

DIPLOMARBEIT

Titel der Diplomarbeit

WIENER- UND KALMAN-FILTER-THEORIE UND DEREN VERWENDUNG IN DER HYDROMETEOROLOGISCHEN DATENVERARBEITUNG

angestrebter akademischer Grad

Magister der Naturwissenschaften (Mag. rer. nat.)

Verfasser:	Emanuel Hecher
Matrikel-Nummer:	0202871
Studienrichtung:	A 405 Diplomstudium Mathematik
Betreuer:	Hans G. Feichtinger

Wien, am 19.8.2009

Inhaltsverzeichnis

Vorwort	5
1 Grundlagen	7
2 Die Wiener-Filter-Theorie	13
2.1 Überblick	13
2.2 Die Funktionsweise des Wiener-Filters	13
2.2.1 Die Berechnung des Schätzfehlers	14
2.2.2 Die Berechnung der Filterkoeffizienten, die Wiener-Hopf-Gleichung und die Wiener-Lösung	15
2.3 Das Orthogonalitätsprinzip	16
2.4 Vektorschreibweise des minimalen Fehlers	17
2.5 Rauschentfernung, Signalprädiktion und Signalglättung	18
2.5.1 Die Signalglättung	18
2.5.2 Die Filterung	19
2.5.3 Die Signalprädiktion	21
2.6 Die Wiener-Filter-Theorie für kontinuierliche Zeit	22
2.6.1 Die nicht-kausale Lösung	23
2.6.2 Die kausale Lösung	24
3 Die Kalman-Filter-Theorie	27
3.1 Überblick	27
3.2 Das "Standard"-Kalman-Filter	28
3.2.1 Problemstellung und die Innovationsfolge	28
3.2.2 Der Minimum-Varianz-Schätzer	32
3.2.3 Die Kalman-Filter-Gleichungen	33
3.3 Das Erweiterte-Kalman-Filter	38
3.3.1 Die Funktionsweise des Erweiterten-Kalman- Filters	38
3.4 Das Ensemble-Kalman-Filter	40

3.4.1	Die Ensemble-Kovarianzmatrizen	40
3.4.2	Die Funktionsweise des Ensemble-Kalman-Filters . . .	41
4	Anwendungen in der hydrometeorologischen Datenverarbeitung	45
4.1	Überblick	45
4.2	Vorhersage von Wasserabflussmengen von Flüssen mit Hilfe der Wiener-Filter-Theorie	46
4.3	Ein lineares Wasserabflussmodell	47
4.3.1	Die Verwendung des "Standard"-Kalman-Filters	47
4.3.2	Ein Vergleich zwischen der Verwendung des "Standard"-Kalman-Filters und des Ensemble- Kalman-Filters	48
4.4	Schätzung des Bodenwassergehalts	50
4.4.1	Modellbeschreibung	50
4.4.2	Versuchsergebnisse	52
	Literaturverzeichnis	55
A	Zusammenfassung	59
B	Abstract	61
C	Lebenslauf	63

Vorwort

Ziel dieser Arbeit ist es, die aus der statistischen Signalverarbeitung stammende Wiener- bzw. Kalman-Filter-Theorie zu beschreiben und deren Verwendung in der hydrometeorologischen Datenverarbeitung zu zeigen.

Bei dem nach seinen Erfinder Norbert Wiener benannten Wiener-Filter, handelt es sich um ein sog. Schätzfilter zur Störungsunterdrückung von Signalen. Im Gegensatz zum klassischen Filterentwurf wird bei der Wiener-Filter-Theorie nicht davon ausgegangen, dass sich, aufgrund unterschiedlicher Frequenzbereiche, Nutz- und Störsignal frequenzmäßig trennen lassen. Bei dem beschriebenen Filter dürfen die Signale im selben Frequenzbereich liegen, da das verwendete Unterscheidungskriterium auf den unterschiedlichen statistischen Eigenschaften der Signale basiert. Die von Wiener entwickelte Filtertechnik hat die bestmögliche Annäherung des Ausgangssignals $\hat{x}(t)$ an das ursprüngliche Meßsignal $x(t)$ zum Ziel [17]. Die Technik wurde erstmals 1960 im Artikel "Extrapolation, Interpolation and Smoothing of Stationary Time Series" [25] vorgestellt.

Das Kalman-Filter ist eine wichtige Verallgemeinerung des Wiener-Filters. Mithilfe der von Rudolf Kalman in "A New Approach to Linear Filtering and Prediction Problems" [11] vorgestellten Theorie, ist es möglich aus fehlerbehafteten Beobachtungen Systemzustände zu schließen. Das 1960 entwickelte Filter ist der optimale Schätzer für eine große Gruppe von Problemstellungen und ein sehr effektiver und nützlicher Schätzer für eine noch viel größere Gruppe.

Die Kalman-Filter-Theorie bewirkt eine unvorstellbare Genauigkeit vieler Technologien. Die ersten Anwendungsgebiete waren in der Luftfahrt und im Militärbereich. Weiters wurde das Filter auch während des Apollo-Programms verwendet. Im Laufe der Zeit fanden sich immer mehr Einsatzmöglichkeiten in fast allen Ingenieurszweigen. Weitere Anwendungen sind beispielsweise GPS, Zielerfassung von Radar, Beobachtungen der Atmosphäre oder auch automatisierte Medikamentenabgabe.

Im ersten Kapitel werden die für das Verständnis der Arbeit wichtigen grund-

legenden Definitionen und Resultate gezeigt.

Kapitel zwei beschäftigt sich mit der Wiener-Filter-Theorie. Es werden u.a. die Funktionsweise des Wiener-Filters, Rauschentfernung, Signalprädiktion und Signalglättung beschrieben.

Im dritten Kapitel wird das Kalman-Filter als Verallgemeinerung des Wiener-Filters behandelt. Die in dieser Arbeit beschriebenen Varianten des Filters sind das Kalman-Filter für lineare Systeme, das Erweiterte-Kalman-Filter und das Ensemble-Kalman-Filter. Mit den beiden beschriebenen Erweiterungen der Theorie für lineare Systeme, ist es möglich auch Zustände nichtlinearer Systeme zu schätzen. Das Ensemble-Filter wurde 1994 in "Sequential Data Assimilation with a Nonlinear Quasi-Geostrophic Model using Monte Carlo Methods to Forecast Error Statistics" [7] von Geir Evensen vorgestellt.

Abschließend werden in Kapitel vier einige hydrometeorologische Anwendungen gezeigt. Es wird sich mit der Vorhersage von Wasserabflussmengen mittels der Wiener-Filter-Theorie beschäftigt. Weiters werden die Varianten des Kalman-Filters in einem linearen Wasserabflussmodell und zur Schätzung des Bodenwassergehalts verwendet.

Kapitel 1

Grundlagen

Für das Verständnis dieser Arbeit werden einige Definitionen und Resultate aus der Theorie der stochastischen Prozesse benötigt. Die folgenden Definitionen wurden, falls nicht anders zitiert, aus [19] entnommen.

Definitionen:

- Ein *stochastischer Prozess mit diskreter Zeit* ist eine Folge von Zufallsvariablen $x(0), x(1), x(2), \dots$, deren Werte alle in derselben Menge, dem Zustandsraum, liegen. (vgl. [10])

Für die nachfolgenden Punkte sei $t \in \mathbb{R}$:

- Ein *stochastischer Prozess mit kontinuierlicher Zeit* ist eine Familie $x(t)$ realer oder komplexer Zufallsvariablen, deren Werte alle in demselben Wahrscheinlichkeitsraum liegen.
- Mit μ_x wird der *Erwartungswert* des Prozesses $x(t)$ (vorausgesetzt dieser existiert) bezeichnet:

$$E[x(t)] = \mu_x$$

- $r_{xx}(t_1, t_2)$ ist die *Autokorrelation* des Prozesses $x(t)$. Sie ist der Erwartungswert der Produkts $x(t_1)x^*(t_2)$:

$$E[x(t_1)x^*(t_2)] = r_{xx}(t_1, t_2)$$

Bemerkung: Die Existenz der Autokorrelation ist mittels diverser Autokorrelationstests statistisch prüfbar.

- Die *Kreuzkorrelation* $r_{xy}(t_1, t_2)$ zweier Prozesse $x(t)$ und $y(t)$ ist der Erwartungswert des Produkts $x(t_1)y^*(t_2)$:

$$E[x(t_1)y^*(t_2)] = r_{xy}(t_1, t_2)$$

- Die *Kreuzkovarianz* $c_{xy}(t_1, t_2)$ ist die Kovarianz der Zufallsvariablen $x(t_1)$ und $y(t_2)$:

$$E[(x(t_1) - \mu_x(t_1))(y^*(t_2) - \mu_y^*(t_2))] = c_{xy}(t_1, t_2)$$

Es folgt:

$$c_{xy}(t_1, t_2) = r_{xy}(t_1, t_2) - \mu_x(t_1)\mu_y^*(t_2)$$

- Die *Autokovarianz* $c_{xx}(t_1, t_2)$ des Prozesses $x(t)$ ist:

$$c_{xx}(t_1, t_2) = r_{xx}(t_1, t_2) - \mu_x(t_1)\mu_x^*(t_2)$$

Bemerkung: Zu jedem Prozess $x(t)$ kann ein durch Subtraktion des Erwartungswertes ein sog. *erwartungswertfreier Prozess* $x_c(t)$ konstruiert werden:

$$x_c(t) = x(t) - \mu(t)$$

Es gilt:

$$c_{x_c x_c}(t_1, t_2) = r_{x_c x_c}(t_1, t_2) = c_{xx}(t_1, t_2)$$

- Ein stochastischer Prozess heißt *stark stationär*, wenn alle seine statistischen Eigenschaften nicht von der Wahl des Zeitursprungs abhängen.

Folgende Bemerkung wurde aus [24] entnommen:

Bemerkung: Aufgrund der in der Praxis schwierigen Überprüfung des Begriffs der starken Stationarität ist meistens der nun definierte mathematisch etwas schwächere Begriff der *schwachen Stationarität* relevant. Dieser lässt sich anhand der Erwartungswerte zeigen.

- Ein stochastischer Prozess heißt *schwach stationär (wide-sense stationary)*, wenn der Erwartungswert konstant ist d.h. er zu jedem Zeitpunkt dieselben Erwartungswerte hat und seine Autokorrelation nur von der Differenz $\tau = t_1 - t_2$ abhängt:

$$\begin{aligned} E[x(t)] &= \mu_x = \text{konstant} \\ r_{xx}(t_1, t_2) &= r_{xx}(\tau) = E[x(t + \tau)x^*(t)] \end{aligned}$$

Es gilt:

- $r_{xx}(\tau) = E[x(t + \frac{\tau}{2})x^*(t - \frac{\tau}{2})]$
- Jeder stark stationäre Prozess ist auch schwach stationär.

Bemerkung: Die Umkehrung des zweiten Punktes gilt im Allgemeinen nicht [19].

- Sei nun $x(t)$ ein schwach stationärer Zufallsprozess. Die zu dem Zufallsvektor

$$\mathbf{x}(k) = \begin{bmatrix} x(k) & x(k-1) & \cdots & x(k-M) \end{bmatrix}^T$$

gehörige *Autokorrelationsmatrix* ist von der Form

$$\begin{aligned} \mathbf{R}_{xx} &= \\ &= \begin{bmatrix} E[x^2(k)] & E[x(k)x(k-1)] & \cdots \\ E[x(k-1)x(k)] & E[x^2(k-1)] & \cdots \\ E[x(k-2)x(k)] & E[x(k-2)x(k-1)] & \cdots \\ \vdots & \vdots & \ddots \\ E[x(k-M)x(k)] & E[x(k-M)x(k-1)] & \cdots \\ \cdots & E[x(k)x(k-M)] \\ \cdots & E[x(k-1)x(k-M)] \\ \cdots & E[x(k-2)x(k-M)] \\ \ddots & \vdots \\ \cdots & E[x^2(k-M)] \end{bmatrix} \\ &= E[\mathbf{x}(k)\mathbf{x}^T(k)] \end{aligned}$$

Aufgrund der Annahme schwacher Stationarität, können die Elemente der Matrix jeweils als Autokorrelationsfunktionen mit unterschiedlichen Verschiebungen τ identifiziert werden:

$$r_{xx}(\tau) = E[x(k)x(k-\tau)]$$

Damit gilt:

$$\begin{aligned} \mathbf{R}_{xx} &= E[\mathbf{x}(k)\mathbf{x}^T(k)] \\ &= \begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \cdots & r_{xx}(M) \\ r_{xx}(-1) & r_{xx}(0) & \cdots & r_{xx}(M-1) \\ \vdots & \vdots & \ddots & \vdots \\ r_{xx}(-M) & r_{xx}(-M+1) & \cdots & r_{xx}(0) \end{bmatrix} \end{aligned}$$

Bemerkung: Für komplexe Signale wird für die Definition der Autokorrelationsmatrix die *Hermitische Transposition*, $\mathbf{x}^H = (\mathbf{x}^T)^\star$, verwendet. Es gilt also

$$\mathbf{R}_{xx} = E[\mathbf{x}(k)\mathbf{x}^H(k)].$$

Einige Eigenschaften der Autokorrelationsmatrix [18]:

- Bei einem komplexen Signal gilt $\mathbf{R}_{xx}^{\star T} = \mathbf{R}_{xx}^H = \mathbf{R}_{xx}$. Die Matrix ist also *hermetisch*.
Bei einem reellen Signal ist $\mathbf{R}_{xx}$ *symmetrisch*. Es gilt also $\mathbf{R}_{xx}^T = \mathbf{R}_{xx}$, d.h. $r_{xx}(-\tau) = r_{xx}(\tau)$.
- Für ein schwach stationäres Eingangssignal hat die Autokorrelationsmatrix *Toeplitz-Struktur*, d.h. die Elemente der Diagonalen sind jeweils identisch.
- $\mathbf{R}_{xx}$ ist nie negativ definit und fast immer positiv definit.

Beweis: siehe [18] Seite 45

- Bei *weißem Rauschen* (*white noise*) handelt es sich um ein zufälliges Signal (oder einen Zufallsprozess) mit konstanter Amplitude im Leistungsdichtespektrum, d.h. die Rauschenergie verteilt sich konstant über alle Frequenzen. Aufgrund seiner unendlichen Signalenergie ist weißes Rauschen nur als ein theoretisches Modell anzusehen. In der Praxis wird ein Signal als *white* bezeichnet, wenn es über ein definiertes Frequenzband ein flaches Spektrum hat.

Die mathematische Definition von weißem Rauschen sieht folgendermaßen aus:

- Ein endlich-dimensionaler Zufallsvektor $\mathbf{x}(k)$ heißt *white* genau dann, wenn

$$\begin{aligned}\mu_x &= E[\mathbf{x}(k)] = 0 \\ \mathbf{R}_{xx} &= E[\mathbf{x}(k)\mathbf{x}^T(k)] = \sigma^2 \mathbf{I}\end{aligned}$$

gilt.

- Ein stochastischer Prozess mit kontinuierlicher Zeit $x(t)$ mit $t \in \mathbb{R}$ heißt *white noise process* genau dann, wenn

$$\begin{aligned}\mu_x(t) &= E[x(t)] = 0 \\ r_{xx}(t_1, t_2) &= E[x(t_1)x(t_2)^T] = H\delta(t_1 - t_2)\end{aligned}$$

gilt. Die Autokorrelationsfunktion eines weißen Rauschens ist also ein Dirac-Impuls bei Null mit dem Faktor H . Weißes Rauschen ist stark stationär.

- Bei *gaussian noise* handelt es sich um Rauschen mit einer Verteilung der Dichtewahrscheinlichkeit die einer Normalverteilung entspricht.

Gaussian white noise hat Erwartungswert 0 und die Kovarianzfunktion $(t_1, t_2) \mapsto \delta(t_1 - t_2)$.

Bemerkung: Gauß'sches weißes Rauschen wird oft in mathematischen Problemstellungen verwendet, da es sich sehr gut zur Approximation von realen Situationen eignet.

Kapitel 2

Die Wiener-Filter-Theorie

Die Wiener-Filter-Theorie wurde in den 1940er Jahren hauptsächlich von Norbert Wiener und seinen Kollegen entwickelt [25]. Sie bildet die Grundlage für die Entwicklung signalabhängiger Filter.

2.1 Überblick

Einführend wird in 2.2, als Grundlage für die weiteren Abschnitte, die Funktionsweise des Wiener-Filters wie sie in [14] beschrieben wird, gezeigt. Im Abschnitt 2.2.1 wird die Berechnung des Schätzfehlers gezeigt und anhand eines Beispiels aus [20] näher erklärt. Die Berechnung der Filterkoeffizienten, die Wiener-Hopf-Gleichung und die Wiener-Lösung sind die Hauptbestandteile von 2.2.2. In 2.3 wird das wichtige Orthogonalitätsprinzip beschrieben. Die verwendeten Quellen dieses Abschnittes sind [19], [18], [14] und [20]. Unterkapitel 2.4 zeigt die Vektorschreibweise des minimalen Fehlers und dient hauptsächlich der Vorbereitung für den darauf folgenden Abschnitt. In 2.5 werden die Anwendungsgebiete Rauschentfernung (2.5.2), Signalprädiktion (2.5.3) und Signalglättung (2.5.1) behandelt. Diese und auch 2.4 wurden zum größten Teil aus [13] entnommen. Das Unterkapitel 2.6 beschreibt die mathematisch interessantere Filterung für kontinuierliche Zeit. Es wird zwischen der nicht-kausalen (2.6.1) und der kausalen Lösung (2.6.2) unterschieden. Als Quellen für diesen Abschnitt dienten hauptsächlich [19] und [5].

2.2 Die Funktionsweise des Wiener-Filters

Seien $x_i(n)$ die Eingangssignale, $f(n)$ das Zielsignal und $\hat{f}(n)$ die Schätzung des Zielsignals. Die Aufgabe des Filters ist es, die Eingangssignale zu einer

möglichst guten Schätzung des Zielsignals zu verarbeiten. Der dabei entstehende Schätzfehler wird mit $e(n)$ bezeichnet. Der Schätzwert ist von der Form

$$\hat{f}(n) = \sum_{i=0}^M w_i x_i(n) = \mathbf{w}^T \mathbf{x}(n),$$

wobei $\mathbf{w} = [w_0 \ w_1 \ \cdots \ w_M]^T$ der Filterkoeffizientenvektor und $\mathbf{x}(n) = [x_0(n) \ x_1(n) \ \cdots \ x_M(n)]^T$ der Eingangssignalvektor ist [14]. Es sollen nun die Gewichtungsfaktoren w_i der Summe der Eingangssignale so gewählt werden, dass der Schätzfehler im Mittel möglichst klein bleibt. Dies geschieht mittels einer Minimierung des mittleren quadratischen Fehlers (mean-square error).

2.2.1 Die Berechnung des Schätzfehlers

Der Schätzfehler wird folgendermaßen berechnet[14]:

$$\begin{aligned} E[e^2(n)] &= E[(f(n) - \mathbf{w}^T \mathbf{x}(n))^2] \\ &= E[f^2(n) - 2\mathbf{w}^T \mathbf{x}(n)f(n) + \mathbf{w}^T \mathbf{x}(n)\mathbf{x}^T(n)\mathbf{w}] \\ &= E[f^2(n)] - 2\mathbf{w}^T E[\mathbf{x}(n)f(n)] + \mathbf{w}^T E[\mathbf{x}(n)\mathbf{x}^T(n)]\mathbf{w} \\ &= \sigma_f^2 - 2\mathbf{w}^T \mathbf{r}_{\mathbf{x}f} + \mathbf{w}^T \mathbf{R}_{xx} \mathbf{w}. \end{aligned}$$

σ_f^2 ist die Leistung des Zielsignals (Varianz) und $\mathbf{r}_{\mathbf{x}f}$ der Kreuzkorrelationsvektor zwischen Eingangssignalen und Zielsignal. Bei der Matrix $\mathbf{R}_{xx}$ handelt es sich um die Kovarianzmatrix der Eingangssignale, die von folgender Form ist:

$$\mathbf{R}_{xx} = \begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \cdots & r_{xx}(M) \\ r_{xx}(-1) & r_{xx}(0) & \cdots & r_{xx}(M-1) \\ \vdots & \vdots & \ddots & \vdots \\ r_{xx}(-M) & r_{xx}(-M+1) & \cdots & r_{xx}(0) \end{bmatrix}$$

Beispiel 2.1: Angenommen die Einträge dieser Matrix ($r_{xx}(0) = 1.0$, $r_{xx}(1) = 0$) sind bereits gefunden. Außerdem sind die Varianz des Zielsignals ($\sigma_f^2 = 24.40$) und der Kreuzkorrelationsvektor ($\mathbf{r}_{\mathbf{x}f} = [2 \ 4.5]^T$) bekannt.

Einsetzen in obige Gleichung ergibt:

$$\begin{aligned} E[e^2(n)] &= 24.40 - 2 [w_0 \ w_1] \begin{bmatrix} 2 \\ 4.2 \end{bmatrix} + [w_0 \ w_1] \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} w_0 \\ w_1 \end{bmatrix} \\ &= 24.40 - 4w_0 - 9w_1 + w_0^2 + w_1^2 \end{aligned}$$

Da für die Minimierung des Fehlers der mean-square error (MSE) verwendet wurde, ist diese quadratische Gleichung für beliebige Filterkoeffizienten w_i definiert.

Man beachte, dass für das Finden der optimalen Wiener-Filterkoeffizienten das Zielsignal benötigt wird.

Nun wird die durch Einsetzen verschiedener Werte von w_0 und w_1 in die Funktion $E[e^2(n)]$ erhaltene Fläche betrachtet. Die optimalen Wiener Filterkoeffizienten sind genau jene Gewichtungsfaktoren, die im Vektor des Tiefpunktes der so entstandenen parabolischen nach oben konkaven Fläche enthalten sind. Die Gewichtungsfaktoren, die einen möglichst kleinen Schätzfehler erreichen, sind im betrachteten Beispiel $w_0 = 2$ und $w_1 = 4.5$ (vgl.[20]).

2.2.2 Die Berechnung der Filterkoeffizienten, die Wiener-Hopf-Gleichung und die Wiener-Lösung

Aufgrund der in **Beispiel 2.1** beschriebenen Form der erhaltenen Fläche muss am Extremum von $E[e^2(n)]$ die erste Ableitung den Wert Null und die zweite Ableitung einen positiven Wert haben [20].

Also

$$\frac{\partial E[e^2(w_0, w_1)]}{\partial w_0} = 0 \quad \frac{\partial E[e^2(w_0, w_1)]}{\partial w_1} = 0$$

und

$$\frac{\partial^2 E[e^2(w_0, w_1)]}{\partial^2 w_0} > 0 \quad \frac{\partial^2 E[e^2(w_0, w_1)]}{\partial^2 w_1} > 0.$$

Übertragen auf $M+1$ Filterkoeffizienten, ergibt sich durch den verschwindenden Gradienten

$$\begin{aligned} \nabla_{\mathbf{w}} E[e^2(n)] &= -2\mathbf{r}_{\mathbf{x}f} + 2\mathbf{R}_{xx}\mathbf{w}_{opt} \stackrel{!}{=} 0 \\ \mathbf{r}_{\mathbf{x}f} &= \mathbf{R}_{xx}\mathbf{w}_{opt} \quad (\text{Wiener-Hopf-Gleichung}) \\ \mathbf{w}_{opt} &= \mathbf{R}_{xx}^{-1}\mathbf{r}_{\mathbf{x}f} \quad (\text{Wiener-Lösung}). \end{aligned}$$

Wie in [14] ergibt sich der minimale Fehler

$$E[e^2(n)] = \sigma_e^2 = \sigma_f^2 - \mathbf{w}_{opt}^T \mathbf{r}_{\mathbf{x}f}.$$

Um wieder kurz auf das vorherige Beispiel zurückzukommen: Die Fehlerleistung für den optimalen Koeffizientenvektor entspricht der Distanz zwischen dem Tiefpunkt und der $\mathbf{w}$ -Ebene[20].

2.3 Das Orthogonalitätsprinzip

Das für Approximationsaufgaben wichtige Orthogonalitätsprinzip behält auch in der Wiener-Filter-Theorie seine Gültigkeit. Es besagt im Wesentlichen, dass das Fehlersignal im Optimum immer senkrecht zum Eingangssignal steht (vgl. [19] und [20]).

Für $i = 0, \dots, M$ gilt

$$\frac{\partial E[e^2(n)]}{\partial w_i} = \frac{\partial}{\partial w_i} E[e(n)e(n)] = 2E[e(n)\frac{\partial e(n)}{\partial w_i}] = 0 \quad (*_{2.1})$$

Aber aus

$$e(n) = f(n) - \sum_{i=0}^M w_i x_i(n)$$

folgt

$$\frac{\partial e(n)}{\partial w_i} = -x_i(n).$$

Einsetzen in (*2.1) ergibt

$$E[e_{min.}(n)x_i(n)] = 0 \quad i = 1, \dots, n \quad (*_{2.2})$$

bzw.

$$E[e_{min.}(n)\mathbf{x}(n)] = \mathbf{0}$$

weitere gilt

$$E[e_{min.}^2(n)] = \sigma_{e_{min.}}^2 = E[y(n) \cdot e_{min.}(n)] \quad (*_{2.3})$$

wobei

$$\begin{aligned} e_{min.}(n) &= f(n) - (w_0^{opt.} x_0(n) + \dots + w_M^{opt.} x_M(n)) \\ &= f(n) - \mathbf{w}_{opt.}^T \mathbf{x}(n). \end{aligned}$$

Die Gültigkeit von (*2.3) wird folgendermaßen gezeigt [14]:

Aus $\mathbf{r}_{\mathbf{x}f} = \mathbf{R}_{xx} \mathbf{w}_{opt}$ und $E[x_i x_j] = r_{x_i x_j} = r_{x_j x_i} = E[x_j x_i]$ folgt

$$r_{x_i f} = \mathbf{w}_{opt.}^T \begin{bmatrix} r_{x_i x_0} \\ r_{x_i x_1} \\ \vdots \\ r_{x_i x_M} \end{bmatrix} = \mathbf{w}_{opt.}^T E[\mathbf{x} \cdot x_i] = \mathbf{w}_{opt.}^T r_{\mathbf{x}x_i}.$$

Weiters gilt

$$\begin{aligned}
E[e_{min.}^2(n)] &= E[e_{min.}(n)(f(n) - \mathbf{w}_{opt.}^T \mathbf{x}(n))] \\
&= E[e_{min.}(n)f(n)] - E[e_{min.}(n)\mathbf{w}_{opt.}^T \mathbf{x}(n)] \\
&= E[e_{min.}(n)f(n)] - \mathbf{w}_{opt.}^T E[e_{min.}(n)\mathbf{x}(n)].
\end{aligned}$$

Wegen (*2.2) fällt der zweite Term in der letzten Gleichung weg und (*2.3) ist gezeigt.

Bemerkung: Im vorherigen Abschnitt wurde das Wiener-Filter bestimmt, indem als Gütefunktion der mittlere quadratische Fehler gewählt und mit Hilfe des Gradienten das Minimum der Fehlerfunktion $E[e^2(n)]$ ermittelt wurde. Das Wiener-Filter kann auch über das Orthogonalitätsprinzip hergeleitet werden. Beide Wege führen zum gleichen optimalen Filter $\mathbf{w}_{opt.}$, da in beiden Fällen der MSE optimiert wird. Die Herleitung über das Orthogonalitätsprinzip kann in [18] betrachtet werden.

Bemerkung: Geometrisch betrachtet ist der optimale Approximationsvektor eine orthogonale Projektion des Zielvektors in den Raum der Approximationsvektoren [20].

2.4 Vektorschreibweise des minimalen Fehlers

Für die im nächsten Abschnitt behandelten Anwendungsgebiete benötigen wir eine Vektorschreibweise des minimalen Fehlers. Sie ist eine direkte Folgerung der skalaren Darstellung.

Sei

$$\hat{f}(n) = \sum_{i=0}^M w_{ni} x_i(n)$$

für $n = 1, \dots, p$.

Die optimalen Gewichtungsfaktoren werden wie zuvor beschrieben gefunden. Die einzelnen $\hat{f}(n)$ bilden den Vektor $\hat{\mathbf{f}}$ und mit obiger Vorgehensweise ergibt sich

$$\mathbf{M}_{\hat{\mathbf{f}}} = \mathbf{C}_{ff} - \mathbf{C}_{fx} \mathbf{C}_{xx}^{-1} \mathbf{C}_{xf},$$

wobei $\mathbf{C}_{ff}$ die $p \times p$ Kovarianzmatrix ist. Die $p \times (M+1)$ Matrix $\mathbf{C}_{fx}$ beschreibt die Kreuzkorrelation [13].

Bemerkung: Bei den Diagonaleinträgen der $p \times p$ Matrix $\mathbf{M}_{\hat{\mathbf{f}}}$ handelt es sich genau um die einzelnen minimalen Fehler $E[e^2(n)]$ (vgl. [13]).

2.5 Rauschentfernung, Signalprädiktion und Signalglättung

Sind die verschiedenen Eingangssignale $x_i(n)$ abhängig vom Zielsignal oder werden sie aus einem einzigen Eingangssignal durch Zeitverzögerung gewonnen, d.h. $x_i(n) = x(n - i)$, ergeben sich für Wiener-Filter u.a. mit der Signalglättung, der Rauschentfernung und der Signalprädiktion, drei wichtige Anwendungsgebiete [14].

Es sei vorausgesetzt, dass diese Eingangssignale, sie werden aus Übersichtlichkeitsgründen nun mit $x(0), x(1), \dots, x(M)$ bezeichnet, schwach stationäre und mittelwertfreie Prozesse sind.

Da es sich nun ausschließlich um reelle Signale handelt, d.h. es gilt $r_{xx}(j) = r_{xx}(-j)$, ist die Kovarianzmatrix $\mathbf{C}_{xx}$ nun von symmetrischer Toeplitz-Struktur. Sie ist also von der Form:

$$\mathbf{C}_{xx} = \begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \cdots & r_{xx}(M) \\ r_{xx}(1) & r_{xx}(0) & \cdots & r_{xx}(M-1) \\ \vdots & \vdots & \ddots & \vdots \\ r_{xx}(M) & r_{xx}(M-1) & \cdots & r_{xx}(0) \end{bmatrix} = \mathbf{R}_{xx}$$

wobei $r_{xx}(j)$ die Autokorrelationsfunktion des Prozesses $x(k)$ ist [13].

2.5.1 Die Signalglättung

Für $k = 0, \dots, M$ beschreibt die Gleichung

$$x(k) = f(k) + \eta(k)$$

das durch Rauschen $\eta(k)$ gestörte Eingangssignal $x(k)$. Das Wiener-Filter berechnet die Schätzung des Nutzsignals $y = f(k)$ aus den Daten $x(0), \dots, \dots, x(M)$. Im Gegensatz zu der in Abschnitt 2.5.2 beschriebenen Filterung, ist es bei der Signalglättung erlaubt, zukünftige Daten zu verwenden. Während die Glättung beispielsweise zur Schätzung von $f(1)$ die Daten $\{x(0), \dots, \dots, x(M)\}$ zulässt, verwendet die Filterung nur $x(0)$ und $x(1)$. Das Glätten kann folglich nur durchgeführt werden, wenn bereits alle Eingangswerte bekannt sind.

Insgesamt soll also der Vektor der Nutzsignale,

$$\mathbf{y} = \mathbf{f} = [f(0) \quad f(1) \quad \cdots \quad f(M)]^T,$$

aus dem Vektor der Eingangssignale,

$$\mathbf{x} = [x(0) \quad x(1) \quad \cdots \quad x(M)]^T,$$

geschätzt werden.

Es gilt

$$r_{xx}(j) = r_{ff}(j) + r_{\eta\eta}(j)$$

da keine Korrelation zwischen Nutzsignal und Rauschen angenommen wird.

Es folgt

$$\mathbf{C}_{xx} = \mathbf{R}_{xx} = \mathbf{R}_{ff} + \mathbf{R}_{\eta\eta}.$$

und

$$\mathbf{C}_{yx} = E[\mathbf{f}\mathbf{x}^T] = E[\mathbf{f}(\mathbf{f} + \boldsymbol{\eta})^T] = \mathbf{R}_{ff}.$$

Bei $\mathbf{C}_{yx}$ handelt es sich wieder um die Kreuzkorrelation zwischen Eingangssignalen und Nutzsignalen.

Es wird die optimale Schätzung

$$\hat{\mathbf{f}} = \mathbf{R}_{ff}(\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta})^{-1}\mathbf{x}$$

erhalten.

Die $(M+1) \times (M+1)$ Matrix

$$\mathbf{W} = \mathbf{R}_{ff}(\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta})^{-1}$$

wird *Wiener-Glättungs-Matrix* genannt.

Es ergibt sich der minimale Fehler

$$\mathbf{M}_{\hat{f}} = \mathbf{R}_{ff} - (\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta})^{-1}.$$

2.5.2 Die Filterung

Es soll nun nur mit Hilfe der vergangenen und der aktuellen Daten das Signal vom Rauschen getrennt werden. Das Ziel ist es also $y = f(k)$ aus den Daten $\mathbf{x} = [x(0) \ x(1) \ \cdots \ x(k)]^T$ zu schätzen. Dies muss für alle Werte $k = 0, 1, \dots, M$ wiederholt werden.

Die Eingangssignale setzen sich wieder aus den unkorrelierten Prozessen Nutzsignal und Rauschen zusammen. Es gilt wieder

$$\mathbf{C}_{xx} = \mathbf{R}_{ff} + \mathbf{R}_{\eta\eta},$$

wobei $\mathbf{R}_{ff}$ und $\mathbf{R}_{\eta\eta}$ Autokorrelationsmatrizen sind.

Es folgt

$$\begin{aligned} \mathbf{C}_{yx} &= E[f(k) [x(0) \ x(1) \ \cdots \ x(k)]] \\ &= E[f(k) [f(0) \ f(1) \ \cdots \ f(k)]] \\ &= [r_{ff}(k) \ r_{ff}(k-1) \ \cdots \ r_{ff}(0)] \\ &= \hat{\mathbf{r}}_{ff}^T. \end{aligned}$$

Wie im Abschnitt über die Funktionsweise des Wiener-Filters beschrieben, wird $\hat{f}(k) = \mathbf{r}_{ff}^T \mathbf{R}_{xx}^{-1} \mathbf{x}$ berechnet.

Für den Gewichtevektor $\mathbf{w} = [w_0 \ w_1 \ \dots \ w_M]^T$ gilt

$$\mathbf{w} = (\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta})^{-1} \mathbf{r}_{ff}.$$

Definition: Eine *zeitveränderliche Impulsantwort* $h^{(k)}(j)$ ist folgendermaßen definiert:

$$h^{(k)}(j) = a_{k-j}, \quad j = 0, 1, \dots, k.$$

$h^{(k)}(j)$ ist die Antwort des Filters zum Zeitpunkt k auf einen j samples zuvor angewandten Impulses.

Der Vektor $\mathbf{h} = [h^{(k)}(0) \ h^{(k)}(1) \ \dots \ h^{(k)}(k)]^T$ ist folglich nur der auf den Kopf gestellte Vektor $\mathbf{w}$.

Einsetzen ergibt ein *time varying FIR-Filter* (**f**inite **i**mpulse **r**esponse)

$$\hat{f}(k) = \sum_{j=0}^k h^{(k)}(j) x(k-j).$$

Für $\mathbf{r}_{ff} = [r_{ff}(0) \ r_{ff}(1) \ \dots \ r_{ff}(k)]^T$ gilt aufgrund der symmetrischen Toeplitz-Struktur von $\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta}$

$$(\mathbf{R}_{ff} + \mathbf{R}_{\eta\eta}) \mathbf{h} = \mathbf{r}_{ff}$$

bzw.

$$\begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \dots & r_{xx}(k) \\ r_{xx}(1) & r_{xx}(0) & \dots & r_{xx}(k-1) \\ \vdots & \vdots & \ddots & \vdots \\ r_{xx}(k) & r_{xx}(k-1) & \dots & r_{xx}(0) \end{bmatrix} \begin{bmatrix} h^{(k)}(0) \\ h^{(k)}(1) \\ \vdots \\ h^{(k)}(k) \end{bmatrix} = \begin{bmatrix} r_{ff}(0) \\ r_{ff}(1) \\ \vdots \\ r_{ff}(k) \end{bmatrix}$$

für $r_{xx}(j) = r_{ff}(j) + r_{\eta\eta}(j)$. Dies sind die *Wiener-Hopf Filtergleichungen*. Sie müssen für jeden Wert k gelöst werden.

Ist k groß genug, wird das Filter zeitinvariant. Für $k \rightarrow \infty$, wird die zeitabhängige Impulsantwort $h^{(k)}$ durch ihre zeitinvariante Version ersetzt.

Der Versuch $f(k)$ auf der Basis von Daten der Gegenwart und der unendlichen Vergangenheit zu schätzen, ergibt dasselbe Gleichungssystem (*unendlicher Wiener-Filter*).

2.5.3 Die Signalprädiktion

Ziel der Prädiktion ist es $y = x(M+l)$ für $l \geq 1$ aus $\mathbf{x} = [x(0) \ \cdots \ x(M)]^T$ zu schätzen. Der resultierende Schätzer wird *l-step linear predictor* genannt.

Verwendet wird

$$y = \mathbf{C}_{yx} \mathbf{C}_{xx}^{-1} \mathbf{x}$$

mit

$$\mathbf{C}_{xx} = \mathbf{R}_{xx},$$

wobei $\mathbf{R}_{xx}$ eine $(M+1) \times (M+1)$ Matrix ist und

$$\begin{aligned} \mathbf{C}_{yx} &= E[x(M+l) [x(0) \ x(1) \ \cdots \ x(M)]^T] \\ &= [r_{xx}(M+l) \ r_{xx}(M-1+l) \ \cdots \ r_{xx}(l)]^T \\ &= \mathbf{r}_{xx}^T. \end{aligned}$$

Es folgt

$$\hat{x}(M+l) = \mathbf{r}_{xx}^T \mathbf{R}_{xx}^{-1} \mathbf{x}$$

und mit

$$\mathbf{w} = \mathbf{R}_{xx}^{-1} \mathbf{r}_{xx}$$

ergibt sich

$$\hat{x}(M+l) = \sum_{j=0}^M w_j x(j).$$

Es wird nun w_j durch $h(M+1-k)$ ersetzt. Dieser Schritt erlaubt einen Vergleich mit der zuvor beschriebenen Signalfilterung. Aus

$$\begin{aligned} \hat{x}(M+l) &= \sum_{j=0}^M h(M+1-j) x(j) \\ &= \sum_{j=1}^{M+1} h(j) x(M+1-j) \end{aligned}$$

wird erkannt, dass die Vorhersage der Output eines Filters mit Impulsantwort $h(k)$ ist.

Unter Berücksichtigung, dass $\mathbf{h}$ wieder der auf den Kopf gestellte Vektor $\mathbf{w}$ ist, gilt

$$\mathbf{R}_{xx} \mathbf{h} = \mathbf{r}_{xx}$$

mit $\mathbf{r}_{xx} = [r_{xx}(l) \ r_{xx}(l+1) \ \cdots \ r_{xx}(M+l)]^T$.

Es ergeben sich die Gleichungen

$$\begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \cdots & r_{xx}(M) \\ r_{xx}(1) & r_{xx}(0) & \cdots & r_{xx}(M-1) \\ \vdots & \vdots & \ddots & \vdots \\ r_{xx}(M) & r_{xx}(M-1) & \cdots & r_{xx}(0) \end{bmatrix} \begin{bmatrix} h(0) \\ h(1) \\ \vdots \\ h(M) \end{bmatrix} = \begin{bmatrix} r_{xx}(l) \\ r_{xx}(l+1) \\ \vdots \\ r_{xx}(M+l) \end{bmatrix},$$

die sog. *Wiener-Hopf Prädiktionsgleichungen* für den *l-step linear predictor* basierend auf $M+1$ (vergangenen) samples.

Die für den Spezialfall $l=1$ resultierenden Gleichungen sind identisch mit den *Yule-Walker Gleichungen*.

Der sich ergebende minimale quadratische Fehler ist von der Form

$$\mathbf{M}_{\hat{x}} = r_{xx}(0) - \mathbf{r}_{xx}^T \mathbf{R}_{xx}^{-1} \mathbf{r}_{xx}$$

bzw.

$$\begin{aligned} \mathbf{M}_{\hat{x}} &= r_{xx}(0) - \mathbf{r}_{xx}^T \mathbf{w} \\ &= r_{xx}(0) - \sum_{j=0}^M w_j r_{xx}(M+l-j) \\ &= r_{xx}(0) - \sum_{j=0}^M h(M+1-j) r_{xx}(M+l-j) \\ &= r_{xx}(0) - \sum_{j=0}^{M+1} h(j) r_{xx}(j+(l-1)). \end{aligned}$$

Bemerkung: Als Quelle für die soeben beschriebenen Anwendungsgebiete diene zum größten Teil [13].

2.6 Die Wiener-Filter-Theorie für kontinuierliche Zeit

Als Quellen für diese Sektion dienten hauptsächlich [5] und [19].

Sei $f(t)$ ein unbekanntes sample eines stochastischen Prozesses mit kontinuierlicher Zeit. Es gilt

$$x(t) = f(t) + \eta(t),$$

wobei die spektralen Leistungsdichten der stationären Prozesse $f(t)$ und $\eta(t)$ bekannt sind.

2.6.1 Die nicht-kausale Lösung

Ziel ist es, den Prozess $f(t)$ in Form einer Linearkombination aus vergangenen und zukünftigen Daten zu schätzen. Dies ist gleichbedeutend damit, dass die Impulsantwort $h(t)$ eines linearen, nicht-kausalen Filters gesucht wird, der folgendes erfüllt:

Für den Input $x(t)$ ist

$$\hat{f}(t) = \int_{-\infty}^{\infty} x(t - \alpha)h(\alpha)d\alpha$$

die Schätzung von $f(t)$ mit minimalem quadratischem Fehler e .

$\hat{f}(t)$ ist also eine Linearkombination der Daten $x(t - \tau)$, wobei τ alle Werte zwischen $-\infty$ und ∞ annimmt. Aus dem Orthogonalitätsprinzip für unendliche Summen folgt

$$E[(f(t) - \int_{-\infty}^{\infty} x(t - \alpha)h(\alpha)d\alpha)x(t - \tau)] = 0$$

für alle τ .

Es ergibt sich

$$r_{fx}(\tau) = \int_{-\infty}^{\infty} r_{xx}(\tau - \alpha)h(\alpha)d\alpha$$

für alle τ .

Die Impulsantwort für das optimale System wird durch das Lösen dieser Gleichung erhalten.

Mit der Substitution

$$r_{fx}(\tau) \leftrightarrow S_{fx}(\omega) \quad r_{xx}(\tau) \leftrightarrow S_{xx}(\omega) \quad h(t) \leftrightarrow H(\omega)$$

folgt aus dem Faltungstheorem und der obigen Gleichung

$$S_{fx}(\omega) = S_{xx}(\omega)H(\omega).$$

Die für optimale Filter passende Funktion $H(\omega)$ ist daher gegeben durch

$$H(\omega) = \frac{S_{fx}(\omega)}{S_{xx}(\omega)}.$$

Da für alle τ , $f(t) - \hat{f}(t)$ orthogonal zu $x(t - \tau)$ ist, folgt

$$E[(f(t) - \hat{f}(t))\hat{f}(t)] = 0$$

und für den minimalen quadratischen Fehler ergibt sich:

$$\begin{aligned}
e &= E[(f(t) - \hat{f}(t))^2] \\
&= E[(f(t) - \int_{-\infty}^{\infty} x(t - \alpha)h(\alpha)d\alpha)f(t)] \\
&= r_{ff}(0) - \int_{-\infty}^{\infty} r_{fx}(\alpha)h(\alpha)d\alpha.
\end{aligned}$$

Aus Substitution und der Verwendung der Parseval'schen Formel für Integrale folgt

$$e = \frac{1}{2\pi} \int_{-\infty}^{\infty} (S_{ff}(\omega) - S_{fx}^*(\omega)H(\omega))d\omega.$$

Spezialfall:

Sind, wie in den vorhergehenden Abschnitten angenommen, das Nutzsinal $f(t)$ und das Rauschen $\eta(t)$ orthogonal, also unkorreliert, d.h. $r_{f\eta}(\tau) = 0$, ergibt sich

$$S_{fx}(\omega) = S_{ff}(\omega) \quad S_{xx}(\omega) = S_{ff}(\omega) + S_{\eta\eta}(\omega).$$

Es gilt daher

$$H(\omega) = \frac{S_{ff}(\omega)}{S_{ff}(\omega) + S_{\eta\eta}(\omega)}$$

und

$$e = \frac{1}{2\pi} \int_{-\infty}^{\infty} \frac{S_{ff}(\omega)S_{\eta\eta}(\omega)}{S_{ff}(\omega) + S_{\eta\eta}(\omega)}d\omega.$$

2.6.2 Die kausale Lösung

Es können nun nur die vergangenen Daten verwendet werden. Der Schätzer ist gegeben durch

$$\hat{f}(t) = \int_0^{\infty} x(t - \alpha)h(\alpha)d\alpha.$$

Das heißt, er ist der Output eines kausalen Filters mit Input $x(t)$. Folglich ist $\hat{f}(t)$ also eine Linearkombination der Daten $x(t - \tau)$, wobei τ alle Werte von 0 bis ∞ annimmt.

Um den mittleren quadratischen Fehler zu minimieren, ist es notwendig die passende Impulsantwort $h(t)$ zu finden, sodass

$$E[(f(t) - \int_0^{\infty} x(t - \alpha)h(\alpha)d\alpha)x(t - \tau)] = 0 \quad \text{für } \tau > 0$$

erfüllt ist.

Es ergibt sich die *Wiener-Hopf-Gleichung* für Integrale

$$r_{fx}(\tau) - \int_0^\infty r_{xx}(\tau - \alpha)h(\alpha)d\alpha = 0,$$

die nur für $\tau > 0$ gültig ist.

Sei nun

$$\gamma(\tau) = r_{fx}(\tau) - \int_0^\infty r_{xx}(\tau - \alpha)h(\alpha)d\alpha. \quad (*_{2.4})$$

Es gilt

$$\gamma(\tau) = 0 \quad \text{für } \tau > 0$$

und

$$h(\tau) = 0 \quad \text{für } \tau < 0.$$

Es genügt also, eine kausale Funktion $h(t)$ und eine antikausale Funktion $\gamma(\tau)$ zu finden, die zusammen $(*_{2.4})$ erfüllen. Seien die Laplace-Transformationen von $\gamma(\tau)$ und $h(t)$ mit $Y(s)$ bzw. $H(s)$ bezeichnet. Wie in [19] gezeigt, sind $Y(s)$ für $\operatorname{Re} s < 0$ und $H(s)$ für $\operatorname{Re} s > 0$ analytisch.

Das Integral in $(*_{2.4})$ ist die Faltung von r_{xx} mit $h(\tau)$. Anwendung der Laplace-Transformation auf beiden Seiten ergibt

$$Y(s) = S_{fx}(-js) - S_{xx}(-js)H(s). \quad (*_{2.5})$$

Die Funktionen $S_{fx}(-js)$ und $S_{xx}(-js)$ sind Laplace-Transformationen von $r_{fx}(\tau)$ bzw. $r_{xx}(\tau)$.

Es stellt sich nun die Aufgabe, zwei Funktionen $Y(s)$ und $H(s)$ zu finden, sodass $(*_{2.5})$ gilt und den oben beschriebenen Bedingungen bezüglich der Analytizität nicht widersprochen wird.

Folgende Lösung dieses Problems und der nachfolgende Beweis sind [19] entnommen:

- Zuerst wird die Funktion $S_{xx}(-js)$ als Produkt

$$S_{xx}(-js) = A^+(s)A^-(s)$$

dargestellt. Hierbei sind der Faktor $A^+(s)$ und dessen inverse Funktion $\frac{1}{A^+(s)}$ für $\operatorname{Re} s > 0$ bzw. der zweite Faktor $A^-(s)$ und dessen Inverse $\frac{1}{A^-(s)}$ für $\operatorname{Re} s < 0$ analytisch. Ist das Spektrum $S_{xx}(\omega)$ rational in ω wird, um $A^+(s)$ und $A^-(s)$ zu determinieren, $S_{xx}(-js)$ faktorisiert und die Pole und Nullstellen in der linken Halbebene des Pol-Nullstellen-Diagramms werden $A^+(s)$ zugeordnet bzw. die Pole und Nullstellen in der rechten Halbebene werden $A^-(s)$ zugeteilt.

- Nun wird $\frac{S_{fx}(-js)}{A^-(s)}$ gebildet und folgendermaßen als Summe ausgedrückt:

$$\frac{S_{fx}(-js)}{A^-(s)} = B^+(s) + B^-(s),$$

wobei $B^+(s)$ für $\operatorname{Re} s > 0$ und $B^-(s)$ für $\operatorname{Re} s < 0$ analytisch ist. Im Falle, dass das Spektrum $S_{fx}(\omega)$ rational in ω ist, wird um $B^+(s)$ und $B^-(s)$ zu determinieren, mittels Partialbruchzerlegung eine Summe aus linearen Termen gebildet. Die Terme deren Pole sich auf der linken Seite des Pol-Nullstellen-Diagramms befinden werden $B^+(s)$ zugeordnet und die restlichen Terme werden $B^-(s)$ zugeteilt.

- Die gesuchte Funktion $H(s)$ ist nun gegeben durch

$$H(s) = \frac{B^+(s)}{A^+(s)}$$

Beweis: Da die Funktionen $B^+(s)$ und $\frac{1}{A^+(s)}$ so konstruiert sind, dass sie für $\operatorname{Re} s < 0$ analytisch sind, ist die Funktion $H(s)$ für $\operatorname{Re} s > 0$ offensichtlich ebenfalls analytisch. Um den Beweis zu vervollständigen, ist es notwendig zu zeigen, dass die durch Einsetzen von $H(s)$ in (*2.5) resultierende Funktion $Y(s)$ für $\operatorname{Re} s < 0$ analytisch ist.

Unter Verwendung der drei Gleichungen der einzelnen Punkte der Lösung ergibt sich

$$Y(s) = (B^+(s) + B^-(s))A^-(s) - \frac{A^+(s)A^-(s)B^+(s)}{A^+(s)} = A^-(s)B^-(s).$$

Die Funktionen $A^-(s)$ und $B^-(s)$ sind nach Konstruktion für $\operatorname{Re} s > 0$ analytisch. □

Es ergibt sich, aus denselben Überlegungen wie im nicht-kausalen Fall, nach der Anpassung der Integrationsgrenzen der minimale quadratische Fehler

$$e = r_{ff}(0) - \int_0^\infty r_{fx}(\alpha)h(\alpha)d\alpha.$$

Bemerkung: Die soeben in Unterkapitel 2.6 erhaltenen Resultate können mit geringem Aufwand auf zufällige Folgen ausgeweitet bzw. übertragen werden.

Kapitel 3

Die Kalman-Filter-Theorie

Das Kalman-Filter wurde 1960 von Rudolf Kalman erstmals für zeitdiskrete lineare Systeme entwickelt [11]. Es beschreibt einen rekursiven Filter, der den stochastischen Zustand von dynamischen Systemen aus Messungen die von Rauschen gestört sind, schätzt. Es wird dabei der mittlere quadratische Fehler minimiert.

Die wichtigsten Varianten des Kalman-Filters sind das Erweiterte-Kalman-Filter [4], das für Anwendungen in der Meteorologie wichtige Ensemble-Kalman-Filter und das Kalman-Bucy-Filter (nachzulesen in [12] und [23]). Bei letzterem handelt es sich um die Erweiterung des Filters auf zeitkontinuierliche Systeme. Mit dem in dieser Arbeit behandeltem Erweiterten-Kalman-Filter ist es möglich auch Zustände nichtlinearer Systeme zu schätzen. Auf Basis dieses Filters wurde in den 90er Jahren das, in einem eigenen Abschnitt behandelte, Ensemble-Kalman-Filter entwickelt.

3.1 Überblick

Unter Zuhilfenahme von [11], [2] und [4] wird in 3.2 die Kalman-Filter-Theorie zur Zustandsschätzung linearer stochastischer Systeme beschrieben. Um im weiteren Verlauf der Arbeit Verwechslungen mit anderen Varianten des Kalman-Filters vorzubeugen, wird das Kalman-Filter für lineare Systeme, das "Standard"-Kalman-Filter genannt. In 3.2.1 wird die Problemstellung, wie sie in [4] angegeben ist, gezeigt und die sog. *Innovationsfolge* definiert. Dieser Unterabschnitt enthält unter anderem auch Lemmata worin die wichtigsten Eigenschaften der neu definierten Begriffe beschrieben werden. 3.2.2 beschäftigt sich mit dem Schätzvektor $\hat{\mathbf{x}}_k$, dessen Ziel es ist $\mathbf{x}_k$ so zu schätzen, dass die Varianz des Schätzfehlers minimal ist. Im abschließenden Unterkapitel der "Standard"-Kalman-Filter-Theorie 3.2.3 werden nun

die sog. *Kalman-Filter-Gleichungen* hergeleitet. Neben der oben genannten Literatur wird hier zusätzlich [28] verwendet. In diesem Abschnitt werden auch die *Kalman-Gain-Matrizen* und die *Prädiktor-Korrektur-Gleichung* eingeführt. Am Ende des Kapitels werden die im letzten Unterabschnitt erzielten Ergebnisse in übersichtlicher Form zu den *Kalman-Filter-Gleichungen* zusammengefasst. Diese wichtigen Gleichungen werden dann im Laufe der Arbeit des Öfteren verwendet.

In Abschnitt 3.3 wird das Erweiterte-Kalman-Filter [4], die wohl wichtigste Variante des Kalman-Filters, behandelt. Die beschriebene Theorie erlaubt es, auch Zustände nichtlinearer Systeme zu schätzen.

Abschnitt 3.4 zeigt eine Alternative zum Erweiterten-Kalman-Filter, das sogenannte Ensemble-Kalman-Filter. Die Theorie, in der die Kovarianzinformationen durch Fortpflanzung eines Ensembles erhalten werden, wurde erstmals 1994 von Geir Evensen [7] beschrieben. Die *Ensemble-Kovarianzmatrizen* werden in 3.4.1, der Ablauf des Filters in 3.4.2 beschrieben. Zusätzlich zu [7] wurden in diesem Abschnitt noch [8] und [9] als Quellen verwendet.

3.2 Das "Standard"-Kalman-Filter

Als Quellen für diese Sektion dienten hauptsächlich [4], [2] und [28].

3.2.1 Problemstellung und die Innovationsfolge

Es wird ein lineares stochastisches System betrachtet:

$$\begin{aligned}\mathbf{x}_{k+1} &= \mathbf{A}_k \mathbf{x}_k + \Gamma_k \xi_k \\ \mathbf{v}_k &= \mathbf{C}_k \mathbf{x}_k + \eta_k,\end{aligned}$$

wobei es sich bei $\mathbf{A}_k$ um eine $n \times n$ -, bei Γ_k eine $n \times p$ - und bei $\mathbf{C}_k$ um eine $q \times n$ -Matrix handelt ($1 \leq p, q \leq n$). Die drei Matrizen sind bekannt und konstant. Die erste Gleichung heißt *Systemgleichung* und zweite wird *Beobachtungsgleichung* genannt.

Weiters gilt:

$$\begin{aligned}E[\xi_k] &= 0, & E[\xi_k \xi_l^T] &= \mathbf{Q}_k \delta_{kl}, \\ E[\eta_k] &= 0, & E[\eta_k \eta_l^T] &= \mathbf{R}_k \delta_{kl}, \\ E[\xi_k \eta_l^T] &= 0, & E[\mathbf{x}_0 \xi_k^T] &= 0, & E[\mathbf{x}_0 \eta_k^T] &= 0,\end{aligned}$$

für alle $k, l = 0, 1, \dots$. Die Matrizen $\mathbf{Q}_k$ und $\mathbf{R}_k$ sind positiv definit und symmetrisch.

Definition: Sei $\mathbf{x}$ ein n -dimensionaler und $\mathbf{w}$ ein q -dimensionaler Zufallsvektor. Das "innere Produkt" $\langle \mathbf{x}, \mathbf{w} \rangle$ ist definiert durch:

$$\langle \mathbf{x}, \mathbf{w} \rangle = E[(\mathbf{x} - E(\mathbf{x}))(\mathbf{w} - E(\mathbf{w}))^T].$$

Es handelt sich hierbei um die Kovarianz von $\mathbf{x}$ und $\mathbf{w}$.

Bemerkung: $\langle \mathbf{x}, \mathbf{w} \rangle$ ist eine $n \times q$ -Matrix.

Definition: $\| \mathbf{w} \|_q$ ist die positive Quadratwurzel von $\langle \mathbf{w}, \mathbf{w} \rangle$. Somit ist $\| \mathbf{w} \|_q$ eine nicht-negativ definite $q \times q$ -Matrix mit

$$\| \mathbf{w} \|_q^2 = \| \mathbf{w} \|_q \| \mathbf{w} \|_q^T = \langle \mathbf{w}, \mathbf{w} \rangle.$$

Gleichermaßen ist $\| \mathbf{x} \|_n$ die positive Quadratwurzel von $\langle \mathbf{x}, \mathbf{x} \rangle$.

Seien nun $\mathbf{w}_0, \dots, \mathbf{w}_r$ q -dimensionale Zufallsvektoren. Mit $Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ wird die lineare Hülle (linear span) bezeichnet. Es gilt also

$$\begin{aligned} Y(\mathbf{w}_0, \dots, \mathbf{w}_r) &= \\ &= \{ \mathbf{y} : \mathbf{y} = \sum_{i=0}^r P_i \mathbf{w}_i, \text{ wobei } P_0, \dots, P_r \text{ konstante } n \times q \text{ Matrizen sind} \}. \end{aligned}$$

Definition: Sei F_k folgendermaßen definiert:

$$F_k := \min \{ \text{spur} \| \mathbf{x}_k - \hat{\mathbf{y}} \|_n^2 : \mathbf{y} \in Y(\mathbf{w}_0, \dots, \mathbf{w}_r) \}.$$

Es stellt sich nun die Aufgabe ein $\hat{\mathbf{y}}$ in $Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ zu finden, so dass

$$\text{spur} \| \mathbf{x}_k - \hat{\mathbf{y}} \|_n^2 = F_k$$

gilt.

Das gesuchte $\hat{\mathbf{y}}$ wird folgendermaßen charakterisiert:

Lemma: $\hat{\mathbf{y}} \in Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ erfüllt $\text{spur} \| \mathbf{x}_k - \hat{\mathbf{y}} \|_n^2 = F_k$ genau dann, wenn

$$\langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_j \rangle = O_{n \times q}$$

für alle $j = 0, 1, \dots, r$ gilt. Außerdem ist $\hat{\mathbf{y}}$ eindeutig, d.h. es ist

$$\text{spur} \| \mathbf{x}_k - \hat{\mathbf{y}} \|_n^2 = \text{spur} \| \mathbf{x}_k - \tilde{\mathbf{y}} \|_n^2$$

genau dann, wenn $\hat{\mathbf{y}} = \tilde{\mathbf{y}}$ gilt.

Beweis: Angenommen es gilt $\text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 = F_k$, aber $\langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_{j_0} \rangle = C \neq 0_{n \times q}$ für einige j_0 mit $0 \leq j_0 \leq r$. Folglich ist $\mathbf{w}_{j_0} \neq 0$, so dass $\parallel \mathbf{w}_{j_0} \parallel_q^2$ eine positiv definite, symmetrische Matrix ist. Die Existenz der inversen Matrix $\parallel \mathbf{w}_{j_0} \parallel_q^{-2}$ ist dadurch gezeigt. Es ist daher $C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} C^T \neq 0_{n \times n}$ und eine nicht-negativ definite, symmetrische Matrix. Es kann gezeigt werden, dass

$$\text{spur}\{C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} C^T\} > 0 \quad (*_{3.1})$$

gilt.

Der Vektor $\hat{\mathbf{y}} + C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} \mathbf{w}_{j_0}$ ist nun in $Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ und wegen $(*_{3.1})$ gilt

$$\begin{aligned} & \text{spur} \parallel \mathbf{x}_k - (\hat{\mathbf{y}} + C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} \mathbf{w}_{j_0}) \parallel_n^2 \\ = & \text{spur}\{\parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 - \langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_{j_0} \rangle (C \parallel \mathbf{w}_{j_0} \parallel_q^{-2})^T \\ & - C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} \langle \mathbf{w}_{j_0}, \mathbf{x}_k - \hat{\mathbf{y}} \rangle + C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} \parallel \mathbf{w}_{j_0} \parallel_q^2 (C \parallel \mathbf{w}_{j_0} \parallel_q^{-2})^T\} \\ = & \text{spur}\{\parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 - C \parallel \mathbf{w}_{j_0} \parallel_q^{-2} C^T\} \\ < & \text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 = F_k. \end{aligned}$$

Dies widerspricht aber der Definition von F_k .

Umgekehrt, sei $\langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_j \rangle = 0_{n \times q}$ für alle $j = 0, 1, \dots, r$. Angenommen $\mathbf{y}$ sei ein beliebiger n -dimensionaler Zufallsvektor in $Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ und $\mathbf{y}_0 = \mathbf{y} - \hat{\mathbf{y}} = \sum_{j=0}^r P_{0j} \mathbf{w}_j$, wobei die P_{0j} , für $j = 0, 1, \dots, r$, konstante $n \times q$ -Matrizen sind.

Es gilt:

$$\begin{aligned} & \text{spur} \parallel \mathbf{x}_k - \mathbf{y} \parallel_n^2 \\ = & \text{spur} \parallel (\mathbf{x}_k - \hat{\mathbf{y}}) - \mathbf{y}_0 \parallel_n^2 \\ = & \text{spur}\{\parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 - \langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{y}_0 \rangle - \langle \mathbf{y}_0, \mathbf{x}_k - \hat{\mathbf{y}} \rangle + \parallel \mathbf{y}_0 \parallel_n^2\} \\ = & \text{spur}\{\parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 - \sum_{j=0}^r \langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_j \rangle P_{0j}^T \\ & - \sum_{j=0}^r P_{0j} \langle \mathbf{x}_k - \hat{\mathbf{y}}, \mathbf{w}_j \rangle^T + \parallel \mathbf{y}_0 \parallel_n^2\} \\ = & \text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 + \text{spur} \parallel \mathbf{y}_0 \parallel_n^2 \\ \geq & \text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2, \end{aligned}$$

so dass $\text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{y}} \parallel_n^2 = F_k$. Gleichheit wird genau dann erreicht, wenn $\text{spur} \parallel \mathbf{y}_0 \parallel_n^2 = 0$ oder $\mathbf{y}_0 = 0$ (also $\mathbf{y} = \hat{\mathbf{y}}$) gilt. $\square$

Definition: Gegeben sei eine Datenfolge, bestehend aus den Zufallsvektoren $\{\mathbf{v}_j\}, j = 0, \dots, k$, der Dimension q . Die *Innovationsfolge* $\{\mathbf{z}_j\}, j = 0, \dots, k$, von $\{\mathbf{v}_j\}$ ist definiert durch

$$\mathbf{z}_j = \mathbf{v}_j - \mathbf{C}_j \hat{\mathbf{y}}_{j-1}, j = 0, \dots, k,$$

mit $\hat{\mathbf{y}}_{-1} = 0$ und

$$\hat{\mathbf{y}}_{j-1} = \sum_{i=0}^{j-1} \hat{\mathbf{P}}_{j-1,i} \mathbf{v}_i \in Y(\mathbf{v}_0, \dots, \mathbf{v}_{j-1}), j = 1, \dots, k,$$

wobei die $q \times n$ -Matrizen $\mathbf{C}_j$ die Beobachtungsmatrizen, wie sie am Anfang des Kapitels verwendet werden, sind und die $n \times q$ -Matrizen $\hat{\mathbf{P}}_{j-1,i}$ so gewählt werden, dass $\hat{\mathbf{y}}_{j-1}$ das zuvor beschriebene Minimierungsproblem löst. $Y(\mathbf{w}_0, \dots, \mathbf{w}_r)$ aus obiger Definition wird dabei durch $Y(\mathbf{v}_0, \dots, \mathbf{v}_{j-1})$ ersetzt.

Folgendes Lemma beschreibt die Korrelationseigenschaften der Innovationsfolge:

Lemma: Die Innovationsfolge $\{\mathbf{z}_j\}$ von $\{\mathbf{v}_j\}$ erfüllt folgende Eigenschaft:

$$\langle \mathbf{z}_j, \mathbf{z}_l \rangle = (R_l + \mathbf{C}_l \|\mathbf{x}_l - \hat{\mathbf{y}}_{l-1}\|_n^2 \mathbf{C}_l^T) \delta_{jl},$$

wobei $R_l = \text{Var}(\eta_l) > 0$ gilt.

Bemerkung: Aus Übersichtlichkeitsgründen wird

$$\hat{\mathbf{e}}_j = \mathbf{C}_l(\mathbf{x}_l - \hat{\mathbf{y}}_{l-1})$$

gesetzt.

Beweis: Aus obiger Definition, der Bemerkung und der Beobachtungsgleichung $\mathbf{v}_k = \mathbf{C}_k \mathbf{x}_k + \eta_k$ (vgl. S.28) folgt:

$$\mathbf{z}_j = \hat{\mathbf{e}}_j + \eta_j,$$

wobei die Folge $\{\eta_k\}$ das Rauschen beschreibt und

$$\langle \eta_l, \hat{\mathbf{e}}_j \rangle = O_{q \times q} \quad \forall \quad l \geq j \quad (*_{3.2})$$

gilt. Bei $O_{q \times q}$ handelt es sich um die $q \times q$ -Nullmatrix.

- Für $j = l$ folgt nun

$$\begin{aligned} \langle \mathbf{z}_l, \mathbf{z}_l \rangle &= \langle \hat{\mathbf{e}}_l + \eta_l, \hat{\mathbf{e}}_l + \eta_l \rangle \\ &= \langle \hat{\mathbf{e}}_l, \hat{\mathbf{e}}_l \rangle + \langle \eta_l, \eta_l \rangle \\ &= \mathbf{C}_l \|\mathbf{x}_l - \hat{\mathbf{y}}_{l-1}\|_n^2 \mathbf{C}_l^T + R_l. \end{aligned}$$

- Sei nun $j \neq l$. Wegen $\langle \hat{\mathbf{e}}_l, \hat{\mathbf{e}}_j \rangle^T = \langle \hat{\mathbf{e}}_j, \hat{\mathbf{e}}_l \rangle$ kann o.B.d.A. $j > l$ angenommen werden. Durch Verwendung des vorhergehenden Lemmas ergibt sich

$$\begin{aligned}
\langle \mathbf{z}_j, \mathbf{z}_l \rangle &= \langle \hat{\mathbf{e}}_j, \hat{\mathbf{e}}_l \rangle + \langle \hat{\mathbf{e}}_j, \eta_l \rangle + \langle \eta_j, \hat{\mathbf{e}}_l \rangle + \langle \eta_j, \eta_l \rangle \\
&= \langle \hat{\mathbf{e}}_j, \hat{\mathbf{e}}_l + \eta_l \rangle \\
&= \langle \hat{\mathbf{e}}_j, \mathbf{z}_l \rangle \\
&= \langle \hat{\mathbf{e}}_j, \mathbf{v}_l - \mathbf{C}_l \hat{\mathbf{y}}_{l-1} \rangle \\
&= \langle \mathbf{C}_j (\mathbf{x}_j - \hat{\mathbf{y}}_{j-1}), \mathbf{v}_l - \mathbf{C}_l \sum_{i=0}^{l-1} \hat{\mathbf{P}}_{l-1,i} \hat{\mathbf{v}}_i \rangle \\
&= \mathbf{C}_j \langle \mathbf{x}_j - \hat{\mathbf{y}}_{j-1}, \mathbf{v}_l \rangle - \mathbf{C}_j \sum_{i=0}^{l-1} \langle \mathbf{x}_j - \hat{\mathbf{y}}_{j-1}, \mathbf{v}_i \rangle \hat{\mathbf{P}}_{l-1,i}^T \mathbf{C}_l^T \\
&= O_{q \times q}.
\end{aligned}$$

□

Aus dem Lemma und der Eigenschaft $\mathbf{R}_j > 0$ folgt, dass $\{\mathbf{z}_j\}$ eine "orthogonale" Folge von nicht-Nullvektoren ist, d.h.

$$\mathbf{e}_j = \|\mathbf{z}_j\|_q^{-1} \mathbf{z}_j.$$

$\{\mathbf{e}_j\}$ ist eine "orthonormale" Folge, im Sinne, dass $\langle \mathbf{e}_i, \mathbf{e}_j \rangle = \delta_{ij} \mathbf{I}_q$ für alle i und j gilt.

3.2.2 Der Minimum-Varianz-Schätzer

Ziel ist es nun, den Vektor $\mathbf{x}_k$ so zu schätzen, dass die Varianz des Schätzfehlers minimal ist. Dieser neue Schätzvektor wird mit $\hat{\mathbf{x}}_k$ bezeichnet und ist von der Form

$$\hat{\mathbf{x}}_k = \sum_{i=0}^k \langle \mathbf{x}_k, \mathbf{e}_i \rangle \mathbf{e}_i.$$

Bemerkung: $\hat{\mathbf{x}}_k$ ist die Fourier-Entwicklung von $\mathbf{x}_k$ bezüglich der Folge $\{\mathbf{e}_j\}$.

Lemma: $\hat{\mathbf{x}}_k$ ist der Minimum-Varianz-Schätzer von $\mathbf{x}_k$.

Beweis: Da

$$\langle \hat{\mathbf{x}}_k, \mathbf{e}_j \rangle = \sum_{i=0}^k \langle \mathbf{x}_k, \mathbf{e}_i \rangle \langle \mathbf{e}_i, \mathbf{e}_j \rangle = \langle \mathbf{x}_k, \mathbf{e}_j \rangle,$$

gilt, ergibt sich

$$\langle \mathbf{x}_k - \hat{\mathbf{x}}_k, \mathbf{e}_j \rangle = O_{n \times q}, \quad j = 0, 1, \dots, k. \quad (*_{3.3})$$

Es folgt

$$\langle \mathbf{x}_k - \hat{\mathbf{x}}_k, \mathbf{v}_j \rangle = O_{n \times q}, \quad j = 0, 1, \dots, k,$$

und aus dem Lemma auf S.29 wird

$$\text{spur} \parallel \mathbf{x}_k - \hat{\mathbf{x}}_k \parallel_n^2 = \min \{ \text{spur} \parallel \mathbf{x}_k - \mathbf{y} \parallel_T^2 : \mathbf{y} \in Y(\mathbf{v}_0, \dots, \mathbf{v}_k) \}.$$

erhalten.

3.2.3 Die Kalman-Filter-Gleichungen

Mit Hilfe der im vorhergehenden Abschnitt beschriebenen Resultate wird nun die Herleitung der Kalman-Filter-Gleichungen gezeigt.

Bemerkung: Es sei vorausgesetzt, dass $\{\xi_k\}$ und $\{\eta_k\}$ mittelwertfreie Gauß'sche white noise Folgen sind, sodass $\text{Var}(\xi_k) = \mathbf{Q}_k$ und $\text{Var}(\eta_k) = \mathbf{R}_k$ positiv definite Matrizen sind und $E(\xi_k \eta_l) = 0$ für alle k und l gilt. Weiters wird der Anfangszustand $\mathbf{x}_0$ als unabhängig von ξ_k und η_k , im Sinne von $E(\mathbf{x}_0 \xi_k^T) = 0$ und $E(\mathbf{x}_0 \eta_k^T) = 0$ für alle k , angenommen.

Aus den Annahmen ist bekannt, dass

$$\langle \xi_{k-1}, \mathbf{e}_j \rangle = O_{n \times q}, \quad j = 0, 1, \dots, k-1,$$

so dass folgendes gilt:

$$\begin{aligned} \hat{\mathbf{x}}_k &= \sum_{j=0}^k \langle \mathbf{x}_k, \mathbf{e}_j \rangle \mathbf{e}_j \\ &= \sum_{j=0}^{k-1} \langle \mathbf{x}_k, \mathbf{e}_j \rangle \mathbf{e}_j + \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k \\ &= \sum_{j=1}^{k-1} \{ \langle A_{k-1} \mathbf{x}_{k-1}, \mathbf{e}_j \rangle \mathbf{e}_j + \langle \Gamma_{k-1} \xi_{k-1}, \mathbf{e}_j \rangle \mathbf{e}_j \} + \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k \\ &= A_{k-1} \sum_{j=0}^{k-1} \langle \mathbf{x}_{k-1}, \mathbf{e}_j \rangle \mathbf{e}_j + \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k \\ &= A_{k-1} \hat{\mathbf{x}}_{k-1} + \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k. \end{aligned}$$

Definition: Der Ausdruck $A_{k-1}\hat{\mathbf{x}}_{k-1}$ wird nun mit $\hat{\mathbf{x}}_{k|k-1}$ bezeichnet, d.h. also

$$A_{k-1}\hat{\mathbf{x}}_{k-1} = \hat{\mathbf{x}}_{k|k-1},$$

wobei $\hat{\mathbf{x}}_{k-1} := \hat{\mathbf{x}}_{k-1|k-1}$ gilt.

Infolgedessen gilt

$$\hat{\mathbf{x}}_k = \hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k.$$

Um nun die sog. Prädiktor-Korrektur des Kalman-Filters zu erhalten, ist es nun offensichtlich notwendig, die Existenz einer konstanten $n \times q$ -Matrix $\mathbf{G}_k$, mit

$$\langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k = \mathbf{G}_k(\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}),$$

zu zeigen.

Um dies zu erreichen wird der Zufallsvektor $(\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1})$ näher betrachtet. Nun wird folgende Eigenschaft erhalten:

Lemma: Für $j = 0, 1, \dots, k$ gilt

$$\langle \mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}, \mathbf{e}_j \rangle = \|\mathbf{z}_k\|_q \delta_{kj}.$$

Beweis: Aus obiger Definition, (*3.3) und den Orthogonalitätsbedingungen

$$\langle \hat{\mathbf{y}}_j, \mathbf{z}_k \rangle = O_{n \times q}, \quad j = 0, 1, \dots, k-1,$$

folgt

$$\begin{aligned} & \langle \mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}, \mathbf{e}_k \rangle \\ &= \langle \mathbf{v}_k - \mathbf{C}_k(\hat{\mathbf{x}}_{k|k} - \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k), \mathbf{e}_k \rangle \\ &= \langle \mathbf{v}_k, \mathbf{e}_k \rangle - \mathbf{C}_k \{ \langle \hat{\mathbf{x}}_{k|k}, \mathbf{e}_k \rangle - \langle \mathbf{x}_k, \mathbf{e}_k \rangle \} \\ &= \langle \mathbf{v}_k, \mathbf{e}_k \rangle - \mathbf{C}_k \langle \hat{\mathbf{x}}_{k|k} - \mathbf{x}_k, \mathbf{e}_k \rangle \\ &= \langle \mathbf{v}_k, \mathbf{e}_k \rangle \\ &= \langle \mathbf{z}_k + \mathbf{C}_k \hat{\mathbf{y}}_{k-1}, \|\mathbf{z}_k\|_q^{-1} \mathbf{z}_k \rangle \\ &= \langle \mathbf{z}_k, \mathbf{z}_k \rangle \|\mathbf{z}_k\|_q^{-1} + \mathbf{C}_k \langle \hat{\mathbf{y}}_{k-1}, \mathbf{z}_k \rangle \|\mathbf{z}_k\|_q^{-1} \\ &= \|\mathbf{z}_k\|_q. \end{aligned}$$

Andererseits ergibt sich aus der Definition, (*3.2) und (*3.3)

$$\begin{aligned} & \langle \mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}, \mathbf{e}_j \rangle \\ &= \langle \mathbf{C}_k \mathbf{x}_k + \eta_k - \mathbf{C}_k(\hat{\mathbf{x}}_{k|k} - \langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k), \mathbf{e}_j \rangle \\ &= \mathbf{C}_k \langle \mathbf{x}_k - \hat{\mathbf{x}}_{k|k}, \mathbf{e}_j \rangle + \langle \eta_k, \mathbf{e}_j \rangle + \mathbf{C}_k \langle \mathbf{x}_k, \mathbf{e}_k \rangle \langle \mathbf{e}_k, \mathbf{e}_j \rangle \\ &= O_{q \times q}, \end{aligned}$$

für $j = 0, 1, \dots, k-1$. □

Der q -dimensionale Zufallsvektor $(\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1})$ kann für einige konstante $q \times q$ -Matrizen $\mathbf{M}_i$ (siehe [4] S.39), als Summe $\sum_{i=0}^k \mathbf{M}_i \mathbf{e}_i$ ausgedrückt werden.

Aus dem gerade formulierten Lemma folgt nun für $j = 0, 1, \dots, k$,

$$\left\langle \sum_{i=0}^k \mathbf{M}_i \mathbf{e}_i, \mathbf{e}_j \right\rangle = \|\mathbf{z}_k\|_q \delta_{kj},$$

so dass $\mathbf{M}_0 = \mathbf{M}_1 = \dots = \mathbf{M}_{k-1} = 0$ und $\mathbf{M}_k = \|\mathbf{z}_k\|_q$ gilt.

Es ist daher

$$\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1} = \mathbf{M}_k \mathbf{e}_k = \|\mathbf{z}_k\|_q \mathbf{e}_k.$$

Weiters wird

$$\langle \mathbf{x}_k, \mathbf{e}_k \rangle \mathbf{e}_k = \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1})$$

mit

$$\mathbf{G}_k := \langle \mathbf{x}_k, \mathbf{e}_k \rangle \|\mathbf{z}_k\|_q^{-1}$$

erhalten.

Die Matrizen $\mathbf{G}_k$ werden *Kalman-Gain-Matrizen* genannt.

Zusammen mit der nach der Definition erhaltenen Darstellung von $\hat{\mathbf{x}}_k$ ergibt sich die *Prädiktor-Korrektur-Gleichung*:

$$\hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}).$$

Bemerkung: $\hat{\mathbf{x}}_{k|k}$ ist unter gewissen Anfangsbedingungen eine gute Schätzung von $\mathbf{x}_k$. Dies wird folgendermaßen gezeigt:

Ausgehend von

$$\begin{aligned} & \mathbf{x}_k - \hat{\mathbf{x}}_{k|k} \\ = & \mathbf{A}_{k-1} \mathbf{x}_{k-1} + \Gamma_{k-1} \xi_{k-1} - \mathbf{A}_{k-1} \hat{\mathbf{x}}_{k-1|k-1} - \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \mathbf{A}_{k-1} \hat{\mathbf{x}}_{k-1|k-1}), \end{aligned}$$

erhält man unter Verwendung von

$$\mathbf{v}_k = \mathbf{C}_k \mathbf{x}_k + \eta_k = \mathbf{C}_k \mathbf{A}_{k-1} \mathbf{x}_{k-1} + \mathbf{C}_k \Gamma_{k-1} \xi_{k-1} + \eta_k$$

die Gleichung

$$\begin{aligned} & \mathbf{x}_k - \hat{\mathbf{x}}_{k|k} \\ = & (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} (\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}) + (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \Gamma_{k-1} \xi_{k-1} - \mathbf{G}_k \eta_k. \end{aligned}$$

Aus der Voraussetzung, dass die Folgen $\{\xi_k\}$ und $\{\eta_k\}$ (zur Erinnerung sei erwähnt, dass diese Folgen das Rauschen beschreiben) mittelwertfrei sind, folgt

$$E(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}) = (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} E(\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}),$$

so dass

$$E(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}) = (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} \cdots (\mathbf{I} - \mathbf{G}_1 \mathbf{C}_1) \mathbf{A}_0 E(\mathbf{x}_0 - \hat{\mathbf{x}}_{0|0})$$

gilt.

Wird nun

$$\hat{\mathbf{x}}_{0|0} = E(\mathbf{x}_0)$$

gesetzt, gilt $E(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}) = 0$ oder $E(\hat{\mathbf{x}}_{k|k}) = E(\mathbf{x}_k)$ für alle k . $\square$

Nächstes Ziel ist es eine rekursive Formel für $\mathbf{G}_k$ zu finden:

Unter Verwendung von (*3.3) und den Eigenschaften

$$\langle \mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}, \hat{\mathbf{x}}_{k-1|k-1} \rangle = O_{n \times n}$$

bzw.

$$\begin{aligned} \langle \mathbf{x}_k, \xi_k \rangle &= O_{n \times n}, & \langle \hat{\mathbf{x}}_{k|k}, \xi_j \rangle &= O_{n \times n} \\ \langle \mathbf{x}_k, \eta_j \rangle &= O_{n \times q}, & \langle \hat{\mathbf{x}}_{k-1|k-1}, \eta_k \rangle &= O_{n \times q}, \end{aligned}$$

für $j = 0, \dots, k$, folgt

$$\begin{aligned} 0 &= \langle \mathbf{x}_k - \hat{\mathbf{x}}_{k|k}, \mathbf{v}_k \rangle \\ &= \langle (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} (\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}) + (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \Gamma_{k-1} \xi_{k-1} - \mathbf{G}_k \eta_k, \\ &\quad \mathbf{C}_k \mathbf{A}_{k-1} ((\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}) + \hat{\mathbf{x}}_{k-1|k-1}) + \mathbf{C}_k \Gamma_{k-1} \xi_{k-1} + \eta_k \rangle \\ &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} \|\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}\|_n^2 \mathbf{A}_{k-1}^T \mathbf{C}_k^T \\ &\quad + (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T \mathbf{C}_k^T - \mathbf{G}_k \mathbf{R}_k. \quad (*3.4) \end{aligned}$$

Definition: Es ist

$$\mathbf{P}_{k,k} = \|\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}\|_n^2$$

und

$$\mathbf{P}_{k,k-1} = \|\mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1}\|_n^2.$$

Es gilt dann:

$$\begin{aligned} \mathbf{P}_{k,k-1} &= \|\mathbf{A}_{k-1} \mathbf{x}_{k-1} + \Gamma_{k-1} \xi_{k-1} - \mathbf{A}_{k-1} \hat{\mathbf{x}}_{k-1|k-1}\|_n^2 \\ &= \mathbf{A}_{k-1} \|\mathbf{x}_{k-1} - \hat{\mathbf{x}}_{k-1|k-1}\|_n^2 \mathbf{A}_{k-1}^T + \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T \end{aligned}$$

oder

$$\mathbf{P}_{k,k-1} = \mathbf{A}_{k-1} \mathbf{P}_{k-1,k-1} \mathbf{A}_{k-1}^T + \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T. \quad (*3.5)$$

Andererseits wird aus (*3.4)

$$\begin{aligned} & (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{A}_{k-1} \mathbf{P}_{k-1,k-1} \mathbf{A}_{k-1}^T \mathbf{C}_k^T \\ & + (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T \mathbf{C}_k^T - \mathbf{G}_k \mathbf{R}_k = 0 \end{aligned}$$

erhalten.

Umformen ergibt

$$\begin{aligned} & \mathbf{G}_k [\mathbf{R}_k + \mathbf{C}_k (\mathbf{A}_{k-1} \mathbf{P}_{k-1,k-1} \mathbf{A}_{k-1}^T + \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T) \mathbf{C}_k^T] \\ & = [\mathbf{A}_{k-1} \mathbf{P}_{k-1,k-1} \mathbf{A}_{k-1}^T + \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T] \mathbf{C}_k^T \\ & = \mathbf{P}_{k,k-1} \mathbf{C}_k^T. \end{aligned}$$

Für $\mathbf{G}_k$ wird also folgendes Ergebnis erhalten:

$$\check{\mathbf{G}}_k = \mathbf{P}_{k,k-1} \mathbf{C}_k^T (\mathbf{R}_k + \mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T)^{-1}, \quad (*3.6)$$

wobei $\mathbf{R}_k$ positiv definit und $\mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T$ nicht-negativ definit ist. Es kann gezeigt werden, dass deren Summe positiv definit ist.

Nächstes Ziel ist es nun, $\mathbf{P}_{k,k}$ durch $\mathbf{P}_{k,k-1}$ auszudrücken und zusammen mit (*3.5) eine rekursive Darstellung zu erhalten.

Dies wird folgendermaßen erreicht:

$$\begin{aligned} \mathbf{P}_{k,k} &= \| \mathbf{x}_k - \hat{\mathbf{x}}_{k|k} \|^2_n \\ &= \| \mathbf{x}_k - (\hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1})) \|^2_n \\ &= \| \mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1} - \mathbf{G}_k (\mathbf{C}_k \mathbf{x}_k + \eta_k) + \mathbf{G}_k \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1} \|^2_n \\ &= \| (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) (\mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1}) - \mathbf{G}_k \eta_k \|^2_n \\ &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \| \mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1} \|^2_n (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k)^T + \mathbf{G}_k \mathbf{R}_k \mathbf{G}_k^T \\ &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1} (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k)^T + \mathbf{G}_k \mathbf{R}_k \mathbf{G}_k^T. \end{aligned}$$

Wegen (*3.6) gilt aber

$$\begin{aligned} & (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1} (\mathbf{G}_k \mathbf{C}_k)^T \\ &= \mathbf{P}_{k,k-1} \mathbf{C}_k^T \mathbf{G}_k^T - \mathbf{G}_k \mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T \mathbf{G}_k^T \\ &= \mathbf{G}_k \mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T \mathbf{G}_k^T + \mathbf{G}_k \mathbf{R}_k \mathbf{G}_k^T - \mathbf{G}_k \mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T \mathbf{G}_k^T \\ &= \mathbf{G}_k \mathbf{R}_k \mathbf{G}_k^T \end{aligned}$$

und es folgt

$$\begin{aligned}\mathbf{P}_{k,k} &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1} (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k)^T + (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1} (\mathbf{G}_k \mathbf{C}_k)^T \\ &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1}.\end{aligned}$$

Aus den in diesem Abschnitt gezeigten Ergebnissen erhält man nun die *Kalman-Filter-Gleichungen*:

$$\begin{aligned}\mathbf{P}_{0,0} &= \|\mathbf{x}_0 - \hat{\mathbf{x}}_{0|0}\|_n^2 = \text{Var}(\mathbf{x}_0) \\ \mathbf{P}_{k,k-1} &= \mathbf{A}_{k-1} \mathbf{P}_{k-1,k-1} \mathbf{A}_{k-1}^T + \Gamma_{k-1} \mathbf{Q}_{k-1} \Gamma_{k-1}^T \\ \mathbf{G}_k &= \mathbf{P}_{k,k-1} \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T + \mathbf{R}_k)^{-1} \\ \mathbf{P}_{k,k} &= (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1} \\ \hat{\mathbf{x}}_{0|0} &= E(\mathbf{x}_0) \\ \hat{\mathbf{x}}_{k|k-1} &= \mathbf{A}_{k-1} \hat{\mathbf{x}}_{k-1|k-1} \\ \hat{\mathbf{x}}_{k|k} &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}) \\ k &= 1, 2, \dots\end{aligned}$$

3.3 Das Erweiterte-Kalman-Filter

3.3.1 Die Funktionsweise des Erweiterten-Kalman-Filters

In diesem Abschnitt wird nun ein nichtlineares Modell betrachtet:

$$\begin{aligned}\mathbf{x}_{k+1} &= f_k(\mathbf{x}_k) + H_k(\mathbf{x}_k) \xi_k \\ \mathbf{v}_k &= g_k(\mathbf{x}_k) + \eta_k,\end{aligned}$$

wobei f_k und g_k vektorwertige Funktionen mit Bildbereichen in $\mathbb{R}^n$ bzw. $\mathbb{R}^q$ für $1 \leq q \leq n$ sind und H_k eine matrixwertige Funktion mit Bildbereich in $\mathbb{R}^n \times \mathbb{R}^q$ ist. Es sind für jedes k die ersten partiellen Ableitungen, nach allen Komponenten von $\mathbf{x}_k$, von $f_k(\mathbf{x}_k)$ und $g_k(\mathbf{x}_k)$ stetig.

Angenommen, es handelt sich bei $\{\xi_k\}$ und $\{\eta_k\}$ um mittelwertfreie Gaussian white noise Folgen mit Bildbereichen in $\mathbb{R}^p$ bzw. $\mathbb{R}^q$ für $1 \leq p, q \leq n$. Weiters sei:

$$\begin{aligned}E[\xi_k \xi_l^T] &= \mathbf{Q}_k \delta_{kl}, & E[\eta_k \eta_l^T] &= \mathbf{R}_k \delta_{kl}, \\ E[\xi_k \eta_l^T] &= 0, & E[\xi_k \mathbf{x}_0^T] &= 0, & E[\eta_k \mathbf{x}_0^T] &= 0,\end{aligned}$$

für alle k und l .

Wie in der Vorgehensweise des linearen Modells, sind die anfängliche Schätzung $\hat{\mathbf{x}}_0 = \hat{\mathbf{x}}_{0|0}$ und $\hat{\mathbf{x}}_{1|0}$ folgendermaßen gewählt:

$$\hat{\mathbf{x}}_0 = E(\mathbf{x}_0), \quad \hat{\mathbf{x}}_{1|0} = f_0(\hat{\mathbf{x}}_0).$$

Nun werden die $\hat{\mathbf{x}}_k = \hat{\mathbf{x}}_{k|k}$, für $k = 1, 2, \dots$, mit Hilfe von

$$\hat{\mathbf{x}}_{k+1|k} = f_k(\hat{\mathbf{x}}_k) \quad (*3.7)$$

und

$$\begin{aligned} \mathbf{x}_{k+1} &= \mathbf{A}_k \mathbf{x}_k + \mathbf{u}_k + \Gamma_k \xi_k \\ \mathbf{w}_k &= \mathbf{C}_k \mathbf{x}_k + \eta_k, \end{aligned}$$

fortlaufend bestimmt.

$\mathbf{A}_k, \mathbf{u}_k, \Gamma_k, \mathbf{w}_k$ und $\mathbf{C}_k$ werden wie folgt erhalten:

Vorausgesetzt $\hat{\mathbf{x}}_j$ ist so bestimmt, dass $\hat{\mathbf{x}}_{j+1|j}$, für $j = 0, 1, \dots, k$, so wie in (*3.7) definiert wird. Die Verwendung der linearen Taylor Approximation von $f_k(\mathbf{x}_k)$ auf $\hat{\mathbf{x}}_k$ und von $g_k(\mathbf{x}_k)$ auf $\hat{\mathbf{x}}_{k|k-1}$ ergibt dann

$$\begin{aligned} f_k(\mathbf{x}_k) &\simeq f_k(\hat{\mathbf{x}}_k) + \mathbf{A}_k(\mathbf{x}_k - \hat{\mathbf{x}}_k) \\ g_k(\mathbf{x}_k) &\simeq g_k(\hat{\mathbf{x}}_{k|k-1}) + \mathbf{C}_k(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1}), \end{aligned}$$

wobei

$$\mathbf{A}_k = \left[\frac{\partial f_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_k) \right] \quad \text{und} \quad \mathbf{C}_k = \left[\frac{\partial g_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) \right]$$

ist.

Wird nun

$$\begin{aligned} \mathbf{u}_k &= f_k(\hat{\mathbf{x}}_k) - \mathbf{A}_k \hat{\mathbf{x}}_k \\ \Gamma_k &= \mathbf{H}_k(\hat{\mathbf{x}}_k) \\ \mathbf{w}_k &= \mathbf{v}_k - g_k(\hat{\mathbf{x}}_{k|k-1}) + \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1} \end{aligned}$$

gesetzt, kann das nichtlineare Modell durch das lineare approximiert werden. Diese Linearisierung ist aber nur möglich, nachdem $\hat{\mathbf{x}}_k$ bestimmt wurde. Durch die Kenntnis von $\hat{\mathbf{x}}_0$ ist die Systemgleichung für $k = 0$ gültig. Nun wird $\hat{\mathbf{x}}_1 = \hat{\mathbf{x}}_{1|1}$ als die Schätzung von $\mathbf{x}_1$ mit optimalen Gewichten im linearen Modell definiert. Aus den Resultaten des vorhergehenden Abschnitts (vgl. mit den Kalman-Filter-Gleichungen S.38) wird

$$\begin{aligned} \hat{\mathbf{x}}_k &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k(\mathbf{w}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}) \\ &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k((\mathbf{v}_k - g_k(\hat{\mathbf{x}}_{k|k-1}) + \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}) - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}) \\ &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k(\mathbf{v}_k - g_k(\hat{\mathbf{x}}_{k|k-1})) \end{aligned}$$

erhalten, wobei es sich bei $\mathbf{G}_k$ um die Kalman-Gain-Matrix für das lineare Modell im k -ten Moment handelt.

Der resultierende Filterungsprozess wird *Erweiterter-Kalman-Filter* genannt.

Der Algorithmus kann folgendermaßen zusammengefasst werden:

$$\begin{aligned}
\mathbf{P}_{0,0} &= \text{Var}(\mathbf{x}_0), \quad \hat{\mathbf{x}}_0 = E(\mathbf{x}_0). \\
\text{Für } k &= 1, 2, \dots, \\
\mathbf{P}_{k,k-1} &= \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right] \mathbf{P}_{k-1,k-1} \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right]^T \\
&\quad + \mathbf{H}_{k-1}(\hat{\mathbf{x}}_{k-1}) \mathbf{Q}_{k-1} \mathbf{H}_{k-1}^T(\hat{\mathbf{x}}_{k-1}) \\
\hat{\mathbf{x}}_{k|k-1} &= f_{k-1}(\hat{\mathbf{x}}_{k-1}) \\
\mathbf{G}_k &= \mathbf{P}_{k,k-1} \left[\frac{\partial g_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) \right]^T \\
&\quad \cdot \left[\left[\frac{\partial g_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) \right] \mathbf{P}_{k,k-1} \left[\frac{\partial g_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) \right]^T + \mathbf{R}_k \right]^{-1} \\
\mathbf{P}_{k,k} &= \left[\mathbf{I} - \mathbf{G}_k \left[\frac{\partial g_k}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) \right] \right] \mathbf{P}_{k,k-1} \\
\hat{\mathbf{x}}_{k|k} &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k(\mathbf{v}_k - g_k(\hat{\mathbf{x}}_{k|k-1})).
\end{aligned}$$

3.4 Das Ensemble-Kalman-Filter

3.4.1 Die Ensemble-Kovarianzmatrizen

Die Fehlerkovarianzmatrizen der Schätzwerte, $\mathbf{P}_{k,k-1}$ und $\mathbf{P}_{k,k}$, sind für Kalman-Filter unter Verwendung von Erwartungswerten definiert. Das heißt:

$$\mathbf{P}_{k,k-1} = E[(\hat{\mathbf{x}}_{k|k-1} - \mathbf{x}_k)(\hat{\mathbf{x}}_{k|k-1} - \mathbf{x}_k)^T], \quad (*3.8)$$

$$\mathbf{P}_{k,k} = E[(\hat{\mathbf{x}}_{k|k} - \mathbf{x}_k)(\hat{\mathbf{x}}_{k|k} - \mathbf{x}_k)^T]. \quad (*3.9)$$

(vgl. [4] S.26 und 27).

Im Gegensatz dazu, ist es beim Ensemble-Kalman-Filter praktischer sogenannte *Ensemble-Kovarianzmatrizen* rund um den Mittelwert des Ensembles zu verwenden:

$$\mathbf{P}_{k,k-1} \cong \mathbf{P}_{k,k-1}^e = \overline{(\hat{\mathbf{x}}_{k|k-1} - \hat{\mathbf{x}}_{k|k-1})(\hat{\mathbf{x}}_{k|k-1} - \hat{\mathbf{x}}_{k|k-1})^T}, \quad (*3.10)$$

$$\mathbf{P}_{k,k} \cong \mathbf{P}_{k,k}^e = \overline{(\hat{\mathbf{x}}_{k|k} - \hat{\mathbf{x}}_{k|k})(\hat{\mathbf{x}}_{k|k} - \hat{\mathbf{x}}_{k|k})^T}. \quad (*3.11)$$

Im Vergleich zu den oberen Gleichungen wird der Erwartungswert nun durch den Durchschnitt des Ensembles, der durch einen Balken dargestellt wird, ersetzt.

Der Mittelwert des Ensembles kann demzufolge als beste Schätzung interpretiert werden. Dadurch kann auch die Streuung des Ensembles rund um den Mittelwert durchwegs als die Fehlerkovarianz der besten Schätzung aufgefasst werden.

3.4.2 Die Funktionsweise des Ensemble-Kalman-Filters

Auf Basis der Ensemble-Kovarianzen, die durch die Gleichungen (*3.10) und (*3.11) definiert sind, wird nun das *Ensemble-Kalman-Filter* beschrieben.

Es wird zunächst ein Ensemble mit N Beobachtungen erzeugt, aus dem ein zufälliger Wert gezogen wird:

$$\mathbf{v}_{i,k} = \mathbf{v}_k + \eta_{i,k}, \quad i = 1, \dots, N.$$

Es handelt sich bei $\mathbf{v}_k$ um die aktuelle Beobachtung und bei der Zufallsvariable $\eta_{i,k}$ um den Beobachtungsfehler, der als vordefinierte Beobachtungsvarianz $\mathbf{R}_k$ hat. $\mathbf{v}_{i,k}$ ist eine Zufallsvariable deren Erwartungswert aufgrund von [1], mit dem von $\mathbf{v}_k$ identisch sein soll.

Es werden nun aus den N Beobachtungsfehlern diejenigen aussortiert, deren Mittelwert ungleich 0 ist. Durch die so erreichte Mittelwertfreiheit der stochastischen Störungen werden eventuelle spätere Ungenauigkeiten vermieden.

Definition: Es sei die *Ensemble-Fehlerkovarianzmatrix* definiert durch

$$\mathbf{R}_k^e = \overline{\eta\eta^T}.$$

Bemerkung: Für unbegrenzte Ensemblegröße konvergiert die eben definierte Fehlerkovarianzmatrix gegen die Fehlerkovarianzmatrix $\mathbf{R}_k$ wie sie im "Standard"-Kalman-Filter verwendet wurde.

Die nachfolgenden Resultate sind sowohl unter Verwendung der Fehlerkovarianzmatrix $\mathbf{R}_k$, als auch unter Verwendung der Ensemble-Kovarianzmatrix $\mathbf{R}_k^e$, gültig. Letztere ist immer wieder für Anwendungen des Filters praktisch.

Der Funktionsweise des Ensemble-Kalman-Filters basiert auf Korrekturen jedes einzelnen Ensemblemitglieds, die durch

$$\hat{\mathbf{x}}_{k|k}^i = \hat{\mathbf{x}}_{k|k-1}^i + \mathbf{G}_k^e(\mathbf{v}_{i,k} - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}^i) \quad (*3.12)$$

gegeben sind.

Die verwendete Matrix $\mathbf{G}_k^e$ ist durch

$$\mathbf{G}_k^e = \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^{-1} \quad (*3.13)$$

definiert und identisch mit der Kalman-Gain-Matrix.

Bemerkung: Für endliche Ensemblegröße wird mit der Verwendung der Ensemble-Kovarianzen eine Approximation der wahren Kovarianzen eingeführt.

Lemma: Aus Gleichung (*3.12) folgt

$$\overline{\hat{\mathbf{x}}_{k|k}} = \overline{\hat{\mathbf{x}}_{k|k-1}} + \mathbf{G}_k^e (\overline{\mathbf{v}_k} - \mathbf{C}_k \overline{\hat{\mathbf{x}}_{k|k-1}}), \quad (*3.14)$$

wobei wegen der Mittelwertfreiheit der Störungen nun $\overline{\mathbf{v}_k} = \mathbf{v}_k$ gilt.

Abgesehen von der Verwendung von $\mathbf{G}_k^e$, $\mathbf{P}_{k,k-1}^e$, $\mathbf{P}_{k,k}^e$ und $\mathbf{R}_k^e$ anstatt $\mathbf{G}_k$, $\mathbf{P}_{k,k-1}$, $\mathbf{P}_{k,k}$ und $\mathbf{R}_k$, verhält sich $\overline{\hat{\mathbf{x}}_{k|k}}$ zu $\overline{\hat{\mathbf{x}}_{k|k-1}}$ im Ensemble-Kalman-Filter gleich wie $\hat{\mathbf{x}}_{k|k}$ zu $\hat{\mathbf{x}}_{k|k-1}$ im "Standard"-Kalman-Filter.

Es ist nun gezeigt, dass nach updaten der einzelnen Ensemblemitglieder mit Beobachtungsfehlern, ein neues Ensemble mit "korrekten" Störungen erzeugt wurde.

Aus den Gleichungen (*3.12) und (*3.14) wird

$$\hat{\mathbf{x}}_{k|k}^i - \overline{\hat{\mathbf{x}}_{k|k}} = (\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k) (\hat{\mathbf{x}}_{k|k-1}^i - \overline{\hat{\mathbf{x}}_{k|k-1}}) + \mathbf{G}_k^e (\mathbf{v}_{i,k} - \overline{\mathbf{v}_k}),$$

erhalten.

Es kann nun folgendes Resultat hergeleitet werden:

$$\begin{aligned} \mathbf{P}_{k,k}^e &= \overline{(\hat{\mathbf{x}}_{k|k} - \overline{\hat{\mathbf{x}}_{k|k}})(\hat{\mathbf{x}}_{k|k} - \overline{\hat{\mathbf{x}}_{k|k}})^T} \\ &= \overline{((\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k)(\hat{\mathbf{x}}_{k|k-1}^i - \overline{\hat{\mathbf{x}}_{k|k-1}}) + \mathbf{G}_k^e (\mathbf{v}_{i,k} - \overline{\mathbf{v}_k}))} \\ &\quad \overline{((\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k)(\hat{\mathbf{x}}_{k|k-1}^i - \overline{\hat{\mathbf{x}}_{k|k-1}}) + \mathbf{G}_k^e (\mathbf{v}_{i,k} - \overline{\mathbf{v}_k}))^T} \\ &= (\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k) \overline{(\hat{\mathbf{x}}_{k|k-1}^i - \overline{\hat{\mathbf{x}}_{k|k-1}})(\hat{\mathbf{x}}_{k|k-1}^i - \overline{\hat{\mathbf{x}}_{k|k-1}})^T} (\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k)^T \\ &\quad + \mathbf{G}_k^e \overline{(\mathbf{v}_{i,k} - \overline{\mathbf{v}_k})(\mathbf{v}_{i,k} - \overline{\mathbf{v}_k})^T} (\mathbf{G}_k^e)^T \\ &= (\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k) \mathbf{P}_{k,k-1}^e (\mathbf{I} - \mathbf{C}_k^T (\mathbf{G}_k^e)^T) + \mathbf{G}_k^e \mathbf{R}_k^e (\mathbf{G}_k^e)^T \\ &= \mathbf{P}_{k,k-1}^e - \mathbf{G}_k^e \mathbf{C}_k \mathbf{P}_{k,k-1}^e - \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T (\mathbf{G}_k^e)^T \\ &\quad + \mathbf{G}_k^e (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e) (\mathbf{G}_k^e)^T \\ &= (\mathbf{I} - \mathbf{G}_k^e \mathbf{C}_k) \mathbf{P}_{k,k-1}^e. \quad (*3.15) \end{aligned}$$

Bei (*3.15) handelt es sich um das traditionelle Resultat der "Standard"-Kalman-Filter-Theorie für die minimale Fehlerkovarianz. Demnach impliziert letztere Herleitung, dass sich aus dem Ensemble-Kalman-Filter für unendliche Ensemblegröße im Grenzwert das selbe Resultat wie aus dem "Standard"-Kalman-Filter ergibt.

Die Herleitung zeigt auch, dass es notwendig ist die Beobachtungen $\mathbf{v}_k$ als Zufallsvariablen zu wählen, um die obige Darstellung der Fehlerkovarianzmatrix

$$\mathbf{R}_k^e = \overline{\eta\eta^T} = \overline{(\mathbf{v}_{i,k} - \bar{\mathbf{v}}_k)(\mathbf{v}_{i,k} - \bar{\mathbf{v}}_k)^T}$$

zu erhalten.

In (*3.15) kann eine Fehlerkovarianzmatrix von vollem Rang verwendet werden, doch der Ersatz der Kovarianzmatrix durch ein Ensemble führt zu einer Auslöschung in der vorletzten Zeile von (*3.15). Es ergibt sich also:

$$\begin{aligned} & \mathbf{G}_k^e (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e) (\mathbf{G}_k^e)^T \\ = & \mathbf{G}_k^e (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e) (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^{-1} \mathbf{C}_k \mathbf{P}_{k,k-1}^e \\ = & \mathbf{G}_k^e \mathbf{C}_k \mathbf{P}_{k,k-1}^e. \quad (*3.16) \end{aligned}$$

Demnach wird die Verwendung einer Fehlerkovarianzmatrix von niedrigerem Rang für vorteilhaft erklärt. Die Gründe für die restlichen kleineren Ungenauigkeiten werden in [9] beschrieben. Für endliche Ensemblegröße werden zum Beispiel durch die in obiger Bemerkung eingeführten Approximation, Fehler verursacht.

In der soeben beschriebenen Herleitung wird die Existenz der Inversen im Kalman-Gain angenommen. Die Herleitung gilt jedoch auch dann, wenn die Matrix in der Inversion von niedrigerem Rang ist.

Die Inverse in (*3.13) kann dann durch die Pseudoinverse ersetzt werden. Das Kalman-Gain kann dann wie folgt geschrieben werden:

$$\mathbf{G}_k^e = \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^+. \quad (*3.17)$$

Ist die Matrix in der Inversion von vollem Rang, so ist (*3.17) identisch mit (*3.13). Nach Verwendung von (*3.17) wird die obige Herleitung (*3.16) zu

$$\begin{aligned} & \mathbf{G}_k^e (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e) (\mathbf{G}_k^e)^T \\ = & \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^+ (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e) \\ & (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^+ \mathbf{C}_k \mathbf{P}_{k,k-1}^e \\ = & \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1}^e \mathbf{C}_k^T + \mathbf{R}_k^e)^+ \mathbf{C}_k \mathbf{P}_{k,k-1}^e \\ = & \mathbf{G}_k^e \mathbf{C}_k \mathbf{P}_{k,k-1}^e, \end{aligned}$$

wobei die Eigenschaft $\mathbf{Y}^+ = \mathbf{Y}^+ \mathbf{Y} \mathbf{Y}^+$ der Pseudoinversen verwendet wird.

Kapitel 4

Anwendungen in der hydrometeorologischen Datenverarbeitung

4.1 Überblick

Abschnitt 4.2 beschäftigt sich mit einer hydrometeorologischen Anwendung der Wiener-Filter-Theorie. Wie in [27] und [26] beschrieben, wird ein Wiener-Filter zur Vorhersage von Wasserabflussmengen von Flüssen verwendet. Diese Vorhersagen werden zur Hochwasservorhersage eingesetzt. Unter Berücksichtigung von [6] wird der Ablauf eines Vorhersagemodells kurz beschrieben.

In 4.3 wird ein Wasserabflussmodell mit linearem Wasserabfluss vorausgesetzt. In diesem Modell sind die Bedingungen für den Einsatz eines "Standard"-Kalman-Filters erfüllt. Dieser wird in 4.3.1 unter Zuhilfenahme von [15] und [4] beschrieben. 4.3.2 zeigt wie dasselbe Problem durch die Verwendung der Ensemble-Kalman-Filter-Theorie behandelt wird. Ebenso werden in diesem Unterabschnitt die Filterperformances des "Standard"-Filters und des Ensemble-Filters für dieses spezielle Beispiel direkt miteinander verglichen [15].

In Abschnitt 4.4 wird die Verwendung der Kalman-Filter-Theorie bei der Schätzung des Bodenwassergehalts gezeigt. In dem in [22] und [21] durchgeführten Versuch werden das Erweiterte- und das Ensemble-Kalman-Filter zur Assimilation von Messdaten verwendet. In 4.3.1 wird das im Gegensatz zu Abschnitt 4.2 nun nichtlineare Modell beschrieben. Unterabschnitt 4.3.2 enthält die aus dem Experiment gewonnenen Ergebnisse für das Erweiterte- bzw. das Ensemble-Filter. Auch die Filterperformances für dieses Beispiel werden wieder in einen direkten Vergleich gestellt.

4.2 Vorhersage von Wasserabflussmengen von Flüssen mit Hilfe der Wiener-Filter-Theorie

Kurzzeitprognosen von Wasserabflussmengen von Flüssen spielen eine wichtige Rolle sowohl für den Ablauf von Wasserwirtschaftssystemen, als auch für Hochwasserwarnsysteme. In [27] und [26] wird die Verwendung der Wiener-Filter-Theorie in der Analyse eines hydrologischen Systems behandelt. Mit Hilfe eines Wiener-Mehrkanal-Filters werden Wasserstände an Messstationen, in diesem speziellen Fall entlang des Rheins, vorhergesagt. Das deutsche hydrologische Bundesinstitut in Koblenz berechnet mit Hilfe der Filtertheorie Wasserstandsvorhersagen für 16 Messstellen entlang des Flusses. Die einzelnen Wasserstände werden alle 6 Stunden an allen Pegelstationen, also nicht nur an den zuvor erwähnten Vorhersagepegeln, gemessen. Das verwendete Modell benötigt als Input die jeweiligen 6-Stunden-Differenzen der gemessenen Pegelstände bzw. die Wasserabflusswerte an den Messstellen des Wassereinzugsgebiets des jeweiligen Vorhersagepegels.

An den insgesamt 16 Vorhersagestellen des Rheins werden die Wasserstände in 6, 12, 18, 24, 30 und 36 Stunden vorhergesagt. Die Messwerte der einzelnen Pegelstationen werden täglich um 5.00, 11.00, 17.00 und 23.00 Uhr abgelesen. Die Vorhersagen sind eine Stunde nach der Abfrage der Pegelstände bekannt. Der Ablauf des Vorhersagemodells kann, in groben Zügen, in folgenden Schritten beschrieben werden [6]:

1. Ablesen der Wasserstände aller Pegel.
2. Dateneingabe
3. Berechnung der Abflusswerte der letzten 6 Stunden
4. Berechnung der 6-Stunden-Differenzen der Wasserstände und der Abflusswerte der einzelnen Messstellen
5. Für jeden einzelnen Pegel wird nun eine Impulsantwortfunktion angewendet, die jeweils unter Berücksichtigung der 6-Stunden-Differenzen der letzten 3 Tage gewichtet wird
6. Berechnung der Wasserstandsdifferenzen bzw. der Differenzen der Abflussmengen der einzelnen Pegel in Bezug auf den Vorhersagepegel
7. Berechnung der Gesamtsumme der Wasserstandsdifferenzen bzw. der Differenzen der Abflussmengen

8. Addition der Gesamtdifferenzen zum aktuellen Wasserstand bzw. Abflusswert
9. Vorhersage des Wasserstandes für ein Zeitmessungsintervall (in diesem Fall 6 Stunden)
 - Um Wasserstandsvorhersagen für das nächste Zeitmessungsintervall treffen zu können, muss mit dem soeben vorhergesagten Wert bei Schritt 5 neu begonnen werden.

In [6] werden die für 6, 12, 18 und 24 Stunden vorhergesagten Wasserstände mit den im Nachhinein gemessenen Werten verglichen. Dieser Vergleich bestätigt die hohe Zuverlässigkeit des Vorhersagemodells. Zeitsegmente mit auffällig großen Abweichungen zwischen Vorhersage und Messung können durch abrupten Wechsel des Wasserstandes erklärt werden. Die maximalen Differenzen bewegen sich zwischen ca. 10 cm (bei der 6-Stunden-Vorhersage) und 20 cm (bei der 24-Stunden-Vorhersage).

Da dieses Anwendungsbeispiel wie es in [27], [26] und [6] beschrieben wird, eher für einen Hydrologen als für einen Mathematiker interessant erscheint, wird nun nicht mehr weiter darauf eingegangen.

4.3 Ein lineares Wasserabflussmodell

4.3.1 Die Verwendung des "Standard"-Kalman-Filters

Da die für Hochwasserprognosen notwendigen Messungen von Wasserständen und deren Umrechnung in Durchflüsse immer wieder Ungenauigkeiten beinhalten, wird das Kalman-Filter zur optimalen Zustandsschätzung eingesetzt. Ziel ist es eine möglichst genaue Vorhersage zu treffen. Durch die Verbindung der Modellergebnisse mit den Messdaten unter Abwägung der jeweils eingetragenen Unsicherheiten wird der Prognosefehler möglichst klein gehalten. In dem unten angeführten Beispiel kann, da mit der Linearität des verwendeten Modells alle Voraussetzungen erfüllt sind, das "Standard"-Kalman-Filter verwendet werden. Durch die Kombination unsicherer Größen soll die Gesamtunsicherheit reduziert und anschließend Datenassimilation durchgeführt werden [15].

Den Modelleingangsgrößen aus Abschnitt 3.2 werden nun gemessene Größen zugeordnet. Es handelt sich dabei beispielsweise um Messwerte von Niederschlag, Verdunstung, usw.. $\mathbf{x}_k$ beschreibt die Schätzwerte des Abflusses zum

Zeitpunkt k . Es wird angenommen, dass die Unsicherheiten aus den Modelleingangsgößen stammen.

In der Beobachtungsgleichung

$$\mathbf{v}_k = \mathbf{C}_k \mathbf{x}_k + \eta_k,$$

handelt es sich nun bei $\mathbf{v}_k$ um Durchflussmengen. Zur Erinnerung sei bemerkt, die Matrix $\mathbf{C}_k$ beschreibt Beziehung zwischen Beobachtungen und Zustandsgrößen des Modells und bei η_k handelt es sich um einen zufälligen Fehler mit Beobachtungsvarianz $\mathbf{R}_k$, der klarerweise die Unsicherheit der Beobachtung darstellt.

Im Kalman-Filter wird nun der aktualisierte Schätzwert des Abflusses $\hat{\mathbf{x}}_{k|k}$ als Linearkombination des Schätzwertes der Abflussvorhersage $\hat{\mathbf{x}}_{k|k-1}$ und der Abflussmessung $\mathbf{v}_k$ dargestellt:

$$\hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \mathbf{G}_k (\mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}),$$

wobei es sich bei

$$\mathbf{G}_k = \mathbf{P}_{k,k-1} \mathbf{C}_k^T (\mathbf{C}_k \mathbf{P}_{k,k-1} \mathbf{C}_k^T + \mathbf{R}_k)^{-1}$$

wieder um das Kalman-Gain handelt. $\mathbf{P}_{k,k-1}$ bezeichnet wie schon in 3.2.3 die a priori Schätzvarianz.

Die aktualisierte a posteriori Schätzvarianz ergibt sich, wie schon im theoretischen Teil der Arbeit über das "Standard"-Kalman-Filter, zu

$$\mathbf{P}_{k,k} = (\mathbf{I} - \mathbf{G}_k \mathbf{C}_k) \mathbf{P}_{k,k-1},$$

(vgl. mit den Kalman-Filter-Gleichungen S.38 und [4] S.42).

Neben [4] diene [15] als Quelle für diesen Abschnitt.

4.3.2 Ein Vergleich zwischen der Verwendung des "Standard"-Kalman-Filters und des Ensemble-Kalman-Filters

Im Gegensatz zum "Standard"-Kalman-Filter werden, wie schon in Abschnitt 3.4 beschrieben, beim Ensemble-Kalman-Filter die statistischen Eigenschaften des Schätzwertes des Abflusses durch ein Ensemble von möglichen Schätzwerten dargestellt. Dieses Ensemble von Abflussschätzwerten wird mit einem Ensemble von Modelleingangsgößen ermittelt. Die nun erhaltenen Schätzwerte haben somit auch die passenden Schätzvarianzen.

Die deduktive Schätzung des Abflusses ist nach 3.4 gegeben durch

$$\hat{\mathbf{x}}_{k|k} = \overline{\hat{\mathbf{x}}_{k|k}} = \frac{1}{N} \sum_{i=1}^N \hat{\mathbf{x}}_{k|k}^i$$

(vgl. hierzu die Bemerkung nach (*_{3.11}) über endliche Ensemblegröße).

Wie schon im theoretischen Teil der Arbeit über den Ablauf des Ensemble-Kalman-Filters gezeigt, wird jedes Ensemblemitglied $\hat{\mathbf{x}}_{k|k}^i$ aktualisiert und basierend auf diesem Wert werden die Schätzwerte $\hat{\mathbf{x}}_{k|k-1}$ und die entsprechende Schätzvarianz ermittelt.

In [15] werden die Rechenergebnisse für die Assimilation des Abflusses aus einem einfachen Linearspeicher mittels "Standard"-Kalman-Filter und Ensemble-Kalman-Filter verglichen. Das behandelte Modell ist von der Form

$$\mathbf{x}_{k+1} = \xi \cdot \mathbf{x}_k,$$

wobei $\xi = 0,9$ gilt.

Es wurden bei der Anwendung des Ensemble-Kalman-Filters zur Beschreibung der statistischen Schätzwerte 10 bzw. 100 Ensemblemitglieder gewählt. Für die Modell- und die Beobachtungsvarianz wurde für dieses Beispiel jeweils der Wert $0,5 (m/s)^2$ angenommen. Anhand diskreter Abflussbeobachtungen, die zur Veranschaulichung der Methode nicht für jeden Berechnungsschritt zur Verfügung stehen, wird die Datenassimilation durchgeführt.

Die Schätzung des Anfangswertes der Zustandsgröße, also des Abflusses, beträgt $1 m^3$ mit Fehlervarianz $2,5 (m/s)^2$. Es wird nun Schritt für Schritt das "Standard"- bzw. das Ensemble-Kalman-Filter angewendet.

Bei beiden Algorithmen wächst, durch den Eintrag zusätzlicher Modellunsicherheiten, die Funktion der Schätzvarianz solange an, bis ein beobachteter Abflusswert vorhanden ist. Die Korrektur mittels der Anwendung des "Standard"- bzw. des Ensemble-Kalman-Filters bewirkt dann eine plötzliche Verringerung der Schätzvarianz. Aufgrund der Unsicherheiten steigt die Funktion danach wieder bis zum Moment der nächsten Nachbesserung.

Der zeitliche Verlauf der Schätzvarianz bildet beim "Standard"-Kalman-Filter eine Sägezahnkurve, wobei die Steigung der Funktion zwischen zwei Messwerten von der eingetragenen Modellunsicherheit bei jedem Berechnungsschritt abhängig ist. Die Unsicherheit wird durch die gewählte Modellvarianz bestimmt.

Beim Ensemble-Kalman-Filter spielt klarerweise die Wahl der Ensemblegröße eine wichtige Rolle für die Entwicklung der Schätzvarianzen. Mit einer Größe von 10 Ensemblemitgliedern gibt es deutliche Unterschiede im Vergleich

der Varianzwerte zwischen "Standard"- und Ensemble-Kalman-Filter. Diese verringern sich aber deutlich durch die Erhöhung der Ensemblegröße auf 100 Realisationen. Je höher die Anzahl an Ensemblemitgliedern bei der Ausführung des Algorithmuses, desto besser werden die statistischen Eigenschaften der Schätzvarianz angenähert.

Im Gegensatz zu den Schätzvarianzen weichen die Ergebnisse für die Datenassimilation mit "Standard"-Kalman-Filter und Ensemble-Kalman-Filter beim Verlauf der Zustandsgrößen nur geringfügig voneinander ab. Auch der Unterschied zwischen den Ensemblegrößen (10 bzw. 100 Mitglieder) ist nur geringfügig. Abweichungen in der Schätzvarianz wirken sich dabei auf die Gewichtung zwischen Modellergebnis und Messung aus. Dies geschieht aber ebenfalls nur in einem sehr geringen Maß.

Durch die soeben beschriebene Anwendung wird gezeigt, dass für die Schätzung der Zustandsgrößen mittels Ensemble-Kalman-Filter-Theorie, schon mit einer Ensemblegröße von 10 Realisationen sehr gute Ergebnisse erzielt werden können.

Bemerkung: Abschließend sei nochmals darauf hingewiesen, dass der Vergleich zwischen "Standard"-Kalman-Filter und Ensemble-Kalman-Filter nur aufgrund der Linearität des gewählten Modells möglich ist.

4.4 Schätzung des Bodenwassergehalts

In [22] und [21] werden das Erweiterte- bzw. das Ensemble-Kalman-Filter zur Schätzung des Bodenwassergehalts angewendet. Die Aufgabe der Filter in dem durchgeführten Experiment ist die Assimilation von Daten die durch Messungen des oberflächennahen Wassergehalts gewonnen werden.

4.4.1 Modellbeschreibung

Die Filtertheorie wird nun im Gegensatz zum vorhergehenden Abschnitt auf ein nichtlineares Modell angewendet. Es wird das sogenannte *Catchment Model*, wie es in [16] beschrieben wird, verwendet. Hierbei handelt es sich um ein Modell, das die Bodenoberflächen eines Gebiets und deren hydrologische Eigenschaften beschreibt.

Es werden das Gesamtbodenwasserdefizit, die Bodenwasserüberschüsse der Bodenoberfläche (2 cm Tiefe) und der Wurzelschicht (1 m Tiefe) als Variablen verwendet. Mit Hilfe dieser Variablen werden Bodenwasserwerte diagnostiziert.

Bemerkung: Beim Gesamtbodenwasserdefizit handelt es sich um den Absolutwert des benötigten Wassers, das zur vollständigen Sättigung des Bodens benötigt wird. Die Bodenwasserüberschüsse beschreiben die Abweichungen von einem ausgeglichenen Bodenwasserhaushalt.

Das nichtlineare Modell kann durch Zusammenfassen der Variablen zum Zeitpunkt k in den Zustandsvektor $\mathbf{x}_k$ formuliert werden (bei den Variablen handelt es sich in diesem Fall um die Variablen der Bodenwasserüberschüsse und der Defizite):

$$\mathbf{x}_{k+1} = f_k(\mathbf{x}_k) + \bar{\xi}_k$$

wobei der nichtlineare Operator $f_k(\cdot)$ deterministische Daten wie z.B. den beobachteten Niederschlag beinhaltet.

Unsicherheiten verursacht durch das Aufstellen der Modellgleichung, bzw. Fehler die durch Beobachtung der Daten entstehen, werden durch den Fehlerterm $\bar{\xi}_k$ ausgedrückt.

Durch aneinanderfügen aller Beobachtungsdaten vom Zeitpunkt k zum Vektor $\mathbf{v}_k$ wird die Gleichung

$$\mathbf{v}_k = \bar{\mathbf{C}}_k \mathbf{x}_k + \bar{\eta}_k$$

erhalten. Zur Vereinfachung wird hierbei ein lineares Verhältnis angenommen.

Der Operator $\bar{\mathbf{C}}_k$ beschreibt den Zusammenhang zwischen den Zuständen (in diesem Fall den Variablen des Bodenwasserüberschusses bzw. des Bodenwasserdefizits) und den gemessenen Variablen (also den Variablen des oberflächennahen Wassergehalts). Beim Term η_k handelt es sich wieder um den Fehler.

Es sei wieder angenommen, dass die Zufallsvariablen $\bar{\xi}_k$ bzw. $\bar{\eta}_k$ mittelwertfrei sind und es sich bei deren Kovarianzen um $\mathbf{Q}_k$ bzw. $\mathbf{R}_k$ handelt. Darüber hinaus seien die beiden Fehlervariablen wieder unkorreliert und white.

Weiters sei noch erwähnt, dass es sich bei den im Experiment verwendeten Beobachtungen um sogenannte synthetische Beobachtungen handelt. Diese werden aus "wahren" Messwerten durch Hinzufügen von zufälligen Messfehlern erhalten. Da die Filterperformance aber sehr stark von den Fehlerparametern abhängt, müssen diese sehr gewissenhaft gewählt werden. Um einen fairen Vergleich zwischen dem Erweiterten- und dem Ensemble-Kalman-Filter zu ermöglichen, werden die für den jeweiligen Filter vorteilhaftesten Parameter verwendet.

4.4.2 Versuchsergebnisse

Um die Leistungen der Filter vergleichen zu können, werden die tatsächlichen Fehler betrachtet. Es handelt sich dabei um die Differenzen der Bodenwasserwerte die durch ein Kontrollexperiment erhalten werden und der Schätzwerte des Erweiterten- bzw. Ensemble-Kalman-Filters. Auffällig ist, dass die Fehler für die Bodenoberfläche verhältnismäßig groß sind. Dies ist aber nicht von den Filterperformances verursacht, sondern auf die gewählten Assimilationsabstände von 3 Tagen zurückzuführen. Klarerweise können auf diese Weise die an der Bodenoberfläche nur kurzzeitig auftretenden Variablen nicht zufrieden stellend beachtet werden. Um eine bessere Qualität der Schätzung an der Bodenoberfläche zu erreichen müsste die Datenassimilation häufiger durchgeführt werden.

Im Unterschied zur Bodenoberfläche ist die im Experiment verwendete Häufigkeit für tiefere Bodenschichten ausreichend. Diese Tatsache ermöglicht einen fairen Vergleich zwischen den beiden Filtermethoden:

Der rechnerische Aufwand des Ensemble-Kalman-Filters ist natürlich stark von der gewählten bzw. notwendigen Ensemblegröße abhängig. Für das erweiterte Kalman-Filter kann aus dem numerischen Integrationschema ein Rechenaufwand abgeleitet werden, der ungefähr dem des Ensemble-Kalman-Filters mit $m + 1$ Ensemblemitgliedern entspricht. Bei m handelt es sich hierbei um die Anzahl der artunterschiedlichen Zustandsvariablen. In diesem Fall ist wie oben beschrieben $m = 3$. Der Rechenaufwand des erweiterten Kalman-Filters entspricht also dem des Ensemble-Kalman-Filters mit einer Ensemblegröße von 4 Mitgliedern.

Zusammenfassend kann gesagt werden, dass die beiden Filtermethoden für vergleichbare Aufgaben in etwa gleich aufwändig sind.

Es stellt sich nun die Frage, ab welcher Ensemblegröße das Ensemble-Kalman-Filter zufrieden stellende Ergebnisse liefert. Vergleiche mit dem erweiterten Kalman-Filter zeigen, dass bei 4 Ensemblemitgliedern die tatsächlichen Fehler gleich bzw. bei mehr als 4 Mitgliedern kleiner sind. Die Fehler des Ensemble-Kalman-Filters ändern sich mit der Ensemblegröße aber nur geringfügig und Konvergenz wird schnell erreicht. Diese schnelle Konvergenz ist auf die sehr geringe Größe des Zustandsvektors zurückzuführen. Die Anzahl von nur 3 Freiheitsgraden (die 3 Überschuss- bzw. Defizitvariablen) ermöglicht es, trotz einer kleinen Anzahl an Ensemblemitgliedern, sehr gute Ergebnisse zu erreichen.

Werden anstelle der oben beschriebenen tatsächlichen Fehler nun die Kovarianzen betrachtet, fällt auf, dass das Ensemble-Kalman-Filter nun mehr Ensemblemitglieder benötigt um zufrieden stellende Ergebnisse zu erzielen.

Versuche mit einer Ensemblegröße von 10 oder weniger sind noch zu ungenau, doch die Fehler mit bereits 20 Mitgliedern sind klein genug und kommen denen mit 500 Ensembles schon sehr nahe. Da die notwendige Größe des Ensembles durch die Genauigkeit der Beobachtungen stark beeinflusst wird, können auf diesem Weg Ensemblemitglieder "eingespart" werden. Die Annahme genauerer Messungen würde zum Beispiel eine Aufwertung des 500-Mitglieder-Ensemble-Kalman-Filters bedeuten.

Ein weiteres Instrument um die Filterperformance zu beurteilen ist die Auswertung der Innovationsfolge. Diese Methode ist von großer Bedeutung, da es für ihren Einsatz nicht notwendig ist die zuvor aus einem Kontrollexperiment erhaltenen "wahren" Bodenwasserwerte zu kennen.

Die Innovationsfolge

$$\mathbf{z}_k = \mathbf{v}_k - \mathbf{C}_k \hat{\mathbf{x}}_{k|k-1}$$

beschreibt den Unterschied zwischen den aktuellen Beobachtungen und der Vorhersage. Ist das Problem linear, handelt es sich bei $\mathbf{z}_k$ um eine Gauß'sche white noise Folge mit Kovarianz $E(\mathbf{z}_k \mathbf{z}_k^T) = \mathbf{C}_k \mathbf{P}_{k|k-1} \mathbf{C}_k^T + \mathbf{R}_k$.

Diese für die Filter einfach zu produzierenden Terme werden nun normalisiert und mit der Standardnormalverteilung $N(0, 1)$ verglichen. Der Vergleich zeigt, dass die normalisierten Innovationen etwas weiter sind, wodurch eine Unterschätzung der aktuellen Kovarianzen wiedergegeben wird. Der Grund dafür ist, dass die benutzten Modellfehler die Ungewissheiten der im Experiment verwendeten Parameter nicht vollständig widerspiegeln können.

Das erweiterte Kalman-Filter und das Ensemble-Kalman-Filter unterscheiden sich vor allem darin, wie die Fehlerkovarianzen fortgepflanzt werden. Zur Erinnerung, die Fortpflanzung des erweiterten Kalman-Filters kann aus

$$\begin{aligned} \mathbf{P}_{k,k-1} = & \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right] \mathbf{P}_{k-1,k-1} \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right]^T \\ & + \mathbf{H}_{k-1}(\hat{\mathbf{x}}_{k-1}) \mathbf{Q}_{k-1} \mathbf{H}_{k-1}^T(\hat{\mathbf{x}}_{k-1}), \end{aligned}$$

die des Ensemble-Kalman-Filters aus

$$\mathbf{P}_{k,k-1} = \overline{(\hat{\mathbf{x}}_{k|k-1} - \hat{\mathbf{x}}_{k|k-1})(\hat{\mathbf{x}}_{k|k-1} - \hat{\mathbf{x}}_{k|k-1})^T}.$$

abgelesen werden.

Bemerkung : Bei diesem Experiment wurde beim erweiterten Filter der Einfachheit halber

$$\mathbf{P}_{k,k-1} = \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right] \mathbf{P}_{k-1,k-1} \left[\frac{\partial f_{k-1}}{\partial \mathbf{x}_{k-1}}(\hat{\mathbf{x}}_{k-1}) \right]^T + \mathbf{Q}_{k-1}$$

verwendet.

Die Tatsache der unterschiedlichen Fehlerfortpflanzung verursacht natürlich auch Unterschiede zwischen den Fehlerkovarianzen der Filter. Während die Fehlerstandardabweichung des Bodenoberflächenwasserüberschusses beim Ensemble-Filter sehr rasch schwankt, ergeben sich beim anderen Filter am Anfang des Experiments weitaus glattere Schwankungen. Dies ist darauf zurückzuführen, dass beim erweiterten Kalman-Filter jedes Mal ein konstanter Modellfehler $\mathbf{Q}$ hinzugefügt wurde. Im Gegensatz dazu wurden beim Ensemble-Kalman-Filter die Fehler als Reaktion auf die aktuellen Vorhersagekonditionen hinzugefügt.

Zusätzlich zu den Fehlerstandardabweichungen erzeugen die Filter noch Fehlerkorrelationen zwischen den Zuständen und den gemessenen Variablen, also in diesem Fall dem oberflächennahen Wassergehalt. Diese Korrelationen können aus den nicht-Diagonalelementen der Matrix $\mathbf{P}$ und dem Operator $\bar{\mathbf{C}}$ (beim Erweiterten-Kalman-Filter), oder direkt aus dem Ensemble (beim Ensemble-Kalman-Filter) abgeleitet werden.

Erwartungsgemäß ist die Fehlerkorrelation zwischen Bodenoberflächenwasserüberschuss (bzw. Wurzelschichtwasserüberschuss) und Oberflächenwasser-gehalt meistens positiv. Im Falle der Bodenoberfläche ist die Korrelation sprunghafter.

Abschließend sei bemerkt, dass beide Filter ein geeignetes Hilfsmittel zur Bodenwasserschätzung darstellen. Aufgrund seiner Robustheit und Flexibilität hat das Ensemble-Filter den Vergleich mit dem Erweiterten-Kalman-Filter für sich entschieden.

Literaturverzeichnis

- [1] BURGERS, G.P.; van LEEUWEN, P.J.; EVENSEN, G.: Analysis Scheme in the Ensemble Kalman Filter. - Monthly Weather Review, 126: 1719-1724, 1998.
- [2] CATLIN, D.E.: Estimation, Control, and the Discrete Kalman Filter. - Springer, New York, 1989.
- [3] CHONAVEL, T.: Statistical Signal Processing. - Springer, London, 2002.
- [4] CHUI, C.K.; CHEN, G.: Kalman Filtering with Real-Time Applications. - Springer, Berlin, Heidelberg, 1991.
- [5] DOOB, J. L.: Stochastic processes. - J. Wiley & Sons, New York, 1953.
- [6] ENGEL, H.: Flood warning and operational flood forecasting in the German part of the Rhine River basin. - Destructive Water: Water-Caused Natural Disasters, their Abatement and Control, IAHS Publ. no. 239, 1997.
- [7] EVENSEN, G.: Sequential Data Assimilation with a Nonlinear Quasi-Geostrophic Model using Monte Carlo Methods to Forecast Error Statistics. - Journal of Geophysical Research, 99: 10,143-10,162, 1994.
- [8] EVENSEN, G.: The Ensemble Kalman Filter: Theoretical Formulation and Practical Implementation. - Ocean Dynamics, 53: 343-367, 2003.
- [9] EVENSEN, G.: The Ensemble Kalman Filter for Combined State and Parameter Estimation: Monte Carlo Techniques for Data Assimilation in Large Systems. - preprint, 2009.
- [10] HOFBAUER, F.: Stochastische Prozesse. - Skriptum zur Vorlesung, 2002.
- [11] KALMAN, R.E.: A New Approach to Linear Filtering and Prediction Problems. - Transactions of the ASME-Journal of Basic Engineering, 82 (Series D): 35-45, Baltimore, 1960.
- [12] KALMAN, R.E.; BUCY, R.S.: New Results of Linear Filtering and Prediction Problems. - Journal of Basic Engineering 83: 95-108, 1961.

- [13] KAY, S.M.: Fundamentals of Statistical Signal Processing, Estimation Theory. - Prentice Hall PTR, New Jersey, 1993.
- [14] KÖHLER, B.-U.: Konzepte der statistischen Signalverarbeitung. - Springer, Berlin, Heidelberg, 2005.
- [15] KOMMA, J.; BLÖSCHL, G.; RESZLER, C.: Nachführung mittels Ensemble-Kalmanfilter. - Wiener Mitteilungen Band 199: Hochwasservorhersage - Erfahrungen, Entwicklungen & Realität, Wien, 2006.
- [16] KOSTER, R.D.; SUAREZ, A.; STIEGLITZ, M.; KUMAR, P.: A catchment-based Approach to Modeling Land Surface Processes in a general Circulation Model. I: Model Structure. - Journal of Geophysical Research, 105, 24,809-24,822, 2000.
- [17] LERCH, R.: Elektrische Messtechnik. - Springer, Berlin, Heidelberg, 2007.
- [18] MOSCHYTZ, G.S.; HOFBAUER, M.: Adaptive Filter: Eine Einführung in die Theorie mit Aufgaben und MATLAB-Simulationen auf CD-ROM. - Springer, Berlin, Heidelberg, 2000.
- [19] PAPOULIS, A.: Signal Analysis. - McGraw-Hill Book Company, 3rd printing, 1986.
- [20] POULARIKAS, A.D.; RAMADAN, Z. M.: Adaptive Filtering Primer with Matlab. - CRC Press, Taylor & Francis Group, Boca Raton, 2006.
- [21] REICHLE, R.H.; McLAUGHLIN, D.; ENTEKHABI, D.: Hydrologic Data Assimilation with the Ensemble Kalman Filter. - Monthly Weather Review, Volume 130: 103-114, 2001.
- [22] REICHLE, R.H.; WALKER, J.; KOSTER, R.D.; HOUSER, P.R.: Extending versus Ensemble Kalman Filtering for Land Data Assimilation. - Journal of Hydrometeorology, Volume 3: 728-740, 2002.
- [23] RUYGAART, P.A.; SOONG, T.T.: Mathematics of Kalman-Bucy Filtering. - Springer, Heidelberg, New York, Tokyo, 1985.
- [24] WENDEMUTH, A.: Grundlagen der digitalen Signalverarbeitung, Ein mathematischer Zugang. - Springer, Berlin, Heidelberg, 2005.
- [25] WIENER, N.: Extrapolation, Interpolation and Smoothing of Stationary Time Series. - The Technology Press of The Massachusetts Institute of Technology and John Wiley & Sons, Inc., New York, 4th printing, London, 1960.
- [26] WILKE, K.: Principles of hydrological forecasting by multichannel Wiener filtering. - Mathematical Models in Hydrology and Water Resources System, Bratislava, 1975.

- [27] WILKE, K.; BARTH, F.: Operational River-Flood Forecasting by Wiener and Kalman Filtering. - Hydrology for the water Management of Large River Basins (Proceedings of the Vienna Symposium), 1991.
- [28] WINKLER, G.: Stochastische Prozesse in der statistischen Modellierung. - GSF-Forschungszentrum für Umwelt und Gesundheit GmbH - IBB-Institut für Biomathematik und Biometrie, Oberschleißheim, 2000.

Anhang A

Zusammenfassung

Die Wiener-Filter-Theorie wurde erstmals 1960 in dem von Norbert Wiener verfassten Artikel "Extrapolation, Interpolation and Smoothing of Stationary Time Series" [25] vorgestellt. Das Schätzfilter zur Störungsunterdrückung von Signalen bildet den ersten Hauptteil dieser Arbeit. Einführend wird die Funktionsweise des Filters erklärt. Dazu zählen die Berechnung des Schätzfehlers und der Filterkoeffizienten, die Behandlung der Wiener-Hopf-Gleichung und die Wiener-Lösung. Weiters werden die wichtigen Anwendungsgebiete Signalglättung, Rauschentfernung und Signalprädiktion gezeigt. Bei der Glättung wird aus dem durch Rauschen $\eta(k)$ gestörten Eingangssignal $x(k) = f(k) + \eta(k)$ mit $k = 0, \dots, M$ die Schätzung des Nutzsinal $f(k)$, aus den Daten $x(0), \dots, x(M)$ berechnet. Im Gegensatz zur Filterung ist es erlaubt zukünftige Daten zu verwenden. Bei der Filterung wird nur mit Hilfe der vergangenen und aktuellen Daten das Signal vom Rauschen getrennt. Ziel der Prädiktion ist es, ein Signal $x(M + l)$ für $l \geq 1$ aus $\mathbf{x} = [x(0) \ \dots \ x(M)]^T$ zu schätzen. Im Anschluss an die Anwendungen wird das Wiener-Filter für kontinuierliche Zeit beschrieben. Es wird die kausale und die nicht-kausale Lösung gezeigt.

Bei dem von Rudolf Kalman in "A New Approach to Linear Filtering and Prediction Problems" [11] vorgestellten Kalman-Filter handelt es sich um eine Verallgemeinerung des Wiener-Filters. Einführend in die Theorie wird das Kalman-Filter für lineare Systeme ausführlich erklärt. Es werden die Innovationsfolge, die den Unterschied zwischen den aktuellen Beobachtungen und der Vorhersage ausdrückt und der Minimum-Varianz-Schätzer beschrieben. Ein weiterer Abschnitt beschäftigt sich mit der Herleitung der Kalman-Filter-Gleichungen. Die wohl wichtigste Variante des Filters ist das Erweiterte Kalman-Filter, mit dem es auch möglich ist Zustände nichtlinearer Systeme zu schätzen. Es wird hierbei das nichtlineare Problem durch ein lineares ge-

hert. Bei dem 1994 von Geir Evensen in "Sequential Data Assimilation with a Nonlinear Quasi-Geostrophic Model using Monte Carlo Methods to Forecast Error Statistics" [7] vorgestellten Ensemble-Kalman-Filter handelt es sich um eine Alternative zum traditionellen und insbesondere auch zum Erweiterten-Kalman-Filter. Bei diesem Filter werden die statistischen Eigenschaften der Schätzwerte durch ein Ensemble von möglichen Schätzwerten dargestellt.

Der praktische Teil der Arbeit befasst sich mit hydrometeorologischen Anwendungen der zuvor beschriebenen Theorien. Es werden die Verwendung des Wiener-Filters zur Vorhersage von Wasserabflussmengen von Flüssen und der Einsatz der Kalman-Filter-Theorie in einem linearen Wasserabflussmodell und zur Schätzung des Bodenwassergehalts gezeigt. Vergleiche zwischen den Performances der Unterschiedlichen Kalman-Filter-Varianten werden ebenfalls durchgeführt. Im linearen Wasserabflussmodell wird das Filter für lineare Systeme dem Ensemble-Filter gegenübergestellt. Bei der Schätzung des Bodenwassergehalts wird das Erweiterte- mit dem Ensemble-Filter verglichen.

Anhang B

Abstract

The Wiener-Filter-Theory has first been explained in an article by Norbert Wiener ("Extrapolation, Interpolation and Smoothing of Stationary Time Series" [25]) and presents the first part of this master's thesis. The introduction of the chapter deals with the functionality of the filter, which includes the calculation of estimation errors and filtering coefficients, the handling with the Wiener-Hopf-Equation, and the Wiener-Solution. Additionally, the most important fields of application as smoothing, filtering and signal prediction are shown. For smoothing, $f(k)$ is to be estimated for $k = 0, \dots, M$ based in the data set $x(0), \dots, x(M)$, where $x(k) = f(k) + \eta(k)$. With smoothing it is allowed to use future data, in contrast to filtering, where the signal sample is estimated based on the present and past data only. The objective of prediction is to estimate a signal $x(M+l)$ for $l \geq 1$ from $\mathbf{x} = [x(0) \ \dots \ x(M)]^T$. The explanation of applications is followed by a description of the Wiener-Filter for continuous time, which includes a demonstration of causal and noncausal estimation.

In "A New Approach to Linear Filtering and Prediction Problems" [11] by Rudolf Kalman the Kalman-Filter, which describes a generalization of the Wiener-Filter, is introduced. In the current thesis, the Kalman-Filter for linear systems is explained comprehensively. A description for the innovations sequence, that expresses the difference between current observations and predictions, and the Minimum-Variance-Estimator is given. The subsequent paragraph deals with the derivation of Kalman-Filter-Equations. The Extended-Kalman-Filter is supposed to be the most important variation of this filter and allows to estimate states of nonlinear systems. In this connection the nonlinear problem is approximated by a linear one. In an article published 1994 by Geir Evensen, "Sequential Data Assimilation with a Nonlinear Quasi-Geostrophic Model using Monte Carlo Methods to Forecast Error Statistics" [7], the Ensemble-Kalman-Filter, that presents an alternative to

the traditional and the Extended-Kalman-Filter, is introduced. In this version of Kalman-Filter the statistical properties of the estimation values are shown by an ensemble of possible estimation values.

The empirical part of the thesis is about hydrometeorological applications of the above-mentioned theories. The usage of Wiener-Filter for the prediction of discharge of rivers and the application of the Kalman-Filter-Theory in a linear discharge-model and for soil moisture estimation are shown. Another important point is the comparison between single performances from different Kalman-Filter-variants, which is also discussed. In the linear discharge-model the Filter of linear systems is contrasted with the Ensemble-Filter. For soil moisture estimation, the Extended-Filter is compared to the Ensemble-Filter.

Anhang C

Lebenslauf

Persönliche Daten

Name	Emanuel Hecher
Geburtsdaten	28. Oktober 1982 in Mödling
Eltern	Hildegard und Josef Hecher
Familienstand	ledig

Schulische Ausbildung

1989 - 1993	Volksschule Höflein
1993 - 2001	BRG Gröhrmühlgasse Wiener Neustadt

Wehrdienst

2002	Grundwehrdienst
-------------	-----------------

Hochschulbildung

2002 - 2009	Diplomstudium Mathematik, Universität Wien
--------------------	--