Dort sind einmal für die geraden Zeitpunkte k und zum zweiten für ungerades k die Werte $\hat{u}(k)$ als Punkte in der komplexen Ebene dargestellt. An der Skalierung der Achsen erkennt man, wie nahe beieinander diese Punkte für die unterschiedlichen Werte von k liegen. Da man bei einem zyklstationären Approximationsfehler mit $K_\Phi = 2$ Konfidenzellipsen für die Messwerte $\hat{u}(k)$ erhält, die sich periodisch wiederholen, wurde für die geraden und die ungeraden Zeitpunkte k je eine Konfidenzellipse bei einem Messwert eingetragen. Die beiden Konfidenzellipsen unterscheiden sich doch erheblich hinsichtlich ihrer Neigung, was ein erster Hinweis auf die Zyklstationarität der Momente des Approximationsfehlers ist.

Die beiden oberen Teilbilder des Bildes 9.3 zeigen die beiden Impulsantworten $\hat{h}_\kappa(k)$ des periodisch zeitvarianten Teilmodellsystems $\mathcal{S}_{lin}$, das von der *nicht* konjugierten Erregung $\mathbf{v}(k)$ gespeist wird. Hier ist die Antwort $\hat{h}_0(k)$ des Systems auf den Impuls $\gamma_0(k)$ mit kleinen Kreisen, und die Antwort $\hat{h}_1(k)$ des Systems auf den verschobenen Impuls $\gamma_0(k-1)$ mit kleinen Kreuzchen dargestellt. Das oberste Teilbild zeigt die ersten 100 Werte des Realteils, während darunter die entsprechenden Imaginärteilwerte zu sehen sind. Für alle Messwerte der Impulsantworten, die nach Kapitel 4.1 aus den Messwerten der bifrequenten Übertragungsfunktion berechnet wurden, ergab sich eine empirische Messwertstreuung zwischen $3,3 \cdot 10^{-7}$ und $3,4 \cdot 10^{-7}$. Die Imaginärteile der Messwerte der Impulsantworten liegen somit in der Größenordnung der Messwertstreuung. Daher kann das von $\mathbf{v}(k)$ erregte Teilsystem im Rahmen der Messgenauigkeit als reellwertig angesehen werden. Die in Bild 2.3 eingezeichneten Teilsysteme $h_{\beta,\kappa}(k)$ und $h_{\gamma,\kappa}(k)$ liefern also einen vernachlässigbar kleinen Beitrag zu der periodisch zeitvarianten Impulsantwort des von $\mathbf{v}(k)$ erregten Teilsystems, so dass bei diesem Teilsystem kein nennenswertes Übersprechen zwischen den beiden Stereokanälen zu beobachten ist. Bei den Realteilen der beiden Impulsantworten ist im obersten Teilbild lediglich eine Verschiebung um einen Takt zu beobachten, wie man dies auch bei einem zeitinvarianten System erwarten würde. Die Unterschiede der beiden Impulsantworten $\hat{h}_0(k)$ und $\hat{h}_1(k+1)$ sind so gering, dass sie bei dieser Art der Darstellung nicht zu erkennen sind.

Die Beträge der Abtastwerte der bifrequenten Übertragungsfunktion sind in den unteren beiden Teilbildern des Bildes 9.3 in halblogarithmischer Darstellung zu sehen. Die dicke Kurve im vorletzten Teilbild zeigt die Hauptdiagonale der bifrequenten Übertragungsfunktion mit dem gleichen Argument μ in beiden Variablen. Im Bereich niedriger normierter Kreisfrequenzen mit $|\Omega| \leq 15/19 \cdot \pi$ liegt diese Kurve so nahe bei 0 dB, dass sie in diesem Bereich konstant erscheint. Deshalb wurde auch eine um den Faktor 1000 gespreizte Darstellung als dünne Kurve in dieses Teilbild eingezeichnet. Man erkennt nun das typi-

sche Verhalten der Tschebyscheff-Tiefpässe, die bei der Erzeugung des Stereo-Multiplex-Signals verwendet wurden. Die Beträge der Übertragungsfunktionen der FIR-Filter, die zur Interpolation der Abtastwerte des Stereo-Multiplex-Signals eingesetzt werden, liegen innerhalb der Nutzbandbreite bis 15 kHz offensichtlich so nahe bei eins, dass nur die Übertragungsfunktionen der identischen Tschebyscheff-Tiefpässe im linken und rechten Stereokanal im Betrag der Hauptdiagonale der bifrequenten Übertragungsfunktion des komplexen Signals ihre Wirkung zeigen. Als letzte Kurve ist in diesem Teilbild auch noch die Varianz dieser Messwerte eingezeichnet. Sie liegt im gesamten Frequenzbereich unterhalb von $4 \cdot 10^{-11} \approx 104 \text{ dB}$ und nimmt ab, je näher man zur halben Abtastfrequenz $19 \text{ kHz} \hat{=} \Omega = \pi$ kommt. Hier wirkt sich aus, dass auch der Approximationsfehlerprozess die beiden ausgangsseitigen FIR-Interpolationsfilter durchläuft, die zugleich den Pilotton bei 19 kHz unterdrücken.

Das unterste Teilbild in Bild 9.3 zeigt den Betrag des Anteils der bifrequenten Übertragungsfunktion, bei dem das zweite Argument um π größer ist als das erste Argument, was bei der gewählten Abtastung der bifrequenten Übertragungsfunktion zu einer Argumentverschiebung von $M/2$ führt. Dieser Anteil spiegelt das periodisch zeitvariante Verhalten des Systems wider, und wäre bei einem zeitinvarianten System konstant null. In der gewählten halblogarithmischen Darstellung liegt dieser Anteil in einem weiten Frequenzbereich etwa bei -74 dB . Im Bereich niedriger Frequenzen fällt dieser Anteil bis in die Größenordnung der ebenfalls dargestellten Messwertstreuung ab. Man erkennt, dass in diesem Frequenzbereich die Messkurve des Betrags der bifrequenten Übertragungsfunktion stark verrauscht erscheint, während sie im restlichen Frequenzbereich doch relativ glatt verläuft.

Die beiden Impulsantworten $\hat{h}_{*,\kappa}(k)$ des periodisch zeitvarianten Teilmodellsystems $\mathcal{S}_{*,lin}$, das von der konjugierten Erregung $\mathbf{v}(k)^*$ gespeist wird, sowie die Beträge der beiden Anteile der bifrequenten Übertragungsfunktion dieses Teilmodellsystems sind in Bild 9.4 dargestellt. Die Streuung der Messwerte der Impulsantworten liegt auch hier bei allen Messwerten zwischen $3,3 \cdot 10^{-7}$ und $3,4 \cdot 10^{-7}$. Da nun der Imaginärteil der beiden Impulsantworten um Größenordnungen größer ist als die Messwertstreuung, bewirkt dieses Teilsystem ein Übersprechen zwischen den beiden Stereokanälen. Bei den beiden Graphiken der Beträge der beiden Anteile der bifrequenten Übertragungsfunktion wurde jeweils als untere Kurve wieder die Messwertstreuung eingezeichnet.

Mit Hilfe der Gleichungen 2.1 lassen sich aus den Messwerten $\hat{h}_{\kappa}(k)$ und $\hat{h}_{*,\kappa}(k)$ der der zeitvarianten Impulsantworten der beiden Teilsysteme auch Messwerte für die vier Impulsantworten $h_{\alpha,\kappa}(k)$ bis $h_{\delta,\kappa}(k)$ der reellwertigen Systeme in Bild 2.3 berechnen. Diese sind nun ebenfalls periodisch zeitvariant. Auf eine graphische Darstellung der Impulsantworten und der Übertragungsfunktionen dieser reellwertigen Systeme wird verzichtet.

Die beiden oberen Teilbilder der Bilder 9.5 und 9.6 zeigen jeweils den Real- und Imaginärteil der Messwertfolgen $\hat{\phi}_{\mathbf{n}}(k+\kappa, k)$ und $\hat{\psi}_{\mathbf{n}}(k+\kappa, k)$, mit deren Hilfe man die beiden Kovarianzfolgen $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)^*\}$ und $E\{\mathbf{n}(k+\kappa) \cdot \mathbf{n}(k)\}$ abschätzen kann. Die Messwertfolgen wurden aus den Messwerten $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ für die Näherungen des LDS bzw. KLDS mit Hilfe einer zweidimensionalen DFT berechnet, wie dies in Kapitel 4.2 beschrieben ist. Da bei einem zyklstationären Approximationsfehlerprozess die Kovarianzfolgen nicht nur von dem freien Parameter κ der Differenz der beiden betrachteten Zeitpunkte des Prozesses abhängen, sondern auch vom absoluten Zeitpunkt k , wobei sich die Kovarianzfolgen mit K_{Φ} periodisch in k wiederholen, sind in jedem Teilbild die $K_{\Phi}=2$ Folgen dargestellt, die sich mit $k=0$ und $k=1$ ergeben. Erstere sind jeweils mit kleinen Kreisen und letztere mit kleinen Kreuzchen markiert. Von den Kovarianzfolgen wurden nur die Werte mit $|\kappa| \leq 50$ dargestellt. Im obersten Teilbild in Bild 9.5 kann man an der Stelle $\kappa=0$ die beiden Werte der Varianz des Approximationsfehlerprozesses ablesen, die sich in k periodisch abwechseln. Diese liegen deutlich über dem Wert $7,76 \cdot 10^{-11}$, der sich bei einer Abtastung des Stereo-Multiplex-Signals mit einem AD-Wandler mit einer perfekt gleichmäßigen Kennlinie ergeben würde, wenn man die Varianz als $1/12$ des Quadrats der Quantisierungsstufenhöhe berechnet. Somit überwiegt also der Anteil des Approximationsfehlerprozesses, der durch die simulierte Nichtlinearität der AD-Wandlerkennlinie nach Gleichung (9.2) verursacht wird. In den unteren beiden Teilbildern der Bilder 9.5 und 9.6 besteht die obere Kurve jeweils aus den Messwerten für die Stärken der Impulslinien der Näherung des bifrequenten LDS bzw. KLDS, die in Kapitel 2.3 beschrieben sind. Die Kurven in den vorletzten Teilbildern schätzen jeweils die Stärken der Impulslinien mit $\Omega_2 = \Omega_1$ ab, während die Kurven in den untersten Teilbildern Messwerte für die Stärken der Impulslinien mit $\Omega_2 = \Omega_1 + \pi$ sind. Damit die Zuverlässigkeit dieser Messkurven beurteilt werden kann, enthalten diese Teilbilder zusätzlich die Kurven der Messwertvarianzschätzwerte, die immer unterhalb der Kurven der Messwerte liegen. Durch die Mittelung über $L=50000$ Einzelmessung konnte erreicht werden, dass die Beträge der Messwerte des LDS und KLDS immer deutlich größer als die Messwertstreuungen waren, und man somit relativ zuverlässige Messwerte erhalten konnte.

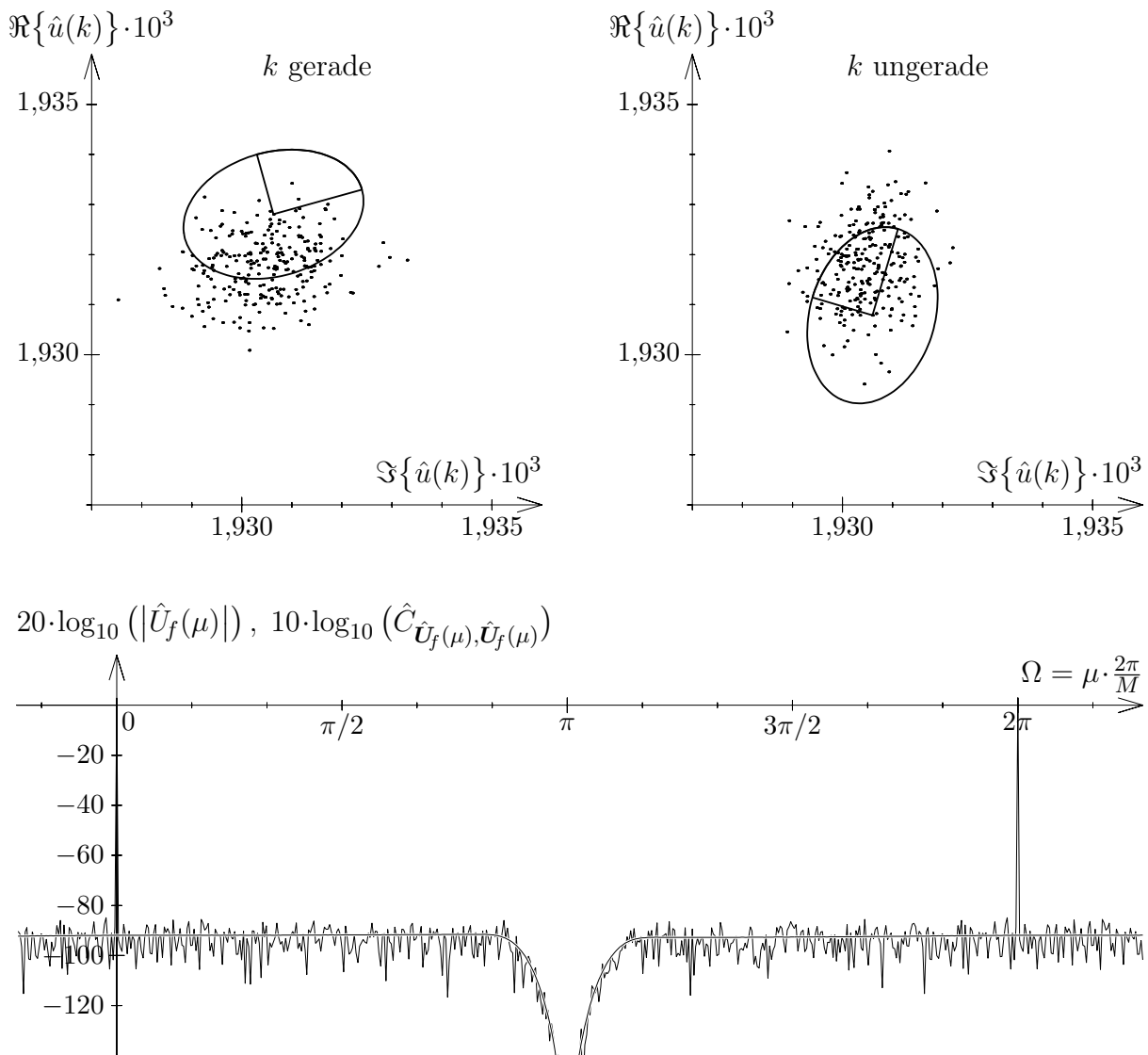
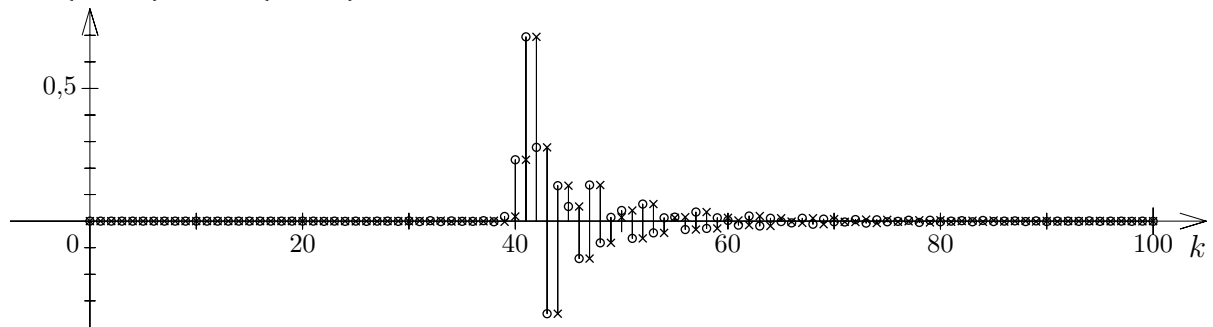


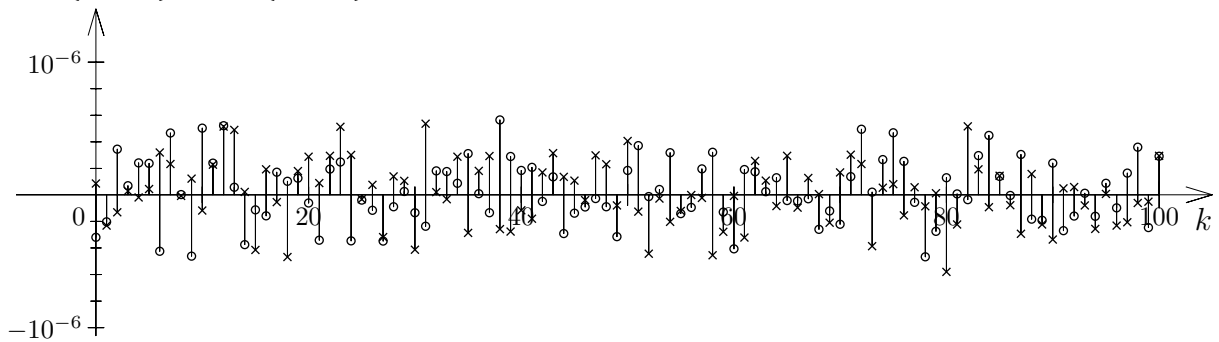
Bild 9.2: Fehlerbehaftete Abtastung eines Stereo-Multiplex-Signals.

Messwerte der deterministischen Störung $u(k)$

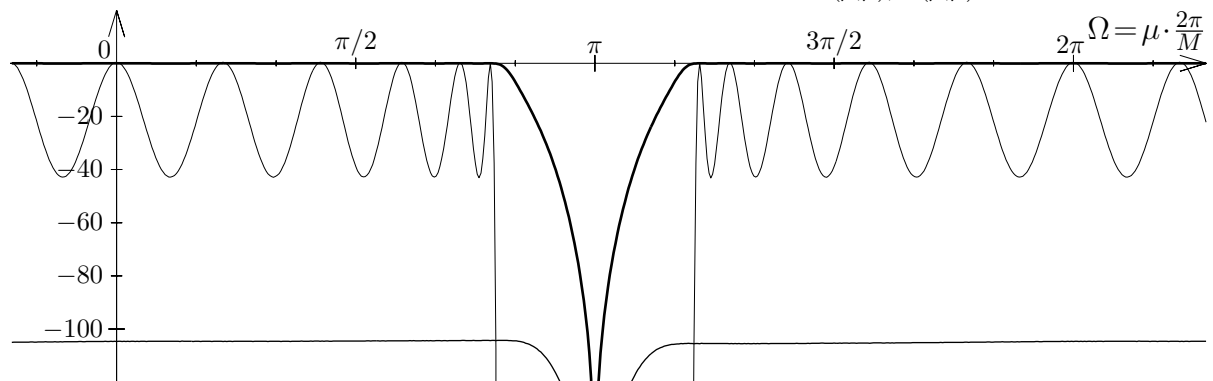
$$\circ \Re\{\hat{h}_0(k)\}, \times \Re\{\hat{h}_1(k)\}$$



$$\circ \Im\{\hat{h}_0(k)\}, \times \Im\{\hat{h}_1(k)\}$$



$$20 \cdot \log_{10}(|\hat{H}(\mu, \mu)|), 20000 \cdot \log_{10}(|\hat{H}(\mu, \mu)|), 10 \cdot \log_{10}(\hat{C}_{\hat{H}(\mu, \mu), \hat{H}(\mu, \mu)})$$



$$20 \cdot \log_{10}(|\hat{H}(\mu, \mu + \frac{M}{2})|), 10 \cdot \log_{10}(\hat{C}_{\hat{H}(\mu, \mu + \frac{M}{2}), \hat{H}(\mu, \mu + \frac{M}{2})})$$

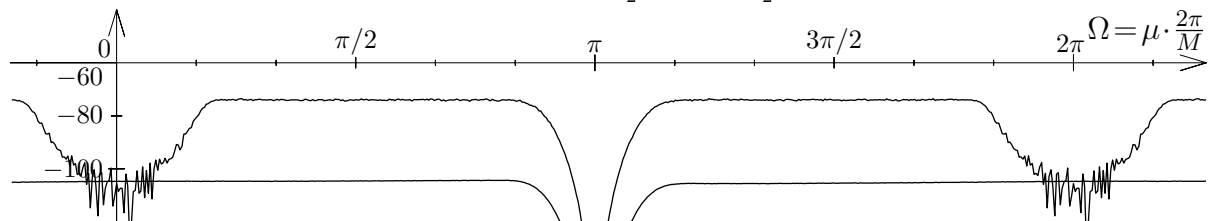


Bild 9.3: Fehlerbehaftete Abtastung eines Stereo-Multiplex-Signals.

Impulsantwort $\hat{h}_\kappa(k)$ und Übertragungsfunktion $\hat{H}(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$

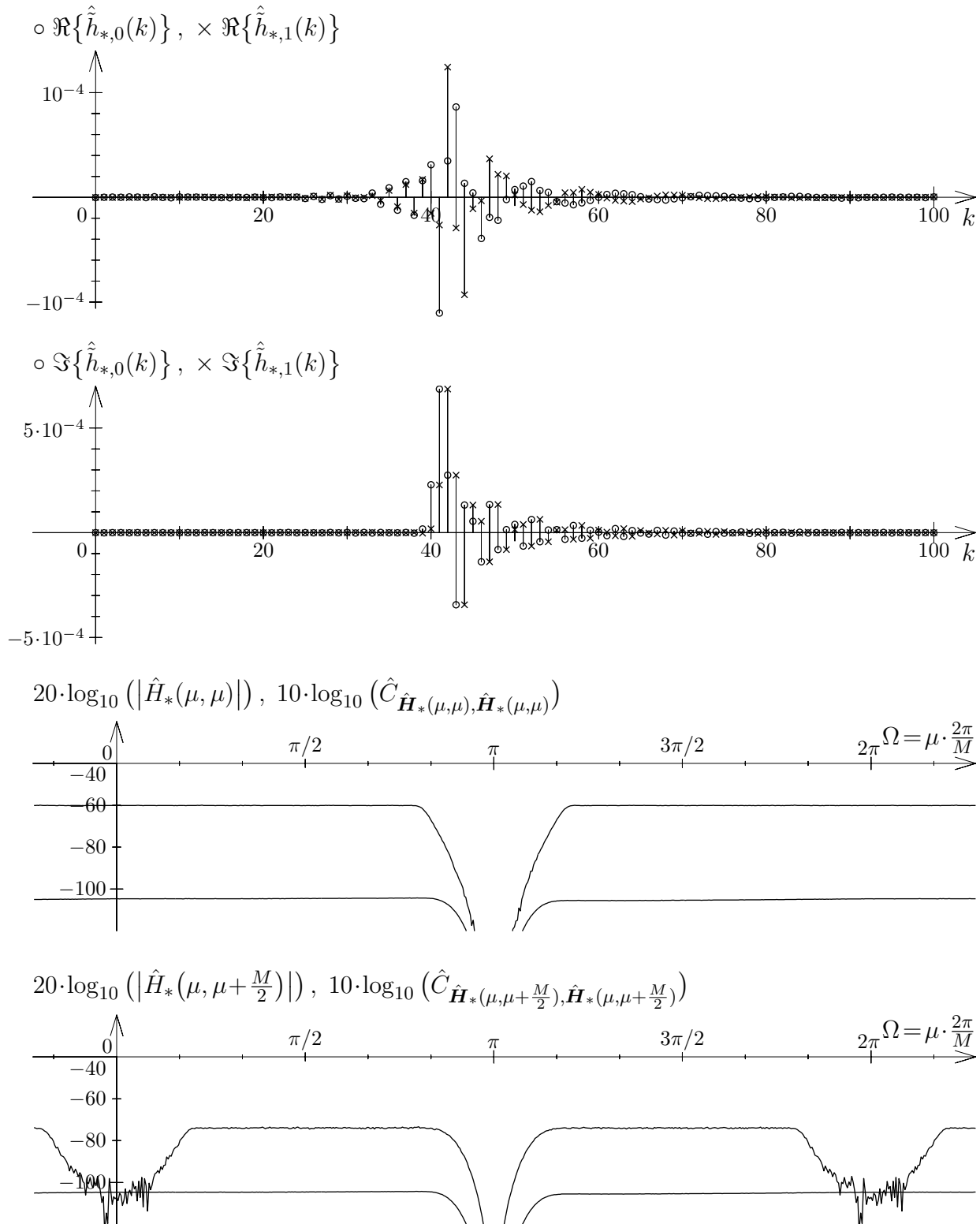


Bild 9.4: Fehlerbehaftete Abtastung eines Stereo-Multiplex-Signals.

 Impulsantwort $\hat{h}_{*,\kappa}(k)$ und Übertragungsfunktion $\hat{H}_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$

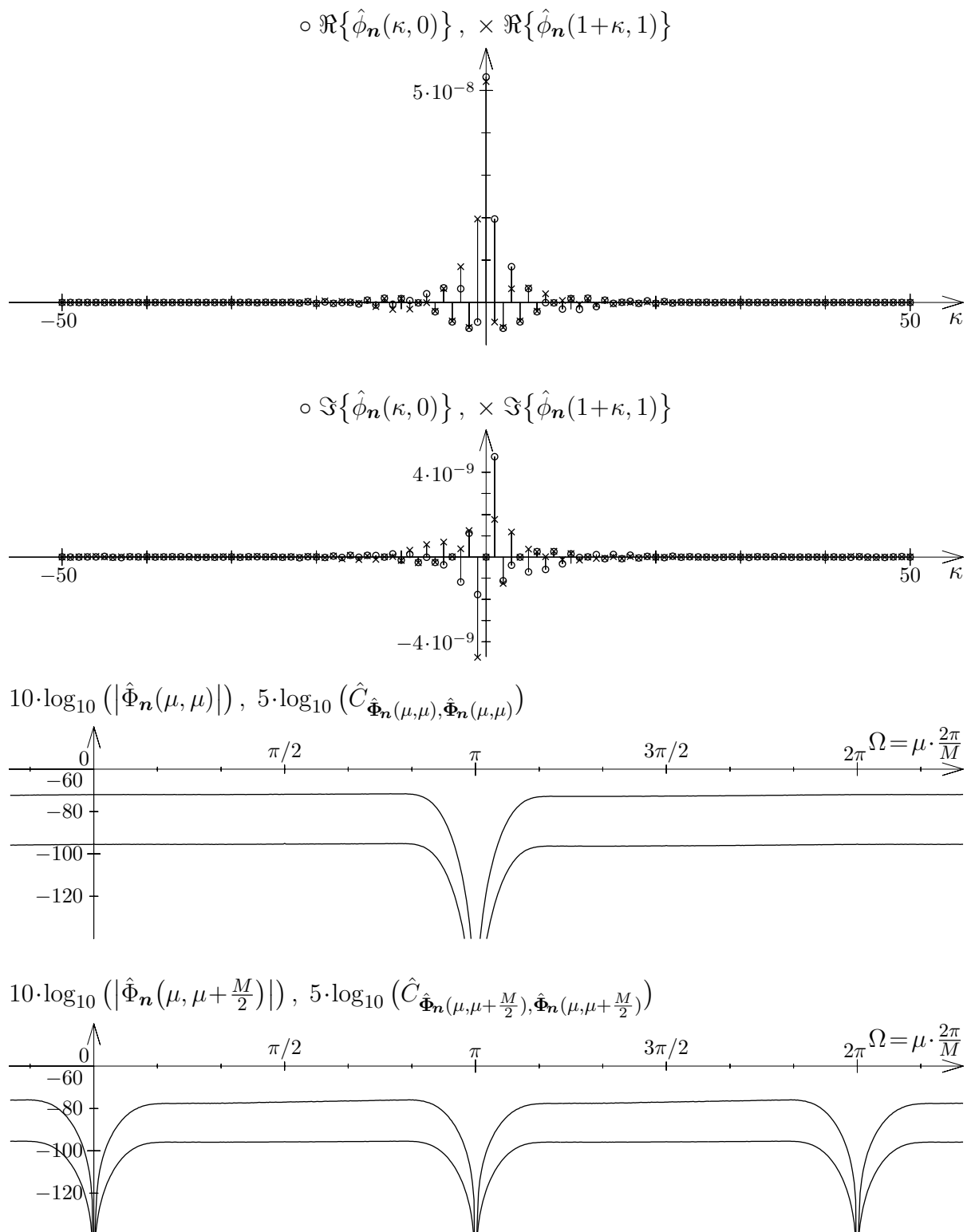


Bild 9.5: Fehlerbehaftete Abtastung eines Stereo-Multiplex-Signals.

Autokovarianzfolge $\hat{\phi}_{\mathbf{n}}(k+\kappa, k)$ und LDS-Näherung $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$

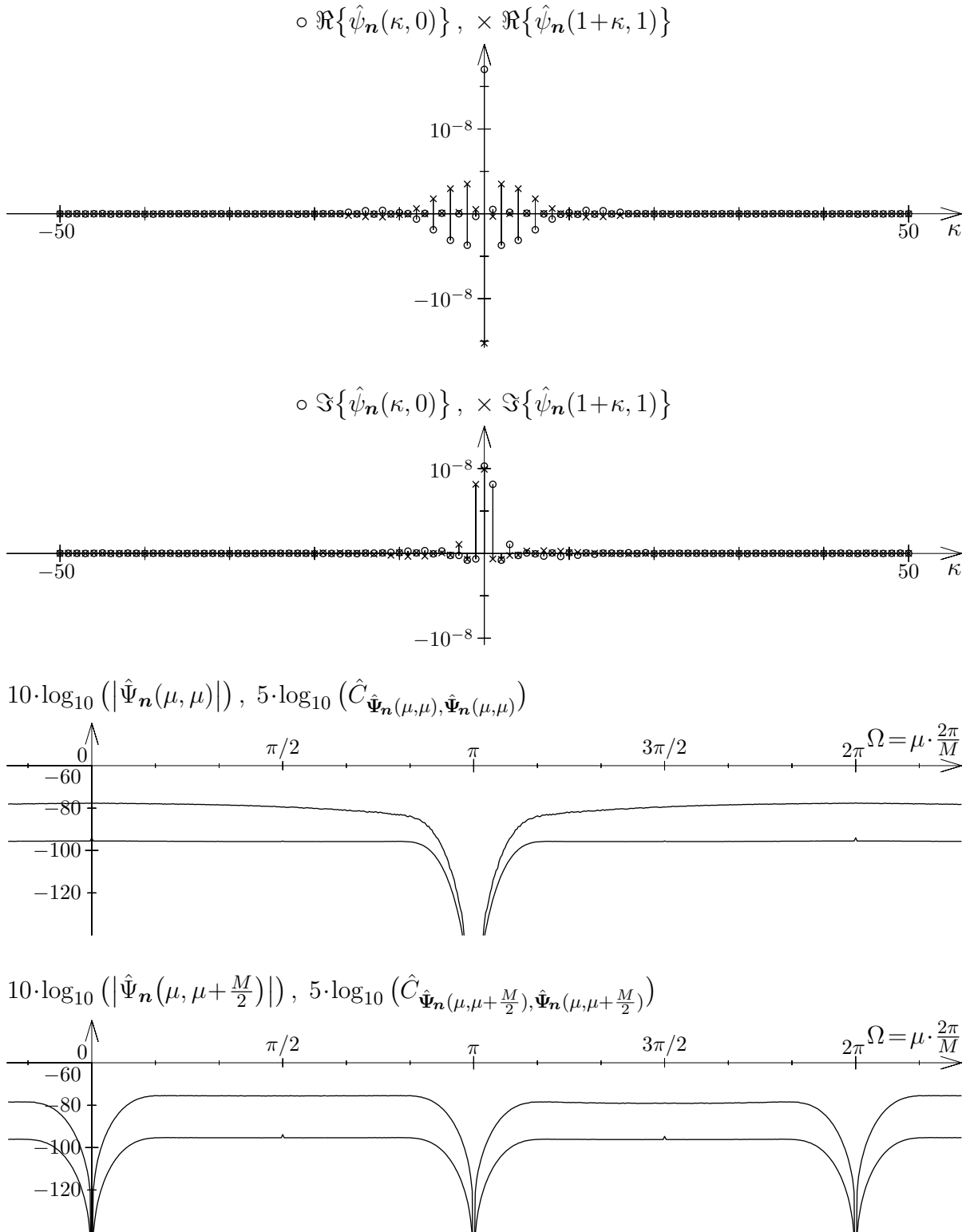


Bild 9.6: Fehlerbehaftete Abtastung eines Stereo-Multiplex-Signals.

Kovarianzfolge $\hat{\psi}_{\mathbf{n}}(k+\kappa, k)$ und KLDS-Näherung $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$

10 Ergänzungen zum Fenster

10.1 Algorithmus für Fenster mit frei wählbaren Nullstellen

In diesem Kapitel wird der in Kapitel [1]:6 vorgestellte Algorithmus zur Konstruktion der Fensterfolge so erweitert, dass es möglich wird, die dort nicht genutzten Freiheitsgrade dazu zu verwenden, weitere Fensterfolgen zu konstruieren, die ebenfalls die Eigenschaften ([1]:2.20) und ([1]:2.27) erfüllen, die den Einsatz der Fensterfolge beim RKM ermöglichen. Auch die hier berechneten Fensterfolgen sind *reell*, so dass an einigen Stellen auf das Konjugieren verzichtet werden kann, und die im Spektrum vorhandene Symmetrie ausgenutzt wird, ohne dass darauf extra hingewiesen wird. Desweiteren sind auch hier die Fensterfolgen in der Regel *nicht* linearphasig, und enthalten meist auch Werte kleiner null.

Die in Kapitel [1]:6.1 beschriebene prinzipielle Konstruktion der Fensterfolge wird im wesentlichen beibehalten und lediglich etwas erweitert. So kann sowohl die zeitdiskrete Fensterfolge wie auch die zugrundeliegende zeitdiskrete, reelle Basisfensterfolge $g(k)$ prinzipiell im gesamten Intervall $k \in [0; F-1]$ von null verschiedene Werte annehmen. Mit dem Parameter A_0 , für den $0 \leq A_0 < N$ gelten muss, kann man festlegen, dass die letzten A_0 Werte innerhalb des Intervalls null sein sollen:

$$\begin{aligned} g(k) = 0 & \quad \forall \quad k < 0 \quad \vee \quad k \geq F - A_0 \\ g(0) \neq 0 & \quad \wedge \quad g(F - A_0 - 1) \neq 0. \end{aligned} \quad (10.1)$$

Die Länge F soll auch hier ein ganzzahliges Vielfaches N der DFT-Länge M beim RKM sein. Bis auf einen konstanten Faktor — den man nicht fest vorgibt, da er nie explizit bestimmt werden muss — wird die Basisfensterfolge durch die nun $F - A_0 - 1$ Nullstellen des Polynoms der Z-Transformierten der Basisfensterfolge festgelegt:

$$z^{F-A_0-1} \cdot G(z) = z^{F-A_0-1} \cdot \sum_{k=0}^{F-A_0-1} g(k) \cdot z^{-k}. \quad (10.2)$$

Dieses Polynom lässt sich als Produkt zweier Polynome darstellen:

$$z^{F-A_0-1} \cdot G(z) = z^{F-N+2 \cdot A_1} \cdot G_1(z) \cdot z^{N-A_0-2 \cdot A_1-1} \cdot G_2(z). \quad (10.3)$$

Die $F-N+2 \cdot A_1$ Nullstellen des ersten Polynoms $z^{F-N+2 \cdot A_1} \cdot G_1(z)$ enthalten die in Kapitel [1]:6.1 genannten, fest vorgegeben Nullstellen am Einheitskreis, sowie weitere $2 \cdot A_1$ Nullstellen, die sich am Einheitskreis im gleichen Frequenzraster daran anschließen:

$$G(e^{j \cdot \frac{\pi}{F} \cdot \nu}) = 0 \quad \text{für} \quad \nu = N+1-2 \cdot A_1 \quad (2) \quad 2 \cdot F-N-1+2 \cdot A_1. \quad (10.4)$$

Die $N-A_0-2 \cdot A_1-1$ Nullstellen des zweiten Polynoms $z^{N-A_0-2 \cdot A_1-1} \cdot G_2(z)$ sind weitestgehend frei wählbar. Sie dürfen jedoch nicht bei $z=0$ oder bei $z=e^{\pm j \cdot \frac{\pi}{F} \cdot (N-1-2 \cdot A_1)}$ liegen, da sie in diesen Fällen gegebenenfalls durch eine Erhöhung von A_0 bzw. A_1 zu berücksichtigen wären. Wird gewünscht, dass irgendwelche der in der letzten Gleichung genannten Nullstellen des ersten Polynoms mit größerer als nur einfacher Vielfachheit auftreten, so sind die dazu benötigten zusätzlichen Nullstellen ebenfalls dem zweiten Polynom zuzuordnen. Desweiteren müssen Nullstellen, die einen von null verschiedenen Imaginärteil besitzen, auf Grund der Reellwertigkeit der Basisfensterfolge $g(k)$ als zueinander konjugiert komplexe Paare auftreten. Die Nullstellen des zweiten Polynoms werden im weiteren mit $z_{0,\rho} = |z_{0,\rho}| \cdot e^{j \cdot \psi_{0,\rho}}$ mit $\rho = 1 \quad (1) \quad N-1-A_0-2 \cdot A_1$ bezeichnet. Die in Kapitel [1]:6 vorgestellte Fensterfolge stellt somit den Spezialfall mit $A_0 = N-1$ und $A_1 = 0$ dar, bei dem keine zusätzlichen Nullstellen frei gewählt werden.

Für das erste Polynom kann man

$$z^{F-N+2 \cdot A_1} \cdot G_1(z) = \frac{z^F + (-1)^N}{\prod_{\nu_2 = \frac{1-N}{2} + A_1}^{\frac{N-1}{2} - A_1} (z - e^{j \cdot \frac{2\pi}{F} \cdot \nu_2})} \quad (10.5)$$

schreiben, wobei für das Produkt im Nenner wieder die in [1] in der Liste der Formelzeichen angegebene Definition verwendet wird, die es zulässt, dass der Laufindex ν_2 auch Werte annehmen kann, die nicht ganzzahlig sind. Das zweite Polynom kann in Form seiner Produktdarstellung als

$$z^{N-1-A_0-2 \cdot A_1} \cdot G_2(z) = \prod_{\rho=1}^{N-1-A_0-2 \cdot A_1} (z - z_{0,\rho}) \quad (10.6)$$

geschrieben werden.

Wie in [1] wird durch die Überlagerung der verschobenen Betragsquadrate des Spektrums $G(e^{j\Omega})$ das Betragsquadrat des Spektrums $F(\Omega)$ gewonnen. Somit überlagern

sich bei den Frequenzen $\Omega = \nu \cdot 2\pi/F$ mit $\nu = N - A_1$ (1) $F - N + A_1$ jeweils die doppelten Nullstellen, die im Polynom $z^{F-N+2\cdot A_1} \cdot G_1(z) \cdot G_1(z^{-1})$ und damit auch im Polynom $z^{F-A_0-1} \cdot G(z) \cdot G(z^{-1})$ vorhanden sind. $z^{F-A_0-1} \cdot |F(-j \cdot \ln(z))|^2$ hat daher bei diesen Frequenzen doppelte Nullstellen. Diese doppelten Nullstellen am Einheitskreis werden bei der Aufspaltung in einen minimalphasigen Anteil $z^{F-A_0-1} \cdot F(-j \cdot \ln(z))$ und den Rest $F(j \cdot \ln(z))$ als einfache Nullstellen jedem der beiden Polynome zugeordnet. Daher gilt die gegenüber Gleichung ([1]:6.15) modifizierte Gleichung:

$$F\left(\nu \cdot \frac{2\pi}{F}\right) = 0 \quad \forall \quad \nu = N - A_1 \text{ (1) } F - N + A_1. \quad (10.7)$$

Die in Gleichung ([1]:6.15) aufgeführten Nullstellen sind hier ebenfalls enthalten.

Auch hier ist die Fensterfolge durch die Werte des Spektrums $F(\Omega)$ bei den Frequenzen $\Omega = \nu \cdot 2\pi/F$ mit $\nu = 0$ (1) $F - 1$ vollständig festgelegt. Nach Gleichung (10.7) sind die meisten Werte des Spektrums bei diesen Frequenzen null. Daher und wegen der Reellwertigkeit von $f(k)$ genügt es das Spektrum der Fensterfolge bei den $N - A_1$ Frequenzen $\Omega = \nu \cdot 2\pi/F$ mit $\nu = 0$ (1) $N - 1 - A_1$ zu berechnen. Dies sind A_1 Werte weniger als bei der Fensterfolge nach Kapitel [1]:6. In Gleichung ([1]:6.16) ist daher in den Grenzen des Summationsindex $N - 1$ durch $N - 1 - A_1$ zu ersetzen.

Die Länge der zu berechnenden Fensterfolge $f(k)$ kann auch hier angegeben werden. Die Länge der Basisfensterfolge ist $F - A_0$. Das Polynom $z^{F-A_0-1} \cdot G(z) \cdot G(z^{-1})$ ist dann vom Grad $2 \cdot (F - A_0 - 1)$. Da die Koeffizienten $d(k)$ des Polynoms $z^{F-A_0-1} \cdot D(z)$ nach Gleichung ([1]:6.10) durch Multiplikation der Fensterautokorrelationsfolge $g_Q(k)$ mit der periodisch fortgesetzten si-Funktion entstanden sind, ist der Grad von $z^{F-A_0-1} \cdot D(z)$ ebenfalls gleich $2 \cdot (F - A_0 - 1)$. Durch Aufspaltung dieses Polynoms in den minimalphasigen Anteil und den Rest entsteht das Polynom $z^{F-A_0-1} \cdot F(-j \cdot \ln(z))$ von halbem Grad. Da die Koeffizienten dieses Polynoms die Werte der Fensterfolge $f(k)$ sind, ist deren Länge gleich der Länge der Basisfensterfolge also $F - A_0$.

Nach Gleichung ([1]:6.7) und Gleichung ([1]:6.13) erhält man aus $|G(z)|^2$ durch Verschiebung und Überlagerung das Betragsquadrat des Spektrums der Fensterfolge. Wenn man für $G_1(z)$ Gleichung (10.5) und für $G_2(z)$ Gleichung (10.6) eingesetzt, erhält man analog zu den Gleichungen ([1]:6.17) und ([1]:6.19) dafür:

$$\begin{aligned} F(-j \cdot \ln(z)) \cdot F(j \cdot \ln(z)) &\sim D(z) = \\ &= \sum_{\nu_1 = \frac{1-N}{2}}^{\frac{N-1}{2}} G_1(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}) \cdot G_1(z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1}) \cdot G_2(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}) \cdot G_2(z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1}) = \end{aligned} \quad (10.8)$$

$$\begin{aligned}
&= \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{z^F \cdot e^{-j \cdot 2\pi \cdot \nu_1} + (-1)^N}{\prod_{\nu_2=\frac{1-N}{2}+A_1}^{\frac{N-1}{2}-A_1} \left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1} - e^{j \cdot \frac{2\pi}{F} \cdot \nu_2} \right)} \cdot \prod_{\rho=1}^{N-1-A_0-2 \cdot A_1} \left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho} \right) \cdot \quad (*) \\
&\quad \cdot \frac{z^{-F} \cdot e^{j \cdot 2\pi \cdot \nu_1} + (-1)^N}{\prod_{\nu_2=\frac{1-N}{2}+A_1}^{\frac{N-1}{2}-A_1} \left(z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} - e^{j \cdot \frac{2\pi}{F} \cdot \nu_2} \right)} \cdot \prod_{\rho=1}^{N-1-A_0-2 \cdot A_1} \left(z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho} \right) = \\
&= \underbrace{\frac{-z^F + 2 - z^{-F}}{\prod_{\nu_3=1-N+A_1}^{N-1-A_1} \left(-z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_3} + 2 - z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_3} \right)}}_{= D_E(z)} \cdot \\
&\quad \cdot \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(-z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_2} + 2 - z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_2} \right) \cdot \\
&\quad \cdot \underbrace{\prod_{\rho=1}^{N-1-A_0-2 \cdot A_1} \left(\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho} \right) \cdot \left(z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho}^* \right) \right)}_{= D_{\bar{E}}(z)}
\end{aligned}$$

Setzt man wieder $z = e^{j \cdot \frac{2\pi}{F} \cdot \nu}$ mit $\nu = 0$ (1) $N-1-A_1$ in die mit (*) gekennzeichnete Form ein, so erhält man für das Betragsquadrat des Spektrums der Fensterfolge bis auf einen konstanten Faktor den gegenüber Gleichung ([1]:6.18) umfangreicheren Ausdruck

$$\begin{aligned}
|F(\nu \cdot \frac{2\pi}{F})|^2 &= F(\nu \cdot \frac{2\pi}{F}) \cdot F(-\nu \cdot \frac{2\pi}{F}) \sim D(e^{j \cdot \frac{2\pi}{F} \cdot \nu}) = \quad (10.9) \\
&F^2 \cdot 4^{1-N+2 \cdot A_1} \cdot \sum_{\nu_1=\max(\nu+A_1,0)-\frac{N-1}{2}}^{\min(\nu-A_1,0)+\frac{N-1}{2}} \prod_{\substack{\nu_2=-\frac{N-1}{2}+A_1 \\ \nu_2 \neq \nu-\nu_1}}^{\frac{N-1}{2}-A_1} \sin\left((\nu-\nu_2-\nu_1) \cdot \frac{\pi}{F}\right)^{-2} \cdot \\
&\quad \cdot \prod_{\rho=1}^{N-1-A_0-2 \cdot A_1} \left((1-|z_{0,\rho}|)^2 + 4 \cdot |z_{0,\rho}| \cdot \sin\left((\nu-\nu_1) \cdot \frac{\pi}{F} - \frac{\psi_{0,\rho}}{2}\right)^2 \right).
\end{aligned}$$

Zum einen sind hier die Abstandsquadrate zu den frei wählbaren Nullstellen $z_{0,\rho}$, deren Nullstellenwinkel mit $\psi_{0,\rho}$ bezeichnet seien, neu hinzugekommen. Zum anderen wurden die Grenzen der Summen- und Produktindizes modifiziert, weil durch die $2 \cdot A_1$ zusätzlichen

Nullstellen am Einheitskreis nun mehr Summanden null sind, und weniger Faktoren in den einzelnen Summanden vorhanden sind. Auch hier braucht der nun modifizierte Vorfaktor $F^2 \cdot 4^{1-N+2A_1}$ wegen der abschließenden Normierung auf $F(0)=M$ nicht explizit berechnet zu werden. Die positiven Wurzeln der mit der letzten Gleichung berechneten Werte sind bis auf einen konstanten Faktor die Beträge der Spektralwerte und somit im wesentlichen die Beträge Fourierreihenkoeffizienten der Fensterfolge.

Um in Gleichung (10.8) von der Form (*) auf die endgültige Form zu kommen, wurde wieder auf den gemeinsamen Hauptnenner aller Summanden erweitert. Dieser wurde zusammen mit dem von ν_1 unabhängigen Zähler vor die Summe gezogen. Der Laufindex ν_2 des in der Summe stehenden Produkts nimmt bei jedem Summanden jeweils wieder die $N-1$ Werte an, die das Produkt entstehen lassen, das zur Erweiterung auf den Hauptnenner benötigt wird. Beim Summanden mit dem Laufindex ν_1 nimmt ν_2 also die Werte $1-N+A_1$ (1) $N-1-A_1$ ohne die Werte $\nu_1+A_1-(N-1)/2$ (1) $\nu_1-A_1+(N-1)/2$ an.

Der Anteil $D_E(z)$ ist mit $z^{F-2\cdot N+2\cdot A_1+1}$ multipliziert ein Polynom aus doppelten Nullstellen am Einheitskreis bei $\Omega = \nu \cdot 2\pi/F$ mit $\nu = N-A_1$ (1) $F-N+A_1$. Jede dieser doppelten Nullstellen tritt als einfache Nullstelle im minimalphasigen Anteil von $D_E(z)$ auf. Die Phase $(F-2\cdot N+1+2\cdot A_1) \cdot \Omega/2$ dieses Anteils ist linear. Die Gruppenlaufzeit ist nun um A_1 größer als bei der in [1] vorgestellten Fensterfolge. Im Bereich der interessierenden Frequenzen $|\Omega| < (N-A_1) \cdot 2\pi/F$ besitzt $D_E(z)$ keine Nullstellen, so dass die Phase des minimalphasigen Anteils von $D_E(z)$ in diesem Intervall keine π -Sprünge hat.

Die Summe $D_{\bar{E}}(z)$ ist auch hier positiv und weist am Einheitskreis keine Nullstellen auf, weil dort alle Summanden nichtnegativ sind, und nie alle N Summanden gleichzeitig null werden. Dies ist dadurch sichergestellt, dass für $z = e^{\pm j \cdot (N-1-2A_1) \cdot \pi/F}$ das Polynom $z^{N-1-A_0-2A_1} \cdot G_2(z)$ von null verschieden ist (siehe Einschränkung für die Wahl von $z_{0,\rho}$ im Anschluss an Gleichung (10.4)), und nur maximal $N-1$ Nullstellen frei gewählt werden können, und somit bei N unterschiedlichen Rotationen nicht bei allen Summanden eine Nullstelle auf demselben Punkt des Einheitskreises liegen kann. Um bei der Berechnung der Phase des minimalphasigen Anteils von $D_{\bar{E}}(z)$ wieder ein Cepstrum brauchbarer Länge zu erhalten empfiehlt es sich auch hier vorher eine Bilineartransformation mit der Substitution von z nach Gleichung ([1]:6.20) durchzuführen.

$$D_{\bar{E}}(z) = \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(-z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_2} + 2 - z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_2} \right) \cdot \prod_{\rho=1}^{N-1-A_0-2A_1} \left((z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho}) \cdot (z^{-1} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} - z_{0,\rho}^*) \right) = \quad (10.10)$$

$$\begin{aligned}
&= \underbrace{\left(\frac{\tilde{z}}{(1+(1-c) \cdot \tilde{z}) \cdot (\tilde{z} + (1-c))} \right)^{2 \cdot (N-1-A_1) - A_0}}_{= \tilde{D}_P(\tilde{z})} \\
&\cdot \underbrace{\sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} (K_{\nu_2} \cdot (\tilde{z} - \tilde{z}_{\nu_2}) \cdot (\tilde{z}^{-1} - \tilde{z}_{\nu_2}^*)) \cdot \prod_{\rho=1}^{N-1-A_0-2A_1} ((\tilde{z} \cdot \tilde{z}_{N,\rho,\nu_1} - \tilde{z}_{Z,\rho,\nu_1}) \cdot (\tilde{z}^{-1} \cdot \tilde{z}_{N,\rho,\nu_1}^* - \tilde{z}_{Z,\rho,\nu_1}^*))}_{= \tilde{D}_N(\tilde{z})}.
\end{aligned}$$

Wieder sind die Nullstellen $\tilde{z}_{\nu_2}$ die Nullstellen, die durch die Bilineartransformation der Nullstellen am Einheitskreis bei $z = e^{j \cdot \frac{2\pi}{F} \cdot \nu_2}$ entstanden sind. Sie berechnen sich weiterhin nach Gleichung ([1]:6.24) mit dem Nullstellenwinkelanteil $\tilde{\psi}_{\nu_2}$ nach Gleichung ([1]:6.26). Auch die Konstante K_{ν_2} berechnet sich unverändert nach Gleichung ([1]:6.25). Neu hinzugekommen ist das Produkt der Faktoren mit den Nullstellen bei $\tilde{z}_{Z,\rho,\nu_1}/\tilde{z}_{N,\rho,\nu_1}$ und $\tilde{z}_{N,\rho,\nu_1}^*/\tilde{z}_{Z,\rho,\nu_1}^*$. Diese Nullstellen berechnen sich als die mit Gleichung ([1]:6.21) Bilineartransformierten der mit dem Drehfaktor $e^{j \cdot \frac{2\pi}{F} \cdot \nu_1}$ multiplizierten Nullstellen $z_{0,\rho}$.

$$\tilde{z}_{Z,\rho,\nu_1} = z_{0,\rho} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} - (1-c) \quad \text{und} \quad \tilde{z}_{N,\rho,\nu_1} = 1 - (1-c) \cdot z_{0,\rho} \cdot e^{j \cdot \frac{2\pi}{F} \cdot \nu_1} \quad (10.11)$$

Im weiteren sei $\tilde{\psi}_{0,\rho,\nu_1}$ der Winkel der Nullstelle $\tilde{z}_{Z,\rho,\nu_1}/\tilde{z}_{N,\rho,\nu_1}$, der sich als Differenz der Winkel des Zählers $\tilde{z}_{Z,\rho,\nu_1}$ und des Nenners $\tilde{z}_{N,\rho,\nu_1}$ der Nullstelle mit Hilfe der 4-Quadranten Arkustangensfunktion berechnen lässt.

In Gleichung (10.10) steht wieder ein Term vor der Summe, dessen Nullstellen alle bei $\tilde{z}=0$, und dessen Polstellen bei $\tilde{z}=c-1$ und $\tilde{z}=1/(c-1)$ liegen. Dieser Anteil liefert für das minimalphasige Polynom wieder einen Phasenbeitrag, der sich nach Gleichung ([1]:6.28) berechnet, und der hier mit der geänderten Vielfachheit $2 \cdot (N-1-A_1) - A_0$ statt $N-1$ in [1] multipliziert zu addieren ist.

Nun müssen wir wieder die Phase des minimalphasigen Anteils des Summenanteils $\tilde{D}_N(\tilde{z})$ in Gleichung (10.10) berechnen. Diese Summe ist für $\tilde{z} = e^{j\tilde{\Omega}}$ wieder positiv reell und lässt sich mit $\tilde{\Omega} = \eta \cdot 2\pi/\tilde{M}$ für alle ganzzahligen Werte η bis auf einen konstanten Faktor in der Form

$$\begin{aligned}
\tilde{D}_N(e^{j\tilde{\Omega}}) &\sim \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(K_{\nu_2} \cdot \sin\left(\frac{\tilde{\Omega}}{2} - \nu_2 \cdot \frac{\pi}{F} - \tilde{\psi}_{\nu_2}\right)^2 \right) \cdot \\
&\quad \cdot \prod_{\rho=1}^{N-1-A_0-2A_1} \left((|\tilde{z}_{N,\rho,\nu_1}| - |\tilde{z}_{Z,\rho,\nu_1}|)^2 + 4 \cdot |\tilde{z}_{N,\rho,\nu_1}| \cdot |\tilde{z}_{Z,\rho,\nu_1}| \cdot \sin\left(\frac{\tilde{\Omega} - \tilde{\psi}_{0,\rho,\nu_1}}{2}\right)^2 \right) =
\end{aligned} \quad (10.12)$$

$$\begin{aligned}
&= \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(K_{\nu_2} \cdot \sin \left(\frac{\pi}{M} \cdot \left(\eta - \widetilde{M} \cdot \left(\frac{\nu_2}{F} + \frac{\tilde{\psi}_{\nu_2}}{\pi} \right) \right) \right)^2 \right) \cdot \\
&\quad \cdot \prod_{\rho=1}^{N-1-A_0-2A_1} \left(\left(|\tilde{z}_{N,\rho,\nu_1}| - |\tilde{z}_{Z,\rho,\nu_1}| \right)^2 + 4 \cdot |\tilde{z}_{N,\rho,\nu_1}| \cdot |\tilde{z}_{Z,\rho,\nu_1}| \cdot \sin \left(\frac{\pi}{M} \cdot \left(\eta - \frac{\widetilde{M}}{2\pi} \cdot \tilde{\psi}_{0,\rho,\nu_1} \right) \right)^2 \right)
\end{aligned}$$

berechnen. Auch hier kann bei jedem Faktor wieder eine Reduktion der ganzen Zahl η um ein Vielfaches von $\widetilde{M}$ in der Art vorgenommen werden, dass der Betrag des Arguments der Sinusfunktion betragslich kleiner als $\pi/2$ bleibt. Es zeigte sich jedoch, dass man auf diese Art der 2π -Reduktion verzichten kann, da dies zu keiner nennenswerten Zunahme der Fehler in der Fensterfolge führt. Der konstanter Faktor, der bisher nicht festgelegt wurde, wird wieder zur Optimierung der Genauigkeit bei der Berechnung des Logarithmus genutzt.

Eine Besonderheit ergibt sich bezüglich der Genauigkeit der Berechnung des Zählers und des Nenners der bilineartransformierten Nullstelle $z_{0,\rho}$. Nach Gleichung (10.11) werden diese Größen als Differenz berechnet. Daher wird der relative Fehler dieser Größen sehr groß, wenn deren Wert selbst sehr klein wird. Dadurch kann auch der Nullstellenwinkel in diesen Fällen nicht mehr exakt berechnet werden. Da jedoch der Beitrag einer frei gewählten Nullstelle $z_{0,\rho}$ zur Summe, die das Betragsquadratspektrum bildet, immer dann klein wird, wenn $\tilde{z}_{Z,\rho,\nu_1}$ oder $\tilde{z}_{N,\rho,\nu_1}$ besonders schlecht berechnet werden kann, wirkt sich diese Besonderheit nicht wesentlich auf die Genauigkeit der Berechnung der Phase aus.

Nach Gleichung ([1]:6.32) liefert uns eine inverse FFT des natürlichen Logarithmus des Terms $\tilde{D}_N(e^{j\eta \cdot \frac{2\pi}{M}})$ das doppelte Cepstrum. Damit das Cepstrum möglichst rasch abklingt, muss der Parameter c geeignet gewählt werden. Da der optimale Parameter c von der Lage der nach der Überlagerung vorhandenen, unbekannten Nullstellen $\hat{z}_0$ abhängig ist, die wiederum von den frei gewählten Nullstellen $z_{0,\rho}$ abhängen, kann hier keine optimale Einstellung empirisch gewonnen werden.

Ich schlage daher vor, den Bilineartransformationsparameter c und die Länge $\widetilde{M}$ der inversen FFT nach den Gleichungen ([1]:6.30) und ([1]:6.31) einzustellen, die Berechnung des Cepstrums durchzuführen, zu kontrollieren ob die Länge der inversen FFT zu kurz war und gegebenenfalls die Länge der inversen FFT zu verdoppeln, den Parameter c zu verändern und die Berechnung der Summe $\tilde{D}_N(e^{j\tilde{\Omega}})$ nach Gleichung (10.12) zu wiederholen. Dies wird solange durchgeführt, bis ein Abbruchkriterium feststellt, dass die Länge der inversen FFT groß genug ist. Um ein Abbruchkriterium zu erhalten, wird der Betrag des berechneten Cepstrums mit einer abfallenden Grenze verglichen, die der Betrag des Cepstrums nicht überschreiten darf. Als Grenze wird die feinere der beiden oberen Abschätzungen nach Ungleichung ([1]:6.38) für den Betrag eines Cepstrums verwendet.

Dabei wird als schlimmster Fall angenommen, dass alle $2 \cdot (N-1-A_1) - A_0$ Nullstellen des minimalphasigen Polynoms, von dem das Cepstrum berechnet wird, den gleichen Betrag haben. Damit ergibt sich für die Grenze $(2 \cdot (N-1-A_1) - A_0)/k \cdot |\tilde{z}_0|^k$. Der Betrag der unbekannten Nullstelle in dieser Grenze wird nun so gewählt, dass die Grenze für $k = \tilde{M}/2$ auf einen Schätzwert $\bar{\sigma}$ für den Mindestwert der Streuung des Rauschsockels abgefallen ist, der bei der Berechnung des Cepstrums mit der inversen FFT unweigerlich entsteht. Für die Grenze ergibt sich daher:

$$|C(k)| \stackrel{!}{<} \frac{2 \cdot N - 2 - 2 \cdot A_1 - A_0}{k} \cdot \left(\frac{\tilde{M} \cdot \bar{\sigma}}{2 \cdot (2 \cdot N - 2 - 2 \cdot A_1 - A_0)} \right)^{\frac{2 \cdot k}{\tilde{M}}}. \quad (10.13)$$

Um einen Schätzwert für den Mindestwert der Streuung des Rauschsockels zu erhalten, wird angenommen, dass der Logarithmus von $\tilde{D}_N(e^{j\tilde{\Omega}})$ für jede Frequenz mit einer Varianz, die sich mit dem „ $Q^2/12$ “-Modell berechnet, verrauscht ist. Als Quantisierungsstufenhöhe wird $\max(|\ln(\tilde{D}_N(e^{j\tilde{\Omega}}))|, 1) \cdot \varepsilon$ eingesetzt (siehe Anhang [1]:A.7). Diese Varianz wird über alle $\tilde{M}$ Frequenzen gemittelt. Die mittlere Rauschleistung, die sich nach einer idealen FFT — ohne Quantisierung der Drehfaktoren und ohne Arithmetikfehler — ergibt, ist um den Faktor $1/\tilde{M}$ kleiner. Die halbe Wurzel der mittleren Rauschleistung nach der idealen FFT wird als Schätzwert für den Mindestwert der Streuung des Rauschsockels im Cepstrum verwendet.

$$\bar{\sigma} = \frac{\varepsilon}{\tilde{M} \cdot \sqrt{48}} \cdot \sqrt{\sum_{\eta=0}^{\tilde{M}-1} \max\left(\ln\left(\tilde{D}_N\left(e^{j \cdot \frac{2\pi}{\tilde{M}} \cdot \eta}\right)\right)^2, 1\right)} \quad (10.14)$$

Dabei berücksichtigt der Faktor $1/2$, dass das Cepstrum sich aus dem Betrag des minimalphasigen Anteils berechnet, und nicht aus den Betragsquadrat. Haben alle unbekannten Nullstellen des Spektrums, dessen Cepstrum berechnet wurde, einen Betrag, der kleiner ist als der für die Berechnung der Grenze verwendete Nullstellenbetrag, so liegt der Betrag des Cepstrums sicher unterhalb der Grenze. Ist jedoch auch nur eine der unbekannten Nullstellen betragsmäßig größer, so wird das Cepstrum die Grenze früher oder später überschreiten, da dann das berechnete Cepstrum einen im wesentlichen exponentiellen Anteil besitzt, der langsamer abklingt als die Grenze. Sofern dieses Überschreiten der Grenze nicht stattfindet, bzw. erst stattfindet, wenn sowohl das Cepstrum als auch die Grenze unterhalb eines Schätzwertes $\hat{\sigma}$ für den maximal möglichen Rauschwert liegen, wird angenommen, dass die Länge der inversen FFT bei der Berechnung des Cepstrums groß genug gewählt war. In diesem Fall kann mit der Berechnung der Phase durch die Auswertung der Sinusreihe wie in [1] fortgefahren werden.

Um einen Schätzwert für dem maximal möglichen Rauschwert im Cepstrum zu erhalten, nehmen wir an, dass der Logarithmus von $\tilde{D}_N(e^{j\tilde{\Omega}})$ bei jeder Frequenz, bei der er berechnet worden ist, einen absoluten Fehler der Größe $\max(|\ln(\tilde{D}_N(e^{j\tilde{\Omega}}))|, 1) \cdot \varepsilon$ aufweist.

Schlimmstenfalls können sich die absoluten Fehler aller $\widetilde{M}$ Frequenzen bei der Berechnung der FFT betragsmäßig addieren. Wenn eine Radix- $\log_2(\widetilde{M})$ FFT durchgeführt wird, entstehen in den einzelnen Stufen zusätzliche Fehler. Als übelsten Fall nehmen wir an, dass sich bei jeder Stufe der FFT ein zusätzlicher absoluter Fehler betraglich addiert, der maximal genauso groß ist, wie die eben genannte Summe der Beträge der absoluten Fehler der Logarithmuswerte aller Frequenzen. Da bei der inversen FFT der Faktor $1/\widetilde{M}$ auftritt, wird sich zum einen maximal noch einmal derselbe Fehler addieren, und zum anderen auch der absolute Fehler entsprechend mit diesem Faktor verringern. Daher kann man abschätzen, dass der absolute Fehler einen Maximalwert von

$$\hat{\sigma} = \frac{\varepsilon \cdot (2 + \log_2(\widetilde{M}))}{2 \cdot \widetilde{M}} \cdot \sum_{\eta=0}^{\widetilde{M}-1} \max\left(\left|\ln\left(\widetilde{D}_N(e^{j \cdot \frac{2\pi}{\widetilde{M}} \cdot \eta})\right)\right|, 1\right) \quad (10.15)$$

nicht überschreiten wird. Wieder berücksichtigt der Faktor $1/2$, dass das Cepstrum halb so groß ist, wie die Werte, die man durch eine FFT aus den Logarithmen der Betragsquadrate der Spektralwerte des minimalphasigen Anteils erhält. Durch diesen relative hoch geschätzten maximal möglichen Wert, der durch Rauschen im Cepstrum entstehen kann, wird vermieden, dass ein „Ausreißer“ im Rauschsockel zu einer unnötigen Verlängerung der FFT-Länge $\widetilde{M}$ führt. Auch ergaben sich dadurch keine Probleme, weil der Betrag des berechneten Cepstrums die Grenze im Falle einer falschen FFT-Länge immer bereits deutlich vor Erreichen des Schätzwertes für das Maximum des Rauschsockels überschritten hat. Selbst bei kritischen FFT-Längen $\widetilde{M}$ war das Cepstrum bei diesem Abbruchkriterium immer bis zum wahren Rauschsockel abgesunken, wenn die FFT-Länge als hinreichend lang detektiert worden ist.

In Bild 10.1 ist links ein Grenzfall des Abklingens des Betrags des Cepstrums abgebildet, bei dem die Länge $\widetilde{M}$ gerade noch für zu klein befunden wurde. Außer dem Betrag des Cepstrums ist noch der Schätzwert für die minimale Streuung des Rauschsockels, sowie die nach unten auf dem Schätzwert für den Maximal möglichen Rauschwert begrenzte Grenze, anhand der entschieden wird ob das Cepstrum schnell genug abklingt, dargestellt. Da in diesem Fall der Betrag des Cepstrums auch in dem punktiert hervorgehobenen Gebiet liegt, wird die FFT-Länge verdoppelt. Bild 10.1 zeigt rechts den sich danach ergebenden Betrag des Cepstrums, der als rasch genug abklingend erkannt wird. Bei der Berechnung dieser Cepstra waren die Parameter $N=5$, $M=256$, $A_0=0$, $A_1=1$ und die zwei frei wählbaren Nullstellen $z_0 = 0,7 \cdot e^{\pm 0,5j}$ eingestellt worden. Zunächst betrug die FFT-Länge $\widetilde{M}=512$ und der Parameter der Bilineartransformation $c=0,06253$. Nach der Verdoppelung der FFT-Länge waren die entsprechenden Werte $\widetilde{M}=1024$ und $c=0,121$.

Wird eine zu kleine FFT-Länge $\widetilde{M}$ festgestellt, so wird $\widetilde{M}$ verdoppelt, und es kann auch gleich der Bilineartransformationsparameter c verändert werden. Eine Veränderung des

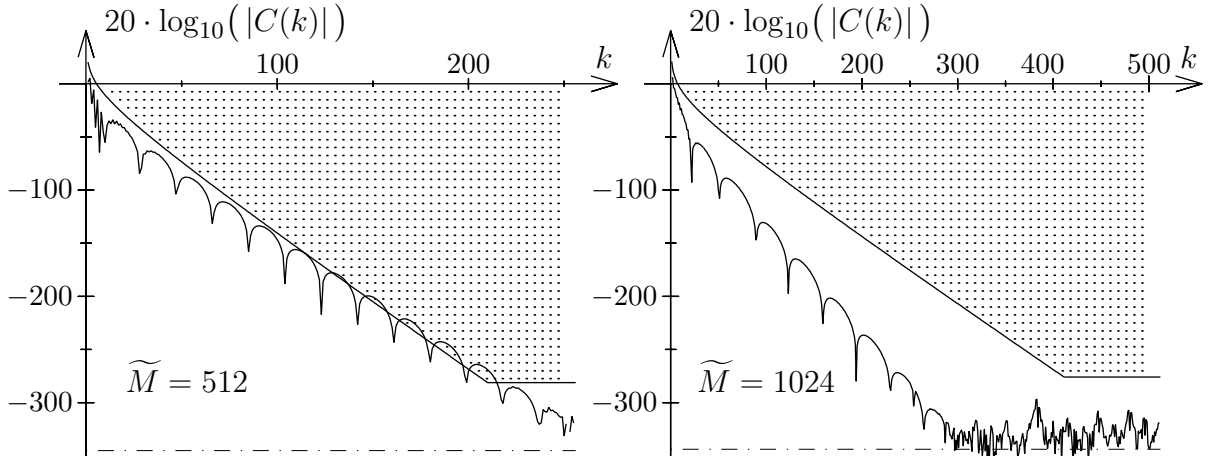


Bild 10.1: Abklingen des Betrags des Cepstrums einer Fensterfolge.

Parameters c entspricht einer weiteren Bilineartransformation der bereits bilineartransformierten Funktion $\tilde{D}_N(\tilde{z})$. Der Parameter dieser weiteren Bilineartransformation sei mit Δc bezeichnet, und entspreche dem Term $1-c$ bei der Bilineartransformation nach Gleichung ([1]:6.20) bzw. ([1]:6.21). Er liegt daher im Bereich zwischen 1 und -1 . Den neuen Wert c_{neu} erhält man aus Δc und aus dem alten Wert c_{alt} gemäß

$$c_{\text{neu}} = \frac{c_{\text{alt}} \cdot (1 - \Delta c)}{1 + (1 - c_{\text{alt}}) \cdot \Delta c} \quad (10.16)$$

Nun wird Δc so gewählt, dass die betragsmäßig größte Nullstelle $\tilde{z}_0$ des Spektrums, dessen Cepstrum mit c_{alt} berechnet wurde, die also den Anteil im Cepstrum hervorruft, der am langsamsten abklingt, weiter entfernt vom Einheitskreis liegt. Der Abstand, den diese Nullstelle vom Einheitskreis hat, ist $1 - \max(|\tilde{z}_0|)$. Nach der Verdoppelung der FFT-Länge und nach der Änderung des Parameters c kann u.U. eine andere Nullstelle den größten Betrag aufweisen. Der Abstand dieser Nullstelle vom Einheitskreis darf dann niemals kleiner sein als der halbe Abstand vor der Änderung des Parameters c , da sonst durch die Verdoppelung der FFT-Länge nichts gewonnen wäre. Man kann näherungsweise zeigen, dass diese Bedingung erfüllt ist, wenn $|\Delta c| < 1/3$ ist. Desweiteren kann man zeigen, dass sich die unbekannten Nullstellen bei einer Änderung des Bilineartransformationsparameters auf Kreisen bewegen, die durch die Punkte $z = \pm 1$ gehen. Somit wird der Abstand einer Nullstelle vom Einheitskreis am größten, wenn die Nullstelle durch die veränderte Bilineartransformation auf die imaginäre Achse abgebildet wird. Außerdem kann man wieder näherungsweise zeigen, dass eine Nullstelle, die sehr nahe am Einheitskreis liegt und einen Winkel $\tilde{\psi}_0$ in der Nähe von $\pi/2$ hat, bei der Wahl von $\Delta c = 1 - \tilde{\psi}_0 \cdot 2/\pi$ auf die imaginäre Achse abgebildet wird. Δc ist daher abhängig vom Winkel der betragsmäßig größten Nullstelle so zu wählen, dass Δc einerseits betragslich auf $1/3$ begrenzt ist, und andererseits für Winkel in der Nähe von $\pi/2$ näherungsweise der lineare Zusammenhang

zwischen Δc und ψ_0 gegeben ist. Die Funktion

$$\Delta c = \frac{1}{3} \cdot \sin\left(\left(\frac{2}{\pi}\right)^2 \cdot \psi_0^3 - \frac{6}{\pi} \cdot \psi_0^2 + \psi_0 + \frac{\pi}{2}\right) \quad (10.17)$$

erfüllt diese Bedingungen.

Um Δc berechnen zu können, muss zunächst der Winkel dieser unbekannten Nullstelle anhand des berechneten Cepstrums bestimmt werden. Das Cepstrum, das mit der inversen FFT berechnet worden ist, ergibt sich als die durch Faltung mit einem Impulskamm mit der Periode $\widetilde{M}$ periodisch fortgesetzte Folge $1/|k| \cdot \sum \tilde{z}_0^{|k|}$. Wenn man aus dieser periodischen Folge einen Ausschnitt mit $\widetilde{M}/4 \leq k < \widetilde{M}/2$ betrachtet, die Überfaltungsfehler vernachlässigt und mit k multipliziert erhält man die endliche Folge

$$\sum \tilde{z}_0^k = \sum \tilde{z}_0^{\frac{\widetilde{M}}{4}} \cdot \tilde{z}_0^{k - \frac{\widetilde{M}}{4}} = \sum \tilde{z}_0^{\frac{\widetilde{M}}{4}} \cdot \tilde{z}_0^{\tilde{k}}, \quad (10.18)$$

also mit der substituierten Zeitvariable $\tilde{k} = k - \widetilde{M}/4$, die Überlagerung abklingender Exponentialfolgen, die mit $\tilde{z}_0^{\widetilde{M}/4}$ gewichtet sind. Die Z-Transformierte dieser Folge mit $\tilde{k} \geq 0$ entspricht der Partialbruchzerlegung $\sum \tilde{z}_0^{\widetilde{M}/4} \cdot z/(z - \tilde{z}_0)$. Für $z = e^{j\Omega}$ wird dieser Ausdruck sein betragliches Maximum in der Gegend der Frequenz haben, deren Polstelle am nächsten am Einheitskreis liegt, da dann sowohl der Vorfaktor $\tilde{z}_0^{\widetilde{M}/4}$ maximal, als auch der Nenner $e^{j\Omega} - \tilde{z}_0$ minimal wird. Dies trifft vor allem dann zu, wenn die Nullstelle mit den größten Betrag deutlich näher am Einheitskreis liegt als alle anderen. Führt man die Z-Transformation nur mit der auf $0 \leq \tilde{k} < \widetilde{M}/4$ zeitlich begrenzten Folge durch, so entspricht das im Spektrum der Faltung mit der periodischen fortgesetzten si-Funktion mit dem Nullstellenabstand $8 \cdot \pi / \widetilde{M}$. Auch das Maximum des gefalteten Spektrums liegt dann etwa bei derselben Frequenz. Für die Frequenzen im Raster $8 \cdot \pi / \widetilde{M}$ lässt sich die Z-Transformation durch eine FFT der Länge $\widetilde{M}/4$ berechnen. Aus dem Ergebnis dieser FFT kann man die Lage des Maximums und damit einen Schätzwert für den Winkel der Nullstelle mit dem größten Betrag bestimmen. Den geschätzten Winkel setzt man in Gleichung (10.17) ein und das Ergebnis wiederum in Gleichung (10.16) und erhält den neuen Wert für c . Mit der doppelten FFT-Länge und dem modifizierten c berechnet man das Cepstrum erneut. Diesen Vorgang wiederholt man solange, bis das Cepstrum die nötige Länge hat.

Nachdem das Cepstrum und damit die Phase analog zu der ersten Variante berechnet worden ist, erhält man nach der Multiplikation mit dem Betrag wieder die Spektralwerte $F(\nu \cdot 2\pi / F)$, aus welchen man wieder die Fensterfolge $f(k)$ gewinnt, indem man die Fourierreihe nach Gleichung ([1]:6.16) auswertet, wobei man die Anteile der Sinus- und Kosinusreihe wieder in umgekehrter Reihenfolge akkumuliert, wie dies in [1] beschrieben

ist. Dabei sind nun weitere $2 \cdot A_1$ Fourierreihenkoeffizienten null, so dass sich eine verkürzte Reihe ergibt. Auch hier empfiehlt es sich *nicht*, die letzten A_0 Werte der Fensterfolge abschließend zu null zu setzen, auch wenn man weiß, dass diese eigentlich exakt null sein müssten, weil dadurch die Nullstellenbedingung ([1]:2.27) schlechter erfüllt werden würde. Ein kommentierter, und auf das wesentliche gekürzter Auszug aus einem Programm, das Fensterfolgen mit dem hier beschriebenen Verfahren mit automatischer Bestimmung der benötigten FFT-Länge $\widetilde{M}$ und des Bilineartransformationsparameters c berechnet, ist in Kapitel 11.1 abgedruckt.

10.2 Berechnung des Spektrums der Fensterfolge

Will man das Spektrum der Fensterfolge für beliebige Frequenzen Ω berechnen, so könnte man entweder die Fourierreihe des Spektrums, deren Koeffizienten die Werte der Fensterfolge sind, für die gewünschten Frequenzen auswerten, oder die periodische si-Funktion mit den bisher berechneten Fourierreihenkoeffizienten $F(\nu \cdot 2\pi/F)/F$ der Fensterfolge falten. Der absolute Fehler ist aber im ersten Fall auf etwa $M \cdot \varepsilon$, und im zweiten Fall auf etwa $\varepsilon/(N \cdot \sin(\Omega/2))$ begrenzt, so dass man in beiden Fällen für große Sperrdämpfung nur mehr einen Rauschsockel berechnen kann, wenn M groß ist und die Sperrdämpfung mit einer höheren Potenz in $\sin(\Omega/2)$ ansteigt. Es ist daher sinnvoller den Betrag und die Phase des Spektrums getrennt zu berechnen. Dies geschieht in ähnlicher Art, wie auch die Fourierreihenkoeffizienten $F(\nu \cdot 2\pi/F)/F$ berechnet worden sind. Dazu berechnet man die Phase wieder aus dem Cepstrum, das mit dem eben genannten Algorithmus berechnet wird, aber diesmal nicht bei den Frequenzen $\Omega = \nu \cdot 2\pi/F$ mit $\nu = 1$ (1) $N-1-A_1$, sondern bei den Frequenzen, für die das Spektrum berechnet werden soll. Der Betrag wird durch vorzeichenrichtiges radizieren des Betragsquadrats gewonnen. Da $F(-j \cdot \ln(z))$ am Einheitskreis nur einfache Nullstellen im äquidistanten Raster aufweist, ist das Vorzeichen bei der Radizierung auf einfache Weise aus Ω bestimmbar. Das Betragsquadrat von $F(\Omega)$ wird durch Überlagerung der verschobenen Betragsquadrate von $G(e^{j\Omega})$ berechnet. Dabei ist jeweils die Polstelle von $G_1(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1})$ nach Gleichung (10.5), die dem Punkt $z = e^{j\Omega}$, für den das Spektrum berechnet werden soll, am nächsten liegt, mit der entsprechenden Nullstelle im Zähler zu kürzen, so dass sich für $z = e^{j\Omega}$ im Zähler die periodisch fortgesetzte, und in der Frequenz entsprechend verschobene si-Funktion ergibt. Deren Betragsquadrat ist mit dem Betragsquadrat des Produktes der Abstände zu den restlichen Polstellen zu dividieren und mit dem Betragsquadrat des Produktes der Abstände zu den verschobenen frei gewählten Nullstellen zu multiplizieren. Da auf diese Weise eine Summe von positiven Produkten, deren Faktoren mit höchstmöglicher relativer Genauigkeit berechnet worden sind, gebildet wird, ist so das Spektrum auch für den Sperrbereich

akkurat berechenbar. In Kapitel 11.2 ist ein Programmkern angegeben, mit dessen Hilfe die Fouriertransformierte bis zu einer Sperrdämpfung von etwa $-10 \cdot \log_{10}(\text{realmin})$ berechnet werden kann. `realmin` ist dabei die kleinste positive Zahl, die am Rechner als Gleitkommazahl mit voller Genauigkeit dargestellt werden kann.

Bild 10.2 zeigt an dem Beispiel mit $M=32$, $N=13$, $A_0=12$ und $A_1=0$ den Betrag des Spektrums der Fensterfolge. Dabei wurde das Spektrum auf drei Arten berechnet. Zuerst wurde eine FFT durchgeführt, dann wurde die diskrete Fouriertransformierte von $f(k)$ durch Auswertung der Summenformel der DFT berechnet, und zuletzt wurde der oben erwähnten Algorithmus zur Berechnung des Spektrums für beliebige Frequenzen eingesetzt, der alle Spektralwerte mit etwa gleicher relativer Genauigkeit berechnet. Bei der Berechnungsvariante mit der FFT entstehen systematische Fehler die nach [7] durch die Quantisierung der Drehfaktoren zum Teil zu erklären sind. Bei der Variante mit der Auswertung der DFT-Summenformel ist der typische Rauschsockel zu erkennen, der teils durch die begrenzte Wortlänge der Werte von $f(k)$ und teils durch die Rundungen der Berechnung selbst verursacht wird. Wird das in der Wortlänge begrenzte Fenster bei dem Vorgang der Fensterung mit einem Signal multipliziert, das auf dieselbe Wortlänge quantisiert ist, so kann man nicht erwarten, dass die hohe theoretische Sperrdämpfung im Bereich von $\Omega=\pi$ zu einer Verbesserung bei der Messung des LDS beiträgt. Man kann abschätzen, dass in unserem Beispiel oberhalb einer Frequenz von etwa $3,5 \cdot 2\pi/M$ der theoretische Betrag des Spektrums unterhalb des Rauschsockels liegt.

10.3 Berechnung der AKF der Fensterfolge

Wenn man die Autokorrelationsfolge $d(k)$ der Fensterfolge berechnen will, hat man mehrere Möglichkeiten. Die erste besteht darin, die Faltungsoperation als Summe für jeden Zeitpunkt $k = A_0 - F$ (1) $F - A_0$ zu berechnen. Diese Möglichkeit kann aber für große Fensterlängen F sehr zeitaufwendig sein. Zum zweiten kann man auch zunächst die $2 \cdot F$ Werte des Betragsquadrats des Spektrums der Fensterfolge für alle Frequenzen im Abstand π/F mit dem im letzten Abschnitt beschriebenen Verfahren bestimmen. Von den $2 \cdot F$ Spektralwerten können höchstens $2 \cdot (F - N + A_1) + 1$ Werte von null verschieden sein. Aus diesen kann man mit einer DFT die Autokorrelationsfolge berechnen. Diese Art der Berechnung hat zum einen den Nachteil, dass bei Verwendung einer FFT die Werte ungenau berechnet werden, und dass zum anderen zuerst zusätzliche Spektralwerte berechnet werden müssen, obwohl die $N - A_1$ Spektralwerte, die zur Berechnung der Fensterfolge benötigt wurden, diese vollständig beschreiben. Bei der dritten Art der Berechnung der

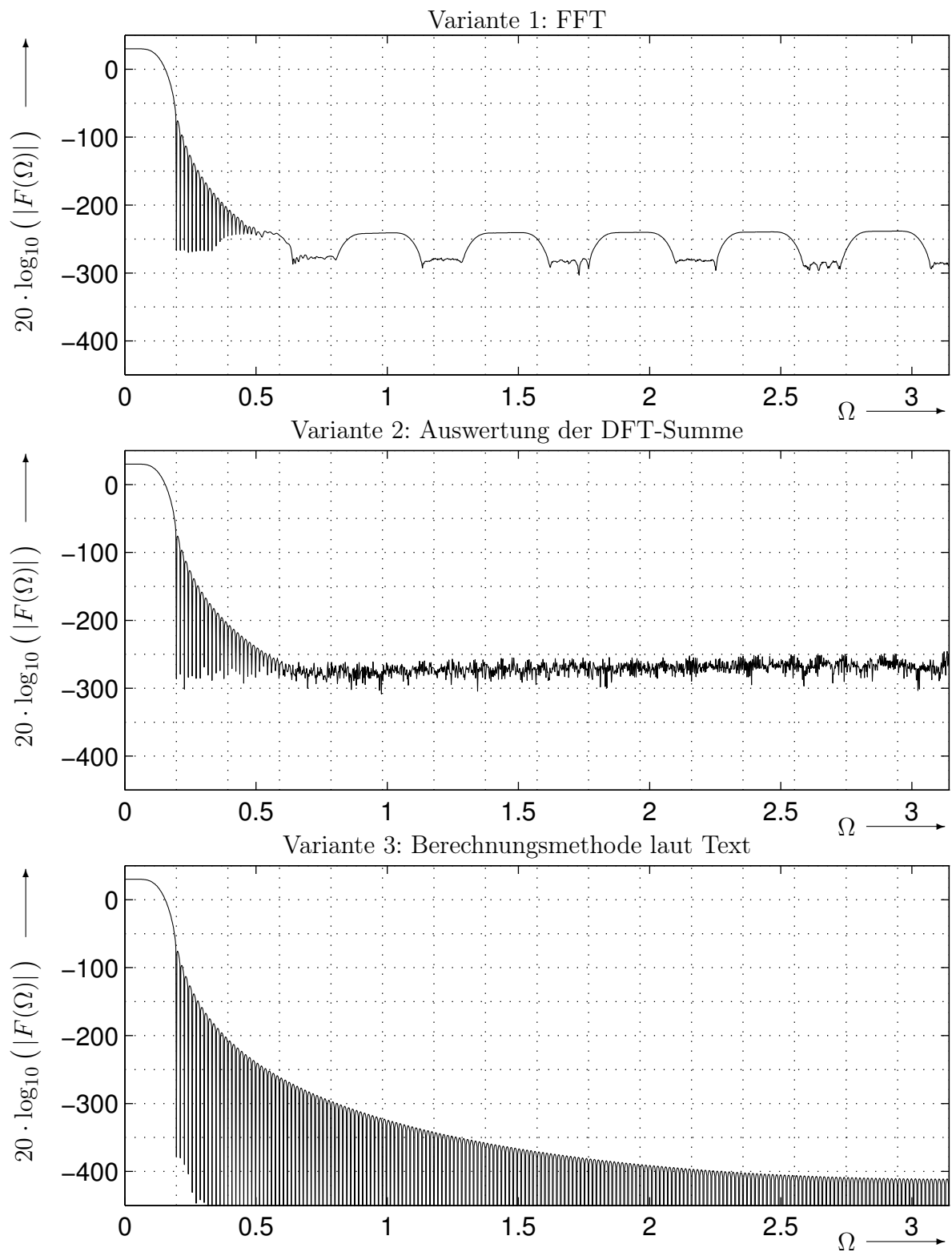


Bild 10.2: Betrag des Spektrums der Fensterfolge mit $M=32$, $N=13$, $A_0=12$ und $A_1=0$.

Autokorrelationsfolge $d(k)$ braucht man nur diese $N - A_1$ Spektralwerte. Die Fensterautokorrelationsfolge lässt sich mit den Sinusreihenkoeffizienten

$$S_\nu = \sum_{\substack{\tilde{\nu}=1+A_1-N \\ \tilde{\nu} \neq \nu}}^{N-A_1-1} \left(\frac{\Re\{F(\nu \cdot \frac{2\pi}{F})^* \cdot F(\tilde{\nu} \cdot \frac{2\pi}{F})\}}{F^2 \cdot M \cdot \tan(\frac{\pi}{F} \cdot (\tilde{\nu} - \nu))} + \frac{\Im\{F(\nu \cdot \frac{2\pi}{F})^* \cdot F(\tilde{\nu} \cdot \frac{2\pi}{F})\}}{F^2 \cdot M} \right) \quad (10.19)$$

für $0 \leq k < F$ nämlich als

$$d(k) = \sum_{\nu=1+A_1-N}^{N-A_1-1} S_\nu \cdot \sin\left(\frac{2\pi}{F} \cdot \nu \cdot k\right) + \frac{F-k}{F^2 \cdot M} \cdot \sum_{\nu=1+A_1-N}^{N-A_1-1} |F(\nu \cdot \frac{2\pi}{F})|^2 \cdot \cos\left(\frac{2\pi}{F} \cdot \nu \cdot k\right) \quad (10.20)$$

schreiben. Für $k \geq F$ ist sie null, und für negatives k ist sie gradesymmetrisch. Im übrigen lässt sich die Autokorrelationsfolge jeder Fensterfolge auf diese Weise berechnen, wobei die Summationsgrenzen entsprechend so zu modifizieren sind, dass alle von null verschiedenen Fourierreihenkoeffizienten (ggf. sind das alle Werte der DFT der Fensterfolge) berücksichtigt werden. Ein kurzer Auszug aus einer Liste eines Programms, das die Fensterautokorrelationsfolge $d(k)$ auf diese Art berechnet, ist in Kapitel 11.3 zu finden.

10.4 Fenster mit weiteren Nullstellen am Einheitskreis im $2\pi/F$ Raster

Nun will ich mich der Frage widmen, wohin die frei wählbaren Nullstellen der Z-Transformierten der Basisfensterfolge $g(k)$ gelegt werden sollten, um die Sperrdämpfung für große Werte von N in der Nähe des Durchlassbereichs auf Kosten der Sperrdämpfung in der Umgebung von $\Omega = \pi$ zu erhöhen.

Der gewünschte Durchlassbereich des Spektrums der Fensterfolge wird durch die zu approximierende mit 2π periodische Rechteckfunktion, die im Intervall $(-\pi; \pi]$ für $|\Omega| < \pi/M$ den Wert M^2 annimmt, und sonst null ist, bestimmt. Das Spektrum der Fensterfolge mit $A_0 = N - 1$ und $A_1 = 0$, die mit dem in Kapitel [1]:6 vorgestellten Algorithmus berechnet worden ist, hat ihre erste Nullstelle am Einheitskreis erst bei der Frequenz $\Omega = 2\pi/M$ also erst bei der doppelten Grenzfrequenz des gewünschten Durchlassbereichs. Daher ist die Sperrdämpfung unmittelbar außerhalb des Durchlassbereichs nicht so hoch, wie dies u. U. gewünscht wird. Das gilt besonders dann, wenn es beim RKM auf eine gute Trennung unmittelbar benachbarter Frequenzpunkte $\mu \cdot 2\pi/M$ ankommt. Stellt man nun fest, dass die Dämpfung der Fensterfolge mit $A_0 = N - 1$ und $A_1 = 0$ bei $\Omega = \pi$ so hoch ist, dass das

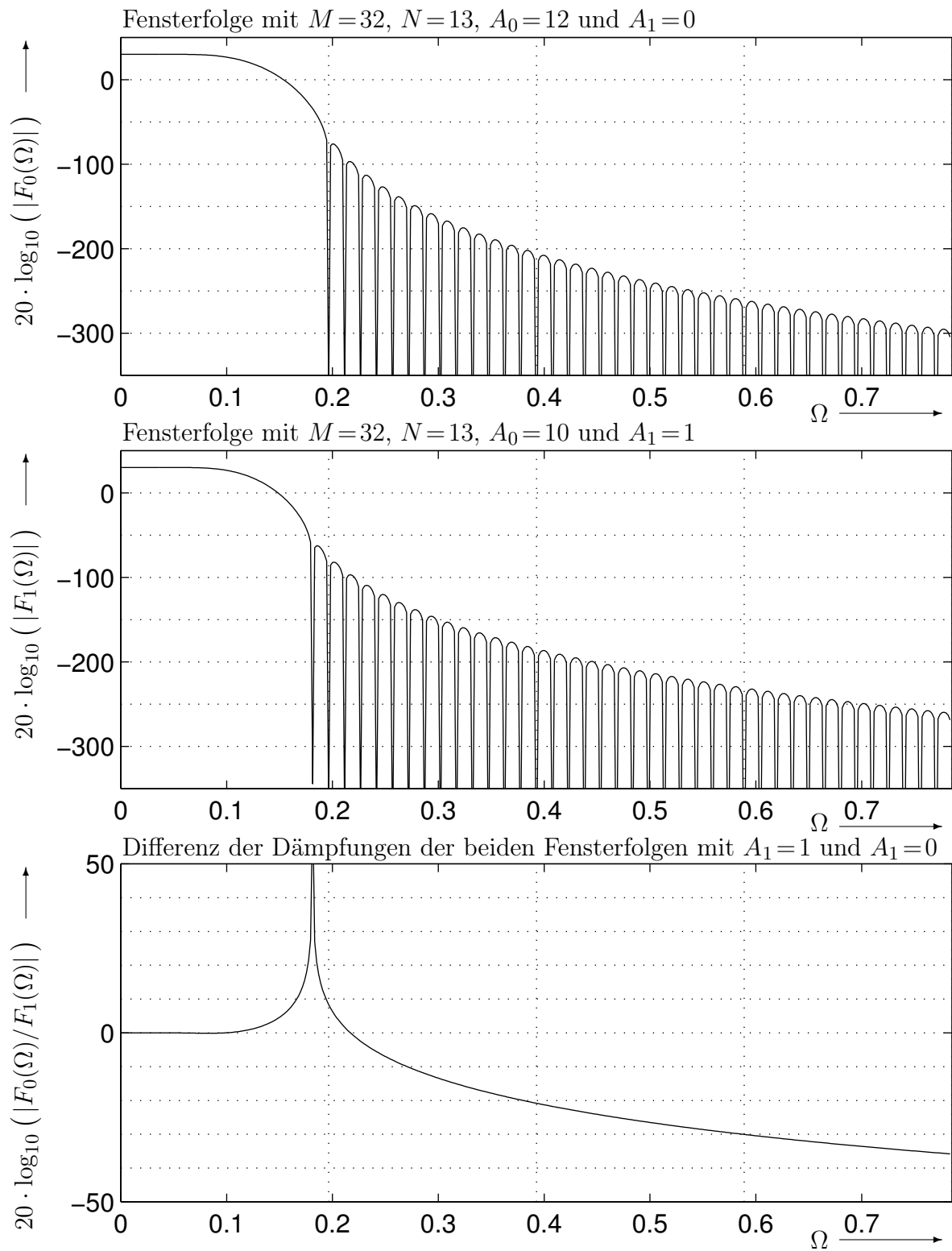


Bild 10.3: Änderung der Sperrdämpfung durch $A_1=1$ statt $A_1=0$.

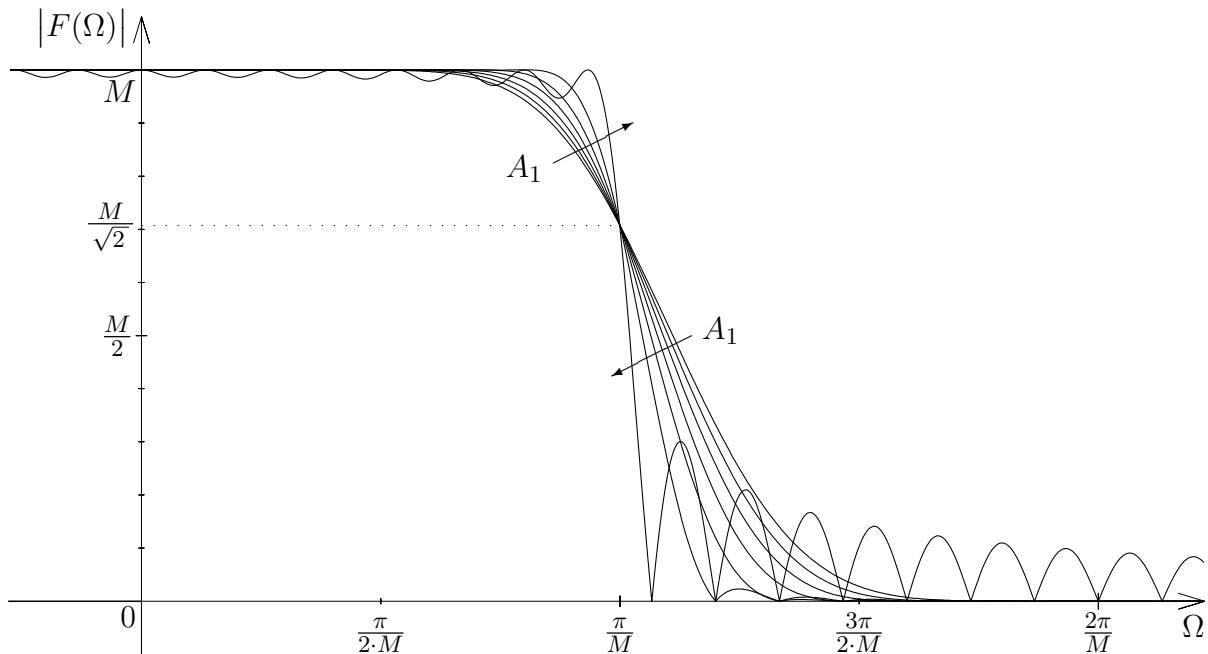


Bild 10.4: Beträge der Spektren der Fensterfolgen mit $M=1024$ und $N=15$.

$A_1 = 1$ (1) 7 als Parameter, $A_0 = 14 - 2 \cdot A_1$.

Spektrum der Fensterfolge in einem großen Frequenzbereich unterhalb des Rauschsockels liegt, so kann man mit $A_0 = N - 3$ und $A_1 = 1$ eine Fensterfolge berechnen, deren Spektrum bei $\Omega = (N - 1) \cdot 2\pi/F$ eine weitere Nullstelle aufweist, die bewirkt, dass das Betragsquadrat des Spektrums im Bereich $\pi/M < |\Omega| < 2 \cdot \pi/M$ kleiner ist als bei der Fensterfolge nach [1]. Man erkaufte sich die höhere Sperrdämpfung in diesem Frequenzbereich jedoch auf Kosten eines Abfalls des Betrags des Spektrums für höhere Frequenzen mit einer Potenz, die um zwei Grade niedriger ist, so dass die Sperrdämpfung oberhalb der Frequenz $\Omega = 2\pi/M$ bald unter die Sperrdämpfung der zuerst berechneten Fensterfolge abfällt. Bild 10.3 zeigt, an dem Beispiel $M=32$, $N=13$, $A_0=10$ und $A_1=1$ wie sich die Sperrdämpfung der Fensterfolge durch die zwei zusätzlichen Nullstellen der Z-Transformierten der Basisfensterfolge am Einheitskreis verändert.

Am Beispiel des Fensters mit $M=1024$ und $N=15$ wird nun gezeigt, wie sich der Betrag des Spektrums verändert, wenn man den Parameter $A_1 = 1$ (1) 7 variiert. Bild 10.4 zeigt die Beträge der Spektren dieser Fensterfolgen im Frequenzbereich zwischen 0 und $2\pi/M$ in linearer Darstellung. Mit zunehmendem Wert des Parameters A_1 entstehen weitere Nullstellen, die sich unterhalb der Frequenz $2\pi/M$ im Abstand $2\pi/F$ an die bereits vorhandenen Nullstellen anschließen. Sie bewirken, dass der Betrag des Spektrums der Fensterfolge mit steigendem A_1 bei der Frequenz π/M immer steiler abfällt. Diese Zunahme der spektralen Flankensteilheit wird jedoch mit einer Abnahme der Sperrdämpfung der

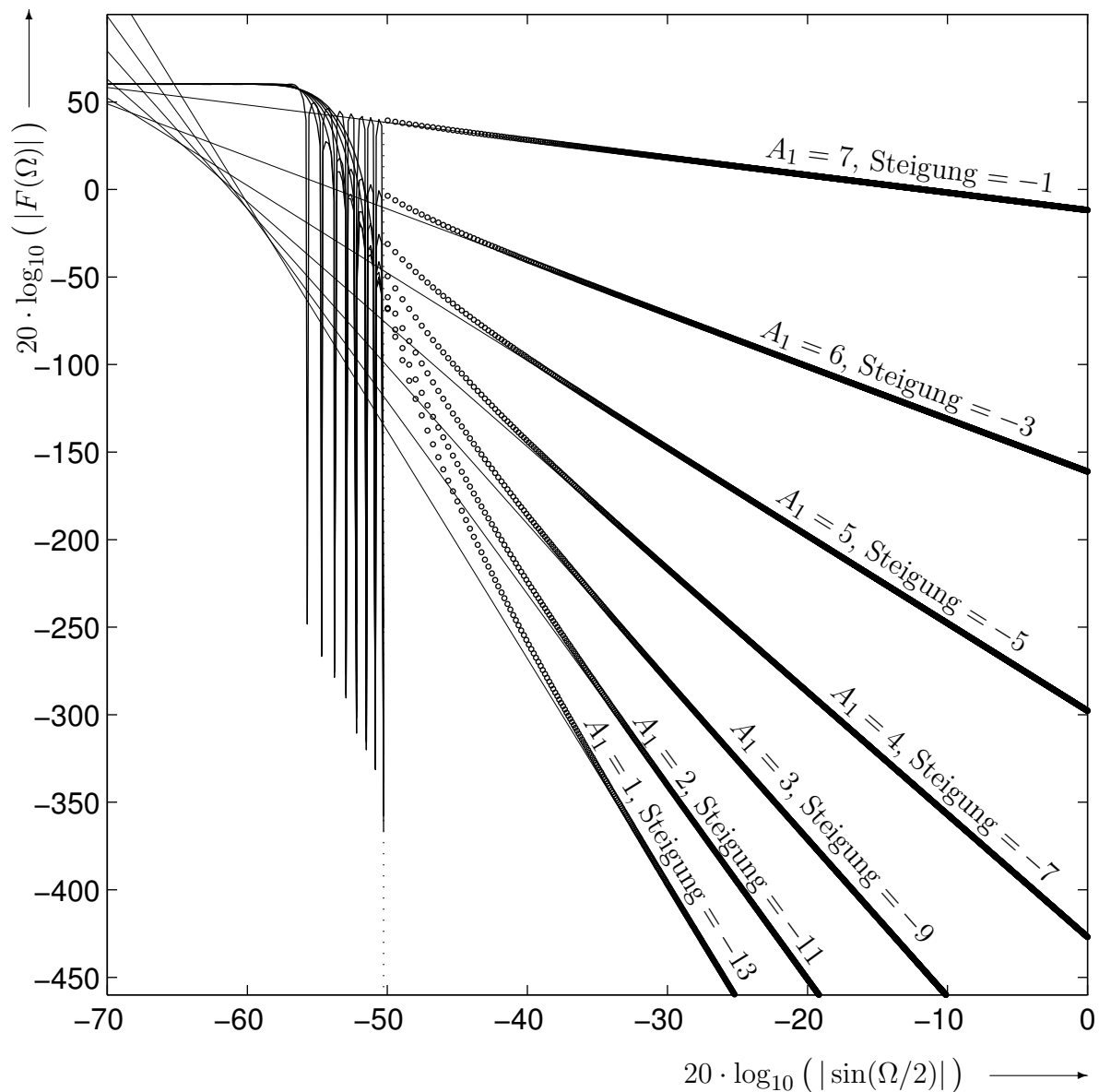


Bild 10.5: Sperrdämpfung der Nebenmaxima der Spektren der Fensterfolgen mit $M = 1024$ und $N = 15$. $A_1 = 1$ (1) 7 als Parameter, $A_0 = 14 - 2 \cdot A_1$

ersten Nebenmaxima und mit einer Abnahme der Potenz des Anstiegs der Sperrdämpfung mit $\sin(\Omega/2)$ erkaufte.

Bild 10.5 zeigt, dass die Sperrdämpfung der Nebenmaxima der Spektren der Fensterfolgen für $\Omega \gg 2\pi/M$ mit der Potenz $N - 2 \cdot A_1$ ansteigt. Bei dieser Darstellung wurde die Ordinate um den Faktor 8 enger skaliert als die Abszisse. Um die Betragsspektren im vollen Dynamikbereich (ca. 500 dB) der Graphik bestimmen zu können, wurden diese nicht mit Hilfe der DFT aus den Fensterfolgen, sondern wie in Kapitel 10.2 beschrieben, berechnet.

10.5 Basisfenster mit sechs zusätzlichen Nullstellen am Einheitskreis

Eine andere Möglichkeit die Sperrdämpfung des Fensters unmittelbar oberhalb des gewünschten Durchlassbereichs zu erhöhen, besteht darin, die frei wählbaren Nullstellen der Z-Transformierten der Basisfensterfolge geeignet zu wählen. Man kann z. B. die Nullstellen des Polynoms $z^{F-N+2\cdot A_1} \cdot G_2(z)$, das nur die frei wählbaren Nullstellen enthält, auf dem Einheitskreis in dem Frequenzbereich verteilen, in dem der Betrag des Spektrums der Fensterfolge mit $A_0 = N-1$ und $A_1=0$ oberhalb des Rauschsockels liegt. Dabei achtet man darauf, dass das Betragsquadrat von $G_2(e^{j\Omega})$ bei niedrigen Frequenzen — also im Durchlassbereich — nicht kleiner sein sollte als in dem Frequenzbereich, in dem man die Sperrdämpfung erhöhen möchte. Durch solch eine Wahl der Nullstellen wird das Spektrum der Basisfensterfolge in dem Frequenzbereich zusätzlich gedämpft, der nach der Überlagerung der verschobenen Betragsquadrate der Spektren den Übergangsbereich ergibt, während das Spektrum der Basisfensterfolge in dem Bereich angehoben wird, aus dem der Bereich des Rauschsockels hervorgeht. Somit wird im Übergangsbereich des Spektrums der Fensterfolge eine höhere Sperrdämpfung erzielt, als bei den Fenster nach [1]. Bild 10.6 zeigt an dem Beispiel mit $M=32$, $N=13$, $A_0=6$, $A_1=0$ und weiteren 6 Nullstellen am Einheitskreis, die das Polynom

$$G_2(z) = (1 - 2 \cdot \cos(0,37) \cdot z^{-1} + z^{-2}) \cdot (1 - 2 \cdot \cos(0,5) \cdot z^{-1} + z^{-2}) \cdot (1 - 2 \cdot \cos(0,6) \cdot z^{-1} + z^{-2})$$

bilden, erstens den Betragsfrequenzgang dieses Polynoms, zweitens die Dämpfung des Spektrums der damit berechneten Fensterfolge und drittens die Veränderung der Dämpfung gegenüber der in Bild 10.2 dargestellten Dämpfung in dem Frequenzbereich, in dem das Spektrum der Fensterfolgen nicht im Rauschsockel versinkt.

Bei der Berechnung der drei Fensterfolgen der letzten beiden Unterkapitel wurden für die Berechnung des Cepstrums folgende FFT-Längen benötigt.

	M	N	A_0	A_1	$\widetilde{M}$
$F_0(\Omega)$	32	13	12	0	256
$F_1(\Omega)$	32	13	10	1	256
$F_2(\Omega)$	32	13	6	0	4096

Eine Kombination der beiden eben beschriebenen Methoden die frei wählbaren Nullstellen einzustellen um im Übergangsbereich einen besseren Dämpfungsfrequenzgang zu erhalten, ist ebenfalls denkbar. Da es für die Wahl der Nullstellen keine allgemeine optimale Lösung geben kann, weil es vom vorgesehenen Einsatz des Fensters abhängt, was unter optimal

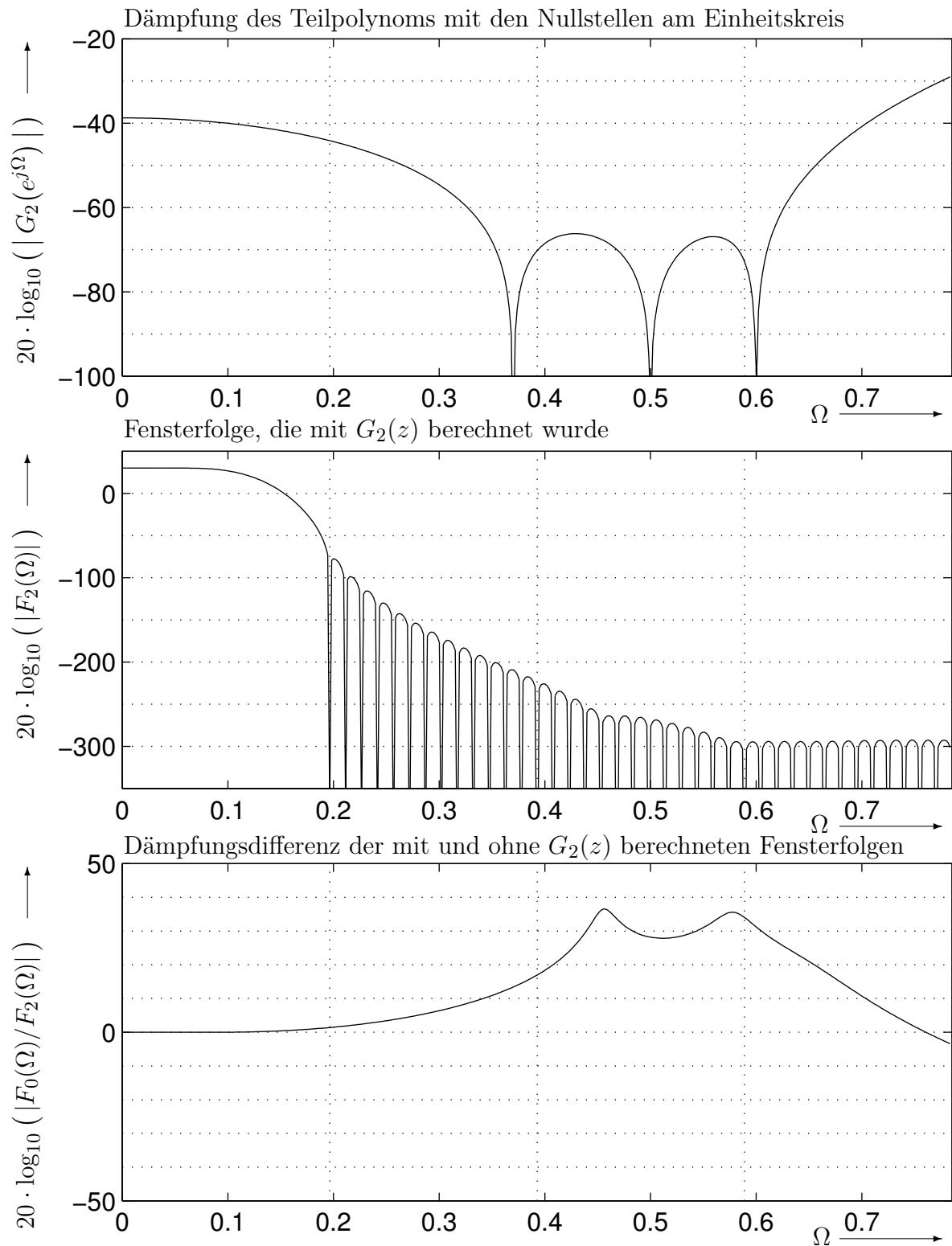


Bild 10.6: Betrag des Spektrums der Fensterfolge mit $A_1=0$ und sechs weiteren Nullstellen am Einheitskreis.

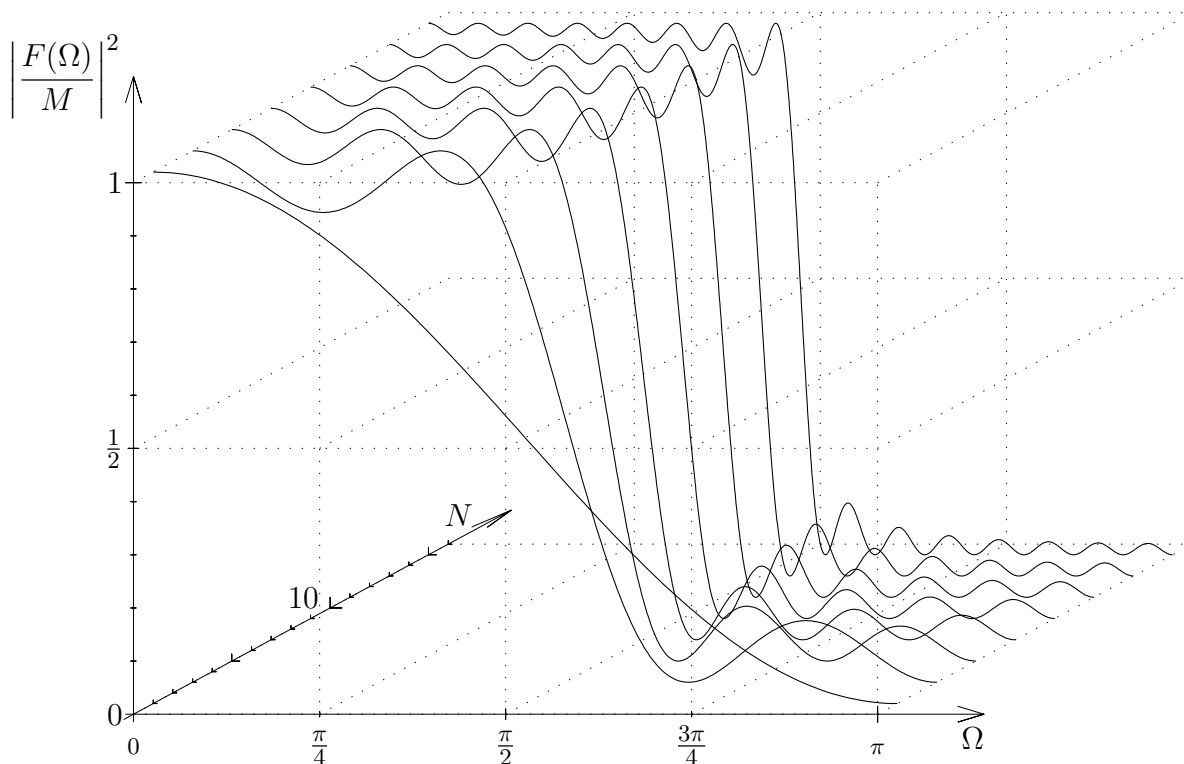


Bild 10.7: Spektren der Autokorrelationsfolgen der Fenster mit $M=2$ (Halbbandfilter). $N = 1$ (2) 15 als Parameter; $A_1 = (N-1)/2$.

zu verstehen ist, erscheint es mir am einfachsten, die Nullstellenlage durch Probieren zu erhalten, also durch iteratives Verschieben der Nullstellen, Berechnen des Spektrums und Beurteilen der Dämpfung im Übergangsbereich. In den meisten Fällen wird man sowieso den Parameter N nicht so groß wählen, dass das Spektrum der Fensterfolge bereits bei niedrigen Frequenzen in den Rauschsockel absinkt, so dass man das in [1] vorgestellte Fenster verwenden kann.

10.6 Halbbandfilter

Bereits in Kapitel [1]:6.3 hatten wir festgestellt, dass die FIR-Filter, die man mit $M=2$ erhält, wenn man die Werte der Fensterautokorrelationsfolge $d(k)$ als Filterkoeffizienten verwendet, Halbbandfilter sind. Man kann nun versuchen die in Bild [1]:6.14 dargestellten Spektralverläufe günstig zu beeinflussen, indem man die bei der Konstruktion der Fensterfolge nach Kapitel 10.1 neu hinzugekommenen Freiheitsgrade ausnutzt. Berechnet man die Filterkoeffizienten als die Werte der Autokorrelationsfolge, die man bei ungeradem N erhält, wenn man A_1 auf den maximal möglichen Wert $A_1 = (N-1)/2$ setzt, so erhält

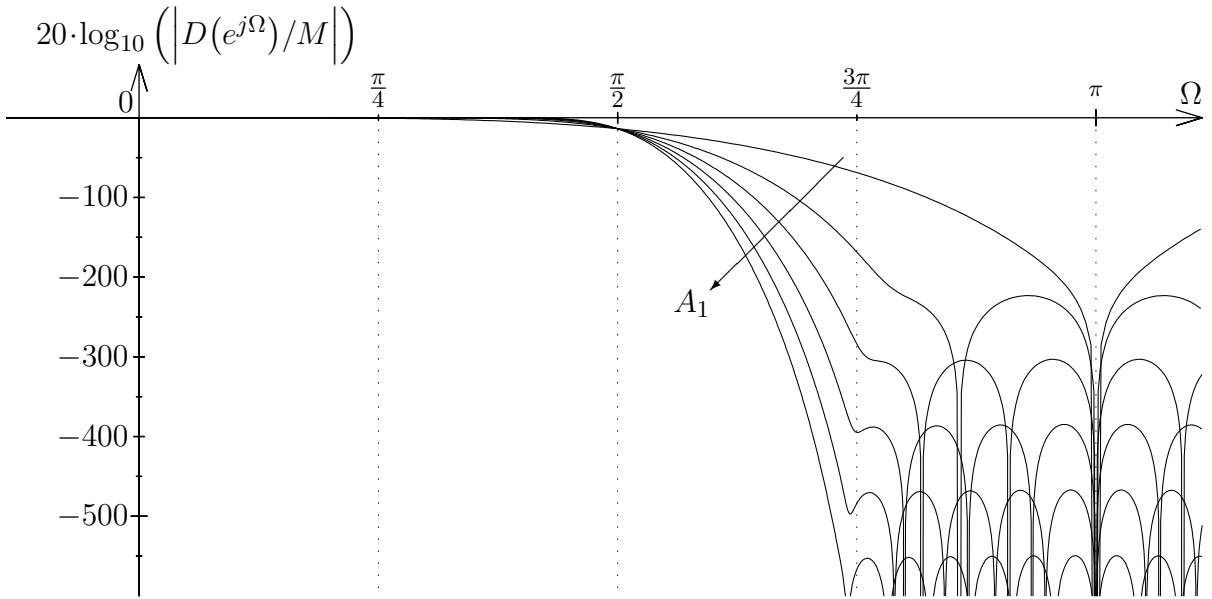


Bild 10.8: Spektren der Autokorrelationsfolgen der Fenster mit $M=2$ (Halbbandfilter).

$A_1 = 0$ (1) 5 als Parameter, $N = 4 \cdot A_1 + 3$ und Nullstellen bei $z_{0,\rho} = e^{(-1)^\rho \cdot j \cdot \frac{2\pi}{F} \cdot (A_1+2)}$ mit $\rho = 1$ (1) $2 \cdot (A_1+1)$.

man die Halbbandfilter, deren Spektren in Bild 10.7 dargestellt sind. Diese zugrundeliegenden Fenster besitzen die Länge $F = M \cdot N = 2 \cdot N$. Dem entsprechend ist die Länge der Fensterautokorrelationsfolge, also des Halbbandfilters, gleich $2 \cdot F - 1 = 4 \cdot N - 1$. Man erkennt, dass im Vergleich zu Bild [1]:6.14 der Übergangsbereich schmäler ist, was jedoch mit einer höheren Welligkeit in Durchlass- und Sperrbereich erkauft wird.

Es erscheint daher sinnvoll, hier den Parameter A_1 nicht maximal zu wählen, sondern auch einige der frei wählbaren Nullstellen $z_{0,\rho}$ dazu zu nutzen, einen Kompromiss aus einer hohen Sperrdämpfung und einer geringen Breite des Übergangsbereichs zu finden. Es zeigte sich, dass man beispielsweise im Frequenzbereich $3\pi/4 < \Omega \leq \pi$ eine fast gleichmäßig hohe minimale Sperrdämpfung erhält, wenn man den Parameter $N = 4 \cdot A_1 + 3$ einstellt, und als frei wählbare Nullstellen $z_{0,\rho}$ das Nullstellenpaar $e^{\pm j \cdot \pi \cdot \frac{A_1+2}{4 \cdot A_1+3}}$ mit der Vielfachheit A_1+1 verwendet. Das Spektrum der Basisfensterfolge besitzt dann nur Nullstellen am Einheitskreis, wobei die niederfrequentesten Nullstellen die Vielfachheit A_1+2 aufweisen, weil die frei gewählten Nullstellen alle auf den sowieso schon vorhandenen einfachen Nullstellen bei diesen Frequenzen zu liegen kommen. Alle weiteren Nullstellen im Abstand $2\pi/F$ sind einfach. Für die Parameterwerte $A_1 = 0$ (1) 5 sind die Betragsfrequenzgänge der entsprechenden Halbbandfilter in Bild 10.8 halblogarithmisch dargestellt. Die Filterlängen der Halbbandfilter sind $2 \cdot F - 1 = 16 \cdot A_1 + 11$. Für $A_1 = 5$ wird bei der Berechnung des Cepstrums eine noch tolerierbare FFT-Länge von $\widetilde{M} = 2^{13}$ benötigt.

10.7 Berechnung einer kontinuierlichen Fensterfunktion

Dieses Unterkapitel behandelt einen Algorithmus, mit dessen Hilfe eine kontinuierliche Fensterfunktion $f_\infty(t)$ berechnet werden kann, die endlich lang ist, und deren Spektrum ähnliche Eigenschaften, wie die in den Gleichungen ([1]:2.20) und ([1]:2.27) beschrieben, besitzt. Es wird — wie sich im weiteren zeigen wird — durch die Art der Konstruktion der Fensterfunktion erreicht, dass erstens die Fensterfunktion zeitlich gemäß

$$f_\infty(t) = 0 \quad \text{für} \quad t < 0 \quad \vee \quad t \geq 1 \quad (10.21)$$

begrenzt ist, dass zweitens das Spektrum $F_\infty(\omega)$ der Fensterfunktion äquidistante Nullstellen auf der imaginären Achse mit Ausnahme des Bereichs niedriger Frequenzen aufweist, so dass

$$F_\infty(2\pi \cdot \nu) = 0 \quad \text{für} \quad \nu \in \mathbb{Z} \quad \text{mit} \quad \nu \neq N-1 \quad (1) \quad N-1 \quad (10.22)$$

gilt, und dass drittens das Betragsquadratspektrum — bei entsprechender Normierung mit einem konstanten Faktor — die Eigenschaft

$$\sum_{\mu=-\infty}^{\infty} |F_\infty(\omega - \mu \cdot N \cdot 2\pi)|^2 = 1. \quad (10.23)$$

besitzt. Die Fensterautokorrelationsfunktion $d_\infty(t) = f_\infty(t) * f_\infty(-t)$ weist daher äquidistante Nullstellen bei Vielfachen von $1/N$ außer bei $t=0$ auf.

Die bisher behandelten zeitdiskreten Fensterfolgen $f(k)$, die von dem Parameter M abhängen, kann man als die Abtastwerte kontinuierlicher Fensterfunktionen für die Abtastzeitpunkte

$$t = \frac{k}{N \cdot M} \quad (10.24)$$

betrachten. Dabei sind die kontinuierlichen Fensterfunktionen, die ebenfalls von M abhängen, außerhalb des in Gleichung (10.21) genannten Zeitintervalls null. Wenn man nun den Grenzübergang $M \rightarrow \infty$ durchführt, so ergibt sich die in k unbegrenzte Fensterfolge, die durch unbegrenzt feine Abtastung der kontinuierlichen Fensterfunktion $N \cdot f_\infty(t)$ entsteht. Wie auch bei den bisher behandelten Fensterfolgen, beschreiben die N Spektralwerte $F_\infty(\nu \cdot 2\pi)$ mit $\nu = 0 \quad (1) \quad N-1$, die kontinuierliche Fensterfunktion vollständig. Die Fensterfunktion lässt sich analog zu Gleichung ([1]:6.16) durch Auswertung der Fourierreihe für jeden beliebigen Zeitpunkt gemäß

$$f_\infty(t) = \begin{cases} \sum_{\nu=1+A_1-N}^{N-1-A_1} F_\infty(\nu \cdot 2\pi) \cdot e^{j \cdot 2\pi \cdot \nu \cdot t} & \text{für} \quad 0 \leq t < 1 \\ 0 & \text{sonst.} \end{cases} \quad (10.25)$$

berechnen. Die Spektralwerte $F_\infty(\nu \cdot 2\pi)$ sind die Grenzwerte der auf M normierten Spektralwerte $F(\nu \cdot \frac{2\pi}{F})$ der zeitdiskreten Fensterfolgen für $M \rightarrow \infty$. Beim Grenzübergang ist die normierte Frequenz Ω , durch die Kreisfrequenz

$$\omega = M \cdot N \cdot \Omega \quad (10.26)$$

zu ersetzen, so dass ein konstanter Wert von ω einem konstanten Produkt aus M und Ω entspricht. Wegen des eben beschriebenen Grenzwertverhaltens, kennzeichne ich die kontinuierliche Fensterfunktion und deren Spektrum mit dem Index ∞ . Die kontinuierliche Fensterfunktion lässt sich natürlich nicht beim RKM einsetzen. Dennoch halte ich es für denkbar, dass auf anderen Gebieten eine kontinuierliche Fensterfunktion mit den eben genannten Eigenschaften von Interesse sein könnte. Daher sei in diesem Kapitel ein Algorithmus aufgeführt, der dem Algorithmus zur Konstruktion der diskreten Fensterfolge mit $A_0 = N-1$ und $A_1 = 0$ sehr ähnlich ist, und mit dessen Hilfe die Spektralwerte $F_\infty(\nu \cdot 2\pi)$ berechnet werden können. Auch für andere Einstellungen der Parameter A_0 und A_1 kann man kontinuierliche Fensterfunktion als Grenzwertlösungen für $M \rightarrow \infty$ erhalten. Der Leser möge die dazu notwendigen Modifikationen selbst herleiten. In dem in Kapitel 11.4 aufgelisteten Programmrumpf ist diese Erweiterung enthalten.

Um die Fensterfunktion $f_\infty(t)$ zu erzeugen gehen wir zunächst wieder von einer reellen Basisfensterfunktion $g_\infty(t)$ aus, die im gleichen Zeitintervall von null verschieden sein kann, wie die Fensterfunktion $f_\infty(t)$. Als Basisfensterfunktion wird ein Ausschnitt einer periodischen Funktion verwendet. Abhängig davon, ob N eine gerade oder ungerade ganze Zahl ist, ist die Periode zwei oder eins. In beiden Fällen lässt sich die Basisfensterfunktion als Fourierreihe schreiben:

$$g_\infty(t) = \begin{cases} \sum_{\nu=\frac{1-N}{2}}^{\frac{N-1}{2}} G_\infty(j \cdot \nu \cdot 2\pi) \cdot e^{j \cdot 2\pi \cdot \nu \cdot t} & \text{für } 0 \leq t < 1. \\ 0 & \text{sonst.} \end{cases} \quad (10.27)$$

Sie ist durch die N Fourierreihenoeffizienten $G_\infty(j \cdot \nu \cdot 2\pi)$ vollständig festgelegt. Das Spektrum der Basisfensterfunktion ist dann

$$\begin{aligned} G_\infty(j\omega) &= \sum_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} G_\infty(j \cdot \nu_2 \cdot 2\pi) \cdot e^{-j \cdot (\frac{\omega}{2} - \nu_2 \cdot \pi)} \cdot \text{si}\left(\frac{\omega}{2} - \nu_2 \cdot \pi\right) = \\ &= e^{-j \cdot \frac{\omega - (N-1) \cdot \pi}{2}} \cdot \sin\left(\frac{\omega - (N-1) \cdot \pi}{2}\right) \cdot \sum_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{2 \cdot G_\infty(j \cdot \nu_2 \cdot 2\pi)}{\omega - \nu_2 \cdot 2\pi}, \end{aligned} \quad (10.28)$$

und weist für $\omega = (2 \cdot \nu - N + 1) \cdot \pi$ mit $\nu \in \mathbb{Z}$ außer für $\nu = 0$ (1) $N - 1$ äquidistante Nullstellen auf. Durch Multiplikation der auf $-1 < t < 1$ zeitlich begrenzten Basisfensterautokorrelationsfunktion

$$g_{Q,\infty}(t) = g_\infty(t) * g_\infty(-t) = \frac{1}{2\pi} \cdot \int_{-\infty}^{\infty} |G_\infty(j\omega)|^2 \cdot e^{j\omega \cdot t} \cdot d\omega \quad (10.29)$$

mit der periodisch fortgesetzten si-Funktion

$$\frac{\sin(N \cdot \pi \cdot t)}{\sin(\pi \cdot t)} \quad (10.30)$$

entsteht die Funktion $d_\infty(t)$. Diese weist dieselben Nullstellen im Abstand $1/N$ auf, wie die periodisch fortgesetzte si-Funktion. Im Frequenzbereich entspricht die Multiplikation der Faltung des Betragsquadratspektrums von $g_\infty(t)$ mit einem zu $\omega = 0$ symmetrischen Impulskamm der konstanten Stärke 2π und der Länge N mit einem Impulsabstand von 2π . Das Faltungsintegral lässt sich daher auch als Summe mit N Summanden schreiben:

$$D_\infty(j\omega) = \sum_{\nu_1 = \frac{1-N}{2}}^{\frac{N-1}{2}} |G_\infty(j\omega - j \cdot \nu_1 \cdot 2\pi)|^2. \quad (10.31)$$

Wie im zeitdiskreten Fall, kann man sich auch im kontinuierlichen Fall anhand der Nullstellenlage der verschobenen, an der Faltung beteiligten Spektren überlegen, dass das so entstandene Spektrum die Gleichung (10.22) erfüllt. Da $|G_\infty(j\omega)|^2$ eine reelle, geradesymmetrische und nichtnegative Funktion in ω ist, und da die Überlagerung ebenfalls symmetrisch zu $\omega = 0$ erfolgt, ist auch $D_\infty(j\omega)$ eine reelle, geradesymmetrische und nichtnegative Funktion in ω . Daher sind alle Nullstellen von $D(s)$ spiegelsymmetrisch zur reellen Achse, alle Nullstellen auf der imaginären Achse sind von gerader Vielfachheit und alle Nullstellen, die nicht auf der imaginären Achse liegen, sind spiegelsymmetrisch zur imaginären Achse. Es handelt sich somit bei $D_\infty(j\omega)$ um das Betragsquadratspektrum einer reellen Fensterfunktion. $d_\infty(t)$ lässt sich daher als $f_\infty(t) * f_\infty(-t)$ schreiben. Durch Wahl eines geeigneten konstanten Faktors kann für das Spektrum der Basisfensterfunktion

$$\sum_{\nu = \frac{1-N}{2}}^{\frac{N-1}{2}} |G_\infty(j \cdot 2\pi \cdot \nu)|^2 = 1 \quad (10.32)$$

erzwungen werden. Dadurch wird erreicht, dass auch die Gleichung (10.23) erfüllt wird.

Im weiteren möchte ich mich auf den Fall beschränken, dass als Basisfensterfunktion die $N - 1$ -te Potenz einer halben Periode der Sinusfunktion $\sin(\pi \cdot t)$ mit geeigneter Amplitude

verwendet wird:

$$g_{\infty}(t) = \begin{cases} 2^{N-1} \cdot \left(\frac{2 \cdot N-2}{N-1}\right)^{-\frac{1}{2}} \cdot \sin(\pi \cdot t)^{N-1} & \text{für } 0 \leq t < 1 \\ 0 & \text{sonst.} \end{cases} \quad (10.33)$$

Bei der in Kapitel [1]:6 verwendeten diskreten Basisfensterfolge waren die letzten $N-1$ Werte null. Wenn wir diese Werte als die mit steigendem M immer enger beieinanderliegenden Abtastwerte einer kontinuierlichen Basisfensterfunktion betrachten, rücken diese $N-1$ Nullstellen mit zunehmendem Parameter M immer enger zusammen, so dass man grenzwertig die in der Basisfensterfunktion $g_{\infty}(t)$ vorhandene $N-1$ -fache Nullstelle erhält. Die der Basisfensterfunktion zugrundeliegende periodische Funktion besitzt die Fourierreihenkoeffizienten

$$G_{\infty}(j \cdot \nu \cdot 2\pi) = \left(\frac{2 \cdot N-2}{N-1}\right)^{-\frac{1}{2}} \cdot j^{-2 \cdot \nu} \cdot \left(\frac{N-1}{\frac{N-1}{2} + \nu}\right) \quad \text{für } \nu = \frac{1-N}{2} (1) \frac{N-1}{2}. \quad (10.34)$$

Wenn man dies und die Partialbruchzerlegung

$$\frac{(N-1)! \cdot \pi^{N-1}}{\prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} \left(\frac{\omega}{2} - \nu_2 \cdot \pi\right)} = j^{N-1} \cdot \sum_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{2 \cdot \left(\frac{N-1}{2} + \nu_2\right)}{\omega - \nu_2 \cdot 2\pi} \cdot j^{-2 \cdot \nu_2} \quad (10.35)$$

in Gleichung (10.28) einsetzt, ergibt sich für des Spektrum der Basisfensterfunktion:

$$\begin{aligned} G_{\infty}(j\omega) &= \left(\frac{2 \cdot N-2}{N-1}\right)^{-\frac{1}{2}} \cdot e^{-j \cdot \frac{\omega \cdot (N-1) \cdot \pi}{2}} \cdot \sin\left(\frac{\omega \cdot (N-1) \cdot \pi}{2}\right) \cdot \sum_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{2 \cdot \left(\frac{N-1}{2} + \nu_2\right)}{\omega - \nu_2 \cdot 2\pi} \cdot j^{-2 \cdot \nu_2} = \\ &= \frac{\pi^{N-1} \cdot (N-1)! \cdot e^{-j \cdot \frac{\omega}{2}} \cdot \sin\left(\frac{\omega \cdot \pi \cdot (N-1)}{2}\right)}{\left(\frac{2 \cdot N-2}{N-1}\right)^{\frac{1}{2}} \cdot \prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} \left(\frac{\omega}{2} - \pi \cdot \nu_2\right)}. \end{aligned} \quad (10.36)$$

Da es sich bei der Basisfensterfunktion um eine endliche, zeitlich begrenzte und kausale Zeitfunktion handelt konvergiert deren Laplacetransformierte für alle $s \in \mathbb{C}$. Sie ergibt sich aus der Fouriertransformierten indem man $s = j\omega$ substituiert.

$$G_{\infty}(s) = \frac{(j \cdot 2\pi)^N \cdot (N-1)! \cdot e^{-\frac{s}{2}} \cdot \sin\left(\frac{s \cdot j \cdot (N-1) \cdot \pi}{2 \cdot j}\right)}{\pi \cdot \left(\frac{2 \cdot N-2}{N-1}\right)^{\frac{1}{2}} \cdot \prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} (s - j \cdot \nu_2 \cdot 2\pi)} \quad (10.37)$$

Man beachte, dass sich die einfachen Polstellen auf der imaginären Achse mit den einfachen Nullstellen der Sinusfunktion kürzen lassen.

Anmerkung: Bei dieser Basisfensterfunktion tritt im Zähler der Laplacetransformierten kein Polynom in s auf. Dadurch ergibt sich ein asymptotischer Anstieg der Sperrdämpfung mit der Potenz in ω , die dem Grad des Nennerpolynoms — also N — entspricht. Im zeitdiskreten Fall entspricht diese Graddifferenz des Zähler- und des Nennerpolynoms der Wahl $A_0 = N - 1$, da dann auch dort die Z-Transformierte der Basisfensterfolge außer den äquidistanten Nullstellen am Einheitskreis im Raster $2\pi/F$ (entspricht der Sinusfunktion im Zähler) keine weiteren Nullstellen enthält. Hätte man im kontinuierlichen Fall eine Basisfensterfunktion mit anderen Fourierreihenkoeffizienten gewählt, so würde im Zähler ein Polynom in s auftreten dessen Grad maximal $N - 1$ wäre. Dieser Fall entspräche bei der zeitdiskreten Fensterfolge der Wahl von beliebigen Nullstellen $z_{0,\rho}$. Ebenso könnte man also zur Konstruktion der kontinuierlichen Fensterfunktion maximal $N - 1$ Nullstellen in der s -Ebene frei wählen. Die Modifikationen, die sich daraus für den im folgenden angegebenen Algorithmus ergäben, sind ähnlich denen, die sich bei der Konstruktion der zeitdiskreten Fensterfolge durch die Einführung der frei wählbaren Nullstellen ergeben haben. Auf eine Darstellung dieser Modifikationen wird jedoch verzichtet. Das in Kapitel 11.4 aufgelistete Programm enthält die Modifikationen, die es ermöglichen $N - 1$ Nullstellen frei zu wählen.

Wie oben beschrieben erhalten wir das Betragsquadrat des Spektrums der Fensterfunktion durch die Überlagerung der verschobenen Betragsquadrate des Spektrums der Basisfensterfunktion. Für $s = j\omega$ lässt sich Betragsquadrat des Spektrums der Fensterfunktion aus

$$\begin{aligned}
 F_\infty(-j \cdot s) \cdot F_\infty(j \cdot s) &= D_\infty(s) = \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} G_\infty(s - j \cdot \nu_1 \cdot 2\pi) \cdot G_\infty(-s + j \cdot \nu_1 \cdot 2\pi) = \\
 &= \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{(j \cdot 2\pi)^N \cdot (N-1)! \cdot e^{-\frac{s-j \cdot \nu_1 \cdot 2\pi}{2}} \cdot \sin\left(\frac{s-j \cdot (N-1+2 \cdot \nu_1) \cdot \pi}{2 \cdot j}\right)}{\pi \cdot \left(\frac{2 \cdot N-2}{N-1}\right)^{\frac{1}{2}} \cdot \prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} (s - j \cdot (\nu_2 + \nu_1) \cdot 2\pi)} \cdot \\
 &\quad \cdot \frac{(j \cdot 2\pi)^N \cdot (N-1)! \cdot e^{\frac{s-j \cdot \nu_1 \cdot 2\pi}{2}} \cdot \sin\left(\frac{-s-j \cdot (N-1-2 \cdot \nu_1) \cdot \pi}{2 \cdot j}\right)}{\pi \cdot \left(\frac{2 \cdot N-2}{N-1}\right)^{\frac{1}{2}} \cdot \prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} (-s - j \cdot (\nu_2 - \nu_1) \cdot 2\pi)} = \\
 &= \frac{(j \cdot 2\pi)^{2 \cdot N} \cdot (N-1)!^2}{\pi^2 \cdot \left(\frac{2 \cdot N-2}{N-1}\right)} \cdot \sin\left(\frac{s}{2 \cdot j}\right)^2 \cdot \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \frac{1}{\prod_{\nu_2=\frac{1-N}{2}}^{\frac{N-1}{2}} (s - j \cdot (\nu_2 + \nu_1) \cdot 2\pi)^2}
 \end{aligned} \tag{10.38}$$

berechnen. Setzt man $s = j \cdot \nu \cdot 2\pi$ mit $\nu = 0$ (1) $N - 1$ ein, so erhält man für das Betrags-

quadrat der gesuchten Spektralwerte der Fensterfunktion den Ausdruck

$$|F_\infty(2\pi \cdot \nu)|^2 = \frac{\sum_{\nu_1=0}^{N-1-|\nu|} \binom{N-1}{\nu_1}^2}{\binom{2 \cdot N-2}{N-1}}, \quad (10.39)$$

der für $\nu=0$ den für die Gültigkeit von Gleichung (10.23) notwendigen Wert Eins ergibt. In Gleichung (10.38) erhält man durch Erweiterung auf den Hauptnenner

$$\begin{aligned} F_\infty(-j \cdot s) \cdot F_\infty(j \cdot s) &= D_\infty(s) = \\ &= \underbrace{-\frac{(2\pi)^{2 \cdot N} \cdot (N-1)!^2 \cdot \sin\left(\frac{s}{2j}\right)^2}{\pi^2 \cdot \binom{2 \cdot N-2}{N-1} \prod_{\nu_3=1-N}^{N-1} (s - j \cdot \nu_3 \cdot 2\pi)^2}}_{= D_{E,\infty}(s)} \cdot \underbrace{(-1)^{N-1} \cdot \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} (s - j \cdot \nu_2 \cdot 2\pi)^2}_{= D_{\bar{E},\infty}(s)}, \end{aligned} \quad (10.40)$$

wobei der Laufindex ν_2 des in der Summe stehenden Produkts jeweils die $N-1$ Werte annimmt, die das Produkt entstehen lassen, das zur Erweiterung auf den Hauptnenner benötigt wird. Beim Summanden mit dem Laufindex ν_1 nimmt ν_2 die Werte $1-N$ (1) $N-1$ ohne die Werte $\nu_1 + (1-N)/2$ (1) $\nu_1 + (N-1)/2$ an.

Bleibt noch die Phase des Spektrums der Fensterfunktion zu berechnen. Der Anteil $D_{E,\infty}(s)$ ist eine Funktion, deren Nullstellen alle doppelt sind und alle auf der imaginären Achse bei $\omega = 2\pi \cdot \nu$ liegen, wobei ν eine ganze Zahl außer $\nu = 1-N$ (1) $N-1$ ist. Jede dieser doppelten Nullstellen tritt als einfache Nullstelle im minimalphasigen Anteil auf. Die Phase des Anteils der einfachen Nullstellen ist für $|\omega| < N \cdot 2\pi$ null. Um zu erreichen, dass der minimalphasige Anteil von $D_{E,\infty}(s)$ zu einer kausalen Zeitfunktion gehört, ist der Anteil der einfachen Nullstellen noch mit $e^{-\frac{s}{2}}$ zu multiplizieren (vgl. Gleichung (10.37)). Für die Frequenzen $\omega = 2\pi \cdot \nu$ mit $\nu = 1-N$ (1) $N-1$ erhalten wir daher die Phase $\nu \cdot \pi$.

$D_{\bar{E},\infty}(s)$ besitzt keine Nullstellen auf der imaginären Achse, weil niemals alle für $s = j\omega$ reellen Summanden zugleich null werden können, und alle Summanden das gleiche Vorzeichen aufweisen. Die Phase des Anteils von $D_{\bar{E},\infty}(s)$, dessen Nullstellen links der imaginären Achse liegen, wird über das Cepstrum berechnet. Es empfiehlt sich wieder vor der Berechnung des Cepstrums eine Bilineartransformation

$$s = c_\infty \cdot 2\pi \cdot \frac{\tilde{z}-1}{\tilde{z}+1}. \quad (10.41)$$

mit dem Transformationsparameter $c_\infty > 0$ durchzuführen. Nach $\tilde{z}$ aufgelöst erhält man:

$$\tilde{z} = \frac{c_\infty \cdot 2\pi + s}{c_\infty \cdot 2\pi - s} \quad (10.42)$$

Für die Frequenzen ω und $\tilde{\Omega}$ ergibt sich folgender Zusammenhang:

$$\tilde{\Omega} = 2 \cdot \arctan\left(c_{\infty} \cdot \frac{\omega}{2\pi}\right). \quad (10.43)$$

Für die aufzusplittende Summe der Produkte in Gleichung (10.40) erhält man dann durch Substitution von s

$$\begin{aligned} D_{\bar{E},\infty}(s) &= (-1)^{N-1} \cdot \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} (s - j \cdot 2\pi \cdot \nu_2)^2 = \tilde{D}_{\bar{E},\infty}(\tilde{z}) = \\ &= \underbrace{\left(\frac{\tilde{z}}{(1+\tilde{z})^2}\right)^{N-1}}_{=\tilde{D}_{P,\infty}(\tilde{z})} \cdot \underbrace{\sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} (K_{\infty,\nu_2} \cdot (\tilde{z} - \tilde{z}_{\infty,\nu_2}) \cdot (\tilde{z}^{-1} - \tilde{z}_{\infty,\nu_2}^*))}_{=\tilde{D}_{N,\infty}(\tilde{z})}, \end{aligned} \quad (10.44)$$

mit den Nullstellen $\tilde{z}_{\infty,\nu_2}$ die durch die Bilineartransformation aus den Nullstellen auf der imaginären Achse bei $j \cdot 2\pi \cdot \nu_2$ entstanden sind.

$$\tilde{z}_{\infty,\nu_2} = \frac{c_{\infty} + j \cdot \nu_2}{c_{\infty} - j \cdot \nu_2} = e^{j \cdot 2 \cdot \arctan\left(\frac{\nu_2}{c_{\infty}}\right)} \quad (10.45)$$

Für die Konstanten K_{∞,ν_2} ergeben sich die Werte

$$K_{\infty,\nu_2} = (2\pi)^2 \cdot (c_{\infty}^2 + \nu_2^2). \quad (10.46)$$

Sie lassen sich mit einem relativen Fehler von etwa ε berechnen. Der Nullstellenwinkel $2 \cdot \arctan(\nu_2/c_{\infty})$ liegt im Intervall $[-\pi; \pi]$. Im gesamten Intervall liegt der absolute Fehler des Nullstellenwinkels maximal in der Größenordnung von ε .

Die Z-Transformierte $\tilde{D}_{P,\infty}(\tilde{z})$ hat eine $2 \cdot N - 2$ -fache Polstelle bei $\tilde{z} = -1$. Der Anteil $\tilde{D}_{P,\infty}(\tilde{z})$ wird nun zurücktransformiert. Man erhält bis auf einen konstanten Faktor den Term

$$D_{P,\infty}(s) = (s + c_{\infty} \cdot 2\pi)^{N-1} \cdot (s - c_{\infty} \cdot 2\pi)^{N-1} \quad (10.47)$$

Dieser Term hat ein zur imaginären Achse symmetrisches Nullstellenpaar der Vielfachheit $N-1$ bei $s = \pm c_{\infty} \cdot 2\pi$. Jeder Polynomfaktor $s + c_{\infty} \cdot 2\pi$ einer einfachen Nullstelle links der imaginären Achse liefert den Phasenbeitrag

$$- \arctan\left(c_{\infty} \cdot \frac{\omega}{2\pi}\right), \quad (10.48)$$

der bei der Berechnung der Phase von $F_{\infty}(\nu \cdot 2\pi)$ entsprechend mit der Vielfachheit $N-1$ zu addieren ist.

Nun müssen wir noch die Phase des minimalphasigen Anteils des Summenanteils $\tilde{D}_{N,\infty}(\tilde{z})$ in Gleichung (10.44) berechnen. Diese Summe ist für $\tilde{z} = e^{j\tilde{\Omega}}$ positiv reell und lässt sich

mit $\tilde{\Omega} = \eta \cdot 2\pi / \tilde{M}$ für alle ganzzahligen Werte η in der Form

$$\begin{aligned} \tilde{D}_{N,\infty}(e^{j\tilde{\Omega}}) &= \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(4 \cdot K_{\infty,\nu_2} \cdot \sin\left(\frac{\tilde{\Omega}}{2} - \arctan\left(\frac{\nu_2}{c_\infty}\right)\right) \right)^2 = \\ &= \sum_{\nu_1=\frac{1-N}{2}}^{\frac{N-1}{2}} \prod_{\nu_2} \left(4 \cdot K_{\infty,\nu_2} \cdot \sin\left(\frac{\pi}{\tilde{M}} \cdot \left(\eta - \frac{\tilde{M}}{\pi} \cdot \arctan\left(\frac{\nu_2}{c_\infty}\right)\right)\right) \right)^2 \end{aligned} \quad (10.49)$$

berechnen. Verwendet man die $\tilde{M}$ ganzzahligen η -Werte

$$\tilde{M} \cdot \left(\frac{1}{\pi} \cdot \arctan\left(\frac{\nu_2}{c_\infty}\right) - \frac{1}{2} \right) < \eta \leq \tilde{M} \cdot \left(\frac{1}{\pi} \cdot \arctan\left(\frac{\nu_2}{c_\infty}\right) + \frac{1}{2} \right) \quad \text{mit } \eta \in \mathbb{Z}, \quad (10.50)$$

so ist das Argument der Sinusfunktion betragslich kleiner als $\pi/2$, und die Fehler bei der Berechnung des Quadrats der Sinusfunktion werden klein gehalten. Da die η -Werte ganze Zahlen sind, die am Rechner exakt darstellbar sind, lässt sich die Differenz aus η und dem durch die endliche Wortlänge quantisierten und normierten Winkel $\tilde{M}/\pi \cdot \arctan(\nu_2/c_\infty)$ mit einem relativen Fehler von ε berechnen. Für diese Werte von η kann daher auch das Quadrat der Sinusfunktion mit derselben relativen Genauigkeit berechnet werden. In der Praxis hat sich gezeigt, dass es auch hier nicht notwendig ist, die Werte η in der eben dargelegten Art um ganzzahlige Vielfache von $\tilde{M}$ zu reduzieren, da die ohne die Reduktion berechneten Fourierreihenkoeffizienten der Fensterfunktion sich nur unwesentlich von den Werten unterscheiden, die man bei Berücksichtigung von Gleichung (10.50) erhält. Die Faktoren K_{∞,ν_2} sind positiv und können ebenfalls mit der gewünschten relativen Genauigkeit berechnet werden. Alle Summanden in Gleichung (10.49) sind nichtnegativ und weisen relative Fehler in der Größenordnung von ε auf. Da bei keiner Frequenz $\tilde{\Omega}$ alle Summanden gleichzeitig null werden ist die Überlagerung aller ν_1 Anteile — also das Betragsquadrat des Spektrums, dessen Phase berechnet werden soll — stets echt positiv und besitzt die geforderte relative Genauigkeit, die es ermöglicht den natürlichen Logarithmus des Betragsquadratspektrums mit der notwendigen absoluten Genauigkeit zu berechnen. Bei der Akkumulation der Summanden in Gleichung (10.49) wird auf den konstanten Faktor $4 \cdot \max(K_{\infty,\nu_2}) = 16 \cdot \pi^2 \cdot (c_\infty^2 + (N-1)^2)$ normiert, um sicherzustellen, dass der am Rechner darstellbare Zahlenbereich nicht überschritten wird. Ein weiterer konstanter Faktor, wird so gewählt, dass das Maximum des Betragsquadratspektrums gleich dem Reziprokwert des Minimums ist, so dass der natürliche Logarithmus mit optimaler Genauigkeit berechnet wird. Eine inverse FFT des Logarithmus von $\tilde{D}_{N,\infty}(e^{j\tilde{\Omega}})$ liefert uns das doppelte, reelle und gradesymmetrische Cepstrum. Der durch die Quantisierungsfehler der FFT entstehende Imaginärteil wird durch eine Realteilbildung unterdrückt. Die gerade Symmetrie nützt man aus, um die Quantisierungsfehler der FFT im Realteil zu

verringern, indem man den Mittelwert des Cepstrums und des gespiegelten Cepstrums bildet. Damit das Cepstrum möglichst rasch abklingt muss der Bilineartransmutationsparameter c_∞ geeignet gewählt werden. c_∞ erhält man aus dem empirisch gewonnenen c nach Gleichung ([1]:6.30) indem man die Bilineartransformation nach Gleichung ([1]:6.21) mit der Näherung $z = e^{j\Omega} \approx 1 + j\Omega = 1 + j\omega/M/N = 1 + s/M/N$ für hinreichend großes M und kleines ω im Durchlassbereich des Spektrums der Fensterfunktion in die Bilineartransformation nach Gleichung (10.42) überführt. Damit ergibt sich:

$$c_\infty = \frac{1}{2\pi} \cdot \lim_{M \rightarrow \infty} c \cdot M \cdot N = \left(\frac{N}{2}\right)^{\frac{4}{3}}. \quad (10.51)$$

Mit $M \rightarrow \infty$ in Gleichung ([1]:6.31) erhält man für die nötige Länge $\widetilde{M}$ der inversen FFT den Mindestwert

$$\widetilde{M} > (\text{Mantissenwortlänge} - 1) \cdot 2^{\ln\left(\frac{N}{3}\right)}, \quad (10.52)$$

der wieder von der Mantissenwortlänge des zur Berechnung der Fensterfunktion verwendeten Rechners abhängt. Wieder wird als Länge der inversen FFT die kleinste Zweierpotenz gewählt, die die Ungleichung (10.52) erfüllt.

Die Werte des Cepstrums für $k > 0$ sind die Sinusreihenoeffizienten der in $\widetilde{\Omega}$ periodischen schiefsymmetrischen Phasenfunktion. Gesucht wird die Phase der Fourierreihenoeffizienten der Fensterfunktion, also die Phase von $F_\infty(\omega)$ bei den Frequenzen $\omega = \nu \cdot 2\pi$ mit $\nu = 1$ (1) $N-1$. Mit Gleichung (10.43) berechnet man die entsprechenden Frequenzen $\widetilde{\Omega}$. Sie sind kleiner als π und lassen sich mit einem relativen Fehler von circa ε berechnen. Nun werden die Werte der in $\widetilde{\Omega}$ periodischen, schiefsymmetrischen Phasenfunktion des minimalphasigen Anteils von $\widetilde{D}_{N,\infty}(\tilde{z})$ durch Auswertung der Sinusreihe bei diesen Frequenzen $\widetilde{\Omega}$ berechnet. Bei der Auswertung der Fouriersinusreihe wird die Reihenfolge der Summanden so gewählt, dass zunächst mit den kleinsten Summanden begonnen wird. Dieses Vorgehen ergibt eine höhere Genauigkeit des Ergebnisses.

Die anschließende Addition dieses Phasenanteils zu dem Phasenanteil des in Gleichung (10.44) vor die Summe gezogenen Terms $\widetilde{D}_{P,\infty}(\tilde{z})$ nach Gleichung (10.48) ergibt die gesuchte Phase von $F_\infty(\nu \cdot 2\pi)$. Der Betrag wurde als die positive Wurzel des Ausdrucks in Gleichung (10.39) berechnet. Damit sind die Fourierreihenoeffizienten der Fensterfunktion berechnet.

Es sei wieder angemerkt, dass es außer für den trivialen Fall mit $N=1$ (Rechteckfenster) auch für $N=2$ (1) 5 geschlossene Lösungen für $F_\infty(\nu \cdot 2\pi)$ gibt, deren Berechnung sich nur für $N=2$ lohnt. Man erhält:

$$\begin{aligned}
F_{\infty}(-2\pi) &= -(1-j)/2 \\
F_{\infty}(0) &= 1 \\
F_{\infty}(2\pi) &= -(1+j)/2.
\end{aligned} \tag{10.53}$$

In Kapitel 11.5 ist ein Hilfsprogramm abgedruckt, das die Werte der Fensterfunktion $f_{\infty}(t)$ für beliebiges t nach Gleichung (10.25) berechnet. Bei der Berechnung von $f_{\infty}(t)$ bilden die Realteile $2 \cdot \Re\{F_{\infty}(\nu \cdot 2\pi)\}$ der berechneten N Werte die Koeffizienten einer Kosinusreihe, die Imaginärteile $2 \cdot \Im\{F_{\infty}(\nu \cdot 2\pi)\}$ bilden die einer Sinusreihe. Beide Reihen bilden zusammen eine zeitkontinuierliche, periodische und reelle Funktion in t . Setzt man die so berechnete Funktion für $t < 0$ und für $t \geq 1$ zu null, so erhält man den gesuchten Wert der kontinuierlichen Fensterfunktion zum Zeitpunkt t . Um ein genaueres Ergebnis zu erhalten wird die Reihenfolge der Summation der einzelnen Reihenglieder wieder vertauscht, da auch hier die ersten Koeffizienten der Reihen die größten sind.

Das Spektrum $F_{\infty}(\omega)$ kann mit dem Hilfsprogramm in Kapitel 11.6 für beliebiges ω hochgenau berechnet werden. Bei der Berechnung des Spektrums der Fensterfunktion für beliebige Frequenzen geht man analog zu dem Fall der diskreten Fensterfolge vor. Wieder berechnet man den Betrag und die Phase des Spektrums getrennt. Die Phase wird wieder aus dem Cepstrum, das mit dem eben genannten Algorithmus berechnet wird, durch Auswertung der Sinusreihe aber diesmal nicht bei den Frequenzen $\omega = \nu \cdot 2\pi$ mit $\nu = 1$ (1) $N-1$, sondern bei den Frequenzen, für die das Spektrum berechnet werden soll, gewonnen. Der Betrag wird durch vorzeichenrichtiges Radizieren des Betragsquadrats gewonnen. Da $F_{\infty}(-j \cdot s)$ auf der imaginären Achse nur einfache Nullstellen im äquidistanten Raster 2π aufweist, ist das Vorzeichen bei der Radizierung auf einfache Weise aus ω bestimmbar.

Das Betragsquadrat von $F_{\infty}(\omega)$ wird durch Überlagerung der verschobenen Betragsquadrate von $G_{\infty}(j\omega)$ berechnet. Dabei ist jeweils die Polstelle von $G_{\infty}(s - j \cdot \nu_1 \cdot 2\pi)$, die dem Punkt $s = j\omega$, für den das Spektrum berechnet werden soll, am nächsten liegt, mit der Sinusfunktion im Zähler von $G_{\infty}(j\omega - j \cdot \nu_1 \cdot 2\pi)$ zu einer in der Frequenz verschobenen si-Funktion in ω zusammenzufassen, auszuwerten, zu quadrieren und mit dem Betragsquadrat des Produktes der Abstände zu den restlichen Polstellen zu dividieren. Da auf diese Weise eine Summe von positiven Produkten, deren Faktoren mit höchstmöglicher relativer Genauigkeit berechnet worden sind, gebildet wird, ist so das Spektrum auch für den Sperrbereich bis zu sehr hohen Frequenzen akkurat berechenbar. Mit dem in Kapitel 11.6 auszugsweise angegebenen Programm kann die Fouriertransformierte bis zu einer Sperrdämpfung von etwa $-10 \cdot \log_{10}(\text{realmin})$ berechnet werden.

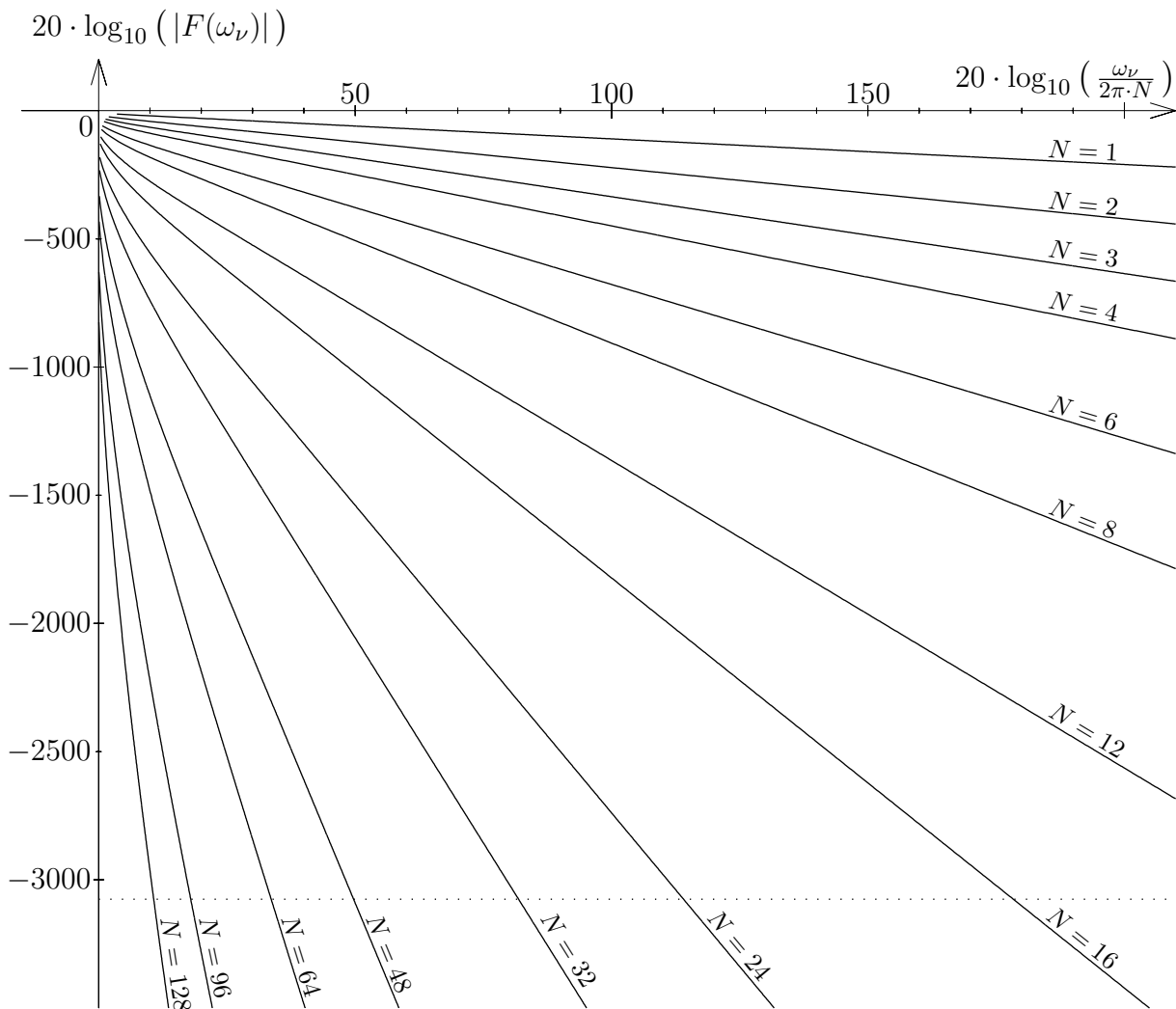


Bild 10.9: Sperrdämpfung der Nebenmaxima der Spektren der kontinuierlichen Fensterfunktionen bei den Frequenzen $\omega_\nu = (2 \cdot \nu + 1) \cdot \pi$ mit N als Parameter.

In Bild 10.9 sind die Nebenmaxima der Beträge der Spektren der kontinuierlichen Fensterfunktionen bei den Frequenzen $\omega_\nu = (2 \cdot \nu + 1) \cdot \pi$ mit ganzzahligem $\nu > N$ doppelt logarithmisch über der auf $2\pi \cdot N$ normierten Frequenz ω aufgetragen. Für jeden der Werte 1, 2, 3, 4, 6, 8, 12, 16, 24, 32, 48, 64, 96 und 128 des Fensterlängenfaktors N wurde eine Kurve eingetragen, wobei einerseits für große Frequenzen nicht alle Nebenmaxima berechnet wurden (es wären einfach zu viele), und andererseits zur besseren Darstellbarkeit die bei den Nebenmaxima berechneten Beträge der Spektralwerte durch gerade Linien miteinander verbunden wurden. Bei dieser Graphik mit der etwas ungewöhnlichen Skalierung der Ordinate sollen nun drei Dinge gezeigt werden. Zum ersten ist zu sehen, dass der asymptotische Anstieg der Sperrdämpfung mit der Potenz N erfolgt. Im Bild 10.9 wurde die Ordinate um den Faktor 20 enger skaliert als die Abszisse und die Beschriftung der

Kurven wurde mit der Neigung eingetragen, die der Potenz ω^{-N} entspricht. Zum zweiten sieht man, dass sich das Quadrat des Betragsfrequenzgangs wirklich bis in die Größenordnung der kleinsten positiven Zahl `realmin`, die am Rechner mit voller Genauigkeit darstellbar ist, berechnen lässt, wenn man diesen über die Summe der verschobenen Betragsquadratspektren berechnet, wobei man deren Summanden jeweils als Produkte der Betragsquadrate der Abstände zu den Nullstellen berechnet. Zum Vergleich wurde der Wert $\sqrt{\text{realmin}}$ als punktierte Linie eingetragen. Drittens wird demonstriert, dass auch extrem große Werte von N kein Problem für den verwendeten Berechnungsalgorithmus darstellen. Lediglich die Dauer der Berechnung nimmt dann — grob geschätzt etwa quadratisch mit N — zu. Bei $N=128$ wurde eine FFT-Länge von $\widetilde{M}=1024$ benötigt.

Mit dem Hilfsprogramm in Kapitel 11.7 kann man die Autokorrelationsfunktion $d_\infty(t)$ des kontinuierlichen Fensters für beliebiges t zumindest theoretisch exakt berechnen. Bei einem kontinuierlichen Fenster kann man das Faltungsintegral, das die Fensterautokorrelationsfunktion definiert, nur näherungsweise und mit einem erheblichen Rechenaufwand berechnen, wenn man hohe Ansprüche an die Genauigkeit der Berechnung stellt. Auch die bei einer diskreten Fensterfolge immer bestehende Möglichkeit, die Fensterautokorrelationsfolge aus endlich vielen Abtastwerten des Betragsquadrats des Fensterspektrums mit einer DFT zu berechnen, ist hier in der Regel nicht gegeben, da wegen der zeitlichen Begrenzung der Autokorrelationsfunktion des Fensters eine unbegrenzte Anzahl spektraler Abtastwerte benötigt werden würde. Mit den Sinusreihenkoeffizienten

$$S_\nu = \sum_{\substack{\tilde{\nu}=1-N \\ \tilde{\nu} \neq \nu}}^{N-1} \frac{\Re\{F_\infty(\nu \cdot 2\pi)^* \cdot F_\infty(\tilde{\nu} \cdot 2\pi)\}}{\pi \cdot (\tilde{\nu} - \nu)} \quad (10.54)$$

kann man jedoch die Fensterautokorrelationsfunktion für $0 \leq t < 1$ als

$$d_\infty(t) = \sum_{\nu=1-N}^{N-1} S_\nu \cdot \sin(2\pi \cdot \nu \cdot t) + (1-t) \cdot \sum_{\nu=1-N}^{N-1} |F_\infty(\nu \cdot 2\pi)|^2 \cdot \cos(2\pi \cdot \nu \cdot t) \quad (10.55)$$

über die bereits bestimmten Abtastwerte des Spektrums der Fensterfunktion berechnen. Die gerade Symmetrie der Fensterautokorrelationsfunktion liefert die Funktionswerte für negatives t . Diese Art der Berechnung der Fensterautokorrelationsfunktion ist übrigens bei allen Fensterfunktionen möglich, die sich als eine Periode einer endlichen Fourierreihe darstellen lassen, wobei die Grenzen der Summationsindizes so zu modifizieren sind, dass alle Fourierreihenkoeffizienten berücksichtigt werden. Bei dem in Kapitel 11.7 enthaltenen Programmauszug sind die Summationsgrenzen so verändert, dass auch für die hier nicht beschriebene freie Wahl weiterer Nullstellen der Laplacetransformierten der Basisfensterfunktion die Berechnung der Fensterautokorrelationsfunktion möglich ist.

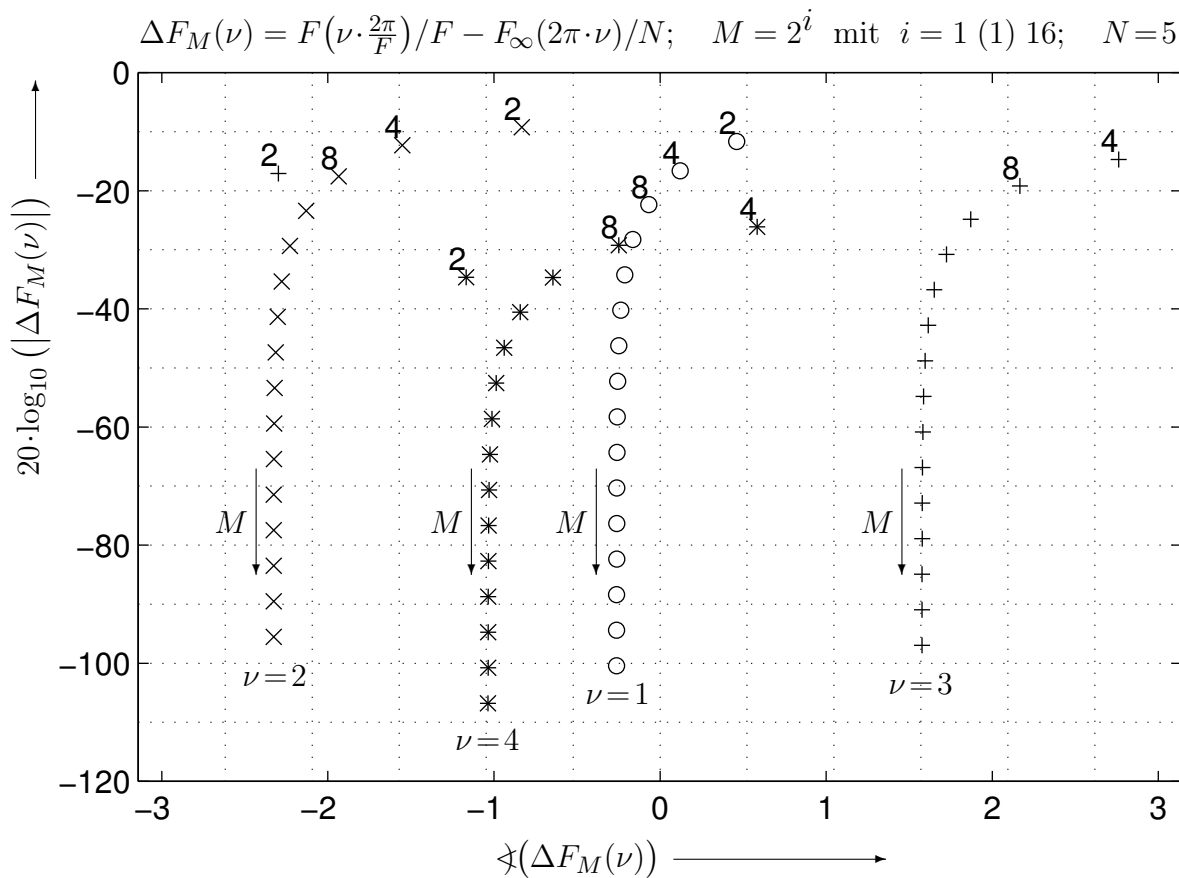


Bild 10.10: Grenzwerte der Fourierreihenkoeffizienten der diskreten Fensterfolgen für $M \rightarrow \infty$ im Vergleich zu den Fourierreihenkoeffizienten der kontinuierlichen Fensterfunktion.

10.8 Die kontinuierlichen Fensterfunktion als Grenzwertlösung

Am Anfang des letzten Unterkapitels wurde gesagt, dass sich die auf N normierten Fourierreihenkoeffizienten der kontinuierlichen Fensterfunktion als die Grenzwerte der Fourierreihenkoeffizienten der diskreten Fensterfolgen nach Kapitel [1]:6 für $M \rightarrow \infty$ ergeben. Um dies zu demonstrieren, sind in Bild 10.10 am Beispiel der Fensterfolgen, die man mit $N=5$ bei dem Algorithmus nach Kapitel [1]:6 erhält, die Abweichungen der Fourierreihenkoeffizienten $F(\nu \cdot \frac{2\pi}{F})/F$ von den auf N normierten Fourierreihenkoeffizienten $F_\infty(2\pi \cdot \nu)$ der kontinuierlichen Fensterfunktion graphisch dargestellt. Dabei wurde diese komplexe Differenz nicht in der komplexen Ebene eingetragen, sondern es wurde der Logarithmus des Betrages der Differenz jeweils über dem Winkel der Differenz aufgetragen. Dadurch kann man auch bei großen Werten von M die dann sehr kleine Differenz noch geeignet darstellen. Für die Werte M wurden alle Zweierpotenzen von 2 bis 2^{16} ausgewählt.

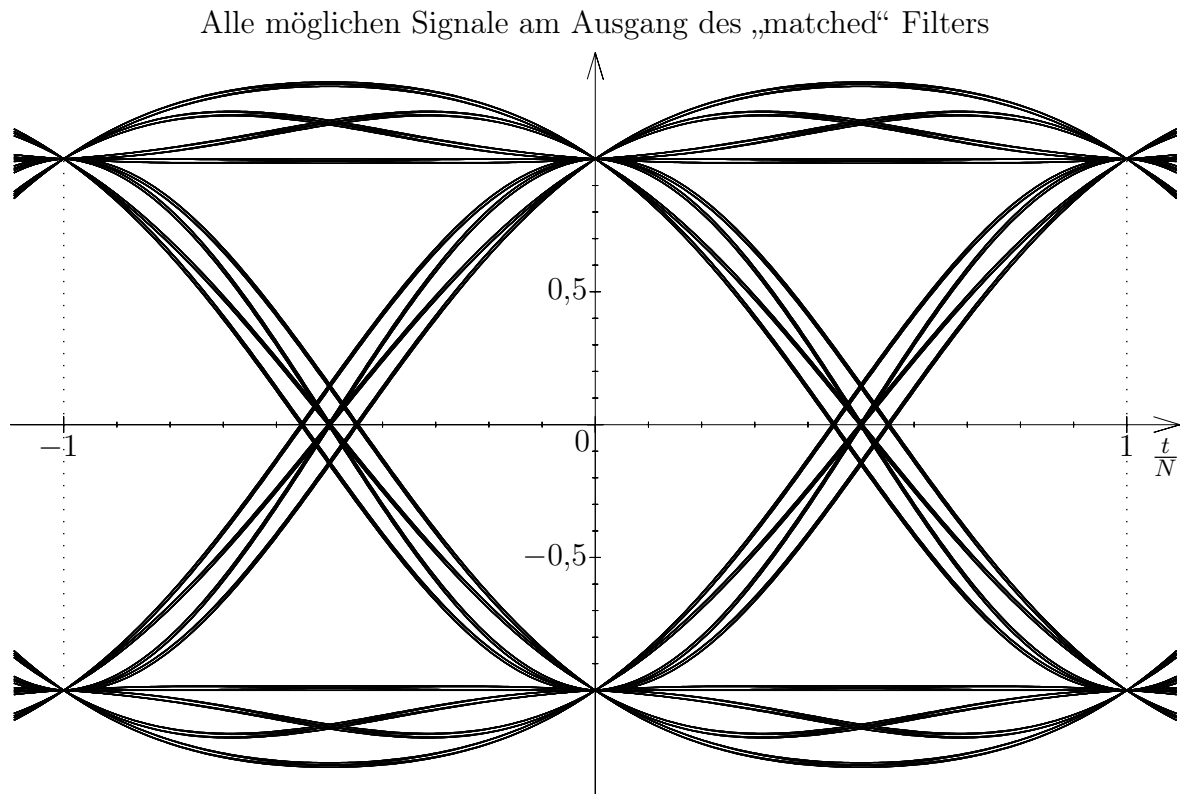


Bild 10.11: Augendiagramm einer ungestörten 2ASK-PAM bei Verwendung der kontinuierlichen Fensterfunktion mit $N=4$ als Sendeimpuls und als „matched“ Filter.

Der Fourierreihenkoeffizient des Gleichanteils wurde nicht eingetragen, da dieser immer mit dem auf N normierten Fourierreihenkoeffizienten der kontinuierlichen Fensterfunktion übereinstimmt. Einerseits erkennt man, dass die Abweichungen mit steigendem M betragsmäßig abnehmen, während der Winkel für große Werte von M etwa konstant bleibt. Andererseits kann man ablesen, dass selbst bei dem relativ großen Wert von $M = 2^{16}$ noch Abweichungen in der fünften Stelle hinter dem Komma auftreten.

10.9 Augendiagramm des AKF der kontinuierlichen Fensterfunktion

Im Anschluss an Gleichung (10.23) wurde festgestellt, dass die Fensterautokorrelationsfunktion $d_{\infty}(t) = f_{\infty}(t) * f_{\infty}(-t)$ äquidistante Nullstellen bei Vielfachen von $1/N$ außer bei $t=0$ aufweist. Sie erfüllt also die erste Nyquist-Bedingung. Bei der Fensterfunktion $f_{\infty}(t)$ handelt es sich somit um einen sog. Wurzel-Nyquist-Impuls. Verwendet man daher die Fensterfunktion als Sendeimpuls bei einer digitalen Übertragung und filtert man das

Empfangssignal mit einem „matched“-Filter, dessen Impulsantwort $f_{\infty}(-t)/N$ genau die gespiegelte und auf N normierte Fensterfunktion ist, so erhält man bei einem verzerrungsfreien ungestörten Kanal am Ausgang des Empfangsfilters die mit der auf N normierten Fensterautokorrelationsfunktion gefalteten Sendesymbole. Bei geeigneter Abtastung kann man dann die digitalen Sendesymbole ohne Intersymbolinterferenzen wiedergewinnen. Um beurteilen zu können, wie gut die Fensterfunktion für solch eine Übertragung geeignet ist, ist in Bild 10.11 das Augendiagramm für eine 2ASK-PAM (**P**uls-**A**mplituden-**M**odulation mit binärer Amplitudenumtastung = **A**mplitude **S**hift **K**eying) aufgetragen. Man erkennt, dass aufgrund fehlender Intersymbolinterferenzen die maximal mögliche vertikale Augenöffnung erreicht werden kann, und dass sich auch eine ganz brauchbare horizontale Augenöffnung ergibt. Möglicherweise kann man die Freiheitsgrade bei der Wahl der Basisfensterfunktion bei der Konstruktion der Fensterfunktion nutzen, um weitere günstige Eigenschaften für die Anwendung als Sendeimpuls bei einer digitalen Übertragung zu erzielen. Die Anwendung der Fensterfunktion als Sendeimpuls wurde im Rahmen dieser Arbeit jedoch nicht weiter untersucht.

Beweis der Formel im Anschluss an das Inhaltsverzeichnis:

Die Nullstellen des Polynoms $z^M - 1$ liegen alle auf dem Einheitskreis im Raster $\frac{2\pi}{M}$. Somit liefert eine Polynomdivision mit $z - 1$:

$$\prod_{\rho=1}^{M-1} \left(z - e^{j \cdot \frac{2\pi}{M} \cdot \rho} \right) = \frac{z^M - 1}{z - 1} = \sum_{\nu=0}^{M-1} z^\nu.$$

Setzt man $z=1$ ein und bildet man auf beiden Seiten der Gleichung den Absolutbetrag, so erhält man:

$$\left| \prod_{\rho=1}^{M-1} \left(1 - e^{j \cdot \frac{2\pi}{M} \cdot \rho} \right) \right| = \prod_{\rho=1}^{M-1} \left| 1 - e^{j \cdot \frac{2\pi}{M} \cdot \rho} \right| = \sum_{\nu=0}^{M-1} 1^\nu = M$$

Mit der Eulerschen Formel können die Beträge der Faktoren des Produkts berechnet werden:

$$\left| 1 - e^{j \cdot \frac{2\pi}{M} \cdot \rho} \right| = \left| e^{j \cdot \frac{\pi}{M} \cdot \rho} - e^{-j \cdot \frac{\pi}{M} \cdot \rho} \right| = 2 \cdot \sin \left(\frac{\pi}{M} \cdot \rho \right).$$

Setzt man dies in die vorletzte Gleichung ein und dividiert man auf beiden Seiten mit 2^{M-1} , so erhält man die gesuchte Formel.

11 MATLAB-Programmauszüge

Hier werden nun die entscheidenden Zeilen der Programme in der Interpretersprache MATLAB angegeben, die für die Berechnung

- der Fourierreihenkoeffizienten der diskreten Fensterfolgen nach Kapitel 10.1 sowie der Werte dieser Fensterfolgen für beliebige Zeitpunkte,
- der Spektren dieser Fensterfolgen für beliebige Frequenzen,
- der Autokorrelationsfolgen dieser Fensterfolgen für beliebige Zeitpunkte,
- der Fourierreihenkoeffizienten der kontinuierlichen Fensterfunktionen nach Kapitel 10.7,
- der Werte dieser Fensterfunktionen für beliebige Zeitpunkte,
- der Spektren dieser Fensterfunktionen für beliebige Frequenzen und
- der Autokorrelationsfunktionen dieser Fensterfunktionen für beliebige Zeitpunkte

benötigt werden. Dabei wird darauf verzichtet, die Teile des Programms abzudrucken, die nicht für die eigentliche Berechnung der Fenster benötigt werden, die aber bei einem guten Programm immer vorhanden sein sollten, wie zum Beispiel eine Überprüfung der Eingabeparameter oder eine adäquate Behandlung von Spezial- und Ausnahmefällen. Zunächst werden jeweils die Programmzeilen in **Schreibmaschinenschrift** aufgelistet, wobei diese durchnummeriert sind, um im folgenden Kommentar auf die Zeilen Bezug nehmen zu können. Es werden dabei jeweils nur die Programmzeilen ausführlich kommentiert, die nicht schon bei einer der vorangegangenen Programmvarianten oder in Kapitel [1]:A.8 kommentiert worden sind. Bei vielen dieser Programme wird vorausgesetzt, dass die MATLAB-interne globale Variable **eps** vorhanden ist, welche die relative Rechnergenauigkeit ε enthält, und die nicht explizit an das jeweilige Programm übergeben werden muss. Gleiches gilt für die Variable **pi**, die — wie der Name schon sagt — π enthält. Die Programmauszüge sind nur stellenweise für eine besonders schnelle Berechnung der Fensterfolge optimiert. Die in [1] durchgeführten Betrachtungen über die Genauigkeit der Berechnung sind hier alle berücksichtigt. Auch wurde weitgehend versucht, bei der Art der Berechnung die in Kapitel 10 beschriebenen Vorgehensweisen beizubehalten, so dass diese Programme auch dazu dienen sollen, dem Leser den konkreten Ablauf zu zeigen, und zu demonstrieren, dass die dort für die theoretische Herleitung angegebenen Formeln nur unwesentlich verändert übernommen werden können.

11.1 Die diskreten Fensterfolgen

Neben dem Fensterlängenfaktor N und dem Abstand M der Nullstellen der Fensterautokorrelationsfolge, werden die Anzahl A_1 der zusätzlichen Nullstellen der Z-Transformierten $G(z)$ der Basisfensterfolge $g(k)$ am Einheitskreis im Raster $2\pi/F$, die Anzahl A_0 der abschließenden Nullwerte der Fensterfolge, sowie die frei wählbaren Nullstellen $z_{0,\rho}$ von $G(z)$ benötigt. Die ganzzahligen Parameter N , M , A_1 und A_0 müssen als skalare Größen (also als 1×1 Matrizen) `N`, `M`, `A_1` und `A_0` beim Programmaufruf angegeben werden. Für diese Parameter müssen die Bedingungen $N > 1$, $M > 1$, $0 \leq A_1 < N/2$ und $0 \leq A_0 < N - 1 - 2 \cdot A_1$ erfüllt sein. Die Nullstellen $z_{0,\rho}$ bilden die Elemente des Vektors `z_0`, der ebenfalls beim Programmaufruf zu übergeben ist. Dieser Vektor muss genau $N - 1 - A_0 - 2 \cdot A_1$ Elemente enthalten. Nullstellen, die nicht reell sind, müssen als zueinander konjugiert komplexe Paare in `z_0` vorhanden sein. Um eine möglichst gute Genauigkeit zu erzielen, sollten Nullstellen außerhalb des Einheitskreises durch die am Einheitskreis gespiegelten ersetzt worden sein, und die Nullstellen sollten nach aufsteigendem Abstand zum Punkt $z = 1$ sortiert sein. Als Ergebnis werden von diesem Programm die Werte der Fensterfolge auf dem Vektor `f_k` sowie die Fourierreihenkoeffizienten für $\nu = 0$ (1) $N - A_1 - 1$ auf dem Vektor `F_nu` zurückgegeben.

```

1: function [ f_k, F_nu ] = fenster( N, M, A_0, A_1, z_0 )
2: F = N * M
3: c = 2 / ( 1 + (N/2)^(M/3/(1-M)) / tan(pi/2/M) )
4: Ms = -log(eps)/log(2) * 2^log(N/3) * 3.6^(1/M)
5: Ms = 2^ceil(log(Ms)/log(2))
6: Ms = max(Ms,16)
7: Ms_OK = 0
8: while ~Ms_OK
9:     eta = 0:Ms/2
10:    F_eta = zeros(1,Ms/2+1)
11:    NF_1 = 1 / ( c^2 + 4 * (1-c) * sin( pi/F * (N-1-A_1) )^2 )
12:    NF_0 = max( [ ( abs( z_0*exp( j*pi/F*(1-N) ) - (1-c) ) + ...
                    abs( 1-(1-c)*z_0*exp( j*pi/F*(1-N) ) ) ) ; ...
                    ( abs( z_0*exp( j*pi/F*(N-1) ) - (1-c) ) + ...
                    abs( 1-(1-c)*z_0*exp( j*pi/F*(N-1) ) ) ) ] ).^(-2)
13:    for nu_1 = (1-N)/2:(N-1)/2
14:        F_eta_1 = ones(1,Ms/2+1)
15:        rho = 0

```

```

16:   for nu_2 = [(1-N+A_1):(nu_1+A_1-(N+1)/2), ...
               (nu_1-A_1+(N+1)/2):(N-1-A_1)]
17:     K_1 = NF_1 * ( c^2 + 4 * (1-c) * sin( pi/F * nu_2 )^2 )
18:     Psi_1 = atan( ( (1-c) * sin(2*pi/F*nu_2) ) / ...
                  ( c + 2 * (1-c) * sin(pi/F*nu_2)^2 ) )
19:     F_eta_1 = F_eta_1 .* K_1 .* sin( pi/Ms*eta-pi/F*nu_2-Psi_1 ).^2
20:     rho = rho + 1
21:     if rho < N-0.5-2*A_1-A_0
22:       Z_Z = z_0(rho) * exp( j*2*pi/F*nu_1 ) - (1-c)
23:       Z_N = 1 - (1-c) * z_0(rho) * exp(j*2*pi/F*nu_1)
24:       abs_Z_Z = abs( Z_Z )
25:       abs_Z_N = abs( Z_N )
26:       F_eta_1 = F_eta_1 .* NF_0(rho) .* ...
                ( (abs_Z_N-abs_Z_Z)^2 + 4 * abs_Z_N * abs_Z_Z * ...
                  sin( pi/Ms*eta-(angle(Z_Z)-angle(Z_N))/2 ).^2 )
27:     end
28:   end
29:   F_eta = F_eta + F_eta_1
30: end
31: NF = 1 / sqrt( max(F_eta) * min(F_eta) )
32: F_eta = NF * F_eta
33: L_eta = log(F_eta)
34: Ceps_2 = ifft( [L_eta,L_eta(Ms/2:-1:2)] )
35: Ceps_2 = real( Ceps_2 )
36: Ceps_2 = ( Ceps_2 + Ceps_2([1,Ms:-1:2]) ) / 2
37: sockel = eps / Ms / sqrt(48) * ...
          sqrt( sum( max( [L_eta,L_eta(Ms/2:-1:2)].^2, 1 ) ) )
38: grenze = ( 2 * (2*N-2-2*A_1-A_0) ) ./ [1:Ms/2] .* ...
          ( Ms * sockel/2/(2*N-2-2*A_1-A_0) ).^( 2*[1:Ms/2]/Ms )
39: max_fehl = sum( max( abs([L_eta,L_eta(Ms/2:-1:2)]), 1 ) ) * ...
            eps * ( 2 + log(Ms)/log(2) ) / Ms
40: grenze = max( grenze, max_fehl )
41: if any( abs(Ceps_2(2:Ms/2+1)) > grenze )
42:   kriterium = abs( fft( Ceps_2(Ms/4+1:Ms/2) .* [Ms/4:Ms/2-1] ) )
43:   [dummy,womax] = max( kriterium(1:Ms/8+1) )
44:   delta_c = 1-(womax-1)*16/Ms
45:   delta_c = sin( delta_c*pi*(1-delta_c^2/2) )/3
46:   c = c * (1-delta_c) / ( 1 + (1-c) * delta_c )
47:   Ms = 2*Ms
48: else
49:   Ms_OK = 1
50: end
51: end

```

```

52: nu = 1:N-A_1-1
53: Omega = 2*pi/F*nu
54: Omega_s = Omega + 2 * atan( ( (1-c) * sin(Omega) ) ./ ...
                                ( c + 2 * (1-c) * sin(Omega/2).^2 ) )
55: phi = zeros(1,N-A_1-1)
56: for nu_i = nu
57:   phi(nu_i) = sin(Omega_s(nu_i)*[Ms/2-1:-1:1]) * Ceps_2(Ms/2:-1:2).';
58: end
59: phi = phi - (N-1-A_1-A_0/2) * Omega_s + (F-1-A_0) * Omega/2
60: if N == 2*A_1+1
61:   F_nu = ones(1,N-A_1)/N
62: else
63:   nu = 0:N-A_1-1
64:   F_nu = zeros(1,N-A_1)
65:   NF_1 = exp(2*log(8/pi*F)-sum(log([5:2:4*N-8*A_1]))/(N-1-2*A_1))
66:   NF_0 = 1 ./ max( [ abs( z_0 * exp(j*pi/F*(N-1)) - 1 ).^2 ; ...
                      abs( z_0 * exp(j*pi/F*(1-N)) - 1 ).^2 ] )
67:   for nu_1 = (1-N)/2:(N-1)/2
68:     F_nu_1 = ( ( nu < nu_1+N/2-A_1 ) & ( nu > nu_1-N/2+A_1 ) ) / eps^2
69:     rho = 0
70:     for nu_2 = (1-N)/2+A_1:(N-3)/2-A_1
71:       rho = rho + 1
72:       if rho < N-0.5-2*A_1-A_0
73:         abs_z_0 = abs(z_0(rho))
74:         F_nu_1 = F_nu_1 .* NF_0(rho) .* ( ( 1-abs_z_0 )^2 + ...
                                             4 * abs_z_0 * sin( pi/F*(nu-nu_1)-angle(z_0(rho))/2 ).^2 )
75:       end
76:       nu_3 = nu - nu_1 - nu_2 - ( nu <= nu_1 + nu_2 )
77:       F_nu_1 = F_nu_1 ./ ( NF_1 .* sin( pi/F*nu_3 ).^2 )
78:     end
79:     F_nu = F_nu + F_nu_1
80:   end
81:   F_nu = sqrt(F_nu)
82:   F_nu = F_nu / ( N * F_nu(1) )
83: end
84: F_nu(2:N-A_1) = F_nu(2:N-A_1) .* exp(-j*phi)

```

```

85: f_k = zeros(1,F)
86: si_k = [ 0:2:(F+0.5)/2,...
            F-[2*floor((F+0.5)/4)+2:2:(3*F+0.5)/2], ...
            [2*floor((3*F+0.5)/4)+2:2:2*F-1]-2*F      ]
87: si_k = sin( (pi/F) * si_k )
88: co_k = [ F-[0:4:2*F-1], [4*floor((2*F-1)/4)+4:4:4*F-2]-3*F ]
89: co_k = sin( (pi/(2*F)) * co_k )
90: for nu = N-A_1-1:-1:1
91:     k_nu = rem( [0:F-1]*nu, F ) + 1
92:     f_k = f_k + 2 * real( F_nu(nu+1) ) * co_k(k_nu) - ...
                2 * imag( F_nu(nu+1) ) * si_k(k_nu)
93: end
94: f_k = f_k + 1/N

```

Durch die bei diesem Programm freigestellte Wahl der nicht festgelegten Nullstellen der Z-Transformierten der Basisfensterfolge ergeben sich zwei wesentliche Modifikationen gegenüber dem im Anhang von [1] angegebenen Programm. Zum einen kann nun die für die Berechnung des Cepstrums benötigte FFT-Länge $\widetilde{M}$ genauso wie der optimale Wert des Bilineartransformationsparameters c nicht zu Beginn festgelegt werden. Die hier notwendige iterative Bestimmung dieser Parameter wird im Anschluss an die zweite wesentliche Modifikation kommentiert. Diese besteht darin, dass sich der Betrag des Spektrums sowohl bei der Berechnung der Beträge der Fourierreihenoeffizienten der Fensterfolge $f(k)$, als auch bei der Berechnung des Betragsquadrats der bilinear Z-Transformierten bei der Berechnung des Cepstrums nun umfangreicher berechnet.

In der **for**-Schleife zur Berechnung der kumulativen Produkte in Gleichung (10.12), die in Zeile 16 beginnt und in Zeile 28 endet, sind die Elemente des Vektors **F_eta_1**, mit deren Hilfe die kumulativen Produkte für alle Frequenzpunkte $\widetilde{\Omega} = \eta \cdot 2\pi / \widetilde{M}$ mit $\eta = 0 (1) \widetilde{M}/2$ nach und nach berechnet werden, noch mit den Betragsquadraten der Abstände zu den neu hinzugekommenen Nullstellen zu multiplizieren. In Gleichung (10.12) ist dies das kumulative Produkt mit dem Laufindex ρ . Die Berechnung der Nullstellenabstände erfolgt abwechselnd, d. h. es wird immer zuerst eine der Nullstellen verarbeitet, die zur Erweiterung auf den Hauptnenner erforderlich war, und dann eine der Nullstellen des Vektors **z_0**. Dazu wird in Zeile 15 **rho** zunächst auf null gesetzt, um dann bei jedem Schleifendurchlauf mit einem neuen Wert von **nu_2** in Zeile 20 um eins erhöht zu werden. In Zeile 21 wird festgestellt, ob schon alle Nullstellen des Vektors **z_0** verarbeitet wurden. Wenn nicht, werden in den Zeilen 22 bis 26 die mit einem geeignet gewählten Normierungsfaktor **NF_0(rho)** multiplizierten Betragsquadrate der Abstände der nächsten um $e^{j \cdot \nu_1 2\pi / F}$ rotierten Nullstelle zu den Punkten $\widetilde{z} = e^{j \cdot \eta \cdot 2\pi / \widetilde{M}}$ für alle η -Werte zugleich berechnet, wie dies in Gleichung (10.12) angegeben ist. Da zur Erweiterung auf den Hauptnenner genau $N - 1$

Nullstellen benötigt werden, wird die Schleife mit dem Index `nu_2`, der alle im Anschluss an Gleichung (10.9) genannten Werte annehmen kann, genau $N-1$ mal durchlaufen. Da der Vektor `z_0` genau $N-1-A_0-2\cdot A_1$ Nullstellen enthält, kann es nicht vorkommen, dass die Schleife beendet wird, bevor nicht alle Nullstellen des Vektors `z_0` verarbeitet worden sind. Bei dieser Programmvariante wurde für jede der frei wählbaren Nullstellen genau ein Normierungsfaktor neu eingeführt. Dies sind die Elemente des Vektors `NF_0`, der in Zeile 12 berechnet wird. Da die Normierungsfaktoren bei allen Frequenzen $\tilde{\Omega} = \eta \cdot 2\pi / \tilde{M}$ und bei allen Summanden der Summe über alle verschobenen Betragsquadratspektren der Basisfensterfolge mit dem Index `nu_1` dieselben sind, entspricht diese Art der Normierung lediglich einer Skalierung aller Spektralwerte des Vektors `F_eta` mit einer gemeinsamen Konstante, nämlich mit dem Produkt aller Elemente des Vektors `NF_0`. Wenn man sich in Gleichung (10.12) in der von $\tilde{\Omega}$ abhängigen Form einen Faktor des kumulativen Produkts mit dem Index ρ betrachtet, so stellt man fest, dass der Faktor den Maximalwert $(|\tilde{z}_{N,\rho,\nu_1}| + |\tilde{z}_{Z,\rho,\nu_1}|)^2$ nicht übersteigt. Für alle möglichen Werte von ν_1 liegen die Nullstellen, die durch Rotation mit dem Drehfaktor $e^{j\cdot\nu_1 2\pi/F}$ aus der Nullstelle $z_{0,\rho}$ hervorgehen, entweder auf dem Einheitskreis, oder auf einem Kreis, der innerhalb des Einheitskreises liegt. Außerhalb des Einheitskreises können diese Nullstellen nicht liegen, da wir vorausgesetzt haben, dass die Nullstellen $z_{0,\rho}$ durch die am Einheitskreis gespiegelten Nullstellen zu ersetzen sind, wenn sie sich zuvor außerhalb des Einheitskreises befinden. Auch nach der Bilineartransformation liegen die Nullstellen entweder auf dem Einheitskreis, oder auf einem Kreis innerhalb des Einheitskreises. Man kann sich nun überlegen, dass in den Fällen, bei denen die richtige Normierung besonders wichtig ist — in der Regel ist das der Fall, wenn M besonders groß ist —, die rotierten und bilineartransformierten Nullstellen mit dem betraglichen Maximalwert von ν_1 besonders nahe am Einheitskreis liegen. Daher wird in Zeile 12 für jede Nullstelle des Vektors `z_0` für die beiden Werte $\nu_1 = N-1$ und $\nu_1 = 1-N$ der maximal mögliche Wert des Faktors des kumulativen Produkts in Gleichung (10.12) bestimmt, und von diesen beiden Werten jeweils der größere genommen, um dessen Reziprokwert als Normierungsfaktor zu verwenden.

Bei dieser Programmvariante sind nun um A_1 weniger Fourierreihenoeffizienten der Fensterfolge $f(k)$ zu berechnen, als in der in [1] vorgestellten Variante. Deshalb hat der Vektor `nu`, der in Zeile 63 festgelegt wird, nun entsprechend weniger Elemente. Die Berechnung der Fourierreihenoeffizienten erfolgt nach Gleichung (10.9) in den Zeilen 60 bis 83. Hier ist nun eine Fallunterscheidung angebracht. Wenn alle frei wählbaren Nullstellen auf den Einheitskreis gelegt werden, so dass A_1 seinen Maximalwert $(N-1)/2$ annimmt, bleibt in Gleichung (10.9) von der Summe mit dem Laufindex ν_1 nur ein Summand mit dem Wert Eins übrig. Dieser Summand entspricht dem Summanden, bei dem die doppelte Polstelle bei $z = e^{j\cdot\nu\cdot 2\pi/F}$ in Gleichung (10.8) mit der doppelten Nullstelle des Zählers gekürzt wurde, und sich somit der Wert F^2 der quadrierten si-Funktion an der Stelle Null ergab, der in

Gleichung (10.9) vor der Summe steht. Daher ist, wenn dieser Fall eintritt, was in Zeile 60 abgefragt wird, der Betrag aller Fourierreihenkoeffizienten von $f(k)$ gleich $1/N$, wie dies in Zeile 61 eingesetzt wird. Anderenfalls werden die Beträge aller Fourierreihenkoeffizienten von $f(k)$ nach Gleichung (10.9) berechnet, wobei auch hier die Nullstellen und die Polstellen abwechselnd verarbeitet werden. Die Reihenfolge der Bearbeitung wird — wie bei der Berechnung des Cepstrums — mit Hilfe der Programmzeilen 69, 71, 72 und 75 realisiert. Warum diese Bearbeitungsreihenfolge gewählt wurde wird im Kommentar des Programms erläutert, mit dessen Hilfe das Spektrum der Fensterfolge für beliebige Frequenzen Ω berechnet werden kann, weil dort Einhaltung dieser Reihenfolge von besonderer Bedeutung ist. Auch die in Zeile 66 berechneten Normierungsfaktoren $NF_0(\rho)$ werden dort näher erläutert. Die Normierungskonstante NF_1 wird in Zeile 65 so gewählt, dass sich ohne Berücksichtigung der Nullstellen des Vektors z_0 vor der abschließenden Normierung auf $(N \cdot F_nu(0))^2$ bei $\nu=0$ ein Wert in der Größenordnung von eins ergibt. Dass dem so ist, wird im Kommentar zum Programm im Anhang von [1] erläutert, wobei zu beachten ist, dass nun um $2 \cdot A_1$ weniger doppelte Polstellen vorliegen. Da nun um A_1 weniger Fourierreihenkoeffizienten zu berechnen sind, wird auch der Vektor F_nu , mit dessen Hilfe die Summe in Gleichung (10.9) realisiert wird, in Zeile 64 mit entsprechend weniger Elementen bereitgestellt. Der Vektor F_nu_1 zur Berechnung der kumulativen Produkte ist ebensolang, und berücksichtigt bei seiner Initialisierung in Zeile 68, die modifizierten Summengrenzen in Gleichung (10.9), indem die Elemente zu null initialisiert werden, die Summanden entsprechen würden, die nicht in der Summe in Gleichung (10.9) vertreten sind. Dass die anderen Elemente dieses Vektors nicht auf eins sondern auf ε^{-2} initialisiert werden, ist Teil der Normierung, die bei dem Programm zur Berechnung des Spektrums der Fensterfolge kommentiert wird. Die Berechnung der Nullstellenabstandsquadrate in den Zeilen 73 und 74 erfolgt bis auf die Normierung mit dem in Gleichung (10.9) angegebenen Term. Die Berechnung der Polstellenabstandsquadrate erfolgt in den Zeilen 76 und 77 wie dies in [1] ausführlich erläutert ist. Die vorläufige Normierung wird in Zeile 82 endgültig korrigiert, nachdem in Zeile 81 aus den Betragsquadraten der Fourierreihenkoeffizienten durch Wurzelziehen die Beträge berechnet worden sind. Danach werden diese in Zeile 84 noch mit der Exponentialfunktion der Phase multipliziert, und es kann daraus die Fensterfolge wie in dem in [1] angegebenen Programm berechnet werden. Es sind hier lediglich um A_1 weniger Fourierreihenkoeffizienten vorhanden, so dass die `for`-Schleife in Zeile 90 nun mit einem entsprechend niedrigeren Wert von `nu` startet.

Bei der Berechnung der Phase der Fourierreihenkoeffizienten in den Zeilen 52 bis 59 ist zum einen ebenfalls zu berücksichtigen, dass nun um A_1 weniger Fourierreihenkoeffizienten berechnet werden müssen, so dass die Vektoren `nu` und `phi` in den Zeilen 52 und 55 entsprechend kürzer ausfallen. Zum anderen ist bei der Addition des Phasenanteils des Teils $\tilde{D}_P(\tilde{z})$ der Bilineartransformierten, der die bekannten Polstellen enthält, nun ebenso

wie bei dem linearphasigen Anteil $D_E(z)$ die geänderte Vielfachheit der Pol- bzw. Nullstellen zu beachten, so dass sich in Zeile 59 modifizierte Vorfaktoren bei den Termen in Ω_{ms} und $\Omega/2$ ergeben.

Nun kommen wir zur dynamischen Anpassung der Länge $\widetilde{M}$ des Cepstrums und des Bilineartransformationsparameters c . Mit der Variable `Ms_OK`, die in Zeile 7 auf null initialisiert wird, wird angezeigt, dass das Cepstrum bisher noch nicht erfolgreich berechnet werden konnte. Daher wird die `while`-Schleife, die in Zeile 8 beginnt und in Zeile 51 endet, beim ersten mal auf jeden Fall durchlaufen. Sollte mit den in den Zeilen 3 bis 6 gewählten Werten $\widetilde{M}$ und c festgestellt worden sein, dass das Cepstrum mit einer genügend großen Länge berechnet worden ist, so wird in Zeile 49 die Variable `Ms_OK` auf eins gesetzt, so dass es nicht zu einem erneuten Schleifendurchlauf kommt. Um festzustellen, ob die FFT-Länge $\widetilde{M}$ bei der Berechnung des Cepstrums groß genug war, ist der Betrag des Cepstrums mit der in Gleichung (10.13) genannten Grenze zu vergleichen. In Zeile 41 wird dieser Vergleich durchgeführt, wobei hier die zweifache Grenze herangezogen wird, weil im Programm auch das doppelte Cepstrum `Ceps_2` verwendet wird. Um die doppelte Grenze in Zeile 38 berechnen zu können, muss zunächst der Mindestwert des Rauschsockels nach Gleichung (10.14) in Zeile 37 bestimmt werden. Um zu vermeiden, dass der Rauschsockel im Cepstrum verhindert, dass eine ausreichende FFT-Länge $\widetilde{M}$ erkannt werden kann, wird diese Grenze noch auf das doppelte des in Gleichung (10.15) genannten Schätzwertes für den maximal zu erwartenden Berechnungsfehler im Cepstrum nach unten begrenzt. Diese Begrenzung erfolgt in den Programmzeilen 39 und 40. Nur falls eine zu kurze FFT-Länge $\widetilde{M}$ detektiert worden ist, wird in Zeile 47 $\widetilde{M}$ verdoppelt, und in Zeile 46 ein neuer Wert des Bilineartransformationsparameters c nach Gleichung (10.16) berechnet. Dazu muss der Wert Δc mit Gleichung (10.17) aus dem Schätzwert des Winkels der Nullstelle, die dem Einheitskreis am nächsten liegt, berechnet werden. Diese Abschätzung erfolgt, indem man bei der Folge, die man durch eine FFT der Länge $\widetilde{M}/4$ aus einem Abschnitt des mit k multiplizierten Cepstrums gewinnt, die Lage des betraglichen Maximums bestimmt, wie dies im Anschluss an Gleichung (10.17) beschrieben ist, und in den Zeilen 42 und 43 ausgeführt wird. Da der Wert des Maximums selbst nicht interessiert, sondern nur die Lage des Maximums, ist dafür die im weiteren nicht mehr verwendete Variable `dummy` vorgesehen. Die Variable `womax` liefert uns in Zeile 43 den Index des Elementes des Vektors `kriterium`, bei dem das Maximum erreicht wird. Da die Elemente dieses Vektors die diskreten Fouriertransformierten des Abschnitts des mit k multiplizierten Cepstrums bei den Frequenzen $\tilde{\eta} \cdot 8\pi/\widetilde{M}$ mit $\tilde{\eta} = 0(1)\widetilde{M}/4-1$ enthalten, ist der um eins reduzierte Index `womax` proportional zum Winkel des Punktes am Einheitskreis, wo das Spektrum sein betragliches Maximum erreicht. Wenn man in Zeile 45 auf der rechten Seite $\text{delta_c} = 1 - 2 \cdot \psi_0/\pi$ einsetzt, erhält man die in Gleichung (10.17) dargestellte Form für Δc , so dass die beiden Zeilen 44 und 45 lediglich eine zweistufige Berechnung der Gleichung (10.17) darstellen,

die vom numerischen Standpunkt günstiger erscheint. Die Abschätzung des Winkels ψ_0 ist nur dann sinnvoll, wenn man auch schon bei dem Startwert von $\widetilde{M}$ entscheiden kann, ob der Parameter c verringert, vergrößert oder beibehalten werden soll. Dazu muss es bei dem Startwert von $\widetilde{M}$ wenigstens die drei möglichen Schätzwerte $\psi_0 \in \{0; \pi/2; \pi\}$ für den unbekannten Nullstellenwinkel geben. $\widetilde{M}/4$ sollte daher wenigstens 4 sein. Dies wird erzwungen, indem der Startwert von $\widetilde{M}$ in Zeile 6 auf den Mindestwert 16 gesetzt wird, falls sich bei der Berechnung in den Zeilen 4 und 5 ein kleinerer Wert ergibt.

Wie bei der in [1] vorgestellten Variante kann man auch hier die Fensterfolge durch Auswertung der Fourierreihe nach Gleichung ([1]:6.16) berechnen. Dies erfolgt in den Zeilen 85 bis 94. Prinzipiell kann man die Fourierreihe nach Gleichung ([1]:6.16) auch für beliebiges reelles k auswerten. Wenn man an nichtganzzahligen Werten k , die auf einem Vektor $\mathbf{k}$ abgespeichert seien, interessiert ist, so ersetzt man den Programmaufruf in Zeile 1 sowie die Programmzeilen ab einschließlich Zeile 85 durch die Programmzeilen:

```

1: function [ f_k, F_nu ] = fenster( N, M, A_0, A_1, z_0, k )
:
85: f_k = zeros(1,length(k))
86: O_k = 2*pi/F*k
87: for nu = N-A_1-1:-1:1
88:   f_k = f_k + 2*real(F_nu(nu+1))*cos(O_k*nu) - ...
           2*imag(F_nu(nu+1))*sin(O_k*nu)
89: end
90: f_k = f_k + 1/N

```

Diese Variante ist natürlich nicht so exakt, wie die Variante für ganzzahliges k .

11.2 Die Spektren der diskreten Fensterfolgen

Neben den im vorigen Unterkapitel genannten Parametern N , M , A_0 , A_1 und dem Vektor der Nullstellen $z_{0,\rho}$, für die die dort genannten Restriktionen gelten sollen, benötigen wir bei diesem Programm noch einen Vektor Ω , der die Frequenzen Ω als Elemente enthält, für die das Spektrum der Fensterfolge berechnet werden soll. Von diesen wird angenommen, dass sie im Bereich zwischen $-\pi$ und π liegen, weil dann die Berechnung mit bestmöglicher Genauigkeit erfolgen kann. Als Ergebnis werden von diesem Programm die Werte des Spektrums der Fensterfolge für die Frequenzen des Vektors Ω auf dem Vektor $\mathbf{F}_\Omega$ zurückgegeben.

```

1: function F_Omega = spektrum( N, M, A_0, A_1, z_0, Omega )
2: F = N * M
3: c = 2 / ( 1 + (N/2)^(M/3/(1-M)) / tan(pi/2/M) )
4: Ms = -log(eps)/log(2) * 2^log(N/3) * 3.6^(1/M)
5: Ms = 2^ceil(log(Ms)/log(2))
6: Ms = max(Ms,16)
7: Ms_OK = 0
8: while ~Ms_OK
9:     eta = 0:Ms/2
10:    F_eta = zeros(1,Ms/2+1)
11:    NF_1 = 1 / ( c^2 + 4 * (1-c) * sin( pi/F * (N-1-A_1) )^2 )
12:    NF_0 = max( [ ( abs( z_0*exp( j*pi/F*(1-N) )-(1-c) ) + ...
                    abs( 1-(1-c)*z_0*exp( j*pi/F*(1-N) ) ) ) ; ...
                    ( abs( z_0*exp( j*pi/F*(N-1) )-(1-c) ) + ...
                    abs( 1-(1-c)*z_0*exp( j*pi/F*(N-1) ) ) ) ] ).^(-2)
13:    for nu_1 = (1-N)/2:(N-1)/2
14:        F_eta_1 = ones(1,Ms/2+1)
15:        rho = 0
16:        for nu_2 = [(1-N+A_1):(nu_1+A_1-(N+1)/2), ...
                     (nu_1-A_1+(N+1)/2):(N-1-A_1)]
17:            K_1 = NF_1 * ( c^2 + 4 * (1-c) * sin( pi/F * nu_2 )^2 )
18:            Psi_1 = atan( ( (1-c) * sin(2*pi/F*nu_2) ) / ...
                           ( c + 2 * (1-c) * sin(pi/F*nu_2)^2 ) )
19:            F_eta_1 = F_eta_1 .* K_1 .* sin( pi/Ms*eta-pi/F*nu_2-Psi_1 ).^2
20:            rho = rho + 1
21:            if rho < N-0.5-2*A_1-A_0
22:                Z_Z = z_0(rho) * exp( j*2*pi/F*nu_1 ) - (1-c)
23:                Z_N = 1 - (1-c) * z_0(rho) * exp(j*2*pi/F*nu_1)
24:                abs_Z_Z = abs( Z_Z )
25:                abs_Z_N = abs( Z_N )
26:                F_eta_1 = F_eta_1 .* NF_0(rho) .* ...
                           ( (abs_Z_N-abs_Z_Z)^2 + 4 * abs_Z_N * abs_Z_Z * ...
                             sin( pi/Ms*eta-(angle(Z_Z)-angle(Z_N))/2 ) ).^2 )
27:            end
28:        end
29:        F_eta = F_eta + F_eta_1
30:    end

```

```

31: NF = 1 / sqrt( max(F_eta) * min(F_eta) )
32: F_eta = NF * F_eta
33: L_eta = log(F_eta)
34: Ceps_2 = ifft( [L_eta,L_eta(Ms/2:-1:2)] )
35: Ceps_2 = real( Ceps_2 )
36: Ceps_2 = ( Ceps_2 + Ceps_2([1,Ms:-1:2]) ) / 2
37: sockel = eps / Ms / sqrt(48) * ...
           sqrt( sum( max( [L_eta,L_eta(Ms/2:-1:2)].^2, 1 ) ) )
38: grenze = ( 2 * (2*N-2-2*A_1-A_0) ) ./ [1:Ms/2] .* ...
           ( Ms * sockel/2/(2*N-2-2*A_1-A_0) ).^( 2*[1:Ms/2]/Ms )
39: max_fehl = sum( max( abs([L_eta,L_eta(Ms/2:-1:2)]), 1 ) ) * ...
           eps * ( 2 + log(Ms)/log(2) ) / Ms
40: grenze = max( grenze, max_fehl )
41: if any( abs(Ceps_2(2:Ms/2+1)) > grenze )
42:     kriterium = abs( fft( Ceps_2(Ms/4+1:Ms/2) .* [Ms/4:Ms/2-1] ) )
43:     [dummy,womax] = max( kriterium(1:Ms/8+1) )
44:     delta_c = 1-(womax-1)*16/Ms
45:     delta_c = sin( delta_c*pi*(1-delta_c^2/2) )/3
46:     c = c * (1-delta_c) / ( 1 + (1-c) * delta_c )
47:     Ms = 2*Ms
48: else
49:     Ms_OK = 1
50: end
51: end
52: len_0 = length(Omega)
53: Omega_s = Omega + 2 * atan( ( (1-c) * sin(Omega) ) ./ ...
                             ( c + 2 * (1-c) * sin(Omega/2).^2 ) )
54: phi = zeros(1,len_0)
55: for nu_i = 1:len_0
56:     phi(nu_i) = sin(Omega_s(nu_i)*[Ms/2-1:-1:1]) * Ceps_2(Ms/2:-1:2).'
57: end
58: phi = phi - (N-1-A_1-A_0/2) * Omega_s + (F-1-A_0) * Omega/2
59: F_Omega = zeros(1,len_0+1)
60: if N == 2*A_1+1
61:     NF_1 = 1
62: else
63:     NF_1 = exp(2*log(8/pi*F)-sum(log([5:2:4*N-8*A_1])))/(N-1-2*A_1))
64: end
65: NF_0 = 1 ./ max( [ abs( z_0 * exp(j*pi/F*(N-1)) - 1 ).^2 ; ...
                    abs( z_0 * exp(j*pi/F*(1-N)) - 1 ).^2 ] )
66: for nu_1 = (1-N)/2:(N-1)/2
67:     Omega_nu_1 = [Omega,0]/(2*pi) - nu_1/F
68:     Omega_nu_1 = (2*pi) * ( Omega_nu_1 - round(Omega_nu_1) )

```

```

69: nu_4 = 2 * round( F/(2*pi) * Omega_nu_1 + (N-1)/2 ) - N + 1
70: nu_4 = ( 1-N+2*A_1 ) .* ( nu_4 < 2*A_1-N ) + ...
      nu_4 .* ( abs(nu_4) < N-2*A_1 ) + ...
      ( N-1-2*A_1 ) .* ( nu_4 > N-2*A_1 )
71: nu_4 = nu_4/2
72: d_Omega = Omega_nu_1 - 2*pi/F * nu_4
73: F_Omega_1 = ( abs(d_Omega) < eps/F )
74: F_Omega_1 = F_Omega_1 / eps^2 + (~F_Omega_1) .* ...
      ( sin(F/2*d_Omega) ./ ...
      ( F*eps*sin(d_Omega/2) + F_Omega_1 ) ).^2
75: rho = 0
76: for nu_2 = (1-N)/2+A_1:(N-3)/2-A_1
77:     rho = rho + 1
78:     if rho < N-0.5-2*A_1-A_0
79:         abs_z_0 = abs(z_0(rho))
80:         F_Omega_1 = F_Omega_1 .* NF_0(rho) .* ( ( 1-abs_z_0 )^2 + ...
            4*abs_z_0*sin((Omega_nu_1-angle(z_0(rho)))/2).^2 )
81:     end
82:     nu_3 = nu_2 + ( nu_2 >= nu_4 )
83:     F_Omega_1 = F_Omega_1 ./ (NF_1.*sin( Omega_nu_1/2-pi/F*nu_3).^2 )
84: end
85: F_Omega = F_Omega + F_Omega_1
86: end
87: Omega_nu_1 = abs( F / (2*pi) * [Omega,0] ) - N + A_1 + 0.5
88: Omega_nu_1 = 0.5 * ( Omega_nu_1 > 0 ) .* Omega_nu_1 + 0.25
89: Omega_nu_1 = 1 - 2 * ( ( Omega_nu_1 - floor(Omega_nu_1) ) > 0.5 )
90: F_Omega = Omega_nu_1 .* sqrt(F_Omega)
91: F_Omega = ( M / F_Omega(len_0+1) ) * F_Omega(1:len_0)
92: F_Omega = F_Omega .* exp(-j*phi)

```

Auch wenn man die Werte des Spektrums der Fensterfolge $f(k)$ bei dieser Programmvariante nicht mehr nur für die Frequenzen $\Omega = \nu \cdot 2\pi/F$ mit $\nu = 0 \dots N-1$ berechnet, weil man hier nicht an den Koeffizienten der Fourierreihenentwicklung der Fensterfolge interessiert ist, muss man das Cepstrum genauso berechnen, wie dies im vorigen Unterkapitel beschrieben ist. Die Zeilen 2 bis 51 sind daher identisch mit den Zeilen des Programms in Unterkapitel 11.1. In den Zeilen 53 bis 58 wird wieder die Phase der zu bestimmenden Spektralwerte durch Auswertung der Sinusreihe bestimmt, deren Koeffizienten die Cepstralwerte sind. Die Auswertung erfolgt nun für alle Frequenzen des Vektors `Omega`, dessen Elementanzahl in Zeile 52 bestimmt wird. Daher ist der Vektor `phi` nun in Zeile 54 mit einer veränderten Länge bereitzustellen, und die `for`-Schleife der Zeilen 55 bis 57 ist nun für jede der Frequenzen des Vektors `Omega` einmal zu durchlaufen.

Ab Zeile 59 werden die Beträge der zu bestimmenden Spektralwerte berechnet. Da dies nun nicht mehr nur die Spektralwerte für niedrige Frequenzen im Raster der Nullstellen des Spektrums der Fensterfolge auf dem Einheitskreis sind, und da das Spektrum für große Werte von M und N im Bereich hoher Frequenzen $\Omega \approx \pi$ wegen des potenzmäßigen Anstiegs der Sperrdämpfung extrem klein werden kann, ist nun besonderes Augenmerk auf eine geeignete Normierung und eine geeignete Reihenfolge der Verarbeitung der Pol- und Nullstellen zu legen, wenn man die Berechnung der Spektralwerte so durchführen will, dass alle Spektralwerte, die betraglich größer als `realmin` sind, mit dem kleinstmöglichen relativen Fehler berechnet werden können. Das Betragsquadrat der Fensterfolge erreicht ihr absolutes Maximum M^2 immer bei der Frequenz $\Omega = 0$, weil bei dieser Frequenz bei der Überlagerung der um Vielfache von $2\pi/M$ verschobenen Betragsquadrate der Fensterspektren in Gleichung ([1]:2.20), die von der hier zu berechnenden Fensterfolge erfüllt wird, nur ein einziger der stets positiven Summanden übrigbleibt. Wenn man die Normierungskonstanten `NF_1` und `NF_0(rho)` nun so wählt, dass das Betragsquadrat des Spektralwertes bei der Frequenz $\Omega = 0$ vor der endgültigen Normierung größer als M^2 ist, wird sich der relative Fehler der bis dahin berechneten Spektralwerte durch die abschließende Normierung nur für die Werte erhöhen, die nach der abschließenden Normierung kleiner als `realmin` sind. Um dieses Ziel zu erreichen, wird in Zeile 74 bei der Initialisierung des Vektors, mit dessen Hilfe die kumulativen Produkte berechnet werden, im Nenner jeweils der Faktor `eps`² verwendet. Von diesem Wert kann erwartet werden, dass er wesentlich größer als M^2 , und doch deutlich kleiner als die größte am Rechner darstellbare Zahl ist, so dass es für beliebige Werte Ω nicht zu einem Überlauf kommen kann, wenn man bei der Berechnung der Zwischenergebnisse darauf achtet, dass diese ebenfalls nicht größer als die größte am Rechner darstellbare Zahl werden. Um die endgültige Normierung einfach realisieren zu können, wird wieder auch der Spektralwert bei der Frequenz Null ermittelt, da man von diesem weiß, dass er nach der endgültigen Normierung M sein muss.

Bevor die konkrete Wahl der Normierungskonstanten `NF_1` und `NF_0(rho)` in den Zeilen 60 bis 65 näher erläutert wird, wird die im Programm realisierte Reihenfolge der Verarbeitung der Pol- und Nullstellen analysiert und begründet. Wenn man in Gleichung (10.8) bei der mit (*) gekennzeichneten Form $z = e^{j \cdot \Omega}$ einsetzt, kann man für jede der Frequenzen des Vektors `Omega` und jeden Summanden die um $\nu_1 \cdot 2\pi/F$ verringerte Frequenz berechnen, und bei beiden Nennern jeweils die Polstellen suchen, die dem Punkt $e^{j \cdot (\Omega - \nu_1 2\pi/F)}$ am nächsten liegen. Da die beiden Nenner gleiche Nullstellen aufweisen, sind diese beiden Polstellen identisch. Diese doppelte Polstelle ist diejenige, die jeweils bei jedem Summanden und bei jeder Frequenz des Vektors `Omega` als erste verarbeitet wird. In den Zeilen 67 und 68 werden für jeden Summanden die um $\nu_1 \cdot 2\pi/F$ verringerten Frequenzen `Omega_nu_1` modulo 2π berechnet. In den Zeilen 69 bis 71 werden die auf $2\pi/F$ normierten Winkel der jeweils nächstgelegenen doppelten Polstelle bestimmt. Die Winkeldifferenz wird in Zeile 72

als Vektor `d_Omega` berechnet. Nun kann man wieder die beiden Zähler in Gleichung (10.8) der Form (*) mit der doppelten Polstelle zu einer quadrierten, periodisch fortgesetzten si-Funktion zusammenfassen. Das Hauptmaximum der periodisch fortgesetzten si-Funktion liegt jeweils bei dem Polstellenwinkel, der der gerade betrachteten, verschobenen Frequenz $\Omega - \nu_1 \cdot 2\pi/F$ am nächsten liegt. Wenn die verschobene Frequenz sehr nahe bei dem Hauptmaximum liegt, was in Zeile 73 festgestellt wird, kann man im Rahmen der Rechengenauigkeit die periodisch fortgesetzte si-Funktion durch den Konstanten Wert des Hauptmaximums ersetzen, anderenfalls kann die periodisch fortgesetzte si-Funktion als Quotient zweier Sinusfunktionen sowieso mit der gewünschten Genauigkeit berechnet werden. In der Zeile 74 werden die Werte der quadrierten, periodisch fortgesetzten si-Funktion abhängig von den Werten des Arguments auf eine der beiden Weisen berechnet, um den Vektor `F_Omega_1`, mit dessen Hilfe die kumulativen Produkte berechnet werden, zu initialisieren. Dabei wird zugleich die bereits angesprochene Normierung auf ε^2 vorgenommen. Im weiteren werden wir uns auf den kritischen Fall beschränken, dass M sehr groß ist, und eine reale Chance besteht, dass der am Rechner darstellbare Zahlenbereich verlassen wird. In diesem Fall sinken nun die Werte auf dem Vektor `F_Omega_1` mit steigendem Ω mit $\sin(\Omega/2)^{-2}$ ab. Nach der Initialisierung des Vektors `F_Omega_1` werden die Nullstellen des Vektors `z_0` und die nicht gekürzten Polstellen der Nenner in Gleichung (10.8)(*) abwechselnd verarbeitet. In der Zeile 80 werden zuerst die normierten Quadrate der Nullstellenabstände für alle Frequenzen des Vektors `Omega` zugleich mit den bisher berechneten kumulativen Produkten des Vektors `F_Omega_1` multipliziert. Danach werden jeweils in Zeile 83 die Elemente des Vektors `F_Omega_1` durch die normierten Quadrate der Polstellenabstände dividiert. Dabei wird mit Hilfe des Vektors `nu_3` wieder die gekürzte doppelte Polstelle ausgelassen, ganz ähnlich, wie dies bei dem Programm im Anhang [1]:A.8 beschrieben worden ist. Wenn die gerade verarbeitete Nullstelle $z_{0,\rho}$ sehr nahe bei dem Punkt $z=1$ liegt, werden die Quadrate der Nullstellenabstände zu den Punkten $e^{j\Omega}$ mit steigendem Ω mit $\sin(\Omega/2)^2$ zunehmen. Da für große Werte von M auch alle Polstellen sehr nahe bei $z=1$ liegen, wird dadurch der Einfluss der danach verarbeiteten Polstelle auf den Vektor `F_Omega_1` für hohe Frequenzen $\Omega \approx \pi$ insofern weitgehend kompensiert, als dass sich als Quotient der beiden Abstandsquadrate zu den Pol- und Nullstellen näherungsweise der Konstante Quotient $4 \cdot \text{NF_0}(\rho)/\text{NF_1}$ der Normierungskonstanten ergibt. Wenn die Nullstelle $z_{0,\rho}$ *nicht* so nahe bei $z=1$ liegt, wird sich bei dem Quotienten der Abstandsquadrate zu zwei nacheinander verarbeiteten Null- und Polstellen ein Frequenzverlauf ergeben, der bei hohen Frequenzen insgesamt kleiner als der ebengenannte konstante Quotient ist. Da wir vorausgesetzt hatten, dass die Nullstellen des Vektors `z_0` alle innerhalb des Einheitskreises liegen sollen, wenn sie an dieses Programm übergeben werden, wird der in Zeile 65 berechnete Normierungsfaktor `NF_0(rho)` maximal, wenn die Nullstelle $z_{0,\rho}$ genau bei eins liegt. Die Normierungskonstante `NF_1` wird in

Zeile 63 in der Art berechnet, wie dies bei dem Programm in Kapitel [1]:A.8 beschrieben ist. Es wurde nun untersucht, wie groß der Quotient $4 \cdot \text{NF_0}(\rho) / \text{NF_1}$ in Abhängigkeit von N und A_1 für großes M maximal — also für $z_{0,\rho} = 1$ — werden kann. Dabei stellte sich heraus, dass er außer für $N=2$ und $A_1=0$ immer kleiner 1 ist. Wenn M hinreichend groß ist ($M > 10$), ist der Quotient bei steigendem M praktisch konstant. Desweiteren stellte sich heraus, dass Wert des Quotienten bei einem festen Wert von N und variabler Anzahl A_1 um so näher bei eins liegt, je häufiger die Schleife mit dem Index `nu_2` durchlaufen wird. Im eben ausgeschlossenen Fall ergibt sich ein maximaler Wert für den Quotienten, der für steigendes M gegen einen Wert von etwa 2,1875 konvergiert. Da es in diesem Fall nur maximal eine frei wählbare Nullstelle $z_{0,\rho}$ gibt, und nur ein Null-Polstellenpaar zu verarbeiten ist, so dass die Schleife mit dem Index `nu_2` nur einmal durchlaufen wird, besteht hier selbst für extrem große Werte von M kaum die Gefahr, dass der mit dem kleinstmöglichen relativen Fehler darstellbare Zahlenbereich verlassen wird, und somit ist dieser Maximalwert des Quotienten akzeptabel. In allen anderen Fällen, ist es nicht möglich, dass die Anhebung der Werte des Vektors `F_Omega_1` für $\Omega \approx \pi$ durch die Nullstelle $z_{0,\rho}$ nicht durch die danach verarbeitete Polstelle mehr als kompensiert wird. Nach einem Schleifendurchlauf bei der Berechnung der kumulativen Produkte in Gleichung (10.8)(*) werden die Werte des Vektors `F_Omega_1`, bei denen die größte Gefahr besteht, dass der mit dem kleinstmöglichen relativen Fehler darstellbare Zahlenbereich verlassen wird, kleiner sein als vor diesem Schleifendurchlauf. Man kann sich anhand der Anzahl der zu verarbeitenden Pol- und Nullstellen überlegen, dass zuletzt immer eine Polstelle behandelt wird. Es kann also nicht vorkommen, dass ein Wert, der vor einem Schleifendurchlauf zu klein ist, um mit voller Genauigkeit dargestellt werden zu können, oder der gar zu null quantisiert wurde, nach dem Schleifendurchlauf bei ideal fehlerfreier Berechnung in einen Bereich angehoben wird, der wieder mit voller Genauigkeit darstellbar wäre. Damit solches bei der abschließenden Normierung nicht passiert, darf der letzte Wert des Vektors `F_Omega_1`, der nach Zeile 67 den Wert bei der Frequenz Null enthält, der für die endgültige Normierung benötigt wird, nicht kleiner als M^2 werden. Die Normierungskonstante `NF_1` wurde so gewählt, dass der bei der Initialisierung eingestellte Wert ε^{-2} auch nach der vollständigen Berechnung etwa erhalten bleibt, wenn keine Nullstelle $z_{0,\rho}$ angegeben wird. Die Normierung der Abstandsquadrate der Nullstellen wird in Zeile 65 so gewählt, dass sich für $\Omega=0$ bei allen möglichen Rotationen der Nullstellen mit $e^{j \cdot \nu_1 \cdot 2\pi / F}$ ein Abstand zum Punkt $z=1$ ergibt, der den Maximalwert Eins wenigstens bei einem Summanden der Summe in Gleichung (10.8)(*) auch annimmt. Daher, und weil es immer mehr Summanden als Nullstellen gibt, kann man annehmen, dass in der Summe wenigstens einige Summanden auftreten, die trotz des Einflusses der Nullstellen nicht so extrem viel kleiner als ε^{-2} sind, dass der Wert bei $\Omega=0$ vor der abschließenden Normierung gleich unter M^2 abfällt.

11.3 Die diskreten Fensterautokorrelationsfolgen

Dieses Programm benötigt die Zeitpunkte k , für die die Fensterautokorrelationsfolge berechnet werden soll, als Vektor $\mathbf{k}$. Desweiteren muss der Zeilenvektor $\mathbf{F_nu}$ die Fourierreihenoeffizienten $F(\nu \cdot \frac{2\pi}{F})/F$ der Fensterfolge für $\nu = 0$ (1) $N - A_1 - 1$ enthalten. Diesen Vektor erhält man entweder bei dem Programm im Anhang von [1] oder bei dem Programm in Unterkapitel 10.1 als Ergebnis. Als drittes Eingabeargument ist noch die Länge F der Fensterfolge anzugeben.

```

1: function d_k = fenster_akf( k, F_nu, F )
2: N_A = length(F_nu)
3: M = round( F * F_nu(1) )
4: k = abs(k)
5: O_k = 2*pi/F*k
6: d_k = zeros(1,length(k))
7: for nu = N_A:-1:2
8:     F_sin = real( F_nu(nu)' * ...
        sum( [F_nu(N_A:-1:1)';F_nu([2:nu-1,nu+1:N_A]).'] ./ ...
            tan( pi/F*[2-N_A-nu:-1,1:N_A-nu] ).' ) ) / M + ...
        imag( F_nu(nu)' * ...
            sum( [F_nu(N_A:-1:1)';F_nu([2:nu-1,nu+1:N_A]).'] ) ) / M
9:     d_k = d_k + 2 * F_sin * sin(O_k*(nu-1)) + ...
        2 * abs(F_nu(nu))^2 / M * cos(O_k*(nu-1)) .* (F-k)
10: end
11: d_k = d_k + abs(F_nu(1))^2 / M * (F-k)

```

In den Zeilen 2 und 3 werden die benötigten Parameter bestimmt. Die Variable N_A enthält die Differenz $N - A_1$. Da die Fourierreihe der Fensterfolge nur Anteile bis zum $(N - A_1 - 1)$ -fachen der Grundkreisfrequenz enthält, kann der Wert für N_A als die Länge des Vektors $\mathbf{F_nu}$ bestimmt werden. Weil der Fourierreihenoeffizient mit $\nu=0$ immer $1/N$ ist, lässt sich M daraus durch Multiplikation mit der Fensterlänge F berechnen. I. allg. lässt sich der Fourierreihenoeffizient des Gleichanteils $\mathbf{F_nu}(1)$ am Rechner nicht exakt darstellen. Daher wird in Zeile 3 eine Rundung auf die nächste am Rechner exakt darstellbare ganze Zahl vorgenommen, um so die im Produkt $F \cdot \mathbf{F_nu}(1)$ evtl. vorhandenen Rundungsfehler abzuschneiden. Da es sich bei der Fensterautokorrelationsfolge um eine geradesymmetrische Folge handelt, kann man für negatives k ebenso gut die Werte für $-k$ berechnen. Durch Zeile 4 erspart man sich so eine Fallunterscheidung bei der späteren Berechnung der Fensterautokorrelationsfolge. Diese Berechnung wird mit Hilfe der Formeln (10.19) und (10.20) als Überlagerung einer Sinusreihe und einer mit der Dreiecksfunktion multiplizierten Kosinusreihe vorgenommen, wie dies in Kapitel 10.3 beschrieben ist. Dabei wird

bei jedem Reihenglied die mit k multiplizierte Grundkreisfrequenz benötigt, die daher in Zeile 5 für alle Werte des Vektors $\mathbf{k}$ auf dem Vektor $\mathbf{0_k}$ bereitgestellt wird. Die Summe in Gleichung (10.20) wird als `for`-Schleife, die in Zeile 7 beginnt und in Zeile 10 endet, realisiert, indem zunächst in Zeile 6 der Vektor $\mathbf{d_k}$ mit null initialisiert wird, zu dem in Zeile 9 bei jedem Schleifendurchlauf ein Anteil der Sinusreihe und ein Anteil der mit der Dreiecksfunktion multiplizierten Kosinusreihe addiert wird. In Zeile 8 wird der dazu benötigte Sinusreihenkoeffizient nach Gleichung (10.19) berechnet. Die dabei auftretenden Summen werden mit dem MATLAB-Befehl `sum` berechnet.

Mit der hier vorgestellten Variante kann man nicht nur die Fensterautokorrelationsfolge für ganzzahliges k berechnen, sondern auch die Werte der Funktion für beliebiges reelles k . Ist man nur an der Fensterautokorrelationsfolge für ganzzahliges k interessiert, und will man diese Werte aber mit einem geringeren Fehler berechnen, so empfiehlt es sich, bei der Berechnung der Reihe die dann mehrfach auftretenden Werte der Sinus- und Kosinusfunktion wieder außerhalb der `for`-Schleife hochgenau zu berechnen, und sich die jeweils bei einem Schleifendurchlauf benötigten Werte daraus herauszugreifen. Die dazu im eben vorgestellten Programm notwendigen Modifikationen werden hier nicht explizit dargestellt. Man kann sie dem Programm in Unterkapitel 11.1 zur Bestimmung der Fensterfolge $f(k)$ aus den Fourierreihenkoeffizienten (Zeilen 85 bis 94) sinngemäß entnehmen.

11.4 Die Fourierreihenkoeffizienten der kontinuierlichen Fensterfunktionen

Die in Kapitel 10.7 vorgestellte Fensterfunktion stellt den Spezialfall einer Fensterfunktion dar, bei der in Gleichung (10.37) die Laplace-Transformierte der zugrundeliegenden Basisfensterfunktion im Zähler kein Polynom in s aufweist. Das hier vorgestellte Programm berechnet die Fourierreihenkoeffizienten der allgemeineren Fensterfunktion, deren zugrundeliegende Basisfensterfunktion eine Laplace-Transformierte aufweist, bei der in Gleichung (10.37) im Zähler ein Polynom in s vom Grad $N - A_0 - 1$ eingestellt werden kann. Da der Grad des Zählerpolynoms wenigstens um eins kleiner als der Grad des Nennerpolynoms sein muss, wenn die Basisfensterfunktion keine Delta-Distributionen enthalten soll, muss der Graddefekt A_0 ebenfalls zwischen 0 und $N - 1$ gewählt werden. Desweiteren wird das Zählerpolynom durch seine Nullstellen bis auf eine Konstante festgelegt. Gleichung (10.37) zeigt, dass die Laplace-Transformierte der zugrundeliegenden Basisfensterfunktion auf der imaginären Achse immer äquidistante Nullstellen bei $s = j\omega$ mit $\omega = (2 \cdot \nu - N + 1) \cdot \pi$ und mit $\nu \in \mathbb{Z}$ außer für $\nu = 0$ (1) $N - 1$ aufweist. Auch wenn nun ein zusätzliches Zählerpolynom eingefügt wird, ändert sich daran nichts.

Die Nullstellen des hinzugekommenen Zählerpolynoms müssen zur reellen Achse spiegelsymmetrisch liegen, damit sich eine reelle Basisfensterfunktion ergibt. Von den frei einstellbaren Nullstellen des Zählerpolynoms können einige so gelegt werden, dass sie sich im Raster 2π an die eben genannten fest vorgegebenen äquidistanten Nullstellen auf der imaginären Achse anschließen. Somit gilt hier $G_\infty(j \cdot (N-1-2 \cdot A_1+2 \cdot \nu) \cdot \pi) = 0$ für $\nu \in \mathbb{N}$ und $G_\infty(j \cdot (N-1-2 \cdot A_1) \cdot \pi) \neq 0$. Durch die Angabe des Parameters A_1 werden diese Nullstellen mit einfacher Vielfachheit im Zählerpolynom berücksichtigt. Alle weiteren $N-A_0-1-2 \cdot A_1$ frei einstellbaren Nullstellen des Zählerpolynoms werden explizit als $s_{0,\rho}$ angegeben. Die Festlegung der frei wählbaren Nullstellen des hinzugekommenen Zählerpolynoms erfolgt also genau analog zu der Festlegung der frei wählbaren Nullstellen bei der Z-Transformierten der Basisfensterfolge der in Kapitel 10.1 beschriebenen diskreten Fensterfolge. Lediglich liegen nun die Nullstellen in der s -Ebene statt in der z -Ebene und die Nullstellen, die dort im Raster $2\pi/F$ auf dem Einheitskreis liegen, entsprechen nun den Nullstellen auf der imaginären Achse im Raster 2π . Es sei noch erwähnt, dass die in Kapitel 10.7 beschriebene Fensterfunktion den Spezialfall mit $A_0 = N-1$ darstellt.

Neben dem Parameter N , werden die Anzahl A_1 der zusätzlichen Nullstellen der Laplace-Transformierten $G_\infty(s)$ der Basisfensterfunktion $g(k)$ auf der imaginären Achse im Raster 2π , die Graddifferenz A_0 des Nenners und des Zählers von $G_\infty(s)$, sowie die frei wählbaren Nullstellen $s_{0,\rho}$ von $G_\infty(s)$ benötigt. Die ganzzahligen Parameter N , A_1 und A_0 müssen als skalare Größen (also als 1×1 Matrizen) `N`, `A_1` und `A_0` beim Programmaufruf angegeben werden. Für diese Parameter müssen die Bedingungen $N > 1$, $0 \leq A_1 < N/2$ und $0 \leq A_0 < N-1-2 \cdot A_1$ erfüllt sein. Die Nullstellen $s_{0,\rho}$ bilden die Elemente des Vektors `s_0`, der ebenfalls beim Programmaufruf zu übergeben ist. Dieser Vektor muss genau $N-1-A_0-2 \cdot A_1$ Elemente enthalten. Nullstellen, die nicht reell sind, müssen als zueinander konjugiert komplexe Paare in `s_0` vorhanden sein. Um eine möglichst gute Genauigkeit zu erzielen, sollten Nullstellen in der rechten Halbebene durch die an der imaginären Achse gespiegelten ersetzt worden sein, und die Nullstellen sollten nach aufsteigendem Abstand zum Punkt $s=0$ sortiert sein. Als Ergebnis werden von diesem Programm die Fourierreihenoeffizienten der Fensterfunktion für $\nu = 0$ (1) $N-A_1-1$ auf dem Vektor `F_nu` zurückgegeben.

```

1: function F_nu = fenster_koeff( N, A_0, A_1, s_0 )
2: c = (N/2)^(4/3)
3: Ms = -log(eps)/log(2) * 2^log(N/3)
4: Ms = 2^ceil(log(Ms)/log(2))
5: Ms = max(Ms,16)
6: Ms_OK = 0
7: while ~Ms_OK
```

```

8:  eta = 0:Ms/2
9:  F_eta = zeros(1,Ms/2+1)
10:  NF_1 = 1 / ( 1 + ( (N-1-A_1)/c )^2 )
11:  NF_0 = 1 ./ max( [ ( abs( s_0 + j*pi*(1-N) + 2*pi*c ) + ...
                      abs( s_0 + j*pi*(1-N) - 2*pi*c ) ) ; ...
                      ( abs( s_0 + j*pi*(N-1) + 2*pi*c ) + ...
                      abs( s_0 + j*pi*(N-1) - 2*pi*c ) ) ] ).^2
12:  for nu_1 = (1-N)/2:(N-1)/2
13:      F_eta_1 = ones(1,Ms/2+1)
14:      rho = 0
15:      for nu_2 = [(1-N+A_1):(nu_1+A_1-(N+1)/2), ...
                  (nu_1-A_1+(N+1)/2):(N-1-A_1)]
16:          K_1 = NF_1 * ( 1 + ( nu_2 / c )^2 )
17:          Psi_1 = atan( nu_2 / c )
18:          F_eta_1 = F_eta_1 .* K_1 .* sin( pi/Ms*eta - Psi_1 ).^2
19:          rho = rho + 1
20:          if rho < N-0.5-2*A_1-A_0
21:              Z_Z = 2*pi*c + s_0(rho) + j*2*pi*nu_1
22:              Z_N = 2*pi*c - s_0(rho) - j*2*pi*nu_1
23:              abs_Z_Z = abs( Z_Z )
24:              abs_Z_N = abs( Z_N )
25:              F_eta_1 = F_eta_1 .* NF_0(rho) .* ...
                      ( (abs_Z_N-abs_Z_Z)^2 + 4 * abs_Z_N * abs_Z_Z * ...
                      sin( pi/Ms*eta-(angle(Z_Z)-angle(Z_N))/2 ).^2 )
26:          end
27:      end
28:      F_eta = F_eta + F_eta_1
29:  end
30:  NF = 1 / sqrt( max(F_eta) * min(F_eta) )
31:  F_eta = NF * F_eta
32:  L_eta = log(F_eta)
33:  Ceps_2 = ifft( [L_eta,L_eta(Ms/2:-1:2)] )
34:  Ceps_2 = real( Ceps_2 )
35:  Ceps_2 = ( Ceps_2 + Ceps_2([1,Ms:-1:2]) ) / 2
36:  sockel = eps / Ms / sqrt(48) * ...
          sqrt( sum( max( [L_eta,L_eta(Ms/2:-1:2)].^2, 1 ) ) )
37:  grenze = ( 2 * (2*N-2-2*A_1-A_0) ) ./ [1:Ms/2] .* ...
          ( Ms * sockel/2/(2*N-2-2*A_1-A_0) ).^( 2*[1:Ms/2]/Ms )
38:  max_fehl = sum( max( abs([L_eta,L_eta(Ms/2:-1:2)]), 1 ) ) * ...
          eps * ( 2 + log(Ms)/log(2) ) / Ms
39:  grenze = max( grenze, max_fehl )

```

```

40: if any( abs(Ceps_2(2:Ms/2+1)) > grenze )
41:     kriterium = abs( fft( Ceps_2(Ms/4+1:Ms/2) .* [Ms/4:Ms/2-1] ) )
42:     [dummy,womax] = max( kriterium(1:Ms/8+1) )
43:     delta_c = 1-(womax-1)*16/Ms
44:     delta_c = sin( delta_c*pi*(1-delta_c^2/2) )/3
45:     c = c * (1-delta_c) / (1+delta_c)
46:     Ms = 2*Ms
47: else
48:     Ms_OK = 1
49: end
50: end
51: nu = 1:N-A_1-1
52: Omega_s = 2 * atan(nu./c)
53: phi = zeros(1,N-A_1-1)
54: for nu_i = nu
55:     phi(nu_i) = sin(Omega_s(nu_i)*[Ms/2-1:-1:1]) * Ceps_2(Ms/2:-1:2).';
56: end
57: phi = phi - ( N - 1 - A_1 - A_0/2 ) * Omega_s + pi * nu
58: if N == 2*A_1+1
59:     F_nu = ones(1,N-A_1)
60: else
61:     nu = 0:N-A_1-1
62:     F_nu = zeros(1,N-A_1)
63:     NF_1 = exp(2*log(8)-sum(log([5:2:4*N-8*A_1])))/(N-1-2*A_1))
64:     NF_0 = 1 ./ max( [ abs( s_0 + j*pi*(N-1) ).^2 ; ...
                        abs( s_0 - j*pi*(N-1) ).^2 ] )
65:     for nu_1 = (1-N)/2:(N-1)/2
66:         F_nu_1 = ( ( nu < nu_1+N/2-A_1 ) & ( nu > nu_1-N/2+A_1 ) ) / eps^2
67:         rho = 0
68:         for nu_2 = (1-N)/2+A_1:(N-3)/2-A_1
69:             rho = rho + 1
70:             if rho < N-0.5-2*A_1-A_0
71:                 F_nu_1 = F_nu_1 .* NF_0(rho) .* ...
                        abs( j*2*pi*(nu-nu_1) - s_0(rho) ).^2
72:             end
73:             nu_3 = nu - nu_1 - nu_2 - ( nu <= nu_1 + nu_2 )
74:             F_nu_1 = F_nu_1 ./ ( NF_1 .* nu_3 .^ 2 )
75:         end
76:         F_nu = F_nu + F_nu_1
77:     end
78:     F_nu = sqrt(F_nu)
79:     F_nu = F_nu / F_nu(1)
80: end
81: F_nu(2:N-A_1) = F_nu(2:N-A_1) .* exp(-j*phi)

```

Gegenüber der Berechnung der Fourierreihenkoeffizienten der diskreten Fensterfolge nach Kapitel 10.1, die im Programm des Unterkapitels 11.1 enthalten ist, ergeben sich hier einige Änderungen, die nun kommentiert werden. Die Bilineartransformation ist nun nach Gleichung (10.41) und (10.42) definiert. Daher ergeben sich als Startwerte des Parameters c und der für die Berechnung des Cepstrums benötigten FFT-Länge M_s die in den Gleichungen (10.51) und (10.52) angegebenen Werte, die in den Zeilen 2 bis 5 eingestellt werden.

Die Quadrate der Abstände der Punkte des Einheitskreises $e^{j\cdot\tilde{\Omega}}$ von den bilinear transformierten Nullstellen, die zur Erweiterung auf den Hauptnenner eingefügt wurden, werden mit Hilfe des konstanten Normierungsfaktors NF_1 normiert. Dieser wird in Zeile 10 wieder so gewählt, dass der maximal auftretende mit NF_1 normierte Faktor K_1 , der in Zeile 16 berechnet wird, zu eins wird. Dieser Faktor K_1 ist abgesehen von der Normierung gleich der in Gleichung (10.46) angegebenen Konstante K_{∞,ν_2} . In Zeile 17 wird der in Gleichung (10.45) auftretende Winkel der bilineartransformierten Nullstelle berechnet, mit dessen Hilfe in Zeile 18 bei jedem Schleifendurchlauf ein Faktor mit dem bisher berechneten kumulativen Produkt multipliziert werden kann. Auch die Normierungsfaktoren $NF_0(\mathbf{rho})$ werden wieder als das Reziproke des Maximalwertes $(|\tilde{z}_{N,\rho,\nu_1}| + |\tilde{z}_{Z,\rho,\nu_1}|)^2$ des normierten Quadrats des Nullstellenabstands zum Punkt $e^{j\cdot\tilde{\Omega}}$ gewählt. Da hier die um $\nu_1 \cdot 2\pi$ verschobenen Nullstellen $s_{0,\rho}$ anders bilineartransformiert werden, ergeben sich nach der Bilineartransformation die Nullstellen $\tilde{z}_{Z,\rho,\nu_1}/\tilde{z}_{N,\rho,\nu_1}$, deren Zähler und Nenner sich so berechnen lassen, wie dies in den Zeilen 21 und 22 durchgeführt wird. Setzt man dies und die maximal möglichen Verschiebungen $\nu_1 = N-1$ und $\nu_1 = 1-N$ in die maximalen Abstandsquadrate $(|\tilde{z}_{N,\rho,\nu_1}| + |\tilde{z}_{Z,\rho,\nu_1}|)^2$ ein, so erhält man die in Zeile 11 berechneten Normierungsfaktoren $NF_0(\mathbf{rho})$, die den Vektor NF_0 bilden. Damit lassen sich die normierten Abstandsquadrate in Zeile 23 bis 25 exakt genau so berechnen, wie bei der diskreten Fensterfolge.

Bei der evtl. notwendigen Neuberechnung des Bilineartransmutationsparameters c in Zeile 45 ist nun eine leicht modifizierte Formel zu verwenden, die berücksichtigt, dass hier eine andere Art der Bilineartransformation stattfindet. Aus demselben Grund berechnen sich nun die Frequenzen $\tilde{\Omega}$ in Zeile 52 nach Gleichung (10.43). Abgesehen davon, bleibt die Berechnung des über das Cepstrum berechneten Phasenanteils in den Zeilen 53 bis 56 unverändert. Bei dem linearen Phasenanteil in Zeile 57 ergibt sich eine geringfügige Modifikation, da der Anteil, der im Programm des Unterkapitels 11.1 mit $(F-1-A_0) \cdot \Omega/2$ ansteigt, beim Grenzübergang $M \rightarrow \infty$ mit der Substitution nach Gleichung (10.26) den Term $\omega/2$ liefert, der bei den zu bestimmenden Frequenzpunkten gleich $\nu \cdot \pi$ ist.

Die Normierung bei der Berechnung des Betragsquadrates der Fourierreihenkoeffizienten in den Zeilen 63 und 64 weist nun zwei Veränderungen auf. Wenn man berücksichtigt, dass

bei dem Quadrat des Polstellenabstands immer das Quadrat der Sinusfunktion des *halben* Differenzwinkels auftritt, so ergibt sich mit der Substitution nach Gleichung (10.26) bei der im Unterkapitel [1]:A.8 beschrieben Wahl des Normierungsfaktors `NF_1` der in Zeile 63 angegebene Wert, bei dem nun statt $\log(8/\pi \cdot F)$ der Term $\log(8)$ im Exponenten auftritt. Bei den Normierungsfaktoren des Vektors `NF_0` ist in Zeile 64 nun nicht mehr das Quadrat des Abstands der rotierten Nullstellen $z_{0,\rho}$ zum Punkt $z=1$ zu verwenden, sondern das Quadrat des Abstands der verschobenen Nullstellen $s_{0,\rho}$ zum Punkt $s=0$. Dementsprechend wurden bei der Berechnung der Faktoren des kumulativen Produkts in den Zeilen 71 und 74 die Abstandsquadrate, die sich bei der diskreten Fensterfolge mit Hilfe der Quadrate der Sinusfunktion angeben lassen, durch die entsprechenden Abstandsquadrate ersetzt, die sich ohne die Sinusfunktion berechnen. Es sei noch angemerkt, dass sich im Spezialfall mit $A_0 = N-1$ die im Nenner auftretenden Produkte der Polstellenabstände durch Erweiterung mit $(N-1)!$ und bei einer anderen Normierung als mit `NF_1` als die Quadrate von Binomialkoeffizienten schreiben lassen. Dies erklärt die in Gleichung (10.39) angegebene Formel für die Betragsquadrate der Fourierreihenkoeffizienten. Die abschließende Normierung in Zeile 79 liefert die Fourierreihenkoeffizienten in der Art, dass sich für den Gleichanteil `F_nu(1)` der Wert Eins ergibt.

11.5 Die kontinuierlichen Fensterfunktionen

Mit Hilfe der im vorigen Unterkapitel berechneten Fourierreihenkoeffizienten `F_nu` lässt sich die kontinuierliche Fensterfunktion $f(t)$ mit dem folgenden Programm für beliebige Zeitpunkte t des Intervalls $[0; 1)$, die als Vektor `t` zu übergeben sind, berechnen.

```

1: function f_t = fenster( t, F_nu )
2: f_t = zeros(1,length(t))
3: N_A = length(F_nu)
4: O_t = 2*pi*t
5: for nu = N_A-1:-1:1
6:   f_t = f_t + 2 * real(F_nu(nu+1)) * cos(O_t*nu) - ...
           2 * imag(F_nu(nu+1)) * sin(O_t*nu)
7: end
8: f_t = f_t + F_nu(1)

```

Da hier die Berechnung im wesentlichen so abläuft, wie dies bei der zweiten Version des in Unterkapitel 11.1 angegebenen Programms beschrieben ist, bedarf es bei diesem Programm keines weiteren Kommentars.

11.6 Die Spektren der kontinuierlichen Fensterfunktionen

Mit dem folgenden Programm lässt sich das Spektrum der kontinuierlichen Fensterfunktion bis zu sehr hohen Frequenzen mit fast der maximal möglichen Genauigkeit berechnen. Neben den im Unterkapitel 11.4 beschriebenen Parametern benötigt dieses Programm die Frequenzen ω , für die das Spektrum zu berechnen ist, als Zeilenvektor ω .

```

1: function F_omega = spektrum( N, A_0, A_1, s_0, omega )
2: c = (N/2)^(4/3)
3: Ms = -log(eps)/log(2) * 2^log(N/3)
4: Ms = 2^ceil(log(Ms)/log(2))
5: Ms = max(Ms,16)
6: Ms_OK = 0
7: while ~Ms_OK
8:     eta = 0:Ms/2
9:     F_eta = zeros(1,Ms/2+1)
10:    NF_1 = 1 / ( 1 + ( (N-1-A_1)/c )^2 )
11:    NF_0 = 1 ./ max( [ ( abs( s_0 + j*pi*(1-N) + 2*pi*c ) + ...
                        abs( s_0 + j*pi*(1-N) - 2*pi*c ) ) ; ...
                        ( abs( s_0 + j*pi*(N-1) + 2*pi*c ) + ...
                        abs( s_0 + j*pi*(N-1) - 2*pi*c ) ) ] ).^2
12:    for nu_1 = (1-N)/2:(N-1)/2
13:        F_eta_1 = ones(1,Ms/2+1)
14:        rho = 0
15:        for nu_2 = [(1-N+A_1):(nu_1+A_1-(N+1)/2), ...
                     (nu_1-A_1+(N+1)/2):(N-1-A_1)]
16:            K_1 = NF_1 * ( 1 + ( nu_2 / c )^2 )
17:            Psi_1 = atan( nu_2 / c )
18:            F_eta_1 = F_eta_1 .* K_1 .* sin( pi/Ms*eta - Psi_1 ).^2
19:            rho = rho + 1
20:            if rho < N-0.5-2*A_1-A_0
21:                Z_Z = 2*pi*c + s_0(rho) + j*2*pi*nu_1
22:                Z_N = 2*pi*c - s_0(rho) - j*2*pi*nu_1
23:                abs_Z_Z = abs( Z_Z )
24:                abs_Z_N = abs( Z_N )
25:                F_eta_1 = F_eta_1 .* NF_0(rho) .* ...
                    ( (abs_Z_N-abs_Z_Z)^2 + 4 * abs_Z_N * abs_Z_Z * ...
                      sin( pi/Ms*eta-(angle(Z_Z)-angle(Z_N))/2 ) ).^2 )
26:            end
27:        end
28:        F_eta = F_eta + F_eta_1
29:    end

```

```

30: NF = 1 / sqrt( max(F_eta) * min(F_eta) )
31: F_eta = NF * F_eta
32: L_eta = log(F_eta)
33: Ceps_2 = ifft( [L_eta,L_eta(Ms/2:-1:2)] )
34: Ceps_2 = real( Ceps_2 )
35: Ceps_2 = ( Ceps_2 + Ceps_2([1,Ms:-1:2]) ) / 2
36: sockel = eps / Ms / sqrt(48) * ...
           sqrt( sum( max( [L_eta,L_eta(Ms/2:-1:2)].^2, 1 ) ) )
37: grenze = ( 2 * (2*N-2-2*A_1-A_0) ) ./ [1:Ms/2] .* ...
           ( Ms * sockel/2/(2*N-2-2*A_1-A_0) ).^( 2*[1:Ms/2]/Ms )
38: max_fehl = sum( max( abs([L_eta,L_eta(Ms/2:-1:2)]), 1 ) ) * ...
           eps * ( 2 + log(Ms)/log(2) ) / Ms
39: grenze = max( grenze, max_fehl )
40: if any( abs(Ceps_2(2:Ms/2+1)) > grenze )
41:     kriterium = abs( fft( Ceps_2(Ms/4+1:Ms/2) .* [Ms/4:Ms/2-1] ) )
42:     [dummy,womax] = max( kriterium(1:Ms/8+1) )
43:     delta_c = 1-(womax-1)*16/Ms
44:     delta_c = sin( delta_c*pi*(1-delta_c^2/2) )/3
45:     c = c * (1-delta_c) / (1+delta_c)
46:     Ms = 2*Ms
47: else
48:     Ms_OK = 1
49: end
50: end
51: len_o = length(omega)
52: Omega_s = 2 * atan(omega/(2*pi*c))
53: phi = zeros(1,len_o)
54: for nu_i = 1:len_o
55:     phi(nu_i) = sin(Omega_s(nu_i)*[Ms/2-1:-1:1]) * Ceps_2(Ms/2:-1:2).'
56: end
57: phi = phi - ( N - 1 - A_1 - A_0/2 ) * Omega_s + omega/2
58: F_omega = zeros(1,len_o+1)
59: if N == 2*A_1+1
60:     NF_1 = 1
61: else
62:     NF_1 = exp(2*log(8)-sum(log([5:2:4*N-8*A_1])))/(N-1-2*A_1))
63: end
64: NF_0 = 1 ./ max( [ abs( s_0 + j*pi*(N-1) ).^2 ; ...
                    abs( s_0 - j*pi*(N-1) ).^2 ] )

```

```

65: for nu_1 = (1-N)/2:(N-1)/2
66:   omega_nu_1 = [omega,0]/(2*pi) - nu_1
67:   nu_4 = 2 * round( omega_nu_1 + (N-1)/2 ) - N + 1
68:   nu_4 = ( 1-N+2*A_1 ) .* ( nu_4 < 2*A_1-N ) + ...
        nu_4 .* ( abs(nu_4) < N-2*A_1 ) + ...
        ( N-1-2*A_1 ) .* ( nu_4 > N-2*A_1 )
69:   nu_4 = nu_4/2
70:   d_omega = omega_nu_1 - nu_4
71:   F_omega_1 = ( abs(d_omega) < eps )
72:   F_omega_1 = F_omega_1 / eps^2 + (~F_omega_1) .* ...
        ( sin(pi*d_omega) ./ ...
        ( pi*eps*d_omega + F_omega_1 ) ).^2
73:   rho = 0
74:   for nu_2 = (1-N)/2+A_1:(N-3)/2-A_1
75:     rho = rho + 1
76:     if rho < N-0.5-2*A_1-A_0
77:       F_omega_1 = F_omega_1 .* NF_0(rho) .* ...
        abs( j*2*pi*omega_nu_1 - s_0(rho) ).^2
78:     end
79:     nu_3 = nu_2 + ( nu_2 >= nu_4 )
80:     F_omega_1 = F_omega_1 ./ ( NF_1 .* ( omega_nu_1 - nu_3 ) .^ 2 )
81:   end
82:   F_omega = F_omega + F_omega_1
83: end
84: omega_nu_1 = abs( [omega,0] / (2*pi) ) - N + A_1 + 0.5
85: omega_nu_1 = 0.5 * ( omega_nu_1 > 0 ) .* omega_nu_1 + 0.25
86: omega_nu_1 = 1 - 2 * ( ( omega_nu_1 - floor(omega_nu_1) ) > 0.5 )
87: F_omega = omega_nu_1 .* sqrt(F_omega)
88: F_omega = F_omega(1:len_o) / F_omega(len_o+1)
89: F_omega = F_omega .* exp(-j*phi)

```

Die Berechnung des Cepstrums in den Zeilen 2 bis 50 erfolgt identisch, wie bei der in Unterkapitel 11.4 beschriebenen Berechnung der Fourierreihenkoeffizienten des kontinuierlichen Fensters. Wenn man mit der Substitution nach Gleichung (10.26) den Grenzübergang $M \rightarrow \infty$ betrachtet, erhält man aus dem in Unterkapitel 11.2 beschriebenen Programm das hier vorliegende Programm zur Berechnung des Spektrums des kontinuierlichen Fensters für beliebige Frequenzen. Die Kommentare der beiden ebengenannten Unterkapitel lassen sich daher auf dieses Programm sinngemäß übertragen. Eine Besonderheit ist lediglich dadurch gegeben, dass teilweise die auf 2π normierte Kreisfrequenz — also das, was man üblicherweise als Frequenz bezeichnet — verwendet wird.

11.7 Die kontinuierlichen Fensterautokorrelationsfunktionen

Dieses Programm benötigt die Zeitpunkte t , für die die Fensterautokorrelationsfunktion berechnet werden soll, als Vektor $\mathbf{t}$. Desweiteren muss der Zeilenvektor $\mathbf{F_nu}$ die Fourierreihenkoeffizienten $F_\infty(\nu \cdot 2\pi)$ der Fensterfunktion für $\nu = 0$ (1) $N - A_1 - 1$ enthalten. Diesen Vektor erhält man bei dem Programm in Unterkapitel 11.4 als Ergebnis.

```

1: function d_t = fenster_akf( t, F_nu )
2: N_A = length(F_nu)
3: t = abs(t)
4: 0_t = 2*pi*t
5: d_t = zeros(1,length(t))
6: for nu = N_A:-1:2
7:   F_sin = real( F_nu(nu)' * ...
               sum( [F_nu(N_A:-1:1)';F_nu([2:nu-1,nu+1:N_A]).'] ./ ...
               [2-N_A-nu:-1,1:N_A-nu].') ) / pi
8:   d_t = d_t + 2 * F_sin * sin(0_t*(nu-1)) + ...
               2 * abs(F_nu(nu))^2 * cos(0_t*(nu-1)) .* (1-t)
9: end
10: d_t = d_t + abs(F_nu(1))^2 * (1-t)

```

In der Zeile 2 wird der benötigte Parameter N_A als die Differenz $N - A_1$ bestimmt. Da die Fourierreihe der Fensterfunktion nur Anteile bis zum $(N - A_1 - 1)$ -fachen der Grundkreisfrequenz enthält, kann der Wert für N_A als die Länge des Vektors $\mathbf{F_nu}$ bestimmt werden. Da es sich bei der Fensterautokorrelationsfunktion um eine gradesymmetrische Funktion handelt, kann man für negatives t ebenso gut die Werte für $-t$ berechnen. Durch Zeile 3 erspart man sich so eine Fallunterscheidung bei der späteren Berechnung der Fensterautokorrelationsfunktion. Diese Berechnung wird mit Hilfe der Formeln (10.54) und (10.55) als Überlagerung einer Sinusreihe und einer mit der Dreiecksfunktion multiplizierten Kosinusreihe vorgenommen, wie dies am Ende von Kapitel 10.7 beschrieben ist. Dabei wird bei jedem Reihenglied die mit t multiplizierte Grundkreisfrequenz benötigt, die daher in Zeile 4 für alle Werte des Vektors $\mathbf{t}$ auf dem Vektor $\mathbf{0_t}$ bereitgestellt wird. Die Summe in Gleichung (10.55) wird als `for`-Schleife, die in Zeile 6 beginnt und in Zeile 9 endet, realisiert, indem zunächst in Zeile 5 der Vektor $\mathbf{d_t}$ mit null initialisiert wird, zu dem in Zeile 8 bei jedem Schleifendurchlauf ein Anteil der Sinusreihe und ein Anteil der mit der Dreiecksfunktion multiplizierten Kosinusreihe addiert wird. In Zeile 7 wird der dazu benötigte Sinusreihenkoeffizient nach Gleichung (10.54) berechnet. Die dabei auftretende Summe wird mit dem MATLAB-Befehl `sum` berechnet.

Literaturverzeichnis

- [1] Repp, Helmut: *Ein Fenster zur gleichzeitigen Messung der Übertragungsfunktion eines realen Systems und des Leistungsdichtespektrums des überlagerten Rauschens am Systemausgang (Teil 1)*. arXiv:2511.NNNNN
- [2] Fisz, Marek: *Wahrscheinlichkeitsrechnung und mathematische Statistik*. VEB Deutscher Verlag der Wissenschaften, Berlin 1966
- [3] Cramér, Harald: *Mathematical Methodes of Statistics*. Princeton Univ. Press, Princeton 1996
- [4] Schüßler, Hans Wilhelm; Heinle, Frank: *Measuring the Properties of Implemented Systems*. Frequenz, Band 48, S. 3-7,1994
- [5] Heinle Frank: *Analyse und Kompensation von Quantisierungs- und Codierungsfehlern in implementierten Multiraten-Systemen*. Dissertation am Lehrstuhl für Nachrichtentechnik der Universität Erlangen-Nürnberg, Shaker Verlag, Aachen 1998
- [6] Fliege, Norbert: *Multiraten-Signalverarbeitung*. Teubner-Verlag, Stuttgart 1993
- [7] Heute, Ulrich: *Fehler in DFT und FFT*. Ausgewählte Arbeiten über Nachrichtensysteme, Lehrstuhl für Nachrichtentechnik der Universität Erlangen-Nürnberg, 1982
- [8] Autor unbekannt: *What Every Computer Scientist Should Know About Floating-Point-Arithmetic*. Datei mit dem Namen „floating-point.ps“ von SUN Microsystems, Inc., 2550 Garcia Avenue, Mountain View, CA 9403, U.S.A, Part. No.: 800-7895-10, Revision A, June 1992

Anhang

A.1 Identität der Lösungsräume der Gleichungssysteme (2.21) und (2.22)

Die rechten Seiten aller Gleichungen des Gleichungssystems (2.21) kann man umformen zu:

$$\begin{aligned}
 & \mathbb{E}\left\{\left(\mathbf{y}(k) - \mathbb{E}\{\mathbf{y}(k)\}\right) \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}\right)^H\right\} = \tag{A.1} \\
 &= \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H - \mathbf{y}(k) \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H - \mathbb{E}\{\mathbf{y}(k)\} \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H + \mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H\right\} = \\
 &= \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H\right\} - \mathbb{E}\left\{\mathbf{y}(k) \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H\right\} - \\
 &\quad - \mathbb{E}\left\{\mathbb{E}\{\mathbf{y}(k)\} \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H\right\} + \mathbb{E}\left\{\mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H\right\} = \\
 &= \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H\right\} - \mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H - \mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H + \mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H = \\
 &= \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H\right\} - \mathbb{E}\{\mathbf{y}(k)\} \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H = \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H\right\} - \mathbb{E}\left\{\mathbf{y}(k) \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H\right\} = \\
 &= \mathbb{E}\left\{\mathbf{y}(k) \cdot \tilde{\mathbf{V}}(\tilde{\mu})^H - \mathbf{y}(k) \cdot \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}^H\right\} = \mathbb{E}\left\{\mathbf{y}(k) \cdot \left(\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\}\right)^H\right\} \\
 &\quad \forall \quad k = 0 \text{ (1) } F-1 \quad \text{und} \quad \tilde{\mu} = 0 \text{ (1) } M-1.
 \end{aligned}$$

Da auf den linken Seiten aller Gleichungen der Gleichungssysteme (2.21) der Term $e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot k}$ der einzige ist, der von k abhängt, steht auf der linken Seite für jeden festen Wert von $\tilde{\mu}$ eine in k mit M periodische Folge von Zeilenvektoren. Die Anzahl der Elemente dieser Zeilenvektoren ist gleich der Anzahl der Elemente des Zufallsvektors $\tilde{\mathbf{V}}(\tilde{\mu})$ und sei mit R bezeichnet. Wenn nun für diesen festen Wert $\tilde{\mu}$ auf der rechten Seite der Gleichungssysteme (2.21) keine in k mit M periodische Folge von Zeilenvektoren steht, existiert keine Lösung für die Werte der beiden bifrequenten Übertragungsfunktionen. In diesem Fall macht es keinen Sinn, das reale System durch die beiden linearen und periodisch zeitvarianten Modellsysteme zu modellieren, und die Übertragungsfunktionen mit Hilfe des RKM messtechnisch abzuschätzen. Daher wollen wir uns auf den Fall beschränken, dass eine Lösung existiert, was nur sein kann, wenn für diesen festen Wert $\tilde{\mu}$ auch auf den

rechten Seiten der Gleichungssysteme (2.21) eine in k mit M periodische Folge von Zeilenvektoren steht. Alle Gleichungen, bei denen sich k um ein ganzzahliges Vielfaches von M unterscheidet, während $\tilde{\mu}$ gleich ist, sind somit identisch. Wenn wir nun eine Linearkombination dieser identischen Gleichungen bilden, wobei die Summe der Koeffizienten dieser Linearkombination von Null verschieden ist, ändert sich der Lösungsraum des Gleichungssystems nicht. Bevor wir diese Linearkombination durchführen, ersetzen wir zunächst die diskrete Zeitvariable k durch $k = \tilde{k} + \kappa \cdot M$ mit $\tilde{k} = 0 \ (1) \ M-1$, so dass alle Gleichungen mit gleichem $\tilde{k}$ für alle Werte von κ identisch sind, sofern die Zeitpunkte k im Intervall $[0, F-1]$ liegen. Die Linearkombination bilden wir, indem wir die Gleichungen mit den Werten der Fensterfolge $f(\tilde{k} + \kappa \cdot M)$ gewichten, und über alle κ aufsummieren. Da die Fensterfolge für $k \notin [0, F-1]$ nach Gleichung ([1]:2.15) Null ist, sind die Gleichungen mit $k \notin [0, F-1]$ nach der Gewichtung durch die Fensterfolge auch immer erfüllt, so dass $\kappa \in \mathbb{Z}$ gewählt werden kann. Aufgrund der nach Gleichung ([1]:2.27) geforderten Nullstellenlage des Spektrums der Fensterfolge ist die Summe $\sum_{\kappa=-\infty}^{\infty} f(\tilde{k} + \kappa \cdot M)$ immer gleich 1, so dass der Lösungsraum des Gleichungssystems durch diese Modifikation nicht verändert wird. Wir erhalten M^2 Gleichungen, die jeweils auf beiden Seiten einen Zeilenvektor der Dimension $1 \times R$ stehen haben.

$$\begin{aligned}
& \frac{1}{M} \cdot \sum_{\mu=0}^{M-1} \tilde{H}(\mu) \cdot \mathbb{E} \left\{ (\tilde{\mathbf{V}}(\mu) - \mathbb{E}\{\tilde{\mathbf{V}}(\mu)\}) \cdot (\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\})^H \right\} \cdot \\
& \quad \cdot \underbrace{\sum_{\kappa=-\infty}^{\infty} f(\tilde{k} + \kappa \cdot M) \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot \kappa \cdot M} \cdot e^{j \cdot \frac{2\pi}{M} \cdot \mu \cdot \tilde{k}}}_{=1} = \\
& = \mathbb{E} \left\{ \sum_{\kappa=-\infty}^{\infty} f(\tilde{k} + \kappa \cdot M) \cdot \mathbf{y}(\tilde{k} + \kappa \cdot M) \cdot (\tilde{\mathbf{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\mathbf{V}}(\tilde{\mu})\})^H \right\} \\
& \quad \forall \quad \tilde{k} = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ M-1. \tag{A.2}
\end{aligned}$$

Für jeden festen Wert von $\tilde{\mu}$ haben wir ein Gleichungssystem mit je M Gleichungen, die wir nun diskret fouriertransformieren indem wir jeweils die $\tilde{k}$ -te Gleichung mit dem Drehfaktor $e^{-j \cdot \frac{2\pi}{M} \cdot \tilde{\mu} \cdot \tilde{k}}$ multiplizieren und alle Gleichungen mit gleichem $\tilde{\mu}$ über $\tilde{k}$ aufsummieren. Wenn wir für die diskrete Frequenzvariable $\bar{\mu}$ im Exponenten des Drehfaktors $\bar{\mu} = 0 \ (1) \ M-1$ wählen, entspricht das — abgesehen von dem konstanten und von Null verschiedenen Faktor $\sqrt{M}$ — jeweils einer unitären Transformation der M Gleichungen mit gleichem $\tilde{\mu}$, so dass sich die Lösungsräume der Gleichungssysteme mit jeweils festem $\tilde{\mu}$ nicht ändern, und somit auch nicht der Lösungsraum aller M^2 Gleichungssysteme. Wenn man berücksichtigt, dass die dabei auf den linken Seiten der Gleichungen entstehenden Summen über die Exponentialfunktionen nur für $\mu = \bar{\mu}$ von Null verschieden — nämlich M — sind, erkennt man, dass von der Summe über μ nur jeweils nur der Summand mit

$\mu = \bar{\mu}$ übrig bleibt. Wenn man anschließend noch den Index $\bar{\mu}$ durch den Index μ , der dann ja nicht mehr in den Gleichungen auftritt, substituiert ergeben sich wieder M^2 Gleichungssysteme, bei denen jeweils auf beiden Seiten ein Zeilenvektor der Dimension $1 \times R$ steht.

$$\vec{H}(\mu) \cdot \mathbb{E} \left\{ (\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\}) \cdot (\tilde{\vec{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\vec{V}}(\tilde{\mu})\})^H \right\} = \mathbb{E} \left\{ \mathbf{Y}_f(\mu) \cdot (\tilde{\vec{V}}(\tilde{\mu}) - \mathbb{E}\{\tilde{\vec{V}}(\tilde{\mu})\})^H \right\} \\ \forall \quad \mu = 0 \ (1) \ M-1 \quad \text{und} \quad \tilde{\mu} = 0 \ (1) \ M-1. \quad (\text{A.3})$$

Diese Gleichungssysteme enthalten für jeden festen Wert vom μ genau M Gleichungen mit unterschiedlichen Werten von $\tilde{\mu}$. Im Gegensatz dazu weisen die Gleichungssysteme (2.22) für jedes μ genau je *eine* Gleichung mit einem $1 \times R$ Zeilenvektor auf jeder Seite auf. Diese Gleichung, ist eine der M Gleichungen (A.3), nämlich die, die man mit $\tilde{\mu} = \mu$ erhält. Daher kann der Lösungsraum des Gleichungssystems (2.22) allenfalls von einer höheren Dimension sein, wie der Lösungsraum des Gleichungssystems (A.3). Die Dimension des Lösungsraums der Vektorgleichung, die sich im Gleichungssystem (2.22) für den festen Wert vom μ ergibt, ist gleich dem Rang der $R \times R$ Kovarianzmatrix $\underline{C}_{\tilde{\vec{V}}(\mu), \tilde{\vec{V}}(\mu)}$ des Zufallsvektors $\tilde{\vec{V}}(\mu)$ gemäß Gleichung (2.23). Nach [3] liegt der Zufallsvektor $\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\}$ mit der Wahrscheinlichkeit Eins in einem Unterraum des R -dimensionalen komplexen Raums, dessen Dimension gleich dem Rang der Kovarianzmatrix $\underline{C}_{\tilde{\vec{V}}(\mu), \tilde{\vec{V}}(\mu)}$ ist. Da sich jede Kovarianzmatrix mit einer unitären Matrix $\underline{U}$ auf Diagonalform transformieren lässt (Hauptachsentransformation einer hermiteschen quadratischen Form), kann man aus dem Zufallsvektor $\tilde{\vec{V}}(\mu)$ durch eine unitäre Transformation einen neuen Zufallsvektor $\dot{\vec{V}}(\mu) = \underline{U} \cdot \tilde{\vec{V}}(\mu)$ bilden, dessen Kovarianzmatrix

$$\begin{aligned} \underline{C}_{\dot{\vec{V}}(\mu), \dot{\vec{V}}(\mu)} &= \mathbb{E} \left\{ (\dot{\vec{V}}(\mu) - \mathbb{E}\{\dot{\vec{V}}(\mu)\}) \cdot (\dot{\vec{V}}(\mu) - \mathbb{E}\{\dot{\vec{V}}(\mu)\})^H \right\} = \\ &= \mathbb{E} \left\{ (\underline{U} \cdot \tilde{\vec{V}}(\mu) - \mathbb{E}\{\underline{U} \cdot \tilde{\vec{V}}(\mu)\}) \cdot (\underline{U} \cdot \tilde{\vec{V}}(\mu) - \mathbb{E}\{\underline{U} \cdot \tilde{\vec{V}}(\mu)\})^H \right\} = \\ &= \mathbb{E} \left\{ \underline{U} \cdot (\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\}) \cdot (\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\})^H \cdot \underline{U}^H \right\} = \\ &= \underline{U} \cdot \mathbb{E} \left\{ (\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\}) \cdot (\tilde{\vec{V}}(\mu) - \mathbb{E}\{\tilde{\vec{V}}(\mu)\})^H \right\} \cdot \underline{U}^H = \underline{U} \cdot \underline{C}_{\tilde{\vec{V}}(\mu), \tilde{\vec{V}}(\mu)} \cdot \underline{U}^H \end{aligned} \quad (\text{A.4})$$

eine Diagonalmatrix ist, und deren Hauptdiagonalelemente die Varianzen der neuen Zufallsgrößen des transformierten Zufallsvektors $\dot{\vec{V}}(\mu)$ sind. Da der Rang dieser Diagonalmatrix gleich dem Rang der ursprünglichen Kovarianzmatrix ist, ist die Anzahl der transformierten Zufallsgrößen, die die Varianz Null aufweisen, gleich dem Rangdefekt der ursprünglichen Kovarianzmatrix $\underline{C}_{\tilde{\vec{V}}(\mu), \tilde{\vec{V}}(\mu)}$. Die Kovarianz jeder beliebigen Zufallsgröße mit einer der Zufallsgrößen $\dot{\vec{V}}(\mu)$ mit verschwindender Varianz ist immer Null. Nun fassen wir

alle Zufallsgrößen des Spektrums der Erregung, sowie die konjugierten Zufallsgrößen des Spektrums der Erregung, die in den Zufallsspaltenvektoren $\vec{V}(\tilde{\mu})$ für einen festen Wert vom μ für alle $\tilde{\mu} = 0 \ (1) \ M-1$ in den Gleichungssystem (A.3) vorhanden sind, zu einem Zufallsspaltenvektor $\vec{V}(\mu)$ zusammen. Die Anzahl der Elemente dieses Zufallsvektors bezeichnen wir mit r . Da die Gleichung des Gleichungssystems (2.22) für den festen Wert den μ gerade eine der Gleichungen (A.3) — nämlich die mit $\tilde{\mu} = \mu$ — ist, sind in dem Zufallsvektor $\vec{V}(\mu)$ auch die Zufallsgrößen des Zufallsvektors $\vec{V}(\mu)$ vorhanden. Wenn wir nun die $R \times r$ Kovarianzmatrix

$$\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)} = E \left\{ (\vec{V}(\mu) - E\{\vec{V}(\mu)\}) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} \quad (\text{A.5})$$

betrachten, erkennen wir, dass alle Kovarianzen, die mit einer der Zufallsgrößen $\vec{V}(\mu)$ mit verschwindender Varianz gebildet werden, immer Null sind, so dass in den entsprechenden Zeilen von $\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}$ Nullvektoren mit r Elementen zu finden sind. Weil die Anzahl der Nullvektoren gleich dem Rangdefekt der Kovarianzmatrix $\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}$ ist, und weil die Zufallsgrößen des Zufallsvektors $\vec{V}(\mu)$ im Zufallsvektor $\vec{V}(\mu)$ enthalten sind, ist der Rang dieser $R \times r$ Kovarianzmatrix $\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}$ gleich dem Rang der $R \times R$ Kovarianzmatrix $\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}$ des Zufallsvektors $\vec{V}(\mu)$. Wenn man die $R \times r$ Kovarianzmatrix $\underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}$ von links mit dem konjugiert Transponierten $\underline{U}^H$ der unitären Transformationsmatrix multipliziert, erhält man eine $R \times r$ Kovarianzmatrix desselben Rangs:

$$\begin{aligned} \underline{U}^H \cdot \underline{C}_{\vec{V}(\mu), \vec{V}(\mu)} &= \underline{U}^H \cdot E \left\{ (\vec{V}(\mu) - E\{\vec{V}(\mu)\}) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} = \\ &= E \left\{ (\underline{U}^H \cdot \vec{V}(\mu) - E\{\underline{U}^H \cdot \vec{V}(\mu)\}) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} = \\ &= E \left\{ (\vec{V}(\mu) - E\{\vec{V}(\mu)\}) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} = \underline{C}_{\vec{V}(\mu), \vec{V}(\mu)}. \end{aligned} \quad (\text{A.6})$$

Diese Kovarianzmatrix tritt im Gleichungssystem (A.3) auf, wenn man alle M Gleichungen — mit jeweils einem Zeilenvektor der Dimension $1 \times R$ auf beiden Seiten — für einen festen Wert vom μ und für unterschiedliche Werte von $\tilde{\mu}$ jeweils zu einer Gleichung mit einem Zeilenvektor der Dimension $1 \times r$ auf beiden Seiten zusammenfasst, wobei man eventuell mehrfach vorhandene, identische Gleichungen für die Vektorelemente nur einfach berücksichtigt, und die Reihenfolge der Elemente der Zeilenvektoren auf beiden Seiten der Gleichung in geeigneter Weise permutiert, so dass sich dieselbe Reihenfolge ergibt wie bei den Elementen des Zufallsvektors $\vec{V}(\mu)^H$.

$$\begin{aligned} \vec{H}(\mu) \cdot E \left\{ (\vec{V}(\mu) - E\{\vec{V}(\mu)\}) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} &= E \left\{ \mathbf{Y}_f(\mu) \cdot (\vec{V}(\mu) - E\{\vec{V}(\mu)\})^H \right\} \\ \forall \quad \mu &= 0 \ (1) \ M-1 \end{aligned} \quad (\text{A.7})$$

Jede dieser M Gleichungen mit dem Parameter μ besitzt daher denselben Lösungsraum wie die entsprechende Gleichung des Gleichungssystems (2.22) mit demselben Wert μ . Da die Herleitung für alle Frequenzen μ gilt, ist die Identität der Lösungsräume der Gleichungssysteme (2.21) und (2.22) gezeigt.

A.2 Zur Konditionierung der empirischen Kovarianzmatrix

Bei der Berechnung der Messwerte des RKM ist die empirisch gewonnene Kovarianzmatrix einiger Werte des Spektrums der Erregung zu invertieren. Da solch eine Matrix aus einer Stichprobe vom Umfang L der Spektralwerte der Erregung gewonnen wird, hängt deren Konditionierung davon ab, welche konkrete Stichprobe man gerade gezogen hat. Wir wollen daher eine obere Grenze für die Wahrscheinlichkeit herleiten, eine schlecht konditionierte Matrix zu erhalten, und zeigen, dass diese Grenze mit steigendem Umfang L der Stichprobe mindestens indirekt proportional abfällt, wenn man die zufälligen Spektralwerte so wählt, dass deren theoretische Kovarianzmatrix gut konditioniert ist. Es sei darauf hingewiesen, dass die nun folgende Betrachtung auch für die Varianz einer einzigen Zufallsgröße gilt, da es sich dabei um den Fall einer 1×1 Kovarianzmatrix handelt. In Gegensatz zu [1] wird hier der Fall mittelwertbehafteter Zufallsgrößen behandelt.

Gegeben seien R komplexe Zufallsgrößen, die den Zufallsspaltenvektor $\vec{V}$ bilden, und deren Momente bis zur vierten Ordnung alle existieren sollen. Ein Teil der zweiten Momente sind die Elemente der theoretischen $R \times R$ Kovarianzmatrix

$$\underline{C}_{\vec{V}, \vec{V}} = E\left\{(\vec{V} - E\{\vec{V}\}) \cdot (\vec{V} - E\{\vec{V}\})^H\right\} = \underline{C}_{\vec{V}, \vec{V}}^H, \quad (\text{A.8})$$

die immer hermitesch und positiv semidefinit ist. Sie lässt sich daher unitär kongruent auf Diagonalform transformieren, wobei die Diagonalelemente alle nichtnegativ reell sind, und mit dem Zeilenindex (= Spaltenindex) des Diagonalelements monoton fallen. Die Diagonalelemente sind die Singulärwerte s_i mit $i = 1(1)R$ der Kovarianzmatrix, die zugleich deren Eigenwerte sind. Es gilt $s_i \geq s_{i+1}$. Je näher die Konditionszahl bei eins liegt, desto besser ist die Kovarianzmatrix konditioniert, und desto genauer lässt sich deren Inverse berechnen. Jeder beliebige Vektor wird durch die Multiplikation mit der Kovarianzmatrix auf einen Vektor abgebildet, dessen euklidische Norm $\|\dots\|_2$ die Bedingung

$$s_R \cdot \|\vec{x}\|_2 \leq \|\underline{C}_{\vec{V}, \vec{V}} \cdot \vec{x}\|_2 \leq s_1 \cdot \|\vec{x}\|_2 \quad (\text{A.9})$$

erfüllt. Da es sich bei der Spektralnrm um eine mit der euklidischen Vektornorm kompatible Matrixnorm handelt, gilt außerdem für jeden beliebigen Vektor:

$$\|\underline{C}_{\vec{V}, \vec{V}} \cdot \vec{x}\|_2 \leq \|\underline{C}_{\vec{V}, \vec{V}}\|_2 \cdot \|\vec{x}\|_2. \quad (\text{A.10})$$

Die $R \times L$ Matrix $\underline{V}$ enthalte eine konkrete Stichprobe vom Umfang L des Zufallsvektors $\vec{V}$, d. h. jede Spalte dieser Matrix ist ein Element der Stichprobe, also eine konkrete Realisierung des Zufallsvektors $\vec{V}$, und es wurden insgesamt L konkrete Realisierungen dieses Zufallsvektors zu einer Matrix zusammengefasst. Bei der Matrix handelt es sich um eine konkrete Realisierung der Matrix $\underline{V}$ der mathematischen Stichprobe vom Umfang L des Zufallsvektors $\vec{V}$. Dies ist der Fall, wenn die Stichprobenentnahme in der Art erfolgt ist, dass jedes Element der Stichprobe — also jede Spalte der Matrix $\underline{V}$ — von jedem anderen Element unabhängig ist, und die gleiche Verbundverteilung besitzt, wie der Zufallsvektor $\vec{V}$. Aus jeder konkreten Stichprobenmatrix $\underline{V}$ kann man eine empirische Kovarianzmatrix berechnen.

$$\hat{\underline{C}}_{\vec{V},\vec{V}} = \frac{\underline{V} \cdot \underline{1}_{\perp} \cdot \underline{V}^H}{L - 1} = \underline{V} \cdot \frac{L \cdot \underline{E} - \vec{1}^H \cdot \vec{1}}{L \cdot (L - 1)} \cdot \underline{V}^H \quad (\text{A.11})$$

Nun benötigen wir noch eine weitere Matrixnorm $\|\dots\|_F$, nämlich die euklidische Matrixnorm, die auch Frobeniusnorm genannt wird, und die die Wurzel aus der Summe der Betragsquadrate aller Matrixelemente ist. Da auch diese Matrixnorm mit der euklidischen Vektornorm kompatibel ist, gilt auch mit dieser Matrixnorm für jeden beliebigen Vektor:

$$\|(\underline{C}_{\vec{V},\vec{V}} - \hat{\underline{C}}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 = \|(\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 \leq \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F \cdot \|\vec{x}\|_2. \quad (\text{A.12})$$

Zusätzlich benötigen wir noch die Dreiecksungleichung

$$\|\vec{x} + \vec{y}\|_2 \leq \|\vec{x}\|_2 + \|\vec{y}\|_2. \quad (\text{A.13})$$

Damit können wir abschätzen, dass die Norm des Produkts eines beliebigen Vektors mit der empirischen Kovarianzmatrix innerhalb eines bestimmten Intervalls liegen muss.

$$\begin{aligned} \|\hat{\underline{C}}_{\vec{V},\vec{V}} \cdot \vec{x}\|_2 &= \|\underline{C}_{\vec{V},\vec{V}} \cdot \vec{x} + (\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 \leq \\ &\leq \|\underline{C}_{\vec{V},\vec{V}} \cdot \vec{x}\|_2 + \|(\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 \leq \\ &\leq \|\underline{C}_{\vec{V},\vec{V}}\|_2 \cdot \|\vec{x}\|_2 + \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F \cdot \|\vec{x}\|_2 = (s_1 + \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F) \cdot \|\vec{x}\|_2 \end{aligned} \quad (\text{A.14})$$

$$\begin{aligned} \|\hat{\underline{C}}_{\vec{V},\vec{V}} \cdot \vec{x}\|_2 &= \|\underline{C}_{\vec{V},\vec{V}} \cdot \vec{x} - (\underline{C}_{\vec{V},\vec{V}} - \hat{\underline{C}}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 \geq \\ &\geq \|\underline{C}_{\vec{V},\vec{V}} \cdot \vec{x}\|_2 - \|(\underline{C}_{\vec{V},\vec{V}} - \hat{\underline{C}}_{\vec{V},\vec{V}}) \cdot \vec{x}\|_2 \geq \\ &\geq s_R \cdot \|\vec{x}\|_2 - \|\underline{C}_{\vec{V},\vec{V}} - \hat{\underline{C}}_{\vec{V},\vec{V}}\|_F \cdot \|\vec{x}\|_2 = (s_R - \|\underline{C}_{\vec{V},\vec{V}} - \hat{\underline{C}}_{\vec{V},\vec{V}}\|_F) \cdot \|\vec{x}\|_2 \end{aligned} \quad (\text{A.15})$$

Diese beiden Grenzen gelten für beliebige Vektoren, also auch für Vektoren, die auf ihre euklidische Norm normiert worden sind und daher die Länge Eins haben. Der kleinste Singulärwert $\hat{s}_R$ der empirischen Kovarianzmatrix ist die minimale Länge aller Vektoren,

die durch eine Multiplikation mit der empirischen Kovarianzmatrix aus allen Vektoren der Länge Eins entstanden sind. Die Länge des längsten Bildvektors ist entsprechend der größte Singulärwert $\hat{s}_1$ der empirischen Kovarianzmatrix. Falls die Frobeniusnorm der Abweichung der empirischen von der theoretischen Kovarianzmatrix kleiner als s_R ist, kann die Konditionszahl $\hat{K}_{\vec{V}}$ der empirischen Kovarianzmatrix mit den beiden letzten Ungleichungen abgeschätzt werden.

$$\hat{K}_{\vec{V}} = \frac{\hat{s}_1}{\hat{s}_R} \leq \frac{s_1 + \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F}{s_R - \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F} \quad (\text{A.16})$$

Wenn die Frobeniusnorm der Abweichung der empirischen von der theoretischen Kovarianzmatrix außerdem noch kleiner als

$$\|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F < \frac{n-1}{n+K_{\vec{V}}^{-1}} \cdot s_R < s_R \quad \text{mit} \quad n \in \mathbb{R} \quad \text{und} \quad n > 1 \quad (\text{A.17})$$

ist, kann man sicher sein, dass die Konditionszahl der empirischen Kovarianzmatrix höchstens n mal so groß wie die Konditionszahl der theoretischen Kovarianzmatrix ist:

$$\hat{K}_{\vec{V}} \leq \frac{s_1 + \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F}{s_R - \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F} \leq \frac{s_1 + \frac{n-1}{n+K_{\vec{V}}^{-1}} \cdot s_R}{s_R - \frac{n-1}{n+K_{\vec{V}}^{-1}} \cdot s_R} = n \cdot K_{\vec{V}}. \quad (\text{A.18})$$

Falls nun die Wahrscheinlichkeit, dass die Frobeniusnorm der Matrixdifferenz größer als die Schranke in der Ungleichung (A.17) ist, mit steigendem L gegen Null konvergiert, kann man durch eine Erhöhung von L erreichen, dass die Wahrscheinlichkeit, eine empirische Kovarianzmatrix zu erhalten, die n mal schlechter als die theoretische Kovarianzmatrix konditioniert ist, unter einer beliebig kleinen tolerierbaren Schwelle bleibt. Wir wollen daher eine obere Schranke für diese Wahrscheinlichkeit herleiten.

Dazu benötigen wir einen Satz, aus dem sich auch die Tschebyscheffsche Ungleichung ableiten lässt, und den wir aus [2] entnehmen.

Satz: Nimmt eine zufällige Veränderliche $\mathbf{Y}$ nur nichtnegative Werte an, und besitzt sie einen endlichen Mittelwert $E\{\mathbf{Y}\}$, so ist für jede positive Zahl K die Ungleichung

$$P(\mathbf{Y} \geq K) \leq \frac{E\{\mathbf{Y}\}}{K}$$

erfüllt.

Da das Quadrat der Frobeniusnorm einer zufälligen Matrix solch eine zufällige Veränderliche ist, erhalten wir mit

$$\mathbf{Y} = \|\hat{\underline{C}}_{\vec{V},\vec{V}} - \underline{C}_{\vec{V},\vec{V}}\|_F^2 \quad \text{und mit} \quad K = \left(\frac{n-1}{n+K_{\vec{V}}^{-1}} \cdot s_R \right)^2$$

die gesuchte obere Grenze

$$\begin{aligned}
 P\left(\hat{\mathbf{K}}_{\vec{\mathbf{V}}} \geq n \cdot K_{\vec{\mathbf{V}}}\right) &\leq \\
 &\leq P\left(\|\hat{\underline{\mathbf{C}}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}} - \underline{\mathbf{C}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}}\|_{\text{F}} \geq \frac{n-1}{n+K_{\vec{\mathbf{V}}}^{-1}} \cdot s_R\right) = \\
 &= P\left(\|\hat{\underline{\mathbf{C}}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}} - \underline{\mathbf{C}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}}\|_{\text{F}}^2 \geq \left(\frac{n-1}{n+K_{\vec{\mathbf{V}}}^{-1}} \cdot s_R\right)^2\right) \leq \\
 &\leq \left(\frac{n+K_{\vec{\mathbf{V}}}^{-1}}{(n-1) \cdot s_R}\right)^2 \cdot \text{E}\left\{\|\hat{\underline{\mathbf{C}}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}} - \underline{\mathbf{C}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}}\|_{\text{F}}^2\right\}
 \end{aligned} \tag{A.19}$$

für die Wahrscheinlichkeit, dass die Frobeniusnorm der Differenz der empirischen und der theoretischen Kovarianzmatrix oberhalb der zulässigen Schwelle liegt, die gleichzeitig eine obere Grenze für die Wahrscheinlichkeit ist, dass die Konditionszahl der empirischen Kovarianzmatrix höchstens n mal so groß ist, wie die Konditionszahl der theoretischen Kovarianzmatrix. Um diese Grenze angeben zu können, müssen wir den Erwartungswert des Quadrats der Frobeniusnorm der Differenz der empirischen und der theoretischen Kovarianzmatrix berechnen. Nach der Definition der Frobeniusnorm ergibt sich diese als die Wurzel aus der Summe der Betragsquadrate aller Matrixelemente. Der Erwartungswert des Quadrats der Frobeniusnorm ist daher die Summe der Erwartungswerte der Betragsquadrate aller zufälligen Elemente der Matrixdifferenz.

Mit $\vec{\mathbf{V}}_i$ sei die mathematische Stichprobe vom Umfang L der i -ten Zufallsgröße des Zufallsvektors $\vec{\mathbf{V}}$, also die i -te Zeile der Zufallsmatrix $\underline{\mathbf{V}}$ bezeichnet. Das Element in der i -ten Zeile und j -ten Spalte der empirischen Kovarianzmatrix ist

$$\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j} = \vec{\mathbf{V}}_i \cdot \frac{\underline{\mathbf{1}} \cdot \underline{\mathbf{1}}}{L-1} \cdot \vec{\mathbf{V}}_j^{\text{H}}. \tag{A.20}$$

Der Erwartungswert dieses Elements ist

$$\begin{aligned}
 \text{E}\left\{\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j}\right\} &= \text{E}\left\{\vec{\mathbf{V}}_i \cdot \frac{\underline{\mathbf{1}} \cdot \underline{\mathbf{1}}}{L-1} \cdot \vec{\mathbf{V}}_j^{\text{H}}\right\} = \\
 &= \frac{\text{spur}(\underline{\mathbf{1}} \cdot \underline{\mathbf{1}})}{L-1} \cdot \text{E}\{\mathbf{V}_i \cdot \mathbf{V}_j^*\} + \frac{\vec{\mathbf{1}} \cdot \underline{\mathbf{1}} \cdot \vec{\mathbf{1}}^{\text{H}} - \text{spur}(\underline{\mathbf{1}} \cdot \underline{\mathbf{1}})}{L-1} \cdot \text{E}\{\mathbf{V}_i\} \cdot \text{E}\{\mathbf{V}_j^*\} = \\
 &= \text{E}\{\mathbf{V}_i \cdot \mathbf{V}_j^*\} - \text{E}\{\mathbf{V}_i\} \cdot \text{E}\{\mathbf{V}_j^*\} = \mathbf{C}_{\mathbf{V}_i, \mathbf{V}_j},
 \end{aligned} \tag{A.21}$$

und somit gleich dem entsprechenden Element der theoretischen Kovarianzmatrix $\underline{\mathbf{C}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}}$. Die Varianz des Elements in der i -te Zeile und j -ten Spalte der empirischen Kovarianzmatrix ist daher gleich dem Erwartungswert des Betragsquadrats des Elements in der i -te

Zeile und j -ten Spalte der Differenz der empirischen und der theoretischen Kovarianzmatrix. Die Varianz dieses Matrixelements berechnet sich zu

$$\begin{aligned}
 & \mathbb{E}\left\{\left|\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j} - \mathbb{E}\{\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j}\}\right|^2\right\} = \mathbb{E}\left\{\left|\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j} - \mathbf{C}_{\mathbf{V}_i, \mathbf{V}_j}\right|^2\right\} = \quad (\text{A.22}) \\
 & = \mathbb{E}\left\{\left|\vec{\mathbf{V}}_i \cdot \frac{\mathbf{1}_{\perp}}{L-1} \cdot \vec{\mathbf{V}}_j^H - \mathbb{E}\left\{(\mathbf{V}_i - \mathbb{E}\{\mathbf{V}_i\}) \cdot (\mathbf{V}_j - \mathbb{E}\{\mathbf{V}_j\})^*\right\}\right|^2\right\} = \\
 & = \frac{1}{L} \cdot \mathbb{E}\left\{\left|(\mathbf{V}_i - \mathbb{E}\{\mathbf{V}_i\}) \cdot (\mathbf{V}_j - \mathbb{E}\{\mathbf{V}_j\})^* - \mathbb{E}\left\{(\mathbf{V}_i - \mathbb{E}\{\mathbf{V}_i\}) \cdot (\mathbf{V}_j - \mathbb{E}\{\mathbf{V}_j\})^*\right\}\right|^2\right\} + \\
 & \quad + \frac{1}{L \cdot (L-1)} \cdot \left(\left|\mathbb{E}\left\{(\mathbf{V}_i - \mathbb{E}\{\mathbf{V}_i\}) \cdot (\mathbf{V}_j - \mathbb{E}\{\mathbf{V}_j\})\right\}\right|^2 + \right. \\
 & \quad \left. + \mathbb{E}\left\{|\mathbf{V}_i - \mathbb{E}\{\mathbf{V}_i\}|^2\right\} \cdot \mathbb{E}\left\{|\mathbf{V}_j - \mathbb{E}\{\mathbf{V}_j\}|^2\right\}\right).
 \end{aligned}$$

Bei dieser Berechnung wurde zunächst die bilineare Form $\vec{\mathbf{V}}_i \cdot \mathbf{1}_{\perp} \cdot \vec{\mathbf{V}}_j^H$ als Doppelsumme geschrieben, so dass sich durch die Bildung des Betragsquadrats unter anderem eine Vierfachsumme ergibt. Der Erwartungswert dieser Vierfachsumme lässt sich durch eine Fallunterscheidung bezüglich der Indizes der Vierfachsumme — wie in Tabelle [1]:A.1 des Kapitels [1]:A.5 — einigermaßen umfangreich, aber nicht allzu anspruchsvoll, berechnen. Die Varianz der Elemente der empirischen Kovarianzmatrix nimmt also mit steigendem L wenigstens mit α/L ab, wobei α eine positive Konstante ist, die nur von den zweiten und vierten zentralen Momenten des erregenden Zufallsprozessspektrums abhängt. Da die Summe der Varianzen aller Matrixelemente der Erwartungswert des Quadrats der Frobeniusnorm der Differenz der empirischen und der theoretischen Kovarianzmatrix ist, und die Dimension dieser Matrix mit $R \times R$ von L unabhängig ist, nimmt auch der Erwartungswert des Quadrats der Frobeniusnorm mindestens mit der Ordnung $1/L$ ab.

$$\mathbb{E}\left\{\|\hat{\mathbf{C}}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}} - \mathbf{C}_{\vec{\mathbf{V}}, \vec{\mathbf{V}}}\|_{\text{F}}^2\right\} = \sum_{i=1}^L \sum_{j=1}^L \mathbb{E}\left\{\left|\hat{\mathbf{C}}_{\mathbf{V}_i, \mathbf{V}_j} - \mathbf{C}_{\mathbf{V}_i, \mathbf{V}_j}\right|^2\right\} \sim \frac{1}{L} \quad (\text{A.23})$$

Setzt man dies in die Ungleichung (A.19) ein, so sieht man, dass auch die Wahrscheinlichkeit, eine Konditionszahl der empirischen Kovarianzmatrix zu erhalten, die mehr als n mal so groß ist wie die Konditionszahl der theoretischen Kovarianzmatrix, wenigstens mit $1/L$ abnimmt.

Es sei noch angemerkt, dass dies nur eine *obere Schranke* für die Wahrscheinlichkeit ist, dass die empirische Kovarianzmatrix eine mehr als n -fache Konditionszahl der theoretischen Kovarianzmatrix aufweist. Da dieser oberen Schranke eine Vielzahl von Ungleichungen zugrundeliegen, ist anzunehmen, dass die wahre Wahrscheinlichkeit, deren Berechnung — wenn überhaupt — nur bei Kenntnis der gemeinsamen Verbundverteilung aller Elemente des Zufallsvektors $\vec{\mathbf{V}}$ möglich wäre, wesentlich kleiner ist als die angegebene obere

Schranke. Die hier gemachte Herleitung hat den Vorteil, dass keine Aussage über die genaue Verbundverteilung benötigt wird. Es genügt, wenn die theoretische Kovarianzmatrix gut konditioniert ist, um zu gewährleisten, dass auch die empirische Kovarianzmatrix mit großer Wahrscheinlichkeit brauchbar konditioniert ist, wenn man L nur groß genug wählt.

A.3 Zur Wahrscheinlichkeit der empirischen Varianz Null

Für den Fall, dass man wie in [1] in Bild 1.1 nur ein lineares zeitinvariantes Modellsystem mit der Übertragungsfunktion $H(\Omega)$ ansetzt, aber dennoch die deterministische Modellstörung $u(k)$ beibehält, entartet bei der Berechnung der Messwerte der Übertragungsfunktion nach Gleichung (3.8) die Kovarianzmatrix $\hat{\mathbf{C}}_{\vec{\mathbf{V}}(\mu), \vec{\mathbf{V}}(\mu)}$ zu einer 1×1 Matrix $\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$, deren einziges Element die empirische Varianz des Spektralwertes $\mathbf{V}(\mu)$ der Erregung ist. Hier soll nun der Fall untersucht werden, dass diese Varianz zu null wird.

Jedes Element des konkreten Stichprobenvektors $\vec{\mathbf{V}}(\mu)$ wurde zufällig aus der Zufallsgröße $\mathbf{V}(\mu)$ gewonnen, und daher kann $\vec{\mathbf{V}}(\mu)$ als eine konkrete Realisierung des Zufallsvektors $\vec{\mathbf{V}}(\mu)$ der mathematischen Stichprobe von Umfang L betrachtet werden. Wir betrachten also nicht eine konkret durchgeführte Messung mit ihren Messergebnissen, die aus den konkreten Stichprobenvektoren gewonnen wurden, sondern die stochastischen Eigenschaften des Messverfahrens und der Messergebnisse, die man aus den mathematischen Stichprobenvektoren bestimmt. Es ist nicht unmöglich, nämlich dann, wenn alle Elemente der konkreten Stichprobe vom Umfang L der Zufallsgröße $\mathbf{V}(\mu)$ gleich sind, dass ein oder mehrere der dann ebenfalls zufälligen M empirischen Varianzen $\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$ den konkreten Wert Null annehmen. In diesem Fall werden die Messwerte $\hat{H}(\mu)$ und damit auch die Messwerte $\hat{U}_f(\mu)$ undefiniert. Es existiert dann ein komplexer ein- oder mehrdimensionaler Lösungsraum für die Messwerte $\hat{H}(\mu)$ und $\hat{U}_f(\mu)$. Um auch in diesem Fall eindeutige Messwerte zu erhalten gehen wir folgendermaßen vor. Wenn die empirischen Varianz kleiner als eine beliebig kleine aber echt positive Konstante Υ^{-1} ist, verwenden wir statt der Inversen der empirischen Varianz den konstanten endlichen Wert Υ . Die Inverse der empirischen Varianz ist eine Funktion des Zufallsvektors $\vec{\mathbf{V}}(\mu)$ der mathematischen Stichprobe des Spektrums der Erregung, deren Definitionsbereich den Raum, in dem die empirischen Varianz Null wird, nicht mit einschließt. Durch die mit Υ eingeführte Modifikation ist diese Funktion des Zufallsvektors $\vec{\mathbf{V}}(\mu)$ dann auf dem gesamten Raum definiert und endlich. Die Wahrscheinlichkeit, dass wir die Inverse der empirischen Varianz begrenzen müssen, hängt neben der Wahl von Υ auch von der Art — genauer gesagt von der Verteilung — der Zufallsgröße $\mathbf{V}(\mu)$ ab. Im Fall einer diskreten Zufallsgröße $\mathbf{V}(\mu)$, ergibt sich —

wie wir gleich sehen werden — die unerfreuliche Tatsache, dass die Wahrscheinlichkeit, dass die empirische Varianz so klein wird, dass eine Begrenzung ihrer Inversen erforderlich wird, bei einer Vergrößerung von Υ nicht unter einen Minimalwert sinkt, der echt größer als Null ist. Dieser Fall ist der für praktische Anwendungen interessante Fall, weil man üblicherweise die Zufallsgröße $\mathbf{V}(\mu)$ an einem Rechner generiert. Da an Rechnern nur abzählbar endlich viele Zahlen dargestellt werden können, ist der Ergebnisraum der Zufallsgröße des Spektrums der Erregung immer ein abzählbarer endlicher Unterraum der am Rechner darstellbaren Zahlen. Somit ist die Zufallsgröße $\mathbf{V}(\mu)$ diskret. Unabhängig davon, wie groß Υ gewählt worden ist, ist die Wahrscheinlichkeit, dass alle Werte der konkreten Stichprobe vom Umfang L der Zufallsgröße $\mathbf{V}(\mu)$ gleich sind, und somit die empirische Varianz Null wird, so dass ihre Inverse zu begrenzen ist, immer echt positiv. Wir wollen nun die Wahrscheinlichkeit des Eintretens dieses Falls berechnen, wobei wir annehmen wollen, dass Υ so groß gewählt worden ist, dass in keinem anderen Fall, als in dem, dass alle Werte der konkreten Stichprobe vom Umfang L der Zufallsgröße $\mathbf{V}(\mu)$ gleich sind, eine Begrenzung erfolgt. Die Wahrscheinlichkeit $P(\mathbf{V}(\mu)=V)$, dass die Zufallsgröße einen bestimmten Wert des Ergebnisraums annimmt, lässt sich als Grenzwert aus der Verteilung der Zufallsgröße $\mathbf{V}(\mu)$ berechnen.

$$\begin{aligned} P(\mathbf{V}(\mu)=V) &= \\ &= \lim_{\Delta V \rightarrow 0} \left(P(\mathbf{V}(\mu) < V + (1+j) \cdot \Delta V) - P(\mathbf{V}(\mu) < V + \Delta V) - \right. \\ &\quad \left. - P(\mathbf{V}(\mu) < V + j \cdot \Delta V) + P(\mathbf{V}(\mu) < V) \right) \end{aligned} \quad (\text{A.24})$$

Da es sich bei dem Zufallsvektor $\vec{\mathbf{V}}(\mu)$ um eine mathematische Stichprobe handelt, kann dessen Verbundverteilung $P(\vec{\mathbf{V}}(\mu) < \vec{\mathbf{V}}(\mu))$ als die L -te Potenz der Verteilung der Zufallsgröße $\mathbf{V}(\mu)$ geschrieben werden. Die Wahrscheinlichkeit $P_L(\vec{\mathbf{V}}(\mu) = V \cdot \vec{\mathbf{1}})$, dass alle Elemente des Zufallsvektors $\vec{\mathbf{V}}(\mu)$ der mathematischen Stichprobe vom Umfang L zugleich einen bestimmten Wert V annehmen, ist daher die L -te Potenz der Wahrscheinlichkeit $P(\mathbf{V}(\mu)=V)$. Die Wahrscheinlichkeit $P_L(\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}=0)$, dass die empirische Varianz bei einer Anzahl von L Einzelmessungen Null wird, kann durch Summation aus den Wahrscheinlichkeiten $P_L(\vec{\mathbf{V}}(\mu) = V \cdot \vec{\mathbf{1}})$ berechnet und aus der Wahrscheinlichkeit, dass die empirische Varianz bei $L-1$ Einzelmessungen Null wird, abgeschätzt werden.

$$\begin{aligned} P_L(\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}=0) &= \sum_V P(\mathbf{V}(\mu)=V)^L \leq \\ &\leq \max_V \left\{ P(\mathbf{V}(\mu)=V) \right\} \cdot \sum_V P(\mathbf{V}(\mu)=V)^{(L-1)} = \\ &= \max_V \left\{ P(\mathbf{V}(\mu)=V) \right\} \cdot P_{L-1}(\hat{\mathbf{C}}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}=0) \end{aligned} \quad (\text{A.25})$$

Die Summe ist ebenso wie die Maximalwertbestimmung über alle möglichen Werte von V , also über alle Elemente des Ergebnisraums von $\mathbf{V}(\mu)$ zu erstrecken. Durch geeignete Wahl des erregenden Zufallsvektors $\vec{V}$ kann für $L \geq 2$ erreicht werden, dass die Wahrscheinlichkeit die Varianz Null zu erhalten, mit steigender Mittelungsanzahl L rasch absinkt, weil der Faktor zwischen der Wahrscheinlichkeit für die Varianz Null für L und für $L-1$ gemäß der letzten Gleichung mit $\max_V \{P(\mathbf{V}(\mu)=V)\}$ nach oben abgeschätzt werden kann. Diese obere Grenze für den Faktor wird dann besonders klein, wenn die Maximalwahrscheinlichkeit minimal wird. Dann wird die Wahrscheinlichkeit für die empirische Varianz Null besonders rasch abklingen. Wenn man sich die Zufallsgröße $\mathbf{V}(\mu)$ durch Quantisierung aus einer wertkontinuierlichen Zufallsgröße entstanden denkt, die durch ihre Verteilungsdichtefunktion beschrieben wird, kann man eine möglichst kleine Maximalwahrscheinlichkeit besonders einfach erreichen, wenn man die Quantisierung dort besonders fein macht, wo die Verteilungsdichtefunktion der zugrundeliegenden wertkontinuierlichen Zufallsgröße groß ist, und wenn man insgesamt möglichst fein quantisiert.

An einem einfachen Beispiel soll gezeigt werden in welcher Größenordnung die Wahrscheinlichkeit, die Varianz Null zu erhalten, bei einigermaßen realistischer Wahl des Zufallsvektors $\vec{V}$ liegt. Zur Gewinnung der Stichprobe wird ein Zufallszahlengenerator verwendet, der alle möglichen 4-Byte Festkommazahlen mit gleicher Wahrscheinlichkeit unabhängig liefert. Ähnliche Zufallszahlengeneratoren — die allerdings deterministische Zahlenfolgen mit pseudozufälligen Eigenschaften liefern — sind oft in Prozessoren oder Programmen fest eingebaut, oder sind Teil der Programmbibliothek, die man beim Erwerb eines Compilers einer Programmiersprache erhält.

Wir wollen nun annehmen, dass der Real- und der Imaginärteil der Zufallsgrößen $V_\lambda(\mu)$, die man aus je zwei der zufälligen Ausgaben des Zufallszahlengenerators — z. B. mit einer nichtlinearen Kennlinie — berechnet, unabhängig sind, und dass ein eindeutiger Zusammenhang zwischen den Zufallsgrößen $\mathbf{V}_\lambda(\mu)$ und den Paaren von Ausgaben des Zufallszahlengenerators besteht, so dass sich für die Mächtigkeit des Ergebnisraums jeder der Zufallsgrößen $\mathbf{V}_\lambda(\mu)$ der Wert 2^{64} ergibt. Da alle möglichen Werte der Zufallsgröße $\mathbf{V}_\lambda(\mu)$ gleichwahrscheinlich sind, ergibt sich für die Wahrscheinlichkeit, $P(\mathbf{V}(\mu)=V)$ der von V unabhängige konstante Wert 2^{-64} , und für die Wahrscheinlichkeit, dass die Varianz bei *einer* Frequenz μ Null wird, nach der letzten Formel der mit L exponentiell fallende Wert $2^{64 \cdot (1-L)}$. Für alle Frequenzen $\Omega = \mu \cdot 2\pi/M$ werden unabhängige Zufallszahlen dieses Zufallszahlengenerators verwendet. Man kann nun die Wahrscheinlichkeit, dass eine oder mehrere der M Varianzen für die M Frequenzpunkte Null werden, als Funktion von L und M zu $1 - (1 - 2^{64 \cdot (1-L)})^M$ berechnen. Da die Wahrscheinlichkeit, dass die Varianz bei einer Frequenz μ Null wird, hinreichend klein ist, kann die Wahrscheinlichkeit, dass eine oder mehrere der M Varianzen Null werden, mit $M \cdot 2^{64 \cdot (1-L)}$ abgeschätzt werden.

Gibt man sich die maximal zulässige Wahrscheinlichkeit vor, dass eine oder mehrere der M Varianzen Null werden, so kann man die notwendige Mittelungsanzahl L , oder die maximal mögliche Anzahl M der Frequenzpunkte, für die die Messwerte berechnet werden sollen, als Funktion von der jeweils anderen Größe abschätzen. Bei unserem 4-Byte Zufallszahlengenerator ergibt sich mit dieser Abschätzung bereits bei einer Messung mit der kürzest möglichen Mittelungsanzahl von $L = 2$, dass die Wahrscheinlichkeit $\binom{49}{6}^{-2}$ bei zwei Spielen im Lotto „6 aus 49“ beide male einen „Sechser“ zu haben größer ist als die Wahrscheinlichkeit, dass eine oder mehrere der M Varianzen Null werden, wenn man $M \leq 94334$ wählt. Daher wird man die oben beschriebene Begrenzung der Inversen der empirischen Varianzen bei einer praktischen Realisierung des RKM in Form eines Messprogramms nicht wirklich vornehmen, zumal die Messwerte in diesem Fall sowieso keine Aussagekraft besitzen. Sollte man bei einer realen Messung wirklich einmal das außerordentliche „Glück“ haben, auf diesen Fall zu stoßen, so wird man die Messung einfach wiederholen, oder die Mittelungsanzahl L erhöhen. Man wird dies auch dann tun, wenn die Messwertvarianz zu groß wird. Dennoch ist Begrenzung der Inversen der empirischen Varianz für die Berechnung der stochastischen Eigenschaften der Messwerte unumgänglich.

Oft verwendet man Zufallszahlengeneratoren die eine determinierte Abfolge von Zahlen liefern, die z. B. mit Hilfe rückgekoppelter Schieberegister, oder mit Modulo-Arithmetik gewonnen werden, und die sich mit einer extrem langen Periode wiederholen. In diesem Fall kann der singuläre Fall nie eintreten, so dass man bei der Berechnung der Messwerte den singulären Fall nicht extra abprüfen und getrennt behandeln muss. Streng genommen ist die im Hauptteil beschriebene Theorie bei Verwendung eines solchen Zufallszahlengenerators selbst dann nicht mehr gültig, wenn man den Startwert des Zufallszahlengenerators zufällig wählt, weil dann die Unabhängigkeit der einzelnen Elemente der Stichproben nicht mehr gegeben ist. Die Praxis zeigt jedoch, dass die bisher hergeleiteten Ergebnisse auch bei Verwendung solcher Zufallszahlengeneratoren verwendet werden können.

A.4 Grenzen für die Hauptdiagonalelemente der inversen Kovarianzmatrix

Gegeben sei die $R \times R$ Kovarianzmatrix $\underline{C}$. Da diese hermitesch ist, lässt sie sich unitär kongruent auf Diagonalform transformieren.

$$\underline{S} = \underline{U} \cdot \underline{C} \cdot \underline{U}^H \quad \text{mit} \quad \underline{U} \cdot \underline{U}^H = \underline{E} \quad \Leftrightarrow \quad \underline{U}^{-1} = \underline{U}^H \quad (\text{A.26})$$

Die Diagonalelemente der Diagonalmatrix $\underline{S}$ seien mit s_i bezeichnet. Sie sind die Eigenwerte der Kovarianzmatrix $\underline{C}$ und somit zugleich die Inversen der Eigenwerte der inver-

sen Kovarianzmatrix $\underline{C}^{-1}$. Da bei einer hermiteschen Matrix die Eigenwerte gleich ihren Singulärwerten sind, sind s_i zugleich die Singulärwerte von $\underline{C}$ und die Inversen der Singulärwerte von $\underline{C}^{-1}$. Sie sind reell und nichtnegativ. Mit dem $1 \times R$ Einheitszeilenvektor $\vec{E}_n$, dessen n -tes Element eins ist, und der sonst nur Nullen enthält, ergibt die quadratische Form

$$\vec{E}_n \cdot \underline{C}^{-1} \cdot \vec{E}_n^H \quad (\text{A.27})$$

das n -te Hauptdiagonalelement der inversen Kovarianzmatrix. Wenn mit $\vec{U}_n$ der n -te Spaltenvektor und mit $U_{j,n}$ das Element in der j -ten Zeile und in der n -ten Spalte der Matrix $\underline{U}$ bezeichnet ist, kann man für das n -te Hauptdiagonalelement der inversen Kovarianzmatrix auch

$$\begin{aligned} \vec{E}_n \cdot \underline{C}^{-1} \cdot \vec{E}_n^H &= \vec{E}_n \cdot (\underline{U}^H \cdot \underline{S} \cdot \underline{U})^{-1} \cdot \vec{E}_n^H = \vec{E}_n \cdot \underline{U}^{-1} \cdot \underline{S}^{-1} \cdot (\underline{U}^H)^{-1} \cdot \vec{E}_n^H = \\ &= \vec{E}_n \cdot \underline{U}^H \cdot \underline{S}^{-1} \cdot \underline{U} \cdot \vec{E}_n^H = \vec{U}_n^H \cdot \underline{S}^{-1} \cdot \vec{U}_n = \sum_{j=1}^R |U_{j,n}|^2 \cdot s_j^{-1} \end{aligned} \quad (\text{A.28})$$

schreiben. Für diese Summe kann man eine obere und unterer Schranke angeben.

$$\begin{aligned} \max_i (s_i)^{-1} = \min_i (s_i^{-1}) &= \min_i (s_i^{-1}) \cdot \sum_{j=1}^R |U_{j,n}|^2 = \sum_{j=1}^R |U_{j,n}|^2 \cdot \min_i (s_i^{-1}) \leq \\ &\leq \sum_{j=1}^R |U_{j,n}|^2 \cdot s_j^{-1} \leq \\ &\leq \sum_{j=1}^R |U_{j,n}|^2 \cdot \max_i (s_i^{-1}) = \max_i (s_i^{-1}) \cdot \sum_{j=1}^R |U_{j,n}|^2 = \max_i (s_i^{-1}) = \min_i (s_i)^{-1} \end{aligned} \quad (\text{A.29})$$

Die Hauptdiagonalelemente der inversen Kovarianzmatrix sind also alle reell und liegen innerhalb des Intervalls $[\max_i (s_i)^{-1}; \min_i (s_i)^{-1}]$.

A.5 Zur Berechnung des Erwartungswertes einer bilinearen Form

Bei der Herleitung der Messwerte $\hat{\Phi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$ und $\hat{\Psi}_n(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})$ sowie der Varianzen der Messwerte $\hat{H}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$, $\hat{H}_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$, $\hat{u}(k)$ und $\hat{U}_f(\mu)$ tritt immer der Erwartungswert einer zufälligen bilinearen Form auf. Dabei sind die Elemente der zufälligen Matrix der bilinearen Form Funktionen, die nur von dem zufälligen Spektrum der Erregung abhängen. Die Vektoren der bilinearen Form sind ebenfalls zufällig und hängen nur von den Zeit- oder Spektralwerten des gefensterten Approximationsfehlerprozesses ab.

Wegen der Gemeinsamkeit, die alle diese Erwartungswerte aufweisen, werde ich in diesem Unterkapitel die Erwartungswerte

$$\mathbb{E}\left\{\vec{\mathbf{N}}_1 \cdot \underline{\mathbf{A}} \cdot \vec{\mathbf{N}}_2^H\right\} \quad (\text{A.30})$$

einer bilinearen Form mit der allgemeineren Matrix $\underline{\mathbf{A}}$ und den allgemeineren Zufallsvektoren $\vec{\mathbf{N}}_1$ und $\vec{\mathbf{N}}_2$ berechnen, um so eine Vielzahl von Berechnungen mit den im Einzelfall verwendeten Matrizen und Vektoren zu vermeiden. Die Zufallsvektoren sind jedoch nicht beliebig wählbar. Es muss sich bei den L Spalten der Matrix, die die zwei Zufallsvektoren $\vec{\mathbf{N}}_1$ und $\vec{\mathbf{N}}_2$ als Zeilenvektoren enthält, um die Elemente einer mathematischen Stichprobe vom Umfang L des Zufallsspaltenvektors $[\mathbf{N}_1, \mathbf{N}_2]^T$ handeln. Somit sind alle L Spaltenvektoren voneinander unabhängig und haben dieselbe Verbundverteilung, nämlich die Verbundverteilung des Zufallsspaltenvektors, aus dem die Stichprobe entnommen wurde. Wenn mit $\mathbf{N}_{i,\lambda}$ das λ -te Element des Stichprobenvektors $\vec{\mathbf{N}}_i$ bezeichnet ist, gilt somit

$$\mathbb{E}\left\{\mathbf{N}_{i,\lambda_1} \cdot \mathbf{N}_{j,\lambda_2}^*\right\} = \begin{cases} \mathbb{E}\{\mathbf{N}_i \cdot \mathbf{N}_j^*\} & \text{für } \lambda_1 = \lambda_2 \\ \mathbb{E}\{\mathbf{N}_i\} \cdot \mathbb{E}\{\mathbf{N}_j\}^* & \text{für } \lambda_1 \neq \lambda_2 \end{cases} \quad (\text{A.31})$$

$$\forall \quad i, j = 1 \dots 2; \quad \lambda, \lambda_1, \lambda_2 = 1 \dots L.$$

Es sei darauf hingewiesen, dass die hier geforderte Unabhängigkeit für unterschiedliche Werte von λ *nicht* bedeutet, dass die Zufallsgrößen $\mathbf{N}_1$ und $\mathbf{N}_2$ voneinander unabhängig sein müssen. Diese können z. B. zueinander konjugiert oder sogar identisch sein. Damit die weitere Herleitung ihre Gültigkeit behält, müssen die Elemente der Matrix $\underline{\mathbf{A}}$ für alle Werte von λ von den Zufallsspaltenvektoren $[\mathbf{N}_{1,\lambda}, \mathbf{N}_{2,\lambda}]^T$ unabhängig sein.

Wenn mit $\underline{\mathbf{A}}_{\lambda_1,\lambda_2}$ das Element in der λ_1 -ten Zeile und der λ_2 -ten Spalte bezeichnet ist, erhalten wir:

$$\mathbb{E}\left\{\vec{\mathbf{N}}_1 \cdot \underline{\mathbf{A}} \cdot \vec{\mathbf{N}}_2^H\right\} = \mathbb{E}\left\{\sum_{\lambda_1=1}^L \sum_{\lambda_2=1}^L \mathbf{N}_{1,\lambda_1} \cdot \underline{\mathbf{A}}_{\lambda_1,\lambda_2} \cdot \mathbf{N}_{2,\lambda_2}^*\right\} = \dots \quad (\text{A.32a})$$

Der Erwartungswert wird als Summe der Erwartungswerte berechnet.

$$\dots = \sum_{\lambda_1=1}^L \sum_{\lambda_2=1}^L \mathbb{E}\left\{\mathbf{N}_{1,\lambda_1} \cdot \underline{\mathbf{A}}_{\lambda_1,\lambda_2} \cdot \mathbf{N}_{2,\lambda_2}^*\right\} = \dots \quad (\text{A.32b})$$

Die Doppelsumme wird in zwei Teilsummen aufgespalten. Die erste Teilsumme enthält die Hauptdiagonalelemente, während in der zweiten Teilsumme die Nebendiagonalelemente auftreten.

$$\dots = \sum_{\lambda=1}^L \mathbb{E}\left\{\mathbf{N}_{1,\lambda} \cdot \underline{\mathbf{A}}_{\lambda,\lambda} \cdot \mathbf{N}_{2,\lambda}^*\right\} + \sum_{\lambda_1=1}^L \sum_{\substack{\lambda_2=1 \\ \lambda_2 \neq \lambda_1}}^L \mathbb{E}\left\{\mathbf{N}_{1,\lambda_1} \cdot \underline{\mathbf{A}}_{\lambda_1,\lambda_2} \cdot \mathbf{N}_{2,\lambda_2}^*\right\} = \dots \quad (\text{A.32c})$$

Nun wird die vorausgesetzte Unabhängigkeit der Zufallselemente der Matrix $\underline{\mathbf{A}}$ von den Zufallsspaltenvektoren $[\mathbf{N}_{1,\lambda}, \mathbf{N}_{2,\lambda}]^T$ dazu verwendet, den Erwartungswert des Produktes als das Produkt der Erwartungswerte der einzelnen Faktoren zu schreiben.

$$\dots = \sum_{\lambda=1}^L \mathbb{E}\{\mathbf{N}_{1,\lambda} \cdot \mathbf{N}_{2,\lambda}^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda,\lambda}\} + \sum_{\lambda_1=1}^L \sum_{\substack{\lambda_2=1 \\ \lambda_2 \neq \lambda_1}}^L \mathbb{E}\{\mathbf{N}_{1,\lambda_1} \cdot \mathbf{N}_{2,\lambda_2}^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda_1,\lambda_2}\} = \dots \quad (\text{A.32d})$$

Desweiteren berücksichtigen wir nun die Unabhängigkeit der Stichprobenelemente, die in unterschiedlichen Einzelmessungen mit $\lambda_1 \neq \lambda_2$ gewonnen wurden.

$$\dots = \sum_{\lambda=1}^L \mathbb{E}\{\mathbf{N}_{1,\lambda} \cdot \mathbf{N}_{2,\lambda}^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda,\lambda}\} + \sum_{\lambda_1=1}^L \sum_{\substack{\lambda_2=1 \\ \lambda_2 \neq \lambda_1}}^L \mathbb{E}\{\mathbf{N}_{1,\lambda_1}\} \cdot \mathbb{E}\{\mathbf{N}_{2,\lambda_2}^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda_1,\lambda_2}\} = \dots \quad (\text{A.32e})$$

Außerdem ist bei einer mathematischen Stichprobe die Verteilung der Stichprobenelemente gleich der Verteilung der Zufallsgröße, aus der die Stichprobe gewonnen worden ist. Somit ist auch der Erwartungswert der Stichprobenelemente gleich dem Erwartungswert der Zufallsgröße.

$$\dots = \sum_{\lambda=1}^L \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda,\lambda}\} + \sum_{\lambda_1=1}^L \sum_{\substack{\lambda_2=1 \\ \lambda_2 \neq \lambda_1}}^L \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \cdot \mathbb{E}\{\underline{\mathbf{A}}_{\lambda_1,\lambda_2}\} = \dots \quad (\text{A.32f})$$

Die von λ bzw. von λ_1 und λ_2 unabhängigen Erwartungswerte werden vor die Summen gezogen.

$$\dots = \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} \cdot \sum_{\lambda=1}^L \mathbb{E}\{\underline{\mathbf{A}}_{\lambda,\lambda}\} + \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \cdot \sum_{\lambda_1=1}^L \sum_{\substack{\lambda_2=1 \\ \lambda_2 \neq \lambda_1}}^L \mathbb{E}\{\underline{\mathbf{A}}_{\lambda_1,\lambda_2}\} = \dots \quad (\text{A.32g})$$

Nun ergänzen wir die in der Doppelsumme fehlenden Hauptdiagonalelemente.

$$\begin{aligned} \dots &= \left(\mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} - \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \right) \cdot \sum_{\lambda=1}^L \mathbb{E}\{\underline{\mathbf{A}}_{\lambda,\lambda}\} + \\ &\quad + \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \cdot \sum_{\lambda_1=1}^L \sum_{\lambda_2=1}^L \mathbb{E}\{\underline{\mathbf{A}}_{\lambda_1,\lambda_2}\} = \dots \end{aligned} \quad (\text{A.32h})$$

Wieder vertauschen wir die Reihenfolge der Summation und der Erwartungswertbildung.

$$\begin{aligned} \dots &= \left(\mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} - \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \right) \cdot \mathbb{E}\left\{ \sum_{\lambda=1}^L \underline{\mathbf{A}}_{\lambda,\lambda} \right\} + \\ &\quad + \mathbb{E}\{\mathbf{N}_1\} \cdot \mathbb{E}\{\mathbf{N}_2^*\} \cdot \mathbb{E}\left\{ \sum_{\lambda_1=1}^L \sum_{\lambda_2=1}^L \underline{\mathbf{A}}_{\lambda_1,\lambda_2} \right\} = \dots \end{aligned} \quad (\text{A.32i})$$

Mit dem $1 \times L$ Zeilenvektor $\vec{1}$, der nur Einsen enthält, erhalten wir in Matrixschreibweise:

$$\begin{aligned} \dots = & \left(E\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} - E\{\mathbf{N}_1\} \cdot E\{\mathbf{N}_2^*\} \right) \cdot E\left\{ \text{spur}(\underline{\mathbf{A}}) \right\} + \\ & + E\{\mathbf{N}_1\} \cdot E\{\mathbf{N}_2^*\} \cdot E\left\{ \vec{1} \cdot \underline{\mathbf{A}} \cdot \vec{1}^H \right\}. \end{aligned} \quad (\text{A.32j})$$

In dieser Form können wir den gesuchten Erwartungswert der bilinearen Form im Hauptteil weiterverwenden.

A.6 Zur Berechnung der Messwerte des LDS und des KLDS

Wenn man für die Berechnung der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ die gemäß der Gleichungen (3.27) und (3.28) definierten Matrizen $\underline{V}_{\perp}(\mu)$ verwendet, lassen sich alle Messwerte immer durch Akkumulation aus den bei den Einzelmessungen λ verwendeten Spektralwerten der Erregung und des gefensterten Ausgangssignals berechnen, ohne dass dazu die Spektralwerte aller Einzelmessungen abgespeichert werden müssen. Das gilt unabhängig davon, welche Stichprobenvektoren $\vec{V}(\hat{\mu})$, deren Anzahl mit R bezeichnet sei, man in die Matrizen $\check{\underline{V}}(\mu)$ aufnimmt. Hier soll nun erklärt werden, wie sich die Messwerte in einer Form darstellen lassen, die diese Art der Berechnung ermöglicht.

Nach den Gleichungen (3.29) berechnen sich die Messwerte jeweils als die Quotienten einer bilinearen Form und einer M -fachen Matrixspur. Nun dividieren wir Zähler und Nenner durch $L-1$ und beginnen mit der Berechnung des Zählers, wobei wir jeweils die Darstellung des Messwertes wählen, die die Vektoren $\vec{Y}_f(\mu)$ enthält. Die Matrix der bilinearen Form ist jeweils selbst ein Produkt zweier Matrizen $\underline{V}_{\perp}(\mu)$ bei unterschiedlichen Frequenzen, wobei eine der beiden Matrizen noch transponiert und ggf. konjugiert auftritt. Mit Gleichung (3.28) ersetzt man nun die Matrizen dieses Matrixproduktes jeweils durch die Differenz der Matrix $\underline{1}_{\perp}$ und eines anderen Matrixproduktes. Das sich ergebende Produkt der Matrixdifferenzen lässt sich mit Hilfe des Distributivgesetzes der Matrizenmultiplikation unter Beibehaltung der Reihenfolge der Matrizen als vorzeichenbehaftete Summe von Matrixprodukten schreiben. Diese Summe wird mit den Vektoren $\vec{Y}_f(\mu)$ bei unterschiedlichen Frequenzen von links und transponiert und ggf. konjugiert von rechts multipliziert, was der entsprechenden Multiplikation der einzelnen Summanden entspricht. Wenn man nun berücksichtigt, dass die Matrizen $\underline{V}_{\perp}(\mu)$ hermitesch sind, dass die Matrix $\underline{1}_{\perp}$ idempotent ist, und dass der Zähler durch $L-1$ dividiert wurde, erhält man so eine Summe von Produkten von empirischen Kovarianzvektoren und -matrizen und deren Inversen, wobei deren Dimensionen entweder 1 oder R sind. Die Größe der Vektoren und

Matrizen ist somit von der Mittelungsanzahl L unabhängig. Die Elemente der empirischen Kovarianzvektoren und -matrizen sind die empirischen Kovarianzen der zufälligen Spektralwerte der Erregung und des gefensterten Ausgangssignals. Diese kann man, wie zum Beispiel in den Gleichungen (3.10) und (3.11) gezeigt durch Akkumulation aus den bei den Einzelmessungen λ verwendeten Spektralwerten der Erregung und der gemessenen, gefensterten Ausgangssignale berechnen, ohne dass dazu die Spektralwerte aller Einzelmessungen abgespeichert werden müssen. Am Beispiel des Messwertes $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ soll nun die eben verbal beschriebene Vorgehensweise zur Berechnung des Zählers des Messwertes demonstriert werden.

$$\begin{aligned}
 & \frac{\vec{Y}_f(\mu) \cdot \underline{V}_{\perp}(\mu) \cdot \underline{V}_{\perp}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1} = \quad (A.33) \\
 &= \frac{1}{L-1} \cdot \vec{Y}_f(\mu) \cdot \left(\underline{1}_{\perp} - \underline{1}_{\perp} \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_{\perp} \right) \cdot \\
 &\quad \cdot \left(\underline{1}_{\perp} - \underline{1}_{\perp} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^{-1}}{L-1} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \underline{1}_{\perp} \right) \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H = \\
 &= \underbrace{\frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianz}} - \\
 &\quad - \underbrace{\frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianzvektor}} \cdot \hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^{-1} \cdot \underbrace{\frac{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianzvektor}} - \\
 &\quad - \underbrace{\frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H}{L-1}}_{= \text{Kovarianzvektor}} \cdot \hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1} \cdot \underbrace{\frac{\check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianzvektor}} + \\
 &\quad + \underbrace{\frac{\vec{Y}_f(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H}{L-1}}_{= \text{Kovarianzvektor}} \cdot \hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1} \cdot \underbrace{\frac{\check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianzmatrix}} \cdot \\
 &\quad \cdot \hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})}^{-1} \cdot \underbrace{\frac{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \underline{1}_{\perp} \cdot \vec{Y}_f(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})^H}{L-1}}_{= \text{Kovarianzvektor}}
 \end{aligned}$$

Die Nenner der Messwerte der Gleichungen (3.29) sind jeweils die M -fachen Spuren der Matrixprodukte der bilinearen Formen des Zählers. Die dabei auftretenden Matrixspuren sind zwei der in der Tabelle 3.2 aufgeführten Matrixspuren, nämlich die, die mit $\boxed{\text{d}}$ und $\boxed{\text{g}}$ bezeichnet wurden. Wie man diese und auch alle anderen in der Tabelle 3.2 aufgelisteten

Terme $\boxed{\text{a}}$ bis $\boxed{\text{p}}$ durch eine Akkumulation der Produkte der Spektralwerte der bei den Einzelmessungen verwendeten Erregungen berechnen kann, soll nun erläutert werden.

Zunächst ersetzt man dazu wieder die Matrizen des Matrixproduktes, dessen Spur berechnet werden soll, mit Hilfe der Gleichung (3.28) jeweils durch eine Differenz der Matrix $\underline{1}_\perp$ und eines weiteren Matrixproduktes. Wieder lässt sich das Produkt der Matrixdifferenzen als vorzeichenbehaftete Summe von Matrixprodukten schreiben. Die Spur ist dann die vorzeichenbehaftete Summe der Spuren der einzelnen Summanden, also der Matrixprodukte. Jeder einzelne Summand ist entweder die Matrix $\underline{1}_\perp$ oder beginnt mit $\underline{1}_\perp \cdot \check{V}(\dots)^H$ oder $\underline{1}_\perp \cdot \check{V}(\dots)^T$ und endet mit $\check{V}(\dots) \cdot \underline{1}_\perp$ oder $\check{V}(\dots)^* \cdot \underline{1}_\perp$. Die Spur der $L \times L$ Matrix $\underline{1}_\perp$ ist immer $L-1$ und braucht daher nicht durch Akkumulation der Stichprobenelemente berechnet zu werden. Da die Reihenfolge der Faktoren bei einem Produkt zweier Matrizen auf dessen Spur keinen Einfluss hat, kann man die ersten beiden Faktoren jedes einzelnen Matrixproduktes jeweils an das Ende des Matrixproduktes verschieben. Nun nützt man wieder die Tatsache, dass die Matrix $\underline{1}_\perp$ idempotent ist, und erhält dadurch jeweils ein Matrixprodukt, bei dem sich Kovarianzmatrizen und inverse Kovarianzmatrizen abwechseln.

Wie Gleichung (3.10) zeigt, können die Kovarianzmatrizen und somit deren Spuren durch Akkumulation ohne Zwischenspeicherung der Spektralwerte aller Einzelmessungen berechnet werden. Am Beispiel der Matrixspur $\boxed{\text{d}}$ in Tabelle 3.2 sei dies demonstriert.

$$\begin{aligned}
 \boxed{\text{d}} &= \text{spur} \left(\underline{V}_\perp(\mu) \cdot \underline{V}_\perp \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right) \right) = \tag{A.34} \\
 &= \text{spur} \left(\left(\underline{1}_\perp - \underline{1}_\perp \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_\perp \right) \cdot \right. \\
 &\quad \cdot \left. \left(\underline{1}_\perp - \underline{1}_\perp \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1}}{L-1} \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right) \cdot \underline{1}_\perp \right) \right) = \\
 &= \text{spur} \left(\underline{1}_\perp \cdot \underline{1}_\perp - \underline{1}_\perp \cdot \underline{1}_\perp \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1}}{L-1} \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right) \cdot \underline{1}_\perp - \right. \\
 &\quad - \underline{1}_\perp \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_\perp \cdot \underline{1}_\perp + \underline{1}_\perp \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \\
 &\quad \cdot \check{V}(\mu) \cdot \underline{1}_\perp \cdot \underline{1}_\perp \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi})}^{-1}}{L-1} \cdot \check{V} \left(\mu + \tilde{\mu} \cdot \frac{M}{K_\Phi} \right) \cdot \underline{1}_\perp \left. \right) =
 \end{aligned}$$

$$\begin{aligned}
&= \text{spur}(\underline{1}_{\perp}) - \text{spur}\left(\underline{1}_{\perp} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1}}{L-1} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{1}_{\perp}\right) - \\
&\quad - \text{spur}\left(\underline{1}_{\perp} \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_{\perp}\right) + \\
&+ \text{spur}\left(\underline{1}_{\perp} \cdot \check{V}(\mu)^H \cdot \frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1}}{L-1} \cdot \right. \\
&\quad \left. \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{1}_{\perp}\right) = \\
&= \text{spur}(\underline{1}_{\perp}) - \text{spur}\left(\frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1}}{L-1} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{1}_{\perp} \cdot \underline{1}_{\perp} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^H\right) - \\
&\quad - \text{spur}\left(\frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H\right) + \\
&+ \text{spur}\left(\frac{\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1}}{L-1} \cdot \check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right)^H \cdot \frac{\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1}}{L-1} \cdot \right. \\
&\quad \left. \cdot \check{V}\left(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}\right) \cdot \underline{1}_{\perp} \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H\right) = \\
&= \underbrace{\text{spur}(\underline{1}_{\perp})}_{= L-1} - \text{spur}\left(\hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1} \cdot \underbrace{\frac{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}^H}{L-1}}_{= \text{Kovarianzmatrix}}\right) - \\
&\quad - \text{spur}\left(\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1} \cdot \underbrace{\frac{\check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H}{L-1}}_{= \text{Kovarianzmatrix}}\right) + \\
&+ \text{spur}\left(\hat{C}_{\check{V}(\mu), \check{V}(\mu)}^{-1} \cdot \underbrace{\frac{\check{V}(\mu) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}^H}{L-1}}_{= \text{Kovarianzmatrix}} \cdot \hat{C}_{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}), \check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi})}}^{-1} \cdot \right. \\
&\quad \left. \cdot \underbrace{\frac{\check{V}(\mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \underline{1}_{\perp} \cdot \check{V}(\mu)^H}{L-1}}_{= \text{Kovarianzmatrix}}\right)
\end{aligned}$$

Führt man diese Umformungen bei allen in der Tabelle 3.2 aufgeführten Termen durch, so stellt man bei den Termen $\underline{\text{a}}$, $\underline{\text{b}}$, $\underline{\text{c}}$, $\underline{\text{d}}$, $\underline{\text{g}}$, $\underline{\text{i}}$ und $\underline{\text{l}}$ fest, dass die Matrixspuren nur über die Produkte solcher empirischer Kovarianzmatrizen und deren Inverser zu berechnen sind, die bereits bei der Berechnung der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$ und $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot M/K_{\Phi})$

aufgetreten sind. Die Terme $\underline{e}$, $\underline{f}$, $\underline{h}$, $\underline{j}$, $\underline{k}$, $\underline{m}$, $\underline{n}$, $\underline{o}$ und $\underline{p}$ sind nur bei den Frequenzen $\mu = 0 \left(\frac{M}{2 \cdot K_\Phi} \right) M - 1$ des Falles (3.55a) in den Messwert(ko)varianzen zu finden. Bei diesen Frequenzen enthalten die Matrixspuren zusätzlich noch einige weitere Kovarianzmatrizen. Diese hat man bei der Berechnung anderer Messwerte $\hat{\Phi}_n(\mu, \mu + \check{\mu} \cdot M / K_\Phi)$ und $\hat{\Psi}_n(\mu, \mu + \check{\mu} \cdot M / K_\Phi)$ jedoch ebenfalls bereits bestimmt. Der Parameter $\check{\mu}$, der angibt, bei der Berechnung welcher Messwerte die Kovarianzmatrizen bestimmt worden sind, ist dabei eine ganze Zahl, die sich mit den Gleichungen (3.55) aus μ und $\tilde{\mu}$ ermitteln lässt. Alle Kovarianzmatrizen, die man zur Berechnung der Messwert(ko)varianzen benötigt, sind also bereits nach der Berechnung der Messwerte bekannt, so dass man keine zusätzlichen Akkumulatoren vorsehen muss.

A.7 Zu Spur und Rang des Produktes Idempotenter Matrizen

Es treten immer wieder Produkte der idempotenten Matrizen $\underline{V}_1(\mu)$ für unterschiedliche Frequenzen und deren Transponierte auf. Der Rangdefekt aller dieser Matrizen ist auf eine gemeinsame obere Schranke $D_{\max}$ begrenzt. Da der Rang der Transponierten einer Matrix gleich dem Rang der Matrix ist, ist auch deren Rangdefekt auf die gleiche Schranke begrenzt. Die Transponierte einer idempotenten Matrix ist ebenfalls wieder idempotent. Wir wollen uns nun mit dem Produkt $\underline{I}_\Pi = \underline{I}_1 \cdot \underline{I}_2 \cdot \dots \cdot \underline{I}_n$ von n idempotenten $L \times L$ -Matrizen beschäftigen, die alle einen Rangdefekt aufweisen, der bei allen n Matrizen auf $D_{\max}$ begrenzt ist. Die Spur einer idempotenten Matrix ist immer gleich ihrem Rang. Daher ist die Spur jeder Matrix des Produkts auf den Bereich

$$L - D_{\max} \leq \text{spur}(\underline{I}_i) = \text{rang}(\underline{I}_i) \leq L \quad \forall \quad i = 1 \ (1) \ n \quad (\text{A.35})$$

beschränkt. Jeder Zeilenvektor $\vec{I}_1$, der zu allen Eigenvektoren aller n Matrizen $\underline{I}_i$ zum Eigenwert Null orthogonal ist, ist ein Eigenvektor aller n Matrizen $\underline{I}_i$ zum Eigenwert Eins. Wenn also ein Vektor, für den

$$\vec{I}_1 \cdot \underline{I}_i = \vec{I}_1 \quad \forall \quad i = 1 \ (1) \ n \quad (\text{A.36})$$

gilt, an das Matrixprodukt von links heranmultipliziert wird, so ergibt sich:

$$\vec{I}_1 \cdot \underline{I}_1 \cdot \underline{I}_2 \cdot \dots \cdot \underline{I}_n = \vec{I}_1 \cdot \underline{I}_2 \cdot \underline{I}_3 \cdot \dots \cdot \underline{I}_n = \dots = \vec{I}_1 \cdot \underline{I}_n = \vec{I}_1. \quad (\text{A.37})$$

Daher ist jeder dieser Vektoren $\vec{I}_1$ ein Eigenvektor des Matrixprodukts $\underline{I}_\Pi$ zum Eigenwert Eins. Der Raum, der durch alle Vektoren aufgespannt wird, die zu allen Eigenvektoren aller n Matrizen $\underline{I}_i$ zum Eigenwert Null orthogonal sind, hat mindestens die Dimension

$L - n \cdot D_{\max}$, weil es insgesamt höchstens $n \cdot D_{\max}$ Eigenvektoren aller n Matrizen $\underline{\mathbf{I}}_i$ zum Eigenwert Null gibt. Das Matrixprodukt hat daher — wenn L groß genug ist — den Eigenwert Eins mit einer Vielfachheit von mindestens $L - n \cdot D_{\max}$. Multipliziert man einen beliebigen Vektor von links an des Matrixprodukt $\underline{\mathbf{I}}_{\Pi}$, so entspricht dies einer mehrmaligen Projektion des Vektors und dessen euklidische Norm wird dadurch niemals anwachsen. Daher sind die Beträge der restlichen maximal $n \cdot D_{\max}$ Eigenwerte kleiner gleich eins. Der Betrag der Spur des Matrixprodukts $\underline{\mathbf{I}}_{\Pi}$ — also der Betrag der Summe der Eigenwerte — ist daher für hinreichend großes L immer größer als Null und liegt im Bereich

$$L - 2 \cdot n \cdot D_{\max} \leq |\text{spur}(\underline{\mathbf{I}}_{\Pi})| \leq L. \quad (\text{A.38})$$

Wird die Anzahl n der am Matrixprodukt beteiligten Matrizen und der gemeinsame maximale Rangdefekt $D_{\max}$ all dieser Matrizen konstant gehalten, so wächst der Betrag der Spur des Matrixprodukts für hinreichend großes L also näherungsweise linear mit steigendem L .

A.8 Zur Berechnung der Varianzen und Kovarianzen des LDS und KLDS

Mit Hilfe der Gleichung ([1]:A.42) im Anhang [1]:A.5 gelang es uns bei einem stationären Approximationsfehlerprozess die Varianzen und Kovarianzen der Messwerte $\hat{\Phi}_n(\mu)$ und $\hat{\Psi}_n(\mu)$ zu berechnen. Im Fall eines zyklstationären Approximationsfehlerprozesses benötigen wir stattdessen vier Gleichungssysteme, die sich aus Gleichung ([1]:A.42) ableiten lassen. Das erste Gleichungssystem erhalten wir, indem wir die Gleichung ([1]:A.42) dreimal untereinander schreiben, wobei wir in den einzelnen Zeilen die in der Tabelle A.1 aufgeführten Substitutionen vornehmen. Die Zufallsmatrizen $\underline{\mathbf{V}}_1$, $\underline{\mathbf{V}}_2$, $\underline{\mathbf{V}}_3$ und $\underline{\mathbf{V}}_4$ sind dabei wieder idempotente, hermitesche $L \times L$ Matrizen, deren Rangdefekt mit steigender Mittelungsanzahl L auf eine Konstante begrenzt ist. Die Elemente dieser Matrizen sind aus mathematischen Stichproben der Spektralwerte der zufälligen Erregung berechnet worden. Die Zufallsvektoren $\vec{\mathbf{N}}_1$, $\vec{\mathbf{N}}_2$, $\vec{\mathbf{N}}_3$ und $\vec{\mathbf{N}}_4$ enthalten die Elemente der mathematischen Stichproben der als normalverteilt und mittelwertfrei angenommenen Spektralwerte $\mathbf{N}_1$, $\mathbf{N}_2$, $\mathbf{N}_3$ und $\mathbf{N}_4$ des gefensterten Approximationsfehlerprozesses. Auch diese Zufallsvektoren werden in der in Tabelle A.1 angegebenen Art substituiert. Es wird wieder angenommen, dass die in die Elemente der Zufallsmatrizen eingehenden zufälligen Spektralwerte der Erregung unabhängig von den Spektralwerten seien, deren mathematische Stichproben die Zufallsvektoren sind, so dass die im Anhang [1]:A.5 aufgestellten Forderungen alle erfüllt sind. Die drei Gleichungen, die man mit den angegebenen Sub-

in ([1]:A.42):	$\vec{N}_1$	$\vec{N}_2$	$\vec{N}_3$	$\vec{N}_4$	$\underline{A}$	$\underline{B}$	$\underline{c}$
in Zeile 1:	$\vec{N}_1$	$\vec{N}_2$	$\vec{N}_3$	$\vec{N}_4$	$\underline{V}_1 \cdot \underline{V}_2^H$	$\underline{V}_3 \cdot \underline{V}_4^H$	$\underline{c}_1^{-1}$
in Zeile 2:	$\vec{N}_1$	$\vec{N}_3^*$	$\vec{N}_2^*$	$\vec{N}_4$	$\underline{V}_1 \cdot \underline{V}_3^T$	$\underline{V}_2^* \cdot \underline{V}_4^H$	$\underline{c}_2^{-1}$
in Zeile 3:	$\vec{N}_1$	$\vec{N}_4$	$\vec{N}_2^*$	$\vec{N}_3^*$	$\underline{V}_1 \cdot \underline{V}_4^H$	$\underline{V}_2^* \cdot \underline{V}_3^T$	$\underline{c}_3^{-1}$
mit $\underline{c}_1 = \text{spur}(\underline{V}_1 \cdot \underline{V}_2^H) \cdot \text{spur}(\underline{V}_3 \cdot \underline{V}_4^H)$ $\underline{c}_2 = \text{spur}(\underline{V}_1 \cdot \underline{V}_3^T) \cdot \text{spur}(\underline{V}_2^* \cdot \underline{V}_4^H)$ $\underline{c}_3 = \text{spur}(\underline{V}_1 \cdot \underline{V}_4^H) \cdot \text{spur}(\underline{V}_2^* \cdot \underline{V}_3^T)$							

Tabelle A.1: Drei Substitutionen in der Gleichung ([1]:A.42).

stitutionen erhalten hat, lassen sich auch als eine Gleichung

$$\mathbb{E}\{\hat{\vec{K}}\} = \left(\underline{E} + \mathbb{E}\{\underline{D}^{-1} \cdot \underline{S}\} \right) \cdot \tilde{\vec{K}} = \mathbb{E}\{\underline{F}\} \cdot \tilde{\vec{K}} \quad (\text{A.39})$$

schreiben, bei der auf beiden Seiten ein 3×1 Vektor steht. Bei dieser Gleichung enthalten die Vektoren

$$\hat{\vec{K}} = \begin{bmatrix} \underline{c}_1^{-1} \cdot \vec{N}_1 \cdot \underline{V}_1 \cdot \underline{V}_2^H \cdot \vec{N}_2^H \cdot \vec{N}_3 \cdot \underline{V}_3 \cdot \underline{V}_4^H \cdot \vec{N}_4^H \\ \underline{c}_2^{-1} \cdot \vec{N}_1 \cdot \underline{V}_1 \cdot \underline{V}_3^T \cdot \vec{N}_3^T \cdot \vec{N}_2^* \cdot \underline{V}_2^* \cdot \underline{V}_4^H \cdot \vec{N}_4^H \\ \underline{c}_3^{-1} \cdot \vec{N}_1 \cdot \underline{V}_1 \cdot \underline{V}_4^H \cdot \vec{N}_4^H \cdot \vec{N}_2^* \cdot \underline{V}_2^* \cdot \underline{V}_3^T \cdot \vec{N}_3^T \end{bmatrix} \quad (\text{A.40a})$$

als Elemente jeweils das Produkt zweier empirischer Kovarianzen. Die Produkte der theoretischen Kovarianzen derselben Zufallsgrößen sind in dem Vektor

$$\tilde{\vec{K}} = \begin{bmatrix} \mathbb{E}\{\underline{N}_1 \cdot \underline{N}_2^*\} \cdot \mathbb{E}\{\underline{N}_3 \cdot \underline{N}_4^*\} \\ \mathbb{E}\{\underline{N}_1 \cdot \underline{N}_3\} \cdot \mathbb{E}\{\underline{N}_2 \cdot \underline{N}_4\}^* \\ \mathbb{E}\{\underline{N}_1 \cdot \underline{N}_4\} \cdot \mathbb{E}\{\underline{N}_2 \cdot \underline{N}_3\}^* \end{bmatrix} \quad (\text{A.40b})$$

zusammengefasst. Der in Gleichung (A.39) gebildete Erwartungswert des Vektors $\hat{\vec{K}}$ enthält als Elemente somit die theoretischen, zweiten, *nichtzentralen* Momente der als Faktoren auftretenden empirischen Kovarianzen. Die Matrix, deren Erwartungswert diese zweiten Momente mit dem Vektor $\tilde{\vec{K}}$ verknüpft, und die in diesem Unterkapitel des Anhangs als $\underline{F}$ bezeichnet sei, enthält neben der Einheitsmatrix $\underline{E}$ noch die Diagonalmatrix

$$\underline{D} = \begin{bmatrix} \underline{c}_1 & 0 & 0 \\ 0 & \underline{c}_2 & 0 \\ 0 & 0 & \underline{c}_3 \end{bmatrix} \quad (\text{A.40c})$$

und eine weitere Matrix

$$\underline{\mathbf{S}} = \begin{bmatrix} 0 & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \underline{\mathbf{V}}_4^* \cdot \underline{\mathbf{V}}_3^T) & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_4^H) \\ \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_3^T \cdot \underline{\mathbf{V}}_4^* \cdot \underline{\mathbf{V}}_2^H) & 0 & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_3^T \cdot \underline{\mathbf{V}}_2^* \cdot \underline{\mathbf{V}}_4^H) \\ \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_4^H \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_2^H) & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_4^H \cdot \underline{\mathbf{V}}_2^* \cdot \underline{\mathbf{V}}_3^T) & 0 \end{bmatrix}, \quad (\text{A.40d})$$

deren Hauptdiagonale Null ist. Mit Anhang A.7 kann man zeigen, dass sowohl die Matrix $\underline{\mathbf{F}}$ als auch ihr Erwartungswert für hinreichend große Werte der Mittelungsanzahl L immer regulär ist, da die Zufallsmatrizen $\underline{\mathbf{V}}_1$, $\underline{\mathbf{V}}_2$, $\underline{\mathbf{V}}_3$ und $\underline{\mathbf{V}}_4$ idempotent sind und einen auf eine Konstante begrenzten Rangdefekt aufweisen. Die inverse Matrix sei in diesem Unterkapitel des Anhangs als $\underline{\mathbf{G}} = \underline{\mathbf{F}}^{-1}$ bezeichnet.

Da wir an den zweiten *zentralen* Momenten — also den Kovarianzen der empirischen Kovarianzen — interessiert sind, berechnen wir diese, indem wir von den zweiten *nichtzentralen* Momenten jeweils das Produkt der ersten Momente — also den Vektor $\tilde{\mathbf{K}}$ — subtrahieren:

$$\mathbb{E}\{\hat{\mathbf{K}}\} - \tilde{\mathbf{K}} = \mathbb{E}\{\underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}}\} \cdot \tilde{\mathbf{K}}. \quad (\text{A.41})$$

Nach Gleichung (2.41) und (2.51) sind die meisten der theoretischen Kovarianzen des Spektrums des gefensternten Approximationsfehlerprozesses, die im Hauptteil der Abhandlung in den Vektor $\tilde{\mathbf{K}}$ eingesetzt werden, in guter Näherung Null, wenn man eine hoch frequenzselektive Fensterfolge verwendet. Für diese theoretischen Kovarianzen wurden daher auch keine Schätzwerte bestimmt. Daher brauchen wir für die Zeilen des Vektors $\hat{\mathbf{K}}$, der als Faktor solch eine nicht berechnete empirische Kovarianz enthält auch keine eigene Gleichung aufzustellen. Wir erhalten so drei weitere Gleichungen, die sich alle in Vektorschreibweise wie Gleichung (A.39) darstellen lassen, die nun aber andere Vektoren und Matrizen enthalten. Falls eine der beiden theoretischen Varianzen des zweiten Elementes des Vektors $\tilde{\mathbf{K}}$ nach Gleichung (A.40b) Null ist, erhalten wir mit den Substitutionen nach Tabelle A.1 ohne die Zeile 2 die Vektoren und Matrizen

$$\begin{aligned} \hat{\mathbf{K}} &= \begin{bmatrix} \mathbf{c}_1^{-1} \cdot \tilde{\mathbf{N}}_1 \cdot \underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \tilde{\mathbf{N}}_2^H \cdot \tilde{\mathbf{N}}_3 \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_4^H \cdot \tilde{\mathbf{N}}_4^H \\ \mathbf{c}_3^{-1} \cdot \tilde{\mathbf{N}}_1 \cdot \underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_4^H \cdot \tilde{\mathbf{N}}_4^H \cdot \tilde{\mathbf{N}}_2^* \cdot \underline{\mathbf{V}}_2^* \cdot \underline{\mathbf{V}}_3^T \cdot \tilde{\mathbf{N}}_3^T \end{bmatrix}, \quad \tilde{\mathbf{K}} = \begin{bmatrix} \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} \cdot \mathbb{E}\{\mathbf{N}_3 \cdot \mathbf{N}_4^*\} \\ \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_4^*\} \cdot \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_3^*\}^* \end{bmatrix}, \\ \underline{\mathbf{D}} &= \begin{bmatrix} \mathbf{c}_1 & 0 \\ 0 & \mathbf{c}_3 \end{bmatrix} \quad \text{und} \quad \underline{\mathbf{S}} = \begin{bmatrix} 0 & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_4^H) \\ \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_4^H \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_2^H) & 0 \end{bmatrix}, \\ &\text{falls} \quad \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_3\} = 0 \quad \vee \quad \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_4\} = 0. \end{aligned} \quad (\text{A.42})$$

Ist eine der beiden theoretischen Varianzen des dritten Elementes des Vektors $\tilde{\mathbf{K}}$ nach Gleichung (A.40b) Null, so ergeben sich mit den Substitutionen nach Tabelle A.1 ohne

die Zeile 3 die Vektoren und Matrizen

$$\begin{aligned}\hat{\underline{\mathbf{K}}} &= \begin{bmatrix} \underline{\mathbf{c}}_1^{-1} \cdot \vec{\mathbf{N}}_1 \cdot \underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \vec{\mathbf{N}}_2^H \cdot \vec{\mathbf{N}}_3 \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_4^H \cdot \vec{\mathbf{N}}_4^H \\ \underline{\mathbf{c}}_2^{-1} \cdot \vec{\mathbf{N}}_1 \cdot \underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_3^T \cdot \vec{\mathbf{N}}_3^T \cdot \vec{\mathbf{N}}_2^* \cdot \underline{\mathbf{V}}_2^* \cdot \underline{\mathbf{V}}_4^H \cdot \vec{\mathbf{N}}_4^H \end{bmatrix}, \quad \tilde{\underline{\mathbf{K}}} = \begin{bmatrix} \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} \cdot \mathbb{E}\{\mathbf{N}_3 \cdot \mathbf{N}_4^*\} \\ \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_3\} \cdot \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_4\}^* \end{bmatrix}, \\ \underline{\mathbf{D}} &= \begin{bmatrix} \underline{\mathbf{c}}_1 & 0 \\ 0 & \underline{\mathbf{c}}_2 \end{bmatrix} \quad \text{und} \quad \underline{\mathbf{S}} = \begin{bmatrix} 0 & \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \underline{\mathbf{V}}_4^* \cdot \underline{\mathbf{V}}_3^T) \\ \text{spur}(\underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_3^T \cdot \underline{\mathbf{V}}_4^* \cdot \underline{\mathbf{V}}_2^H) & 0 \end{bmatrix}, \\ &\text{falls} \quad \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_4^*\} = 0 \quad \vee \quad \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_3^*\} = 0. \end{aligned} \quad (\text{A.43})$$

Ist sowohl eine der beiden theoretischen Varianzen des zweiten wie auch des dritten Elementes des Vektors $\tilde{\underline{\mathbf{K}}}$ nach Gleichung (A.40b) Null, so ergibt sich mit den Substitutionen der ersten Zeile nach Tabelle A.1 keine echte Vektorgleichung mehr:

$$\begin{aligned}\hat{\underline{\mathbf{K}}} &= \underline{\mathbf{c}}_1^{-1} \cdot \vec{\mathbf{N}}_1 \cdot \underline{\mathbf{V}}_1 \cdot \underline{\mathbf{V}}_2^H \cdot \vec{\mathbf{N}}_2^H \cdot \vec{\mathbf{N}}_3 \cdot \underline{\mathbf{V}}_3 \cdot \underline{\mathbf{V}}_4^H \cdot \vec{\mathbf{N}}_4^H, \quad \tilde{\underline{\mathbf{K}}} = \mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_2^*\} \cdot \mathbb{E}\{\mathbf{N}_3 \cdot \mathbf{N}_4^*\}, \\ \underline{\mathbf{D}} &= \underline{\mathbf{c}}_1 \quad \text{und} \quad \underline{\mathbf{S}} = 0, \end{aligned} \quad (\text{A.44})$$

$$\text{falls} \quad \left(\mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_3\} = 0 \vee \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_4\} = 0 \right) \wedge \left(\mathbb{E}\{\mathbf{N}_1 \cdot \mathbf{N}_4^*\} = 0 \vee \mathbb{E}\{\mathbf{N}_2 \cdot \mathbf{N}_3^*\} = 0 \right).$$

Um Schätzwerte für die Kovarianzen der empirischen Kovarianzen der Zufallsgrößen $\mathbf{N}_1$ bis $\mathbf{N}_4$ zu erhalten berechnen wir zunächst den Erwartungswert des Zufallsvektors $\underline{\mathbf{G}} \cdot \hat{\underline{\mathbf{K}}}$. Der Erwartungswert des i -ten Elementes dieses Zufallsvektors berechnet sich als Erwartungswert der Linearkombination der Elemente $\hat{\underline{\mathbf{K}}}_j$ des Vektors $\hat{\underline{\mathbf{K}}}$ mit den Elementen $\underline{\mathbf{G}}_{i,j}$ der i -ten Zeile der zu $\underline{\mathbf{F}}$ inversen Matrix $\underline{\mathbf{G}}$ als Koeffizienten. Diese Koeffizienten sind unabhängig von den Zufallsgrößen $\mathbf{N}_1$ bis $\mathbf{N}_4$, da diese Funktionen der Elemente der Zufallsmatrizen $\underline{\mathbf{V}}_1$ bis $\underline{\mathbf{V}}_4$, und somit Funktionen der Spektralwerte der zufälligen Erregung sind. Der Erwartungswert dieser Linearkombination ist die Summe der Erwartungswerte der einzelnen Summanden. Mit Hilfe der Gleichung ([1]:A.42) kann man nun den Erwartungswert jedes einzelnen Summanden dieser Linearkombination als eine Summe dreier Summanden berechnen, wobei es je nach Zuordnung der Spektralwerte des gefensterten Approximationsfehlerprozesses zu den Zufallsgrößen $\mathbf{N}_1$ bis $\mathbf{N}_4$ auch vorkommen kann, dass ein oder zwei Summanden verschwinden, wenn sie eine theoretische Kovarianz Null als Faktor enthalten. Um die Gleichung ([1]:A.42) anzuwenden, substituiert man die Zufallsvektoren $\vec{\mathbf{N}}_1$ bis $\vec{\mathbf{N}}_4$ und die Zufallsmatrizen $\underline{\mathbf{A}}$ und $\underline{\mathbf{B}}$ in Gleichung ([1]:A.42) durch die in Tabelle A.1 angegebenen Zufallsvektoren und -matrizen. Der Zufallsfaktor $\underline{\mathbf{c}}$ ist nun jedoch noch mit dem entsprechenden Element $\underline{\mathbf{G}}_{i,j}$ der i -ten Zeile der Zufallsmatrix $\underline{\mathbf{G}}$ zu multiplizieren. Man erhält so für jeden einzelnen Summanden eine Linearkombination der entsprechenden Elemente $\tilde{\underline{\mathbf{K}}}_k$ des Vektors $\tilde{\underline{\mathbf{K}}}$, wobei die Koeffizienten dieser Linearkombination die Erwartungswerte der $\underline{\mathbf{G}}_{i,j}$ -fachen Elemente $\underline{\mathbf{F}}_{j,k}$ der j -ten Zeile der Zufallsmatrix $\underline{\mathbf{F}}$ sind. Als Gleichung geschrieben ergibt sich für den Erwartungswert

des i -ten Elementes des Zufallsvektors $\underline{\mathbf{G}} \cdot \hat{\underline{\mathbf{K}}}$:

$$\begin{aligned} \sum_j \mathbb{E}\{\mathbf{G}_{i,j} \cdot \hat{\mathbf{K}}_j\} &= \sum_j \sum_k \mathbb{E}\{\mathbf{G}_{i,j} \cdot \mathbf{F}_{j,k}\} \cdot \tilde{K}_k = \sum_k \mathbb{E}\left\{ \underbrace{\sum_j \mathbf{G}_{i,j} \cdot \mathbf{F}_{j,k}} \right\} \cdot \tilde{K}_k = \tilde{K}_i. \\ &= \begin{cases} 1 & \text{für } k=i \\ 0 & \text{sonst} \end{cases} \end{aligned} \quad (\text{A.45})$$

Für den Erwartungswert des Zufallsvektors ergibt sich das durchaus nicht triviale Ergebnis

$$\mathbb{E}\{\underline{\mathbf{F}}^{-1} \cdot \hat{\underline{\mathbf{K}}}\} = \mathbb{E}\{\underline{\mathbf{G}} \cdot \hat{\underline{\mathbf{K}}}\} = \tilde{\underline{\mathbf{K}}} = \mathbb{E}\{\underline{\mathbf{F}}\}^{-1} \cdot \mathbb{E}\{\hat{\underline{\mathbf{K}}}\}, \quad (\text{A.46})$$

das besagt, dass sich der Erwartungswert des mit der Zufallsmatrix $\underline{\mathbf{G}}$ abgebildeten Zufallsvektors $\hat{\underline{\mathbf{K}}}$ in ein Produkt der Inversen des Erwartungswertes einer Zufallsmatrix und den Erwartungswert des Vektors faktorisieren lässt, obwohl beide Faktoren von den Spektralwerten der zufälligen Erregung abhängen. Mit dieser Erkenntnis kann man nun zeigen, dass die Zufallsgrößen des Zufallsvektors

$$\underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}} \cdot (\underline{\mathbf{E}} + \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}})^{-1} \cdot \hat{\underline{\mathbf{K}}} \quad (\text{A.47})$$

als Schätzwerte für die Kovarianzen der empirischen Kovarianzen der Zufallsgrößen $\mathbf{N}_1$ bis $\mathbf{N}_4$ dienen können, da ihre Erwartungswerte

$$\begin{aligned} &\mathbb{E}\left\{ \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}} \cdot (\underline{\mathbf{E}} + \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}})^{-1} \cdot \hat{\underline{\mathbf{K}}} \right\} = \\ &= \mathbb{E}\left\{ \left((\underline{\mathbf{E}} + \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}}) - \underline{\mathbf{E}} \right) \cdot (\underline{\mathbf{E}} + \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}})^{-1} \cdot \hat{\underline{\mathbf{K}}} \right\} = \\ &= \mathbb{E}\{\hat{\underline{\mathbf{K}}}\} - \mathbb{E}\left\{ (\underline{\mathbf{E}} + \underline{\mathbf{D}}^{-1} \cdot \underline{\mathbf{S}})^{-1} \cdot \hat{\underline{\mathbf{K}}} \right\} = \mathbb{E}\{\hat{\underline{\mathbf{K}}}\} - \tilde{\underline{\mathbf{K}}} \end{aligned} \quad (\text{A.48})$$

mit den in Gleichung (A.41) angegebenen theoretischen Kovarianzen der empirischen Kovarianzen dieser Zufallsgrößen übereinstimmen.

A.9 Zum Vergleich der Beträge der empirischen Varianz und Kovarianz

Um zu zeigen, dass die konkreten Varianzschätzwerte niemals kleiner sind als die Beträge der entsprechenden konkreten Kovarianzschätzwerte, sind an zwei Stellen etwas aufwendigere Überlegungen anzustellen. Um den Ablauf des Textes im Hauptteil der Abhandlung durch dieses nicht so wichtige und im weiteren nicht mehr verwendete Detail nicht allzu sehr zu unterbrechen, wurden diese beiden Herleitungen in den Anhang verbannt. Zum einen handelt es sich dabei um die Messwert(ko)varianzen des LDS und des KLDS und zum anderen um die Messwert(ko)varianzen der deterministischen Störung.

A.9.1 Messwerte des LDS und des KLDS

Zunächst werden zwei Sätze über die Beträge von Determinanten hergeleitet, die wir für den Beweis der Formel, die wir im Hauptteil der Abhandlung brauchen, benötigen.

Der erste Satz bezieht sich auf den Betrag der Determinante eines Produktes einer Matrix mit ihrer Transponierten. Die Matrix besteht dabei aus einer Auswahl von Zeilen einer unitären Matrix. Da eine Zeilenpermutation weder den Betrag der Determinante noch die Unitaritätseigenschaft einer Matrix beeinflusst, kann man die unitäre Matrix immer so permutieren, dass die ausgewählten Zeilenvektoren die obersten Zeilen der unitären Matrix bilden. Daher wird der folgende Satz ohne Beschränkung der Allgemeinheit für den Fall gezeigt, dass die ausgewählten Zeilenvektoren die ersten Zeilen der unitären Matrix sind.

Gegeben sei eine unitäre $L \times L$ Matrix $\underline{V}$ deren erste R Zeilen die Matrix $\underline{V}_1$ bilden, und deren restliche Zeilen in der Matrix $\underline{V}_2$ zusammengefasst sind.

$$\underline{V} = \left[\begin{array}{c} \overbrace{\underline{V}_1}^{L \text{ Spalten}} \\ \underline{V}_2 \end{array} \right] \}^{R \text{ Zeilen}} \quad \text{mit} \quad \underline{V}^H = \underline{V}^{-1} \quad (\text{A.49})$$

Satz: Der Betrag der Determinante des Produktes $\underline{V}_1 \cdot \underline{V}_1^T$ ist immer kleiner gleich 1.

Wenn man mit s_i die R Singulärwerte der Matrix $\underline{V}_1 \cdot \underline{V}_1^T$ bezeichnet, kann man unter Ausnutzung der Sätze,

- „Die Spektralnorm einer Matrix ist ihr größter Singulärwert“,
- „Die Norm eines Matrixprodukts ist niemals größer als das Produkt der Normen der daran beteiligten Matrizen“,
- „Die Spektralnorm einer unitären Matrix ist 1“,
- „Die Spektralnorm einer Matrix, deren Zeilen beliebig aus den Zeilen einer unitären Matrix ausgewählt worden sind, ist 1“ und
- „Das Produkt der Singulärwerte einer Matrix ist gleich dem Betrag der Determinante einer Matrix“

die Gültigkeit dieses Satzes folgendermaßen zeigen:

$$1 = \|\underline{V}_1\|_2 \cdot \|\underline{V}_1^T\|_2 \geq \|\underline{V}_1 \cdot \underline{V}_1^T\|_2 = \max_i(s_i) \geq \max_i(s_i)^R \geq \prod_{i=1}^R s_i = |\det(\underline{V}_1 \cdot \underline{V}_1^T)|. \quad (\text{A.50})$$

Die Determinante des Produktes $\underline{V}_1 \cdot \underline{V}_1^H$ hingegen ist immer 1, da dieses Produkt wegen der Unitarität der Matrix $\underline{V}$, aus der die Zeilen von $\underline{V}_1$ entnommen wurden, immer die $R \times R$ Einheitsmatrix ist.

Der zweite Satz vergleicht die Determinante eines Produktes einer Matrix mit ihrer konjugiert Transponierten mit dem Betrag der Determinante des Produktes derselben Matrix mit ihrer Transponierten.

Gegeben sei eine $R \times L$ Matrix $\underline{M}$.

Satz: Der Betrag der Determinante des Produktes $\underline{M} \cdot \underline{M}^H$ ist immer größer gleich dem Betrag der Determinante des Produktes $\underline{M} \cdot \underline{M}^T$.

Um dies zu zeigen, betrachten wir zwei Fälle. Im ersten Fall ist $R > L$. Da alle Spaltenvektoren der beiden Matrixprodukte $\underline{M} \cdot \underline{M}^H$ und $\underline{M} \cdot \underline{M}^T$ Linearkombinationen der L Spaltenvektoren der Matrix $\underline{M}$ sind, und diese L Spaltenvektoren schon wegen ihrer zu geringen Anzahl keine Basis der R -dimensionalen Raums sein können, ist die Determinante beider Matrixprodukte immer Null, so dass der Satz mit dem Gleichheitszeichen erfüllt ist.

Im zweiten Fall $R \leq L$ kann man unter Ausnutzung des zuletzt gezeigten Satzes und der Sätze

- „Jede Matrix lässt sich mit einer unitären $R \times R$ Matrix $\underline{U}$, einer unitären $L \times L$ Matrix $\underline{V}$, deren ersten R Zeilen die Matrix $\underline{V}_1$ bilden, der $R \times R$ Diagonalmatrix $\underline{S}$ ihrer reellen, nichtnegativen Singulärwerte und der $R \times (L - R)$ Nullmatrix $\underline{0}$ in der Form

$$\underline{M} = \underline{U} \cdot [\underline{S}, \underline{0}] \cdot \underline{V} = \underline{U} \cdot [\underline{S} \cdot \underline{V}_1, \underline{0}] \quad (\text{A.51})$$

darstellen“,

- „Die Determinante eines Produkts quadratischer Matrizen ist gleich dem Produkt der Determinanten der einzelnen daran beteiligten Matrizen“,
- „Der Betrag der Determinante einer unitären Matrix ist 1“ und
- „Der Betrag eines Produktes komplexer Zahlen ist gleich dem Produkt der Beträge der daran beteiligten komplexen Zahlen“

die Gültigkeit des Satzes folgendermaßen zeigen:

$$\begin{aligned} \det(\underline{M} \cdot \underline{M}^H) &= \\ &= \det\left(\underline{U} \cdot [\underline{S} \cdot \underline{V}_1, \underline{0}] \cdot [\underline{S} \cdot \underline{V}_1, \underline{0}]^H \cdot \underline{U}^H\right) = \\ &= \det\left(\underline{U} \cdot \underline{S} \cdot \underline{V}_1 \cdot \underline{V}_1^H \cdot \underline{S}^H \cdot \underline{U}^H\right) = \end{aligned}$$

$$\begin{aligned}
&= \det(\underline{U}) \cdot \det(\underline{S}) \cdot \det(\underline{V}_1 \cdot \underline{V}_1^H) \cdot \det(\underline{S}^H) \cdot \det(\underline{U}^H) = \\
&= \det(\underline{U} \cdot \underline{U}^H) \cdot \det(\underline{S}) \cdot 1 \cdot \det(\underline{S}^H) = \\
&= |\det(\underline{S})| \cdot |\det(\underline{S})| \geq \\
&\geq |\det(\underline{U})| \cdot |\det(\underline{S})| \cdot |\det(\underline{V}_1 \cdot \underline{V}_1^T)| \cdot |\det(\underline{S}^T)| \cdot |\det(\underline{U}^T)| = \\
&= |\det(\underline{U}) \cdot \det(\underline{S}) \cdot \det(\underline{V}_1 \cdot \underline{V}_1^T) \cdot \det(\underline{S}^T) \cdot \det(\underline{U}^T)| = \\
&= |\det(\underline{U} \cdot \underline{S} \cdot \underline{V}_1 \cdot \underline{V}_1^T \cdot \underline{S}^T \cdot \underline{U}^T)| = \\
&= |\det(\underline{U} \cdot [\underline{S} \cdot \underline{V}_1, \underline{0}] \cdot [\underline{S} \cdot \underline{V}_1, \underline{0}]^T \cdot \underline{U}^T)| = \\
&= |\det(\underline{M} \cdot \underline{M}^T)| \tag{A.52}
\end{aligned}$$

Nun wählen wir $R = 2$ und bezeichnen die beiden Zeilenvektoren der $2 \times L$ Matrix $\underline{M}$ mit $\vec{X}$ und $\vec{Y}^*$:

$$\underline{M} = \begin{bmatrix} \vec{X} \\ \vec{Y}^* \end{bmatrix} \tag{A.53}$$

Setzen wir dies in Ungleichung (A.52) ein, und berechnen wir die Determinanten, so erhalten wir die Ungleichung

$$\begin{aligned}
&\det\left(\begin{bmatrix} \vec{X} \\ \vec{Y}^* \end{bmatrix} \cdot [\vec{X}^H, \vec{Y}^T]\right) = \vec{X} \cdot \vec{X}^H \cdot \vec{Y}^* \cdot \vec{Y}^T - \vec{Y}^* \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^T \geq \\
&\geq \det\left(\begin{bmatrix} \vec{X} \\ \vec{Y}^* \end{bmatrix} \cdot [\vec{X}^T, \vec{Y}^H]\right) = |\vec{X} \cdot \vec{X}^T \cdot \vec{Y}^* \cdot \vec{Y}^H - \vec{Y}^* \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^H| \geq \\
&\geq \left| |\vec{X} \cdot \vec{X}^T \cdot \vec{Y}^* \cdot \vec{Y}^H| - |\vec{Y}^* \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^H| \right| \geq \\
&\geq |\vec{X} \cdot \vec{X}^T \cdot \vec{Y}^* \cdot \vec{Y}^H| - |\vec{Y}^* \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^H|, \tag{A.54}
\end{aligned}$$

wobei wir bei der Determinante auf der rechten Seite mit der Dreiecksungleichung eine untere Grenze angegeben haben. Auf beiden Seiten der Ungleichung addiert man nun die Terme mit dem negativen Vorzeichen. Dann wird weiter umgeformt, wobei man berücksichtigt, dass das Produkt komplexer Zahlen kommutativ ist, dass bei einem Produkt komplexer Zahlen die Faktoren nach belieben konjugiert werden können, ohne dass sich der Betrag des Produkts ändert, dass bei dem Produkt einer komplexen Zahl mit ihrer Konjugierten die Betragsbildung unterbleiben kann, dass sich das Skalarprodukt durch Transponieren nicht ändert, und dass das Konjugieren bei einer reellen Zahl ohne Wirkung bleibt. Wir multiplizieren die Ungleichung dann mit der reellen Zahl a , von der wir festlegen, dass für sie $a \geq 2$ gilt, und addieren auf beiden Seiten den gleichen reellen Term.

Schließlich verwenden wir erneut die Dreiecksungleichung um zu der Form zu gelangen, die im Hauptteil der Abhandlung benötigt wird.

$$\begin{aligned}
\vec{X} \cdot \vec{X}^H \cdot \vec{Y}^* \cdot \vec{Y}^T + |\vec{Y}^* \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^H| &\geq |\vec{X} \cdot \vec{X}^T \cdot \vec{Y}^* \cdot \vec{Y}^H| + \vec{Y}^* \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^T & (A.55) \\
\vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + |\vec{Y} \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^H| &\geq |\vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + |\vec{Y}^* \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^T| \\
\vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + \vec{Y} \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^H &\geq |\vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + |\vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| \\
a \cdot \vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + a \cdot \vec{Y} \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^H &\geq a \cdot |\vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + a \cdot |\vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| \\
a \cdot \vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + a \cdot \vec{Y} \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^H - 2 \cdot |\vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| &\geq \\
&\geq a \cdot |\vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + a \cdot |\vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| - 2 \cdot |\vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| \\
a \cdot \vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + a \cdot \vec{Y} \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^H - 2 \cdot \vec{Y}^* \cdot \vec{X}^H \cdot \vec{X} \cdot \vec{Y}^T &\geq \\
&\geq |a \cdot \vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + |(a-2) \cdot \vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T| \\
a \cdot \vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H + a \cdot \vec{X} \cdot \vec{Y}^H \cdot \vec{Y} \cdot \vec{X}^H - 2 \cdot \vec{X} \cdot \vec{Y}^T \cdot \vec{Y}^* \cdot \vec{X}^H &\geq \\
&\geq |a \cdot \vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T| + (a-2) \cdot \vec{Y} \cdot \vec{X}^T \cdot \vec{X} \cdot \vec{Y}^T \\
- 2 \cdot \vec{X} \cdot \vec{Y}^T \cdot \vec{Y}^* \cdot \vec{X}^H + a \cdot \vec{X} \cdot \vec{Y}^H \cdot \vec{Y} \cdot \vec{X}^H + a \cdot \vec{X} \cdot \vec{X}^H \cdot \vec{Y} \cdot \vec{Y}^H &\geq \\
&\geq |(a-2) \cdot \vec{X} \cdot \vec{Y}^T \cdot \vec{Y} \cdot \vec{X}^T| + a \cdot \vec{X} \cdot \vec{X}^T \cdot \vec{Y} \cdot \vec{Y}^T
\end{aligned}$$

A.9.2 Messwerte der deterministischen Störung

Um zu zeigen, dass die konkreten Varianzschätzwerte der Messwerte $\hat{u}(k)$ des Approximationsfehlerprozesses niemals kleiner sind als die Beträge der entsprechenden konkreten Kovarianzschätzwerte, beginnen wir damit, in den Gleichungen (3.53) die Grenzen der Laufindizes der Doppelsumme zu verändern. Statt die Summe über μ von 0 bis $M-1$ gehen zu lassen, starten wir bei $\hat{\mu} \cdot M / K_\Phi$ und enden bei $M + \hat{\mu} \cdot M / K_\Phi - 1$. Da die Summanden mit M periodisch sind, wird dadurch lediglich die Reihenfolge der Summanden zyklisch permutiert. Wenn wir für alle Werte $\hat{\mu} = 0 \text{ (1) } K_\Phi - 1$ die so entstandenen Doppelsummen addieren, so erhalten wir eine Dreifachsumme, die der mit K_Φ multiplizierten Doppelsumme entspricht. Diese Multiplikation kompensieren wir indem wir die Dreifachsumme auf K_Φ normieren. Anschließend ersetzen wir noch die Laufindizes μ und $\tilde{\mu}$ durch die neuen Laufindizes $\check{\mu} = \mu - \hat{\mu} \cdot M / K_\Phi$ und $\bar{\mu} = \hat{\mu} + \tilde{\mu}$ und wir erhalten, wenn wir wieder die M -Periodizität der Varianzen und Kovarianzen der Messwerte des Spektrums der

deterministischen Störung ausnutzen, folgenden Dreifachsummen.

$$\begin{aligned}
\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\mu=\hat{\mu} \frac{M}{K_\Phi}}^{M+\hat{\mu} \frac{M}{K_\Phi}-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu+\tilde{\mu} \frac{M}{K_\Phi})} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\check{\mu}=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(\check{\mu}+(\hat{\mu}+\tilde{\mu}) \frac{M}{K_\Phi})} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\check{\mu}=0}^{M-1} \sum_{\bar{\mu}=\hat{\mu}}^{K_\Phi+\hat{\mu}-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(\check{\mu}+\bar{\mu} \frac{M}{K_\Phi})} \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu}-\bar{\mu}) \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\check{\mu}=0}^{M-1} \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\bar{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(\check{\mu}+\bar{\mu} \frac{M}{K_\Phi})} \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu}-\bar{\mu}) \cdot k} \\
&\quad \forall \quad k = 0 \ (1) \ F-1
\end{aligned} \tag{A.56a}$$

und

$$\begin{aligned}
\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)^*} &= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\mu=\hat{\mu} \frac{M}{K_\Phi}}^{M+\hat{\mu} \frac{M}{K_\Phi}-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\mu), \hat{U}_f(-\mu-\tilde{\mu} \frac{M}{K_\Phi})^*} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\check{\mu}=0}^{M-1} \sum_{\tilde{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(-\check{\mu}-(\hat{\mu}+\tilde{\mu}) \frac{M}{K_\Phi})^*} \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \tilde{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\check{\mu}=0}^{M-1} \sum_{\bar{\mu}=\hat{\mu}}^{K_\Phi+\hat{\mu}-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(-\check{\mu}-\bar{\mu} \frac{M}{K_\Phi})^*} \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu}-\bar{\mu}) \cdot k} = \\
&= \frac{1}{M^2 \cdot K_\Phi} \cdot \sum_{\check{\mu}=0}^{M-1} \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\bar{\mu}=0}^{K_\Phi-1} \hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \frac{M}{K_\Phi}), \hat{U}_f(-\check{\mu}-\bar{\mu} \frac{M}{K_\Phi})^*} \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu}-\bar{\mu}) \cdot k} \\
&\quad \forall \quad k = 0 \ (1) \ F-1.
\end{aligned} \tag{A.56b}$$

Für die Varianzen und Kovarianzen der Messwerte des Spektrums der gefensterten deterministischen Störung setzen wir die in den Gleichungen (3.48) angegebenen Schätzwerte $\hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \cdot M/K_\Phi), \hat{U}_f(\check{\mu}+\bar{\mu} \cdot M/K_\Phi)}$ und $\hat{C}_{\hat{U}_f(\check{\mu}+\hat{\mu} \cdot M/K_\Phi), \hat{U}_f(-\check{\mu}-\bar{\mu} \cdot M/K_\Phi)^*}$ in der Form ein, die die Vektoren $\hat{\tilde{C}}_U(\dots)$ enthält. Bei diesen Schätzwerten wiederum setzen wir die Messwerte $\hat{\Phi}_{\mathbf{n}}(\dots, \dots)$ und $\hat{\Psi}_{\mathbf{n}}(\dots, \dots)$ gemäß der Gleichungen (3.34) in der Form mit den Vektoren

$\hat{N}_f(\dots)$ ein. Da diese Messwerte mit Matrizen $V_{\perp}(\dots)$ gewonnen wurden, die die Bedingung (3.33) erfüllen, ist dort der Vorfaktor $M^{-1} \cdot (L-1-K(\check{\mu}))^{-1}$ bei allen Werten $\hat{\mu}$ und $\bar{\mu}$ gleich. Er hängt nur von $\check{\mu}$ ab. Somit erhalten wir:

$$\begin{aligned}
 \hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \tag{A.57a} \\
 &= \frac{1}{M \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\check{\mu}=0}^{M-1} \sum_{\hat{\mu}=0}^{K_{\Phi}-1} \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \hat{C}_U(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \hat{C}_U(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \\
 &\quad \cdot \hat{\Phi}_{\mathbf{n}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}}, \check{\mu} + \bar{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} = \\
 &= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \sum_{\hat{\mu}=0}^{K_{\Phi}-1} \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \hat{C}_U(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \hat{C}_U(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \\
 &\quad \cdot \hat{N}_f(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \hat{N}_f(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} \\
 &\quad \forall \quad k = 0 \text{ (1) } F-1
 \end{aligned}$$

und

$$\begin{aligned}
 \hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)^*} &= \tag{A.57b} \\
 &= \frac{1}{M \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\check{\mu}=0}^{M-1} \sum_{\hat{\mu}=0}^{K_{\Phi}-1} \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \hat{C}_U(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_{\Phi}})^* \cdot \hat{C}_U(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \\
 &\quad \cdot \hat{\Psi}_{\mathbf{n}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}}, \check{\mu} + \bar{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} = \\
 &= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \sum_{\hat{\mu}=0}^{K_{\Phi}-1} \sum_{\bar{\mu}=0}^{K_{\Phi}-1} \hat{C}_U(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_{\Phi}})^* \cdot \hat{C}_U(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}})^H \cdot \\
 &\quad \cdot \hat{N}_f(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot \hat{N}_f(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_{\Phi}})^T \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} \\
 &\quad \forall \quad k = 0 \text{ (1) } F-1.
 \end{aligned}$$

Jeder Summand dieser Dreifachsummen ist ein Produkt von zwei Skalarprodukten von jeweils zwei Vektoren mit je L Elementen. Ausführlich geschrieben erhalten wir die Fünffachsummen

$$\begin{aligned}
\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \tag{A.58a} \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\bar{\mu}=0}^{K_\Phi-1} \sum_{\check{\lambda}=1}^L \hat{C}_{U,\check{\lambda}}(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{C}_{U,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi})^* \\
&\quad \cdot \sum_{\check{\lambda}=1}^L \hat{N}_{f,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{N}_{f,\check{\lambda}}(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_\Phi})^* \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} = \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \\
&\quad \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \hat{C}_{U,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{N}_{f,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi}) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \hat{\mu} \cdot k} \cdot \\
&\quad \cdot \sum_{\bar{\mu}=0}^{K_\Phi-1} \hat{C}_{U,\check{\lambda}}(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{N}_{f,\check{\lambda}}(\check{\mu} + \bar{\mu} \cdot \frac{M}{K_\Phi})^* \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot |C_{\check{\lambda},\check{\lambda}}(\check{\mu}, k)|^2 \\
&\quad \forall \quad k = 0 \text{ (1) } F-1
\end{aligned}$$

$$\begin{aligned}
\text{und } \hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)}^* &= \tag{A.58b} \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \sum_{\bar{\mu}=0}^{K_\Phi-1} \sum_{\check{\lambda}=1}^L \hat{C}_{U,\check{\lambda}}(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{C}_{U,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi})^* \\
&\quad \cdot \sum_{\check{\lambda}=1}^L \hat{N}_{f,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi}) \cdot \hat{N}_{f,\check{\lambda}}(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_\Phi}) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot (\hat{\mu} - \bar{\mu}) \cdot k} = \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \\
&\quad \cdot \sum_{\hat{\mu}=0}^{K_\Phi-1} \hat{C}_{U,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{N}_{f,\check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_\Phi}) \cdot e^{j \cdot \frac{2\pi}{K_\Phi} \cdot \hat{\mu} \cdot k} \cdot \\
&\quad \cdot \sum_{\bar{\mu}=0}^{K_\Phi-1} \hat{C}_{U,\check{\lambda}}(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_\Phi})^* \cdot \hat{N}_{f,\check{\lambda}}(-\check{\mu} - \bar{\mu} \cdot \frac{M}{K_\Phi}) \cdot e^{-j \cdot \frac{2\pi}{K_\Phi} \cdot \bar{\mu} \cdot k} = \\
&= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_\Phi} \cdot \sum_{\check{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot C_{\check{\lambda},\check{\lambda}}(\check{\mu}, k) \cdot C_{\check{\lambda},\check{\lambda}}(-\check{\mu}, k) \\
&\quad \forall \quad k = 0 \text{ (1) } F-1.
\end{aligned}$$

Dabei wurde zuletzt jeweils die Abkürzung

$$C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k) = \sum_{\check{\mu}=0}^{K_{\Phi}-1} \hat{C}_{U, \hat{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}})^* \cdot \hat{N}_{f, \check{\lambda}}(\check{\mu} + \hat{\mu} \cdot \frac{M}{K_{\Phi}}) \cdot e^{j \cdot \frac{2\pi}{K_{\Phi}} \cdot \hat{\mu} \cdot k} \quad (\text{A.59})$$

$$\forall \quad k = 0 \ (1) \ F-1, \quad \check{\mu} = 0 \ (1) \ M-1, \quad \hat{\lambda} = 1 \ (1) \ L \quad \text{und} \quad \check{\lambda} = 1 \ (1) \ L$$

eingeführt. In Gleichung (A.58a) ergibt sich dieselbe Summe, wenn wir über $-\check{\mu}$ statt über $\check{\mu}$ summieren, da dies wegen der M -Periodizität der Summanden lediglich die Reihenfolge der Summanden umkehrt. Damit können wir diese Dreifachsumme auch als die Hälfte der Dreifachsumme mit $-\check{\mu}$ und der Dreifachsumme mit $\check{\mu}$, also als die Dreifachsumme der arithmetischen Mittel der Betragsquadrate von $C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k)$ und $C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)$ schreiben. Da das Arithmetische Mittel zweier reeller nichtnegativer Zahlen nie kleiner ist als deren geometrisches Mittel gilt

$$\begin{aligned} \frac{|C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k)|^2 + |C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)|^2}{2} &\geq \sqrt{|C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k)|^2 \cdot |C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)|^2} = \\ &= |C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k)| \cdot |C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)| = |C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k) \cdot C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)| \end{aligned} \quad (\text{A.60})$$

$$\forall \quad k = 0 \ (1) \ F-1, \quad \check{\mu} = 0 \ (1) \ M-1, \quad \hat{\lambda} = 1 \ (1) \ L \quad \text{und} \quad \check{\lambda} = 1 \ (1) \ L.$$

Da der Betrag einer Summe niemals größer ist als die Summe der Beträge der Summanden, kann man nun für $L > K(\check{\mu}) + 1$ beweisen, dass der konkrete Varianzschätzwert des Messwertes $\hat{u}(k)$ des Approximationsfehlerprozesses zum Zeitpunkt k niemals kleiner als der Betrag des konkreten Kovarianzschätzwertes ist.

$$\begin{aligned} \hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)} &= \\ &= \frac{1}{M^2 \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\hat{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot \frac{|C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k)|^2 + |C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)|^2}{2} \geq \\ &\geq \frac{1}{M^2 \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\hat{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \left| \frac{C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k) \cdot C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k)}{L-1-K(\check{\mu})} \right| \geq \\ &\geq \left| \frac{1}{M^2 \cdot (L-1)^2 \cdot K_{\Phi}} \cdot \sum_{\hat{\lambda}=1}^L \sum_{\check{\lambda}=1}^L \sum_{\check{\mu}=0}^{M-1} \frac{1}{L-1-K(\check{\mu})} \cdot C_{\hat{\lambda}, \check{\lambda}}(\check{\mu}, k) \cdot C_{\hat{\lambda}, \check{\lambda}}(-\check{\mu}, k) \right| = \\ &= |\hat{C}_{\hat{\mathbf{u}}(k), \hat{\mathbf{u}}(k)}^*| \\ &\forall \quad k = 0 \ (1) \ F-1 \end{aligned} \quad (\text{A.61})$$

A.10 Anmerkung zur Gauß-Verbundverteilung

Bei der Herleitung der Varianzen und Kovarianzen der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ und $\hat{\Psi}_{\mathbf{n}}(\mu)$ eines stationären Approximationsfehlerprozesses in Kapitel [1]:3.6 bzw. $\hat{\Phi}_{\mathbf{n}}(\mu_1, \mu_2)$ und $\hat{\Psi}_{\mathbf{n}}(\mu_1, \mu_2)$ eines zyklstationären Approximationsfehlerprozesses in Kapitel 3.4 hatten wir vorausgesetzt, dass das Zufallsgrößentupel $[\mathbf{N}_f(\mu), \mathbf{N}_f(-\mu)]^T$ bzw. die Zufallsgrößentupel $[\mathbf{N}_f(\mu_1), \mathbf{N}_f(\mu_2)]^T$ und $[\mathbf{N}_f(\mu_1), \mathbf{N}_f(-\mu_2)]^T$ eine normalverteilte mittelwertfreie Verbundverteilung aufweisen. Die Mittelwertfreiheit ist nicht das entscheidende Kriterium für die Entscheidung, ob man annehmen kann, dass die Forderung erfüllt ist. Wenn die Verbundverteilung des Zufallsgrößentupels eine Normalverteilung ist, so ist sie mittelwertfrei, wenn jede einzelne der beiden Zufallsgrößen des Zufallsgrößentupels mittelwertfrei ist. In [1] hatten wir uns darauf beschränkt, die dort vorgestellte Variante des RKM nur für mittelwertfreie Approximationsfehlerprozesse zu verwenden. Bei der hier vorgestellten RKM-Variante sorgt die Modellierung der deterministischen Störung dafür, dass die Mittelwertfreiheit gegeben ist, wie Gleichung (2.30) zeigt. Die Forderung, dass eine Verbundnormalverteilung vorliegen muss, besagt *nicht*, dass die gemeinsame Verbundverteilung *aller* Spektralwerte $\mathbf{N}_f(\mu)$ für alle $\mu = 0$ (1) $M-1$ eine Normalverteilung sein muss. Es gibt Zufallsvektoren, die *keine* gemeinsame Verbundnormalverteilung aufweisen, bei denen aber alle zweidimensionalen Randverteilungen zweier beliebigen Zufallsgrößen dieses Zufallsvektors Verbundnormalverteilungen sind. Ich möchte dies an einem Beispiel demonstrieren. Um ein anschauliches Beispiel zu erhalten, werde ich dies nicht an einem beliebig dimensionalen Zufallsvektor und an zweidimensionalen Randverteilungen zeigen, sondern ein Beispiel angeben, bei dem beide eindimensionalen Randverteilungen eines zweidimensionalen *nicht* verbundnormalverteilten Zufallsvektors jeweils eindimensionale Normalverteilungen sind. Es bleibt dem Leser überlassen, dieses Beispiel auf einen mehrdimensionalen Zufallsvektor und dessen zweidimensionale Randverteilungen zu übertragen.

Gegeben sei ein zweidimensionaler reeller Zufallsvektor

$$\vec{n} = \begin{bmatrix} n_1 \\ n_2 \end{bmatrix}, \quad (\text{A.62})$$

der die zweidimensionale Verbundverteilungsdichte

$$p_{\vec{n}}(\vec{n}) = \begin{cases} 0 & \text{für } n_2 > n_1 \geq 0 \vee -n_1 < n_2 \leq 0 \vee n_2 < n_1 \leq 0 \vee -n_1 > n_2 \geq 0 \\ \frac{1}{\pi} \cdot e^{-\frac{1}{2} \cdot \|\vec{n}\|^2} = \frac{1}{\pi} \cdot e^{-\frac{1}{2} \cdot (n_1^2 + n_2^2)} & \text{sonst} \end{cases} \quad (\text{A.63})$$

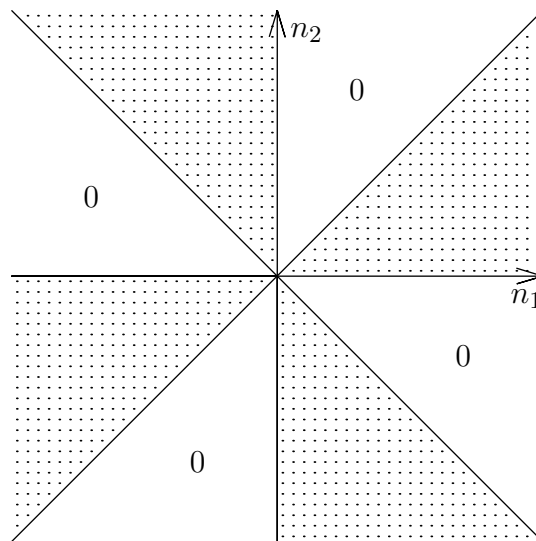


Bild A.1: Gebiet der n_1 - n_2 -Ebene, in dem die Verbundverteilung zweier unkorrelierter, einzeln normalverteilter aber abhängiger Zufallsgrößen von null verschieden ist.

besitze, die in der halben n_1 - n_2 -Ebene null ist, und in der anderen Hälfte der n_1 - n_2 -Ebene doppelt so groß ist, wie die Verbundverteilungsdichte eines zweidimensionalen, unkorrelierten normierten, zentrierten und verbundnormalverteilten Zufallsgrößentupels. In Bild A.1 ist das Gebiet der n_1 - n_2 -Ebene, in dem die Verbundverteilungsdichte von null verschieden ist, punktiert dargestellt. Die Verteilungsdichte nach (A.63) erfüllt alle Bedingungen, die bei einer zweidimensionalen Verbundverteilungsdichte erfüllt sein müssen. Die eindimensionalen Randverteilungsdichten der Zufallsgrößen $\mathbf{n}_1$ und $\mathbf{n}_2$ erhält man, indem man die gemeinsame zweidimensionale Verbundverteilungsdichte über n_1 bzw. über n_2 integriert. Diese Integrale liefern dasselbe, wie wenn man über die Verbundverteilungsdichte eines zweidimensionalen, unkorrelierten normierten, zentrierten und normalverteilten Zufallsgrößentupels integrieren würde, da die Bereiche des Integrals, innerhalb derer die Verbundverteilungsdichte (A.63) null ist, gerade durch die Bereiche kompensiert werden, in denen die Verbundverteilungsdichte doppelt so groß ist wie die zentrierte normierte Verbundnormalverteilungsdichte.

Jede der beiden Zufallsgrößen $\mathbf{n}_1$ und $\mathbf{n}_2$ ist daher eindimensional normalverteilt. Da sich die Verbundverteilungsdichte nach (A.63) jedoch nicht als das Produkt der beiden Randverteilungsdichten schreiben lässt, sind die beiden Zufallsgrößen abhängig. Multipliziert man die Verteilungsdichte nach (A.63) mit $n_1 \cdot n_2$ und integriert man über die gesamte n_1 - n_2 -Ebene, so erhält man die Korrelation der Zufallsgrößen $\mathbf{n}_1$ und $\mathbf{n}_2$. Dieses Integral ergibt dasselbe wie das entsprechende Integral über die zentrierte normierte Verbundnormalverteilungsdichte, nämlich null. Die Zufallsgrößen $\mathbf{n}_1$ und $\mathbf{n}_2$ sind also unkorreliert.

Man hat es also in unserem Beispiel mit zwei unkorrelierten normalverteilten Zufallsgrößen zu tun, die voneinander *abhängig* sind. Dies ist *kein* Widerspruch der Aussage, dass zwei Zufallsgrößen unabhängig sind, wenn sie unkorreliert und verbundnormalverteilt sind. Die Zufallsgrößen unseres Beispiels sind nämlich *nicht* verbundnormalverteilt, sondern lediglich jede für sich ist normalverteilt.

Ebenso ist es vorstellbar, dass es mehrdimensionale Verbundverteilungen von Zufallsvektoren gibt, die sich nicht als das Produkt ihrer ein- oder zweidimensionalen Randverteilungen schreiben lassen, obwohl jede der Zufallsgrößen bzw. Zufallsgrößentupel aus dem Zufallsvektor normalverteilt sind. Die Forderung, dass das Zufallsgrößentupel $[\mathbf{N}_f(\mu), \mathbf{N}_f(-\mu)]^T$, bzw. die Zufallsgrößentupel $[\mathbf{N}_f(\mu_1), \mathbf{N}_f(\mu_2)]^T$ und $[\mathbf{N}_f(-\mu_1), \mathbf{N}_f(\mu_2)]^T$ verbundnormalverteilt sein sollen, ist daher wesentlich weniger streng, als die Forderung, dass *alle* Spektralwerte $\mathbf{N}_f(\mu)$ für *alle* $\mu = 0$ (1) $M-1$ eine gemeinsame Verbundnormalverteilung aufweisen sollen.

Es ist manchmal schwer zu entscheiden, ob man von den Zufallsgrößentupeln, die bei der Berechnung der Messwert(ko)varianzen auftreten, annehmen kann, dass sie verbundnormalverteilt sind, weil die Zufallsgrößen des Tupels durch eine Fensterung und eine anschließende DFT aus den Zufallsgrößen des Approximationsfehlerprozesses $\mathbf{n}(k)$ gewonnen werden. Es sei nun gezeigt, dass auch von den Zufallsgrößen des wahren Approximationsfehlerprozesses $\mathbf{n}(k)$ *nicht* angenommen werden muss, dass sie *alle* — für *alle* $k = 0$ (1) $F-1$ — verbundnormalverteilt sind, um sicherzustellen, dass für die Zufallsgrößen $\mathbf{N}_1, \mathbf{N}_2, \mathbf{N}_3$ und $\mathbf{N}_4$ die bei der Rückführung der vierten Momente auf die zweiten Momente benötigte Beziehung ([1]:A.41) erfüllt ist. Dazu setzen wir dort zunächst für die Zufallsgrößen $\mathbf{N}_i$ ($i = 1$ (1) 4) die zufälligen Spektralwerte $\mathbf{N}_f(\mu_i)$ ein, wobei die gewählten diskreten Frequenzen μ_i auch gleich sein können.

Wir beginnen damit, dass wir im Erwartungswert auf der linken Seite der Gleichung ([1]:A.41) die Zufallsgrößen als Funktionen der Zufallsgrößen $\mathbf{n}(k)$ des Approximationsfehlerprozesses schreiben. Jede der vier Zufallsgrößen ist entweder die diskrete Fouriertransformierte des gefensterten Approximationsfehlerprozesses $\mathbf{n}(k)$ bei der Frequenz μ_i oder deren Konjugierte. Wir erhalten somit den Erwartungswert über ein Produkt von vier Summen (Laufindizes k_i mit $i = 1$ (1) 4), die jeweils eine Linearkombination der Zufallsgrößen $\mathbf{n}(k)$ oder ihrer Konjugierten enthält. Das Produkt der vier Summen kann als Vierfachsumme des Produkts der Summanden geschrieben werden. Nachdem man die Reihenfolge der Summation und der Erwartungswertbildung vertauscht hat, kann man bei jedem Summanden einige bezüglich der Erwartungswertbildung konstante Faktoren vor den Erwartungswert ziehen. Diese Faktoren sind neben den Drehfaktoren $e^{\pm j \cdot \frac{2\pi}{M} \cdot \mu_i \cdot k_i}$ der DFT noch die Werte $f(k_i)$ und $f(k_i)^*$ der verwendeten Fensterfolge. Aus Gleichung ([1]:A.41) erhält man auf diese Weise:

$$\begin{aligned}
& \mathbb{E}\{ \mathbf{N}_1 \cdot \mathbf{N}_2^* \cdot \mathbf{N}_3 \cdot \mathbf{N}_4^* \} = \mathbb{E}\{ \mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^* \cdot \mathbf{N}_f(\mu_3) \cdot \mathbf{N}_f(\mu_4)^* \} = \quad (\text{A.64}) \\
& = \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \cdot \mathbf{n}(k_3) \cdot \mathbf{n}(k_4)^* \} \cdot f(k_1) \cdot f(k_2)^* \cdot f(k_3) \cdot f(k_4)^* \cdot \\
& \quad \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_2 \cdot k_2 + \mu_3 \cdot k_3 - \mu_4 \cdot k_4)}.
\end{aligned}$$

Setzt man nun für den Zufallsvektor $\vec{\mathbf{N}}$ in Gleichung ([1]:A.49) des Anhangs [1]:A.6

$$\vec{\mathbf{N}} = \begin{bmatrix} \mathbf{n}(k_1) \\ \mathbf{n}(k_2) \\ \mathbf{n}(k_3) \\ \mathbf{n}(k_4) \end{bmatrix} \quad \text{mit} \quad k_1, k_2, k_3, k_4 = 0 \ (1) \ F-1 \quad (\text{A.65})$$

ein, so erhält man mit $R=4$ und mit dem Indexquadrupel $[i, j, k, l] = [1, 6, 3, 8]$ aus Gleichung ([1]:A.60) eine Gleichung, die es uns ermöglicht das vierte Moment in Gleichung (A.64) durch die entsprechenden zweiten Momente auszudrücken.

$$\begin{aligned}
\mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \cdot \mathbf{n}(k_3) \cdot \mathbf{n}(k_4)^* \} &= \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \} \cdot \mathbb{E}\{ \mathbf{n}(k_3) \cdot \mathbf{n}(k_4)^* \} + \\
&+ \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_3) \} \cdot \mathbb{E}\{ \mathbf{n}(k_2) \cdot \mathbf{n}(k_4) \}^* + \quad (\text{A.66}) \\
&+ \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_4)^* \} \cdot \mathbb{E}\{ \mathbf{n}(k_2) \cdot \mathbf{n}(k_3)^* \}.
\end{aligned}$$

Dabei genügt es auch beim Approximationsfehlerprozess anzunehmen, dass die Randverteilungen jedes beliebigen Zufallsgrößenquadrupels des Prozesses Verbundnormalverteilungen sind. Somit müssen nicht alle Zufallsgrößen $\mathbf{n}(k)$ für $k = 0 \ (1) \ F-1$ gemeinsam verbundnormalverteilt sein. Desweiteren erkennt man, dass es keine Rolle spielt, welche Fensterfolge man bei der Berechnung der Spektralwerte verwendet, weil die Werte der Fensterfolge vor die Erwartungswertbildung gezogen werden konnten. Ersetzt man die vierten Momente durch die Summe dreier Produkte zweiter Momente, so kann man die Vierfachsumme über je drei Summanden auch als Summe dreier Vierfachsummen schreiben. Nun zieht man die Drehfaktoren der DFT und die Werte der Fensterfolge wieder in die Erwartungswertbildung hinein, wobei man auf die richtige Zuordnung der Indizes achtet. Nach einer eventuellen Vertauschung der Reihenfolge der einzelnen Summationen der Vierfachsummen, kann man aus den innersten beiden Doppelsummen jeweils einen der beiden Erwartungswerte herausziehen. Die innere Doppelsumme ist dann nicht mehr von den Indizes der äußeren Doppelsumme abhängig, und kann somit vor die äußere Doppelsumme gezogen werden, so dass man drei Produkte von jeweils zwei Doppelsummen erhält. Nach Vertauschung der Reihenfolge der Summation und der Erwartungswertbildung kann man die Doppelsumme innerhalb des jeweiligen Erwartungswertes wieder als

das Produkt zweier Einfachsummen schreiben. Jede dieser Einfachsummen ist dann eine der Zufallsgrößen $\mathbf{N}_f(\mu_i)$ ($i = 1$ (1) 4). Man erhält so:

$$\begin{aligned}
& \mathbb{E}\{ \mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^* \cdot \mathbf{N}_f(\mu_3) \cdot \mathbf{N}_f(\mu_4)^* \} = \tag{A.67} \\
&= \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \} \cdot \mathbb{E}\{ \mathbf{n}(k_3) \cdot \mathbf{n}(k_4)^* \} \cdot f(k_1) \cdot f(k_2)^* \cdot f(k_3) \cdot f(k_4)^* \cdot \\
&\quad \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_2 \cdot k_2 + \mu_3 \cdot k_3 - \mu_4 \cdot k_4)} + \\
&+ \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_3) \} \cdot \mathbb{E}\{ \mathbf{n}(k_2) \cdot \mathbf{n}(k_4)^* \} \cdot f(k_1) \cdot f(k_2)^* \cdot f(k_3) \cdot f(k_4)^* \cdot \\
&\quad \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_2 \cdot k_2 + \mu_3 \cdot k_3 - \mu_4 \cdot k_4)} + \\
&+ \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbb{E}\{ \mathbf{n}(k_1) \cdot \mathbf{n}(k_4)^* \} \cdot \mathbb{E}\{ \mathbf{n}(k_2) \cdot \mathbf{n}(k_3)^* \} \cdot f(k_1) \cdot f(k_2)^* \cdot f(k_3) \cdot f(k_4)^* \cdot \\
&\quad \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_2 \cdot k_2 + \mu_3 \cdot k_3 - \mu_4 \cdot k_4)} = \\
&= \mathbb{E}\left\{ \sum_{k_1=-\infty}^{\infty} \sum_{k_2=-\infty}^{\infty} \mathbf{n}(k_1) \cdot \mathbf{n}(k_2)^* \cdot f(k_1) \cdot f(k_2)^* \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_2 \cdot k_2)} \right\} \cdot \\
&\quad \cdot \mathbb{E}\left\{ \sum_{k_3=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbf{n}(k_3) \cdot \mathbf{n}(k_4)^* \cdot f(k_3) \cdot f(k_4)^* \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_3 \cdot k_3 - \mu_4 \cdot k_4)} \right\} + \\
&+ \mathbb{E}\left\{ \sum_{k_1=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \mathbf{n}(k_1) \cdot \mathbf{n}(k_3) \cdot f(k_1) \cdot f(k_3) \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 + \mu_3 \cdot k_3)} \right\} \cdot \\
&\quad \cdot \mathbb{E}\left\{ \sum_{k_2=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbf{n}(k_2)^* \cdot \mathbf{n}(k_4)^* \cdot f(k_2)^* \cdot f(k_4)^* \cdot e^{+j \cdot \frac{2\pi}{M} \cdot (\mu_2 \cdot k_2 + \mu_4 \cdot k_4)} \right\} + \\
&+ \mathbb{E}\left\{ \sum_{k_1=-\infty}^{\infty} \sum_{k_4=-\infty}^{\infty} \mathbf{n}(k_1) \cdot \mathbf{n}(k_4)^* \cdot f(k_1) \cdot f(k_4)^* \cdot e^{-j \cdot \frac{2\pi}{M} \cdot (\mu_1 \cdot k_1 - \mu_4 \cdot k_4)} \right\} \cdot \\
&\quad \cdot \mathbb{E}\left\{ \sum_{k_2=-\infty}^{\infty} \sum_{k_3=-\infty}^{\infty} \mathbf{n}(k_2)^* \cdot \mathbf{n}(k_3) \cdot f(k_2)^* \cdot f(k_3) \cdot e^{+j \cdot \frac{2\pi}{M} \cdot (\mu_2 \cdot k_2 - \mu_3 \cdot k_3)} \right\} = \\
&= \mathbb{E}\{ \mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_2)^* \} \cdot \mathbb{E}\{ \mathbf{N}_f(\mu_3) \cdot \mathbf{N}_f(\mu_4)^* \} + \\
&+ \mathbb{E}\{ \mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_3) \} \cdot \mathbb{E}\{ \mathbf{N}_f(\mu_2) \cdot \mathbf{N}_f(\mu_4) \}^* + \\
&+ \mathbb{E}\{ \mathbf{N}_f(\mu_1) \cdot \mathbf{N}_f(\mu_4)^* \} \cdot \mathbb{E}\{ \mathbf{N}_f(\mu_2) \cdot \mathbf{N}_f(\mu_3)^* \}^*
\end{aligned}$$

Bei der Herleitung der Gleichung ([1]:A.41) im Unterkapitel [1]:A.6 des Anhangs wurde gesagt, dass man die anderen vierten Momente, die die einzelnen Zufallsgrößen evtl. konjugiert enthalten, dadurch erhält, indem man in Gleichung ([1]:A.41) zunächst die gewünschten Zufallsgrößen durch ihre Konjugierten ersetzt, und dann die Indizes entsprechend substituiert. Wenn man die Herleitung von Gleichung (A.64) bis (A.67) mit den

entsprechenden konjugierten Zufallsgrößen durchführt, kann man analog zeigen, dass die formal substituierten Versionen von Gleichung ([1]:A.41) ebenfalls gelten.

Liste der verwendeten Abkürzungen und Formelzeichen

Es werden hier nur die Abkürzungen und Formelzeichen aufgeführt, die nicht in [1] beschrieben sind. Dies gilt auch, wenn die entsprechenden Folgen und Funktionen hier oftmals zweidimensional sind, während sie in [1] nur eindimensional definiert sind.

Allgemeine Formelzeichen und Funktionen

realmin = $2^2 - 2^{\text{Exponentenwortlänge}-1}$ beim IEEE Standard 754. Kleinste positive Zahl, die bei Gleitkommazahlendarstellung am Rechner mit voller Genauigkeit darstellbar ist [8].

$\vec{1}$ Zeilenvektor der nur Einsen enthält.

$\underline{1}_{\perp}$ Idempotente Matrix, die alle Vektoren auf den Nullraum des Einservektors $\vec{1}$ abbildet.

$\vec{E}_n$ Einheitszeilenvektor, dessen n -tes Element eins ist, während alle anderen Elemente null sind.

$\text{kgv}(\dots)$ kleinstes gemeinsames Vielfaches.

$\vec{w}_{K_H, \kappa}$ Vektor der Drehfaktoren der inversen DFT der Länge K_H zum Zeitpunkt κ .

Spezielle Formelzeichen

Die in [1] beschriebenen Konventionen (z. B. Fettdruck $\Rightarrow$ Zufallsgröße, -vektor, -matrix oder -prozess etc.) bezüglich der Variablennamen gelten auch hier, so dass auch in dieser Liste im wesentlichen nur die Formelzeichen angegeben sind, die bei der Beschreibung der Theorie der erweiterten RKM-Varianten vorkommen, und die nicht aus den hier oder in [1] genannten Formelzeichen gemäß dieser Konventionen abgeleitet werden können.

A_0 Mit diesem Parameter legt man fest, dass die letzten A_0 Werte der Fensterfolge innerhalb des Intervalls $k \in [0; F-1]$ null sein sollen.

A_1 Mit diesem Parameter legt man fest, dass das Polynom $z^{F-A_0-1} \cdot G(z)$ zusätzlich zu den in Kapitel [1]:6.1 genannten Nullstellen weitere $2 \cdot A_1$ einfache Nullstellen auf dem Einheitskreis besitzt, die sich im Abstand $2\pi/F$ anschließen.

Δc	Änderung des Parameters der Bilineartransformation bei der Konstruktion der diskreten Fensterfolge ($-1 < c < 1$)
c_∞	Parameter der Bilineartransformation bei der Konstruktion der kontinuierlichen Fensterfunktion ($0 < c_\infty$)
$\hat{\mathbf{C}}_n(\mu)$	Hilfsvektor bei der Berechnung der Kovarianz des n -ten Elementes des Messwertvektors $\hat{\mathbf{H}}(\mu)$ der Übertragungsfunktionen.
$\hat{\mathbf{C}}_U(\mu)$	Hilfsvektor bei der Berechnung der Kovarianz des Spektrums der gefensterten deterministischen Störung.
$\hat{\mathbf{C}}_u(\mu)$	Hilfsvektor bei der Berechnung der Kovarianz der gefensterten deterministischen Störung.
$d_\infty(t)$	$= f_\infty(t) * f_\infty(-t) =$ kontinuierliche Fensterautokorrelationsfunktion.
$D_\infty(s)$	Laplacetransformierte der Fensterautokorrelationsfunktion $d_\infty(t)$.
$D_{E,\infty}(s)$	Anteil der Laplacetransformierten $D_\infty(s)$ der Fensterautokorrelationsfunktion, der nur die Nullstellen auf der imaginären Achse enthält.
$D_{\bar{E},\infty}(s)$	Anteil der Laplacetransformierten $D_\infty(s)$ der Fensterautokorrelationsfunktion ohne die Nullstellen auf der imaginären Achse.
$\tilde{D}_{\bar{E},\infty}(\tilde{z})$	Polynom in $\tilde{z}$, das durch Bilineartransformation aus $D_{\bar{E},\infty}(s)$ entsteht.
$\tilde{D}_{N,\infty}(\tilde{z})$	Anteil von $\tilde{D}_{\bar{E},\infty}(\tilde{z})$, der die Nullstellen in der bilineartransformierten z -Ebene enthält, und dessen minimalphasiger Anteil einen nichtlinearen Phasenbeitrag liefert, der mit Hilfe des Cepstrums berechnet wird.
$\tilde{D}_{P,\infty}(\tilde{z})$	Anteil von $\tilde{D}_{\bar{E},\infty}(\tilde{z})$, der die Polstellen in der bilineartransformierten z -Ebene enthält, und dessen minimalphasiger Anteil nach der bilinearen Rücktransformation einen nichtlinearen Phasenbeitrag liefert, der geschlossen berechnet werden kann.
$D_{P,\infty}(s)$	Bilinear Rücktransformierte von $\tilde{D}_{P,\infty}(\tilde{z})$.
$f_\infty(t)$	Kontinuierliche Fensterfunktion.
$F_\infty(\omega)$	Fouriertransformierte der kontinuierlichen Fensterfunktion $f_\infty(t)$.
$g_\infty(t)$	Basisfensterfunktion, die zur Herleitung der Konstruktion der kontinuierlichen Fensterfunktion $f_\infty(t)$ benötigt wird.
$g_{Q,\infty}(t)$	$= g(t) * g(-t) =$ kontinuierliche Basisfensterautokorrelationsfunktion.
$G_\infty(j\omega)$	Fouriertransformierte der Basisfensterfunktion $g_\infty(t)$.

$G_1(z)$	Anteil von $G(z)$, der die $F - N + 2 \cdot A_1$ einfachen Nullstellen am Einheitskreis im Raster $2\pi/F$ enthält.
$G_2(z)$	Anteil von $G(z)$, der die $N - 1 - A_0 - 2 \cdot A_1$ frei wählbaren Nullstellen enthält.
$h_\kappa(k)$	Zeitvariante Impulsantwort des linearen, komplexwertigen Modellsystems $\mathcal{S}_{lin}$, das vom Eingangssignal $\mathbf{v}(k)$ erregt wird.
$\tilde{h}_\kappa(k)$	Zeitvariante Impulsantwort, die durch periodische Fortsetzung mit der Periode M aus $\tilde{h}_\kappa(k)$ entsteht.
$H(\mu_1, \mu_2)$	Bifrequente Übertragungsfunktion des linearen, komplexwertigen Modellsystems $\mathcal{S}_{lin}$, das vom Eingangssignal $\mathbf{v}(k)$ erregt wird.
$\hat{H}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$	Messwerte der bifrequenten Übertragungsfunktion $H(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$
$h_{*,\kappa}(k)$	Zeitvariante Impulsantwort des linearen, komplexwertigen Modellsystems $\mathcal{S}_{*,lin}$, das vom konjugierten Eingangssignal $\mathbf{v}(k)^*$ erregt wird.
$\tilde{h}_{*,\kappa}(k)$	Zeitvariante Impulsantwort, die durch periodische Fortsetzung mit der Periode M aus $\tilde{h}_{*,\kappa}(k)$ entsteht.
$H_*(\mu_1, \mu_2)$	Bifrequente Übertragungsfunktion des linearen, komplexwertigen Modellsystems $\mathcal{S}_{*,lin}$, das vom konjugierten Eingangssignal $\mathbf{v}(k)^*$ erregt wird.
$\hat{H}_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$	Messwerte der bifrequenten Übertragungsfunktion $H_*(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_H})$
$\vec{H}(\mu)$	Zeilenvektor, der die Werte der sich bei der Systemapproximation theoretisch ergebenden Übertragungsfunktionen $H(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$ und $H_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$ der beiden Modellsysteme zusammenfasst.
$\hat{\vec{H}}(\mu)$	Zeilenvektor, der die Messwerte $\hat{H}(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$ und $\hat{H}_*(\mu, \mu + \hat{\mu} \cdot \frac{M}{K_H})$ der bifrequenten Übertragungsfunktionen der beiden Modellsysteme zusammenfasst.
$\tilde{L}$	Anzahl der Zeitintervalle innerhalb des Zeitraums der Messung, aus denen die L Signalausschnitte für eine RKM-Messung ggf. zufällig entnommen werden können.
K_H	Periode der Zeitvarianz des Modellsystems.
K_Φ	Periode der Zyklostationarität des Approximationsfehlers.
K_S	Kleinstes gemeinsames Vielfaches von K_H und K_Φ .
$K(\mu)$	Anzahl der Zufallsgrößen, die in dem Zufallsvektor $\vec{\mathbf{V}}(\mu)$ zusammengefasst sind.
K_{∞, ν_2}	konstanter Faktor der ν_2 -ten Nullstelle des bilineartransformierten Polynoms mit den Nullstellen am Einheitskreis.

P, Q	Permutationsmatrizen
$\mathcal{S}_{*,lin}$	Eines der beiden linearen Modellsysteme. Dieses Modellsystem wird von dem konjugierten Eingangssignal $\mathbf{v}(k)^*$ erregt, und ergibt sich bei der theoretischen Aufspaltung des gestörten, nichtlinearen realen Systems $\mathcal{S}$ mit Hilfe der wahren, i. Allg. unbekannten Erwartungswerte der anliegenden Prozesse.
S_ν	Sinusreihenoeffizienten der Fensterautokorrelationsfolge bzw. -funktion.
$u(k)$	Deterministisches Störsignal im Modell des realen Systems. Dieses zeitabhängige Störsignal ergibt sich bei der theoretischen Aufspaltung des gestörten, nichtlinearen realen Systems $\mathcal{S}$ mit Hilfe der wahren, i. allg. unbekannten Erwartungswerte der anliegenden Prozesse.
$\hat{u}(k)$	Messwerte des deterministischen Störsignals.
$U_f(\mu)$	Spektrum des gefensterten deterministischen Störsignals.
$\hat{U}_f(\mu)$	Messwerte des Spektrums des gefensterten deterministischen Störsignals.
$\tilde{\mathbf{V}}(\mu)$	Spaltenvektor, der die Zufallswerte $\mathbf{V}(\mu + \hat{\mu} \cdot \frac{M}{K_H})$ und $\mathbf{V}(-\mu - \hat{\mu} \cdot \frac{M}{K_H})^*$ des Spektrums der Erregung zusammenfasst.
$\tilde{\underline{V}}(\mu)$	Stichprobenmatrix der Erregung. Diese $2 \cdot K_H \times L$ Matrix enthält eine konkrete Stichprobe des Zufallsvektors $\tilde{\mathbf{V}}(\mu)$ vom Umfang L .
$\check{\mathbf{V}}(\mu)$	Spaltenvektor, der einen Satz linear unabhängiger Zufallswerte $\mathbf{V}(\mu + \hat{\mu} \cdot \frac{M}{K_S})$ und $\mathbf{V}(-\mu - \hat{\mu} \cdot \frac{M}{K_S})^*$ des Spektrums der Erregung zusammenfasst.
$\check{\underline{V}}(\mu)$	Stichprobenmatrix der Erregung. Diese $K(\mu) \times L$ Matrix enthält eine konkrete Stichprobe des Zufallsvektors $\check{\mathbf{V}}(\mu)$ vom Umfang L .
$\mathbf{x}_*(k)$	Zufallsprozess am Ausgang des Modellsystems $\mathcal{S}_{*,lin}$, der bei Erregung des Modellsystems mit dem konjugierten Zufallsprozess $\mathbf{v}(k)^*$ entsteht.
$z_{0,\rho}$	ρ -te, frei wählbare Nullstelle des Polynoms $z^{(N-1-A_0-2 \cdot A_1)} \cdot G_2(z)$.
$\tilde{z}_{\infty,\nu_2}$	ν_2 -te Nullstelle des Polynoms $\tilde{D}_{N,\infty}(\tilde{z})$ nach der Bilineartransformation am Einheitskreis.
$\tilde{z}_{N,\rho,\nu_1}$	Nenner der ρ -ten Nullstelle des bilineartransformierten Polynoms $\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)^{N-1-A_0-2 \cdot A_1} \cdot G_2\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)$.
$\tilde{z}_{Z,\rho,\nu_1}$	Zähler der ρ -ten Nullstelle des bilineartransformierten Polynoms $\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)^{N-1-A_0-2 \cdot A_1} \cdot G_2\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)$.
ρ	Index der frei wählbaren Nullstelle $z_{0,\rho}$ des Polynoms $z^{(N-1-A_0-2 \cdot A_1)} \cdot G_2(z)$.

- $\bar{\sigma}$ Schätzwert für den Mindestwert der Streuung des Rauschsockels bei der Berechnung des Cepstrums.
- $\hat{\sigma}$ Schätzwert für den maximalen Rauschwert des Rauschsockels bei der Berechnung des Cepstrums.
- $\Phi_{\mathbf{n}}(\Omega_1, \Omega_2)$ Bifrequentes Leistungsdichtespektrum des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.
- $\bar{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Stufenapproximation des Leistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.
- $\tilde{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Näherungswerte der Stufenapproximation des Leistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall, die man mit einer endlich langen Fensterfolge gewinnt.
- $\hat{\Phi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Messwerte des Leistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.
- $\dot{\Phi}_{\mathbf{n}}(\Omega, k)$ Zeitabhängiges Leistungsdichtespektrum des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall, das man aus der Autokorrelationsfolge gewinnt, indem man diese bezüglich der Differenz der Zeitpunkte der an der Autokorrelation beteiligten Zufallsgrößen diskret fouriertransformiert.
- $\dot{\Phi}_{\mathbf{n}}(\Omega, \bar{\mu})$ Bezüglich k invers diskret Fouriertransformierte von $\dot{\Phi}_{\mathbf{n}}(\Omega, k)$.
- $\psi_{0,\rho}$ Winkel der ρ -ten, frei wählbaren Nullstelle des Polynoms $z^{N-1-A_0-2 \cdot A_1} \cdot G_2(z)$.
- $\tilde{\psi}_{0,\rho,\nu_1}$ Winkel der Nullstelle $\tilde{z}_{Z,\rho,\nu_1} / \tilde{z}_{N,\rho,\nu_1}$ des bilineartransformierten Polynoms $\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)^{N-1-A_0-2 \cdot A_1} \cdot G_2\left(z \cdot e^{-j \cdot \frac{2\pi}{F} \cdot \nu_1}\right)$.
- $\overset{?}{\psi}_0$ geschätzter Winkel der betragsmäßig größten, unbekannten Nullstelle $\overset{?}{z}_0$ des bilineartransformierten Polynoms.
- $\Psi_{\mathbf{n}}(\Omega_1, \Omega_2)$ Kreuzleistungsdichtespektrum des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.
- $\bar{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Stufenapproximation des Kreuzleistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.
- $\tilde{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Näherungswerte der Stufenapproximation des Kreuzleistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall, die man mit einer endlich langen Fensterfolge gewinnt.
- $\hat{\Psi}_{\mathbf{n}}(\mu, \mu + \tilde{\mu} \cdot \frac{M}{K_{\Phi}})$ Messwerte des Kreuzleistungsdichtespektrums des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall.

$\dot{\Psi}_{\mathbf{n}}(\Omega, k)$ Zeitabhängiges Kreuzleistungsdichtespektrum des Prozesses $\mathbf{n}(k)$ im zyklstationären Fall, das man aus der Kreuzkorrelationsfolge gewinnt, indem man diese bezüglich der Differenz der Zeitpunkte der an der Kreuzkorrelation beteiligten Zufallsgrößen diskret fouriertransformiert.

$\dot{\Psi}_{\mathbf{n}}(\Omega, \bar{\mu})$ Bezüglich k invers diskret Fouriertransformierte von $\dot{\Psi}_{\mathbf{n}}(\Omega, k)$.

Speicherplatznamen

Die Namen der Variablen, für die bei einer Implementierung am Rechner Speicherbereiche bereitgestellt werden müssen, sind in **Schreibmaschinenschrift** dargestellt. Die Namen der Variablen, die bei den Programmen zur der Berechnung der Fenster, ihrer Spektren und ihrer Autokorrelationsfolgen und -funktionen benötigt werden, sind jeweils im Kommentar der Programme erläutert. Bei der Beschreibung der RKM-Implementierung treten folgende Variablen auf:

$A_1_H(\mu)$ Speicher für die längeren Halbachsen der Konfidenzellipsen der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$.

$A_1_HS(\mu)$ Speicher für die längeren Halbachsen der Konfidenzellipsen der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$.

$A_2_H(\mu)$ Speicher für die kürzeren Halbachsen der Konfidenzellipsen der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$.

$A_2_HS(\mu)$ Speicher für die kürzeren Halbachsen der Konfidenzellipsen der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$.

$A_Phi(\mu)$ Speicher für die halben Konfidenzintervallbreiten der LDS-Messwerte.

$A_1_Psi(\mu)$ Speicher für die längeren Halbachsen der Konfidenzellipsen der Messwerte des KLDS.

$A_2_Psi(\mu)$ Speicher für die kürzeren Halbachsen der Konfidenzellipsen der Messwerte des KLDS.

A_1_u Speicher für die längere Halbachse der Konfidenzellipsen der Messwerte der deterministischen Störung.

A_2_u Speicher für die kürzere Halbachse der Konfidenzellipsen der Messwerte der deterministischen Störung.

$A_1_U(\mu)$ Speicher für die längeren Halbachsen der Konfidenzellipsen der Messwerte des Spektrums der gefensterten, deterministischen Störung.

$A_2_U(\mu)$ Speicher für die kürzeren Halbachsen der Konfidenzellipsen der Messwerte des Spektrums der gefensterten, deterministischen Störung.

- Akku_y(k) Akkumulatorfeld für das gemessene Ausgangssignal $y_\lambda(k)$ des Systems.
- Akku_V(μ) Akkumulatorfeld für die Spektralwerte $V_\lambda(\mu)$ des Eingangssignals des Systems.
- Akku_VQ(μ) Akkumulatorfeld für die Betragsquadrate $|V_\lambda(\mu)|^2$ der Spektralwerte des Eingangssignals des Systems.
- Akku_VV(μ) Akkumulatorfeld für die Produkte $V_\lambda(-\mu) \cdot V_\lambda(\mu)$ der Spektralwerte des Eingangssignals des Systems.
- Akku_VY(μ) Akkumulatorfeld für die Produkte $V_\lambda(-\mu) \cdot Y_{f,\lambda}(\mu)$ der Spektralwerte des Eingangssignals und des gefensterten Ausgangssignals des Systems.
- Akku_YQ(μ) Akkumulatorfeld für die Betragsquadrate $|Y_{f,\lambda}(\mu)|^2$ der Spektralwerte des gefensterten Ausgangssignals des Systems.
- Akku_YV(μ) Akkumulatorfeld für die Produkte $Y_{f,\lambda}(\mu) \cdot V_\lambda(\mu)^*$ der Spektralwerte des Eingangssignals und des gefensterten Ausgangssignals des Systems.
- Akku_YY(μ) Akkumulatorfeld für die Produkte $Y_{f,\lambda}(-\mu) \cdot Y_{f,\lambda}(\mu)$ der Spektralwerte des gefensterten Ausgangssignals des Systems.
- C_HH(μ) Speicher für die Kovarianzen $\hat{C}_{\hat{H}(\mu), \hat{H}(\mu)^*}$ der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$.
- C_HQ(μ) Speicher für die Varianzen $\hat{C}_{\hat{H}(\mu), \hat{H}(\mu)}$ der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$.
- C_HSQ(μ) Speicher für die Varianzen $\hat{C}_{\hat{H}_*(\mu), \hat{H}_*(\mu)}$ der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$.
- C_HSHS(μ) Speicher für die Kovarianzen $\hat{C}_{\hat{H}_*(\mu), \hat{H}_*(\mu)^*}$ der Messwerte der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$.
- C_PhiQ(μ) Speicher für die Varianzen $\hat{C}_{\hat{\Phi}_n(\mu), \hat{\Phi}_n(\mu)}$ der LDS-Messwerte.
- C_PsiQ(μ) Speicher für die Varianzen $\hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)}$ der KLDS-Messwerte.
- C_PsiPsi(μ) Speicher für die Kovarianzen $\hat{C}_{\hat{\Psi}_n(\mu), \hat{\Psi}_n(\mu)^*}$ der KLDS-Messwerte.
- C_uu Speicher für die Kovarianz $\hat{C}_{\hat{u}(k), \hat{u}(k)^*}$ der Messwerte der deterministischen Störung.
- C_uQ Speicher für die Varianz $\hat{C}_{\hat{u}(k), \hat{u}(k)}$ der Messwerte der deterministischen Störung.
- C_UQ(μ) Speicher für die Varianzen $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)}$ der Messwerte des Spektrums der gefensterten, deterministischen Störung.

$C_{UU}(\mu)$	Speicher für die Kovarianzen $\hat{C}_{\hat{U}_f(\mu), \hat{U}_f(\mu)^*}$ der Messwerte des Spektrums der gefensterten, deterministischen Störung.
$C_{VQ}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Varianzen $\hat{C}_{\mathbf{V}(\mu), \mathbf{V}(\mu)}$.
$C_{VV}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Kovarianzen $\hat{C}_{\mathbf{V}(-\mu), \mathbf{V}(\mu)^*}$.
$C_{VY}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Kovarianzen $\hat{C}_{\mathbf{V}(-\mu), \mathbf{Y}_f(\mu)^*}$.
$C_{YQ}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Varianzen $\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{Y}_f(\mu)}$.
$C_{YV}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Kovarianzen $\hat{C}_{\mathbf{Y}_f(\mu), \mathbf{V}(\mu)}$.
$C_{YY}(\mu)$	Speicher für das $L \cdot (L-1)$ -fache der empirischen Kovarianzen $\hat{C}_{\mathbf{Y}_f(-\mu), \mathbf{Y}_f(\mu)^*}$.
$f(k)$	Speicher der F Werte der Fensterfolge.
$H(\mu)$	Speicher der Messwerte $\hat{H}(\mu)$ der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{lin}$.
$HS(\mu)$	Speicher der Messwerte $\hat{H}_*(\mu)$ der Übertragungsfunktion des linearen Teilsystems $\mathcal{S}_{*,lin}$.
$K_{VV}(\mu)$	Speicher für die inversen empirischen Kovarianzmatrizen $\hat{C}_{\check{\mathbf{V}}(\mu), \check{\mathbf{V}}(\mu)}^{-1}$.
$\Phi(\mu)$	Speicher der Messwerte $\hat{\Phi}_{\mathbf{n}}(\mu)$ des LDS.
$\Psi(\mu)$	Speicher der Messwerte $\hat{\Psi}_{\mathbf{n}}(\mu)$ des KLDS.
$u(k)$	Speicher der Messwerte der deterministischen Störung $\hat{u}(k)$.
$U(\mu)$	Speicher der Messwerte des Spektrums $\hat{U}_f(\mu)$ der gefensterten, deterministischen Störung.
$v(k)$	Speicher der Testsignalfolge $v_\lambda(k)$.
$V(\mu)$	Speicher der Spektralwerte $V_\lambda(\mu)$ der Testsignalfolge.
$X_mittel(\mu)$	Speicher für das L -fache der empirischen Mittelwerte der Summe der Spektren der Signale an den Ausgängen der beiden linearen Modellsysteme.
$x_mittel(k)$	Speicher für das L -fache der empirischen Mittelwerte der Summe der Signale an den Ausgängen der beiden linearen Modellsysteme.
$y(k)$	Speicher des gemessenen Ausgangssignals $y_\lambda(k)$.
$Y(\mu)$	Speicher der Spektralwerte $Y_{f,\lambda}(\mu)$ des gemessenen, gefensterten Ausgangssignals.
Y_mittel	Speicher für das L -fache der empirischen Mittelwerte des Spektrums des gefensterten Ausgangssignals.