

Parameterselektion für komplexe Modellierungsaufgaben der Wasserwirtschaft

Moderne CI-Verfahren zur Zeitreihenanalyse

T. Bartz-Beielstein, T. Zimmer und W. Konen

Fakultät für Informatik und Ingenieurwissenschaften

Fachhochschule Köln

Tel.: +49 2261 8196-0

Fax: +49 2261 8196-15

E-Mail: {thomas.bartz-beielstein | wolfgang.konen}@fh-koeln.de

Zusammenfassung

Die Simulation komplexer technischer Vorgänge kann mit unterschiedlichen Modellen durchgeführt werden. Meistens erfolgt die Modellauswahl und anschließende Parametrisierung des Modells nach subjektiven Kriterien. Wir demonstrieren, wie eine systematische Vorgehensweise zur Modellselektion mittels SPO (sequentieller Parameteroptimierung) diesen Vorgang objektivieren kann. Um die einzelnen Schritte nachvollziehbar darzustellen, basiert unsere Darstellung auf einem Beispiel aus der Praxis: der Modellierung von Füllständen in Regenüberlaufbecken.

1 Einleitung

Verfahren der *Computational Intelligence* (CI) haben in den letzten Jahrzehnten einen festen Platz im Repertoire der Simulations- und Optimierungsverfahren erobert. *Neuronale Netze* (NN), *genetische Algorithmen* (GA) oder *genetisches Programmieren* (GP) – um nur einige CI-Verfahren zu nennen – werden tagtäglich in der industriellen Praxis eingesetzt. Die Entscheidung für den Einsatz dieser Verfahren wird dabei häufig von subjektiven Kriterien beeinflusst. Als Kriterien können hier z.B. die Vertrautheit mit einem Verfahren, das persönliche Interesse an beispielsweise bioanalogen Verfahren oder Vorträge auf Fachkonferenzen genannt werden. Nach unserer Erfahrung spielt bei dieser Auswahl der Zufall eine große Rolle. Nur in seltenen Fällen erfolgt die Auswahl allein nach objektiven Maßstäben. Außerdem ist es nicht ausreichend, das beste Verfahren auszuwählen, da der Auswahlprozess selbst gewisse Eigenschaften besitzen sollte¹. Es ist

¹So beschreiben Beyth-Marom et al. [1] diese Situation wie folgt:

... According to the most general normative model, a person facing a decision should (a) list relevant action alternatives, (b) identify possible consequences of those actions, (c) assess the probability of each consequence occurring (if each action were undertaken), (d) establish the relative importance (value or utility) of each consequence, and (e) integrate these values and probabilities to identify the most attractive course of action, following by a defensive decision rule. People who follow these steps are said to behave in a *rational* way. People who do so effectively (e.g., they have accurate probability estimates, they get good courses of action into their list of possibilities) are said to behave *optimally*. Thus, if one does not execute these steps optimally, one can be rational without being very effective at getting what one wants.

also nicht alleine ausreichend, das beste Verfahren zu finden; der Findungsprozess selbst muss optimal gestaltet werden. Dieser Artikel beschäftigt sich zentral mit diesem Findungsprozess und stellt ein Verfahren (SPO) vor, das diesen Prozess verbessern kann.

Uns ist bewusst, dass es einen 100% objektiven Auswahlprozess allein schon aus Zeitgründen nicht geben kann. Stehen aber bereits einzelne Verfahren zur Verfügung und ist ein effizienter Vergleich gewünscht, so kann das in diesem Artikel beschriebene Verfahren hilfreich sein. In diesem Artikel werden drei Zielsetzungen behandelt. Wir betrachten die Modellauswahl und die Bestimmung guter Parameter für ein einzelnes Modell. Diese Zielsetzungen lassen sich mittels eines Meta-Modells in einem gemeinsamen statistischen Framework behandeln. Eine mögliche und bereits in vielen Bereichen erfolgreich angewandte Technik ist die *sequentielle Parameteroptimierung* (SPO) [2, 3, 4]. Als dritte Zielsetzung untersuchen wir verschiedene Einstellung des SPO Meta-Modells.

Zunächst stellen wir in Abschnitt 2 ein Problem aus der Wasserwirtschaft dar. In Abschnitt 3 wird die SPO als Meta-Modellierungsansatz vorgestellt. Abschnitt 4 stellt die verschiedenen Prognoseansätze vor und beschreibt die zugehörigen Schritte zur Datenvor- und -nachbearbeitung. Ebenfalls in diesem Abschnitt wird dargelegt, wie die Parameter des von uns favorisierten Modells eingestellt wurden. Abschnitt 5 beschreibt die Experimente, die in Abschnitt 6 diskutiert werden. Im letzten Abschnitt geben wir eine kurze Zusammenfassung.

2 Füllstandsprognose in Regenüberlaufbecken

In dieser Arbeit vergleichen wir systematisch verschiedene Modellansätze für eine Anwendung aus der Wasserwirtschaft. Dabei steht nicht nur der Vergleich verschiedener Modellierungsansätze im Vordergrund, sondern vor allem die Beschreibung der Verbesserungspotentiale, die sich mit einer gezielten Vorgehensweise erreichen lassen².

Der Vergleich von Modellen zur Prognose der Füllstände in *Regenüberlaufbecken* (RÜB) aufgrund von Regeneintrag und Bodenbeschaffenheit war Gegenstand unserer Untersuchungen. Im Forschungsprojekt KANNST (KANalNetz-Steuerung), aus das von uns benutzte Datenmaterial stammt, wird die Modellierung und Prognose von Füllstandshöhen in RÜB auf Basis einzelner Regenmessungen untersucht [6]. Kanalnetze werden nach dem heutigen Stand der Technik meistens als unregelmäßige Systeme betrieben, bei denen z.B. am Ablauf von Regenüberlaufbecken feste Drosselmengen eingestellt werden [7, 8]. Im Vordergrund einer Weiterentwicklung der Kanalnetzbewirtschaftung steht der Gewässerschutz. In jüngster Zeit werden die hierfür erforderlichen vernetzten Regelungssysteme im Rahmen des Projektes KANNST [6] unter Nutzung moderner Messtechnik und neuer CI-Verfahren entwickelt. Eine vorausschauende Prognose der Füllstände aufgrund von Regeneintrag und Bodenbeschaffenheit ist für diese Aufgaben unerlässlich. Die in [9] dargestellten Ergebnisse basierten auf den mit dem Simulator SWMM [10] erzeugten Regendaten und Füllständen. Hingegen beruhen die in dieser Arbeit präsentierten Ergebnisse auf empirisch erhobenen Daten. Die Anforderungen an die zugrunde liegenden Modellierungsansätze werden dadurch drastisch erhöht, denn die Daten weisen ein intermittierendes und burstartiges Verhalten auf, s. Bild 1. Dieses stellt rückgekoppelte Modelle

²Während in [5] die Prognosemodelle im Vordergrund der Analyse standen, untersuchen wir in diesem Artikel zusätzlich die SPO.

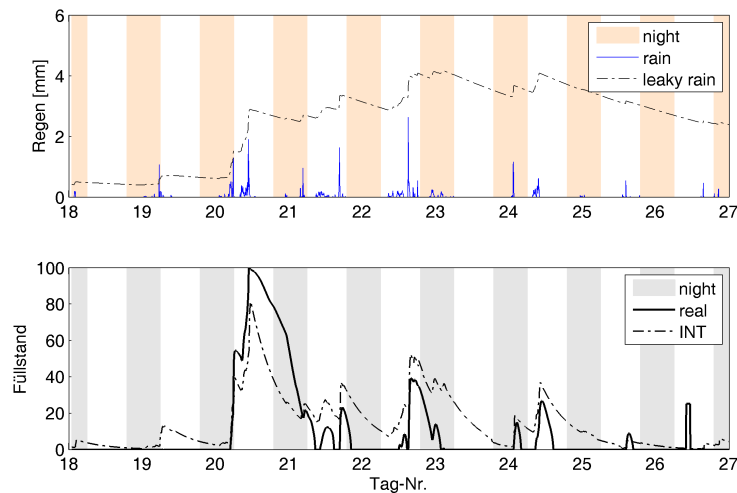


Bild 1: Ausschnitt aus den Kanalnetzdaten. Oben: Die Regendaten haben einen stark burstartigen Charakter. Der 'leaky rain' ist eine gleitende Integration der Regenaktivität. Unten: Die zugehörigen Füllstände im RÜB. Tag Nr. 18 entspricht dem Datum 16.05.2006.

vor starke Probleme. Wir untersuchen, wie verschiedene Prognosemodelle mit solchen Inputdaten arbeiten können. Unterschiedliche Ansätze wie *neuronale Netze* (NN), FIR-Modelle, *Nonlinear AutoRegressive models with eXogenous inputs* (NARX) [11], *Echo State Networks* (ESN) [12], Differentialgleichungen (ODE) oder die Modellierung mit Integralgleichungen (INT) werden dabei von uns eingesetzt.

3 Statistische Meta-Modellierung

Wir unterscheiden exogene Faktoren von endogenen Parametern. Endogene Parameter ändern sich während einer Untersuchung (z.B. Gewichte bei NN während des Lernvorgangs), exogene Faktoren bleiben hingegen konstant (z.B. die Anzahl der Neuronen). Exogene Faktoren werden im weiteren Verlauf unserer Untersuchung betrachtet, endogene Parameter sind nicht Gegenstand unserer Untersuchung. Die exogenen Faktoren werden im Folgenden als *Modellfaktoren* bezeichnet. Zusätzliche Parameter, wie z.B. der Beobachtungszeitraum, sind problemspezifisch. Diese werden im Folgenden als *Problemfaktoren* bezeichnet. Ein Modelldesign beschreibt die Gesamtheit aller modellspezifischen Faktoren, ein Problemdesign beschreibt die problemspezifischen Faktoren. Ein wichtiges Ziel in der Simulation ist die Bestimmung eines optimalen Modelldesigns für ein (oder sogar mehrere) gegebene Problemdesigns. Wir betrachten daher eine typische Situation der Meta-Modellierung, siehe Bild 2. Um eine Vergleichbarkeit zu gewährleisten und zu erreichen, dass jeder Prognoseansatz nach Möglichkeit mit seinem besten Modelldesign operiert, nutzen wir für alle Verfahren die SPO [3, 4], um die bestmögliche Faktoreinstellungen systematisch zu bestimmen. Neu ist in dieser Arbeit auch, dass SPO erstmalig für die Optimierung von ODE-Modellen und Integralgleichungsmodellen (INT) eingesetzt wird. Diese Modelle bieten das Potential, komplexe Anwendungszusammenhänge zu modellieren, die von rein datengetriebenen Verfahren in der Regel nicht erfasst werden. Traditionell haben sie aber das Problem, dass eine größere Anzahl von Parametern einzustellen sind, um quantitativ gute Übereinstimmung zu erzielen. Hier bietet SPO einen

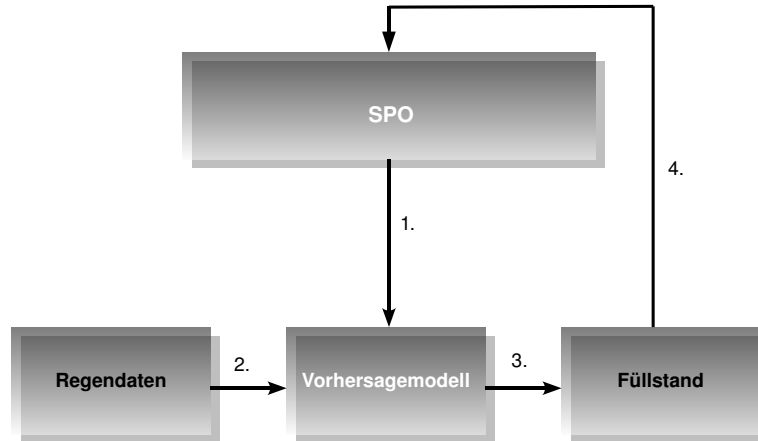


Bild 2: Datenfluss in der SPO Meta-Modellierung. 1. SPO generiert Einstellungen für das Vorhersagemodell (z.B. ein NN). Mit den Regendaten (2.) berechnet das Vorhersagemodell die zugehörigen Füllstände im RÜB (3.) und den Vorhersagefehler. Mit dem Vorhersagefehler (4.) berechnet SPO ein Meta-Modell und generiert weitere Einstellungen für das Vorhersagemodell.

systematischen Ansatz. Ziel unserer Untersuchungen ist festzustellen, ob SPO auch im Fall von ODE und INT gute Optimierungsergebnisse liefert.

4 Material und Methoden

4.1 Daten und Zielfunktion

Die Regendaten wurden minütlich ab dem 28.04.2006 über einen Zeitraum von 108 Tagen erhoben, so dass insgesamt $n = 155.521$ Datensätze zur Verfügung standen. Basierend auf den gemessenen Regendaten r_t ($t = 1, 2, \dots, n$), in den Abbildungen mit „rain“ bezeichnet, sollen die realen Füllstandshöhen in den Becken vorhergesagt werden. Die realen Werte y_t werden mit „real“ und die von unseren Modellen vorhergesagten Werte $\hat{y}_t$ werden mit dem Modellkürzel (z.B. INT) bezeichnet. Die minütlichen Regenmengen wurden zunächst in 5minütigen Intervallen zusammengefasst, die Füllstände in 5minütigen Intervallen gemittelt. Bekannt sind somit r_t und y_t , zu bestimmen ist ein Modell zur Berechnung von $\hat{y}_t$, so dass der Fehler zwischen y_t und $\hat{y}_t$ möglichst gering ist³ ⁴. Als Gü-

³Es ist zu betonen, dass die hier vorgestellten Modelle (bis auf NARX) alle r_t als einzigen zeitveränderlichen Input benutzen.

⁴Auf den ersten Blick mag es erscheinen, dass die Prognose $r_t \rightarrow \hat{y}_t$ gar keine Prognose darstellt, da der Prognosehorizont 0 ist. Jedoch ist die wahre Aufgabe eines Modells, aus den Regeneinträgen bis zum aktuellen Zeitpunkt t eine Prognose über die Bodenbeschaffenheit und deren Auswirkung auf den Füllstand abzugeben. Des Weiteren führt der bis zum Zeitpunkt t gemessene Regeneintrag erst 30-60 min später zu einer Füllstandsänderung, ein Zeitfenster, welches sinnvoll zu Steuerungsmaßnahmen in Kanalnetzen genutzt werden kann.

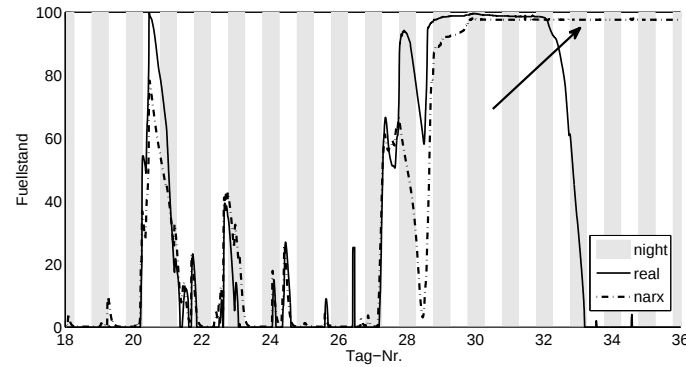


Bild 3: NARX Modellierungsfehler. Der Pfeil kennzeichnet den Bereich, in dem das NARX-Modell nicht „zurück schwingt“, sondern einen konstant hohen Füllstand (größer als 90%) vorhersagt. Dieser Fehler trat systematisch in jeder NARX Simulation auf.

temaß der Prognosen betrachten wir den sog. *Mean Square Error* (MSE): $\frac{1}{n} \sum_t^n (y_t - \hat{y}_t)^2$ und den *Root Mean Square Error* (RMSE): $\sqrt{\frac{1}{n} \sum_t^n (y_t - \hat{y}_t)^2}$.

4.2 Prognoseverfahren

Im Rahmen unserer Analysen habe wir sechs verschiedene Prognoseverfahren betrachtet: NN, ESN, NARX, FIR, ODE und INT. Neuronale Netze finden häufig Anwendung, um komplexe Beziehungen zwischen Ein- und Ausgabegrößen zu modellieren [9]. Dabei ist die zugrunde liegende Abbildungsvorschrift nicht genau bekannt, sondern wird vom NN gelernt. Echo State Networks sind rekurrente neuronale Netze, die ein sog. Reservoir rückgekoppelter Neuronen nutzen, um komplexe dynamische Signale zu generieren [12]. NN und ESN wurden zwar untersucht, aber werden hier nicht weiter dargestellt, da sie im Vergleich mit den anderen Verfahren deutlich schlechtere Ergebnisse lieferten. Der in Bild 3 dargestellte Fehler trat bei NN und ESN bereits nach kurzer Zeit auf und ließ sich nicht durch eine Änderung der Einstellungen der Modellfaktoren ändern. Filter mit endlicher Impulsantwort, engl.: *finite impulse response filter* (FIR-Filter), werden häufig in der Signalverarbeitung eingesetzt. NARX Modelle werden in der Zeitreihenmodellierung und -analyse eingesetzt. Der aktuelle Wert der Zeitreihe wird durch Werte aus der Vergangenheit und zusätzlich aus aktuellen und vergangenen Werten der exogenen Reihe bestimmt [11]. Wir betrachten zudem ein gewöhnliche Differentialgleichungssystem, engl. *ordinary differential equations* (ODE), um die wichtigsten Kausalzusammenhänge des dynamischen Systems Regen-Böden-RÜB zu modellieren. Nachteilig an diesem einfachen ODE-Modell ist, dass es an exponentielles Abklingverhalten des Füllstandes gebunden ist. Dies entspricht nicht den realen Verläufen, weil die Laufzeiten des Wassers durch verschiedene Bodenschichten natürlich unterschiedliche Verzögerungen haben. Deshalb formen wir das ODE-Modell in ein vergleichbares Integralgleichungsmodell (INT) um, bei dem das Abklingverhalten durch Modifikation geeigneter Faltungskerne anpassbar ist.

4.3 SPO-basierte Meta-Modellierung

Die FIR-, NARX-, ODE- und INT-Modelle besitzen zwischen zwei und dreizehn Parametern. In vielen Situationen wird versucht, manuell, d.h. durch Probieren, günstige Parametereinstellungen für das Modell zu finden. Die sog. *one-factor-at-a-time* (OAT) Vorgehensweise findet dabei häufig Anwendung. Allerdings wächst die Anzahl der möglichen Einstellungen mit der Anzahl der Faktoren exponentiell⁵. Daher besitzt die OAT-Vorgehensweise in vielerlei Hinsicht große Nachteile im Vergleich zu einer systematischen Vorgehensweise. Im Bereich der Simulation wurde dies von Kleijnen ausführlich dargestellt [14]. Wir setzten die SPO [4] ein. Diese kombiniert Ansätze aus der klassischen Regressions- und Varianzanalyse mit modernen statistischen Verfahren wie Kriging [15, 16]. SPO ist eine umfassende, auf einem rein experimentellen Ansatz beruhende Methode zur Analyse und Verbesserung deterministischer und stochastischer Simulations- und Optimierungsverfahren [2]. Ein eingangs kleiner Stichprobenumfang mit wenigen Wiederholungen wird im Laufe der SPO vergrößert. Zur Bestimmung neuer Stichproben und der Anpassung der Anzahl der Wiederholungen fließen die bisher gewonnenen Informationen ein, so dass das Modell sequenziell verbessert und die Aussagen immer sicherer werden. Die sequenzielle Vorgehensweise ermöglicht ein Lernen der zugrunde liegenden Abhängigkeiten. SPO kann zudem automatisiert ablaufen und ist für die numerische Simulation und Optimierung praktisch konkurrenzlos. Uns ist zur Zeit kein vergleichbares Verfahren bekannt, das einem Anwender mit geringem Aufwand die Anpassung, den Vergleich und die Analyse verschiedener Modelle ermöglicht. SPO kann direkt zur Optimierung komplexer Probleme eingesetzt werden und wird von den Anwendern, aber auch in der empirisch orientierten Grundlagenforschung an stochastischen Suchverfahren stark nachgefragt⁶. Zudem wird eine frei verfügbare *SPO Toolbox* (SPOT) momentan entwickelt. Mittels SPOT gelang in der hier vorliegenden Arbeit z.B. die Analyse, welcher der Faktoren den größten Einfluss auf die Vorhersagegüte hat, so dass Faktoren mit nur geringem Einfluss im Laufe der SPO-Analyse nicht weiter betrachtet werden müssen.

5 Experimente und Resultate

Es wurden drei ausführliche Studien durchgeführt. Zunächst wurde die Vorhersagegüte auf dem gesamten Datensatz untersucht. Im zweiten Schritt wurde der Datensatz partitio-

⁵Dieses grundlegende Problem ist in der Literatur auch als „Fluch der Dimensionen“ (engl. *curse of dimensionality* [13]) bekannt: Im $\mathbb{R}^1$ reichen 100 Punkte, um das Einheitsintervall im Abstand 0.01 zu belegen. Im $\mathbb{R}^{10}$ sind dafür 10^{20} Punkte erforderlich. Ein Einheitswürfel im $\mathbb{R}^{10}$ ist in diesem Sinne um das 10^{18} -fache größer als im $\mathbb{R}^1$.

⁶SPO wurde beispielsweise in den folgenden Anwendungsfeldern eingesetzt: Maschinenbau: Temperierbohrungen; Luft- und Raumfahrt: Tragflächenoptimierung; Simulation und Optimierung: Fahrstuhlsteuerung; Algorithm Engineering: Graphenalgorithmen; Computational Intelligence: Algorithmische Chemie; Verfahrenstechnik: Entwurf von Destillationsanlagen; Wirtschaftswissenschaft: Modellierung eines Bodenmarktes; Statistik: Selektionsverfahren für Partikelschwarm Verfahren; Informatik: Threshold Selektion und Schrittweitensteuerung für Evolutionsstrategien, Analyse und Anwendung von Partikelschwarm Verfahren; Numerische Mathematik: Vergleich/Analyse klassischer und moderner Optimieralgorithmen; Logistik: Tourenoptimierungs- und Torzuordnungsprobleme; Bioinformatik: Optimierung von Hidden-Markov-Modellen.

Eine Übersicht mit Literaturstellen und eine kurze Einführung in SPO werden in [17] gegeben.

niert, um Training- und Testdaten zur Verfügung zu haben. Im dritten Schritt wurde die Spezifikation des SPO Meta-Modells untersucht.

Für das FIR-Modell mussten fünf Faktoren angepasst werden. Der erste Faktor (s) modelliert die Verdunstung des Niederschlags. Die Gewichte, die nach einer gewissen Verzögerung mit wachsendem t exponentiell abnahmen, benötigen die Einstellung der Verzögerung (d), der Länge des exponentiellen Abklingens (l) und des Abklingfaktors A . Schließlich muss die Skalierung s_2 angepasst werden. Im NARX-Modell sind zwei Faktoren anzupassen: Die Anzahl der Verzögerungen (delay-states, d) und die Anzahl der Neuronen n . Das INT-Modell besitzt 13 zu optimierenden Faktoren und für das ODE-Modell waren sechs Faktoren zu bestimmen. Diese bilden eine Teilmenge der Faktoren des INT-Modells. Um die Übersichtlichkeit zu wahren, haben wir die einzelnen Faktoren in Tab. 1 zusammenfassend dargestellt.

5.1 Erste Untersuchung: Vergleich der Verfahren auf dem gesamten Datensatz

Um einen ersten Vergleich der Prognosegüte der verschiedenen Verfahren zu erhalten, haben wir in diesen Experimenten den gesamten Datensatz (108 Tage) zugrunde gelegt. In diesem Abschnitt beschreiben wir exemplarisch, wie SPO-basiert verbesserte Modellfaktoreinstellungen bestimmt werden können. Diese Vorgehensweise wurde für alle betrachteten Modelle benutzt und wird anhand des INT-Modells (da dies die meisten Faktoren besitzt) im Folgenden skizziert⁷.

I. Vorstudien zur Bestimmung der ROI. Ausgangspunkt waren in Vorstudien [18] ermittelte gute Einstellungen. Diese ersten Einstellungen basierten auf Plausibilitätsüberlegungen. Es wurde eine geringe Anzahl von Testläufen (ohne Designüberlegungen zu berücksichtigen) durchgeführt. Anschließend wurde systematisch ein experimentelles Design entworfen, in dem Intervallgrenzen großzügig gesetzt wurden. Diese Intervallgrenzen werden in der Literatur als *region of interest* (ROI) bezeichnet. Bereits nach drei SPO Runden wurde eine Verbesserung des erwarteten RMSE von 39.25 auf 10.95 erreicht, Überschreitungen des zulässigen Bereichs wurden dabei bewusst in Kauf genommen, so dass SPO einen möglichst großen Bereich explorieren konnte. Nicht zulässige Lösungen erhielten einen Funktionswert, der signifikant über dem in den Vorstudien gefundenen schlechtesten Wert lag⁸.

II. Screening. Im diesem Schritt, der in der Literatur auch als *screening* bezeichnet wird, bestimmten wir die wichtigsten Parameter. Dazu wurde die Laufzeit auf 500 Berechnungen des INT Modells beschränkt, da bereits nach einer kurzen Anfangszeit die wichtigsten Parameter feststehen sollten. Die SPO Analyse zeigte, dass der Parameter α_L den größten Einfluss auf die Vorhersagegüte hat (Bild 5). In diesem Stadium können Ausreißer, die eine valide Modellbildung nicht zulassen, erkannt werden, so dass eine Veränderung der ROI durchgeführt werden kann.

Die baumbasierte Analyse in Bild 4 bestätigt diese Ergebnisse. Sie zeigt, dass der beste Wert wird für $\alpha_L < 0.00357935$ und $3.85311 < \Delta$ erzielt wird.

⁷Eine detaillierte Darstellung dieser Modelle findet sich in [5], generelle Erläuterungen zur SPO-Methodik mit weiteren Beispielen sind in [4] dargelegt.

⁸Es wird ein Minimierungsproblem betrachtet.

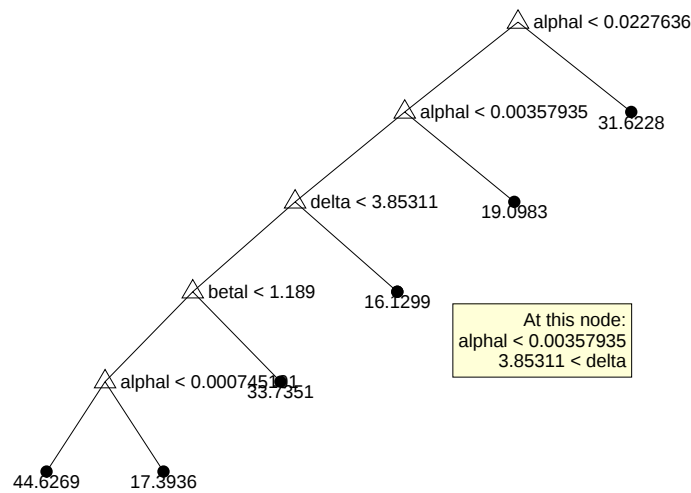


Bild 4: Regressionsbaum. Die Bedingung an jedem Knoten gilt für den linken Teilbaum. Der Faktor mit dem größten Effekt steht an der Wurzel. Der Faktor α_L hat anscheinend den größten Effekt.

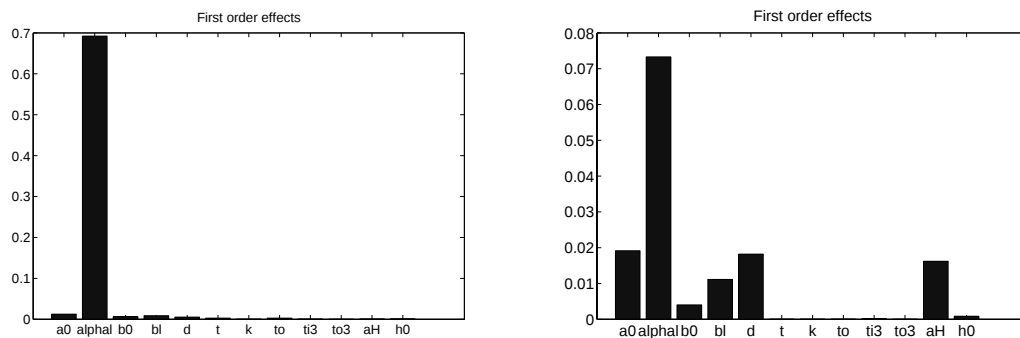


Bild 5: Links: Effekte erster Ordnung (Haupteffekte ohne Interaktionen) nach 500 Auswertungen des INT Modells. Der Faktor α_L hat anscheinend den größten Effekt. Rechts: Haupteffekte im zweiten Schritt der Modellierung. Dadurch, dass passenden Werte für α_L gefunden werden konnten, kommt es zu weniger Ausreißern und die anderen Faktoren besitzen einen größeren Einfluss.

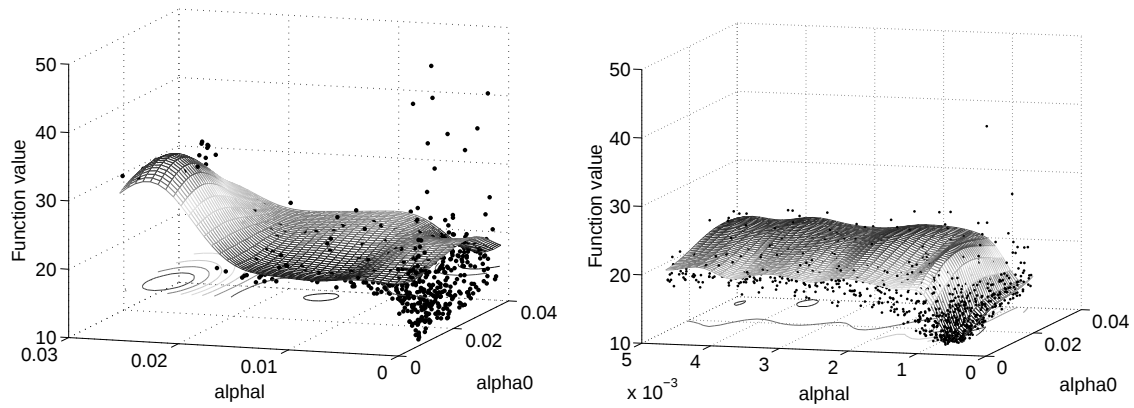


Bild 6: Links: Interaktion der Faktoren α_L und α nach 500 Auswertungen des INT Modells. Die Punkte stellen die Auswertungen dar, sie zeigen starke Fluktuationen im Zielwert (aufgrund der Variation der anderen, hier nicht gezeigten Parameter). Die Fläche ist die SPO-Approximation an die Punkte, die die Basis für weitere Auswertungen ist. Der Faktor α_L hat anscheinend den größten Effekt, kleinere α_L Werte führen zur Verbesserung, aber auch zu höherer Varianz. Rechts: In einer späteren Phase der SPO-Optimierung kann durch Einschränkung der ROI eine wesentliche Reduktion der Varianz erzielt werden, die schwarzen Punkte liegen dichter beieinander.

Die Analysen legen nahe, dass eine Anpassung der ROI für den Faktor α_L bessere Ergebnisse produziert, da dieser mit Abstand den größten Einfluss auf die Vorhersagegüte hat. Im nächsten Schritt beschreiben wir, wie eine Anpassung der ROI durchgeführt werden kann: Für $\alpha_L = 0.0009618$ wurde der beste Funktionswert gefunden. Die Betrachtung des Regressionbaums (Abb. 4) legt nahe, den Wert für α_L kleiner als 0.00357935 und den Wert für Δ größer als 3.85311 zu wählen. Mit der so veränderten ROI werden nun weitere SPO Läufe durchgeführt, wobei die Anzahl n der Simulatoreaufrufe (bzw. Aufrufe des INT-Modells) auf 1000 vergrößert wird. Aus unseren bisherigen Experimenten tragen Werte, die sehr viel größer als 1000 sind, nicht zu einer Verbesserung der Modellfaktoren bei. Im Gegenteil, es können sogar gegenteilige Effekte bei einer Kriging-basierten Modellierung auftreten [19]. Zusätzlich wird durch Bild 6 offensichtlich, dass die Varianz auf diesem Intervall für α_L deutlich zunimmt. Eine ähnliche Betrachtung für Δ zeigte, dass größere Δ Werte vermutlich zu besseren Ergebnissen führen werden, allerdings im Gegensatz zu α_L ohne Veränderung der Varianz.

III. Modellierung und Optimierung. Nachdem in Schritt II von den ursprünglich zwölf Faktoren sechs als signifikant erkannt wurden, nämlich α_L , α , β , β_L , Δ und α_H , beschränkten wir die Suche im weiteren Verlauf auf diese Parameter. Die restlichen Parameter beließen wir auf den bislang besten gefundenen Werten. Mit den so eingestellten Parametern erzielten wir nach 1000 Simulationen den Bestwert RMSE=9.49, siehe Tab. 2.

Der Rechenaufwand bei der unter I.-III. dargestellten Vorgehensweise liegt im Stundenbereich für die FIR-, ODE- und INT-Modelle. Für die NARX-Modellierung war der Aufwand um den Faktor 10 größer, da in jedem Lauf das NARX-Netz trainiert werden musste.

Interessant ist in diesem Zusammenhang auch die Fragestellung, wie weit die von SPO gefundenen Faktoreinstellungen von den manuell bestimmten Werten abweichen, d.h. ist

Tabelle 1: Die 13 Parameter des INT-Vorhersagemodells mit den vom Anwender vorgegebenen zulässigen Bereichen. Das ODE-Modell benutzt daraus nur die 6 Parameter $\alpha, \beta, \tau_{rain}, \Delta, \alpha_L, \beta_L$, die in der Tabelle in hellerem Grau unterlegt sind.

Parameter	Symbol	manuell	Best SPO	Bereich SPO
Abklingkonstante Füllstand (Filter g)	α	0.0054	0.00845722	[0, 0.02]
Abklingkonstante Filter h	α_H	0.0135	0.309797	{0 ... 1}
Abklingkonstante 'leaky rain'	α_L	0.0015	0.000883692	{0 ... 0.0022}
Einkopplung Regen in Füllstand	β	7.0	6.33486	{0 ... 10}
Einkopplung Regen in 'leaky rain'	β_L	0.375	0.638762	{0 ... 2}
Einkopplung K -Term in Füllstand	h_0	0.5	6.87478	{0 ... 10}
Schwelle für 'leaky rain'	Δ	2.2	7.46989	{0 ... 10}
Flankensteilheit aller Filter	κ	1	1.17136	{0 ... 200}
Zeitverzögerung Füllstand zu Regen	τ_{rain}	12	3.82426	{0 ... 20}
Startzeitpunkt Filter h	τ_{in3}	0	0.618184	{0 ... 5}
Endzeitpunkt Filter h	τ_{out3}	80	54.0925	{0 ... 500}
Endzeitpunkt Filter g	τ_{out}	80	323.975	{0 ... 500}
RMSE		12.723	9.48588	

Tabelle 2: Überblick über die Prognosegüte der einzelnen Verfahren mit und ohne SPO auf dem gesamten Datensatz (108 Tage). Das beste Verfahren in jeder Spalte ist weiß unterlegt.

Verfahren	$\langle RMSE \rangle$ zufällige Parametersets	RMSE manuelles Parameterset	RMSE optimiertes Parameterset
FIR	25.42	25.57	20.10
NARX	85.22	75.80	38.15
ODE	39.25	13.60	9.99
INT	31.75	12.72	9.49

der Anwender auch in der Lage, ohne SPO günstige Einstellungen zu finden und welche Faktoren werden manuell völlig falsch eingestellt? Deshalb stellen wir in Tab. 2 drei verschiedene Messungen vor: Erstens haben wir in einem vom Anwender vorgegebenen Bereich (jeweils kleinster und größter zulässiger Wert) 100 zufällige Einstellungen der Modellfaktoren generiert und die zu jeder Einstellung gehörenden Vorhersagefehler berechnet. Zweitens wurde eine vom Anwender nach manuellem Tuning der Modellfaktoren gefundene Einstellung simuliert. Drittens wurde mit einem Budget von $n = 1000$ Modellierungsläufen eine Optimierung der einzelnen Verfahren mittels SPO durchgeführt. Diese Werte sind in Tab. 1 exemplarisch für das INT Modell dargestellt. Die von SPO nach bestmöglicher Einstellung der Modellparameter erzielten RMSE-Werte sind in der vierten Spalte von Tab. 2 dargestellt.

Für die SPO Läufe wurden 5 Wiederholungen berechnet, da NARX ein stochastisches Verfahren ist. Die oben dargestellten Ergebnisse zeigen für das NARX Modell ein eher ernüchterndes Bild: Im Mittel wurde ein RMSE von 35.50 erzielt. Die mit dem NARX Modell erzielten Werte weisen eine hohe Varianz auf. Aber: Das Minimum liegt bei 12.35 und ist durchaus mit den anderen Verfahren vergleichbar. Allerdings ist dies ein zufälliger Wert, der von der Initialisierung des Zufallszahlengenerators (Seed) abhängt.

5.2 Zweite Untersuchung: Generalisierbarkeit der Ergebnisse auf neuen Daten

Die zweite Serie von Experimenten wurde durchgeführt, um festzustellen, ob die gefundenen Einstellungen sich auf neue Daten übertragen lassen. Dazu wurden die Daten in

Tabelle 3: Vergleich der Generalisierbarkeit der einzelnen Verfahren mit und ohne SPO-Optimierung. Die ersten beiden Spalten zeigen den RMSE auf dem Trainingszeitraum (Tage 0...53), die letzten beiden Spalten den RMSE auf dem Testzeitraum (Tage 54...107). Die Zahlen in geschweiften Klammern sind die Testfehler hochgerechnet auf gleiche Regenintensität wie im Trainingszeitraum (siehe Abschnitt 5.2).

Verfahren	RMSE Train (manuell)	RMSE Train (mit SPO)	RMSE Test (manuell)	RMSE Test (mit SPO)
FIR	33.45	20.05	13.73 {19.61}	13.91 {19.86}
NARX	50.86	17.25	20.38 {29.10}	9.91 {14.15}
ODE	13.15	11.03	12.18 {17.40}	8.14 {11.62}
INT	14.99	10.82	9.47 {13.53}	7.71 {11.01}

Trainings- und Testdaten unterteilt. Die erste Hälfte, also 54 Tage, wurde zum Lernen benutzt. Der Fehler wurde dann auf den restlichen Daten bestimmt. Es wurden wiederum die in Abschnitt 5.1 beschriebenen Modellfaktoren optimiert. Tabelle 3 stellt die Ergebnisse zusammen. Es mag zunächst verwundern, dass beim ODE- und INT-Modell die Fehler auf den Trainingsdaten fast immer höher sind als auf den Testdaten. Dies liegt daran, dass der Testzeitraum (Mitte Juni - Mitte August) weniger Niederschläge enthielt als der Trainingszeitraum (Mitte April - Mitte Juni) und damit an trockenen Tagen leichter vorhersagbar war. Korrigiert man aber den Testzeitraum auf gleiche Regenintensität wie den Trainingszeitraum (Faktor $1.428 = 30/21$ Regentage), so zeigen die Zahlen in geschweiften Klammern, dass Trainings- und Testfehler jeweils gleichauf liegen. Es findet also durch die Modelle ODE und INT kein „Auswendiglernen“ statt.

5.3 Dritte Untersuchung: Wahl der initialen Designgröße

Zu Beginn einer SPO wird das zu optimierende Modell $M \in \mathcal{M}$ mehrfach, z.B. $k = 10$ mal, ausgewertet. Anschließend wird ein SPO-Metamodell $M^* \in \mathcal{M}^*$, z.B. ein stochastisches Prozessmodell oder ein Regressionsbaum, bestimmt, um weitere, verbesserte Einstellungen für M vorherzusagen. Um Zeit und Kosten zu sparen, ist es wünschenswert, k möglichst klein zu wählen. Eine untere Schranke ist durch das Metamodell M^* und durch die Anzahl der zu optimierenden Faktoren n im Modell M gegeben. So benötigt ein stochastisches Prozessmodell mit linearem Regressionsterm (also ein Polynom erster Ordnung, siehe [15, 20]) $k = n + 1$ initiale Experimente. Ein quadratisches Metamodell benötigt bereits $k = \frac{1}{2}(n + 1)(n + 2)$ initiale Experimente. Sollen, wie im INT-Modell $n = 13$ verschiedene Faktoren bestimmt werden, dann müssen $k = 91$ initiale Faktoreinstellungen ausgewertet werden. Dies ist in einigen industriellen Anwendungen nicht möglich.

Doch dies ist nicht das einzige Problem. Wir betrachten zudem ein „Henne-Ei Problem“, da die Wahl des optimalen Metamodells (z.B. linear oder quadratisch) von modellinhärenten Strukturen beeinflusst wird [21]: Gibt es keinen linearen Zusammenhang in den Daten (oder dem den Daten zugrunde liegenden Prozess), dann ist auch ein lineares Modell ungeeignet. Zu Beginn eines SPO Laufs liegen allerdings keine Daten vor – diese sollen ja gerade durch SPO möglichst optimal erzeugt werden. Letztendlich ist die Wahl einer passenden initialen Designgröße n nur durch Erfahrungen des Anwenders zu bestimmen. Liegen keine Erfahrungen vor, dann können vorexperimentelle Testläufe erste Hinweise geben.

Algorithm 1 Sequentielle Parameter Optimierung

```
1: procedure SPO( $\mathcal{A}_d, \mathcal{P}_d$ ) ▷ Algorithmus und Problem Design
2:   Select  $p \in \mathcal{P}_d$  and set  $t = 0$  ▷ Probleminstanz
3:    $X_A(t) = \{x^{(1)}, x^{(2)}, \dots, x^{(k)}\}$  ▷ Sample  $k$  initiale Punkte, z.B. mit LHS
4:   repeat
5:      $y_j^{(i)} = Y_j(x^{(i)}, p) \forall x^{(i)} \in X_A(t)$  and  $j = 1, \dots, r(t)$  ▷ Fitnessauswertung
6:      $\bar{Y}^{(i)}(t) = \sum_{j=1}^{r(t)} y_j^{(i)}(t) / r(t)$  ▷ Statistik (Mittelwert) für den  $i$ -ten Punkt
7:      $x_b$  with  $b = \arg \min_i (\bar{y}^{(i)})$  ▷ Bestimmung des besten Punktes
8:      $Y(\cdot) = \mathcal{F}(\beta, \cdot) + \mathcal{Z}(\cdot)$  ▷ Meta-Modell (hier: Kriging)
9:      $X_S = \{x^{(k+1)}, \dots, x^{(k+s)}\}$  ▷ Erzeuge  $s$  neue Samples,  $s \gg k$ 
10:     $y(x^{(i)}), i = 1, \dots, k + s$  ▷ Vorhersage der Fitness mittels Meta-Modell
11:     $I(x^{(i)})$  for  $i = 1, \dots, s + k$  ▷ Bestimme erwartete Verbesserung, cf. [16]
12:     $X_A(t + 1) = X_A(t) \cup \{x^{(k+i)}\}_{i=1}^m$  ▷ Füge  $m$  Punkte mit dem größten  $I(\cdot)$  dem
       Design hinzu
13:    if  $x_b(t) = x_b(t + 1)$  then
14:       $r(t + 1) = 2r(t)$  ▷ Anpassung der Wiederholungen
15:    end if
16:     $t = t + 1; k = k + m$  ▷ Zähler hochsetzen
17:  until Budget exhausted
18: end procedure
```

Eine formale Beschreibung der SPO ist in Alg. 1 dargestellt. Um Floor- oder Ceiling-Effekte zu vermeiden, muss eine geeignete Probleminstanz ausgewählt werden. Es muss sichergestellt werden, dass diese nicht zu leicht, aber auch nicht zu schwierig ist. In diesen beiden Fällen sind die Ergebnisse der Vergleiche der Vorhersagemodelle bedeutungslos. Wir benutzen, abhängig vom verwendeten Meta-Modell ($M^* \in \mathcal{M}^*$), unterschiedliche Initialdesigns wie *Latin Hypercube Designs* (LHD) oder fraktionelle faktorielle Designs. Nachdem das Vorhersagemodell ($M \in \mathcal{M}$) mit den k initialen Faktoreinstellungen gelaufen ist, wird das Meta-Modell M^* berechnet, um weitere, interessante Designpunkte zu bestimmen. Weitere m Punkte mit der größten erwarteten Verbesserung werden dem Design hinzugefügt. Hierbei sollte m im Vergleich zu s relativ klein sein. Die Update-Regel für die Anpassung der Anzahl der Wiederholungen ist nur erforderlich bei stochastischen Verfahren wie z.B. NN. Bei den deterministischen Verfahren wird nur eine Auswertung benötigt ($r(t) = r = 1$). Letztendlich gibt es je nach Optimierungsziel unterschiedliche Terminierungsbedingungen. In der dritten Untersuchung haben wir die Größe des Initialdesigns bei konstantem Budget (500 Berechnung des Vorhersagemodells M) systematisch variiert. Diese Versuchsserie wurde zudem für unterschiedliche Meta-Modelle M^* durchgeführt. Im Einzelnen wurden ein Polynom nullter Ordnung (in der Literatur auch als ordinary Kriging bezeichnet), ein Polynom erster und ein Polynom zweiter Ordnung benutzt. Zusätzlich wurde ein baumbasiertes Meta-Modell untersucht. Die Modelle wurden mit `regpoly0`, `regpoly1`, `regpoly2` und `tree` bezeichnet. Die Ergebnisse sind graphisch in Bild 7 dargestellt.

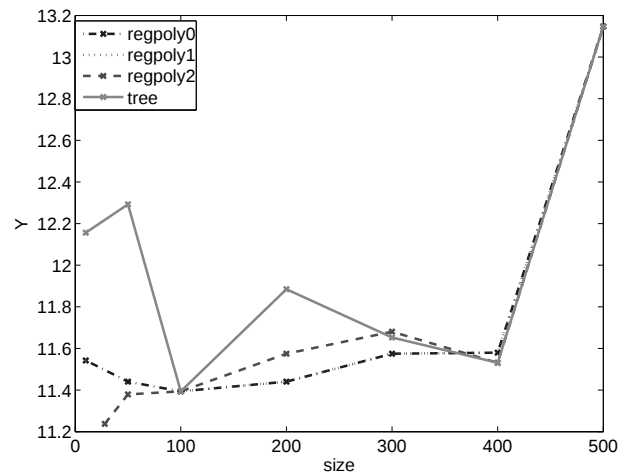


Bild 7: Auswirkungen unterschiedlicher initialer Designgrößen auf das SPO Ergebnis.

6 Diskussion

6.1 Vergleich der Prognosemodelle (M)

Es wurde deutlich, dass Ansätze wie NN, ESN, FIR oder NARX mit den Daten gar nicht oder nur teilweise zurecht kommen. Im Vergleich zu den NN-Modellen lassen sich die FIR-Modelle sehr schnell und einfach implementieren [9]. Die Kernroutine besteht aus ca. 10 Zeilen Programmcode. Die Berechnung der Vorhersagewerte ist entsprechend schnell, FIR-Modelle sind zudem sehr stabil. Interessant ist, dass die FIR-Modelle bessere Ergebnisse erzielen als alle anderen Modelle, wenn die Faktoreinstellungen zufällig (allerdings aus einer vom Anwender als sinnvoll beschriebenen ROI) gewählt werden, siehe Tab. 2. Stellt man jedoch die Faktoren ein, z.B. indem man Erfahrungen aus Testläufen heranzieht oder SPO anwendet, dann verlieren FIR-Modelle im Vergleich zu ODE und INT eindeutig. Damit liefern die FIR-Modelle im Portfolio eine gute Basis für einen Vergleich mit komplexeren Modellen (NN, NARX, ODE und INT). Modelle, die nicht in der Lage sind, die FIR Resultate zu übertreffen, dürfen ohne Weiteres als ungeeignet bezeichnet werden.

In der Anfangsphase der Simulation, z.B. während der ersten 25 Tage (siehe Bild 3), erzielten NARX-Modelle einen guten Vorhersagefehler. Das Modell ging aber dann in einen nahezu gesättigten Zustand (ca. 90% Füllstand) über, der konstant bis zum Ende der Simulation prognostiziert wurde. Ein ähnliches Verhalten zeigte auch ESN. Dieser Fehler trat systematisch auf, so dass wir NARX Modelle für diese Art von Daten als nicht geeignet halten. Zudem erfordern NARX-Modelle durch den Trainingsaufwand eine signifikant längere Laufzeit als alle anderen von uns betrachteten Modelle. Sicherlich gibt es effektivere Verfahren für Verbesserung der Faktoren des NARX Modells. Die von uns gewählte SPO Modellierung ist für zwei Faktoren überdimensioniert. SPO findet normalerweise Anwendung, falls drei oder mehr Faktoren zu optimieren sind. Eine sinnvolle Obergrenze liegt bei 20 Faktoren. Aber auch Modelle mit mehreren hundert oder gar tausend Faktoren können nach einem vorher durchgeführten Screening, z.B. *sequential bifurcation* [19], mit SPO optimiert werden. Um eine Vergleichbarkeit mit den anderen Techniken zu gewährleisten, haben wir uns auch für das NARX-Modell für die Standard-SPO Vorgehensweise

entschieden. Jedes der von uns betrachteten Modelle (FIR, NARX, ODE und INT) wurde mit der gleichen Systematik optimiert. Die modellbasierten Ansätze ODE und INT zeigen hier deutliche Vorteile. Bereits mit den manuell, vor dem SPO-Einsatz gefundenen Parameter erreicht man mit RMSE 13.60 bzw. 12.72 eine starke Verbesserung (61% bzw. 64%) gegenüber FIR und NARX. Der Einsatz von SPO ist wichtig, denn diese bereits guten Werte konnten durch SPO nochmals um 26% gesteigert werden (s. Tab. 2).

6.2 Vergleich der Meta-Modelle (M^*)

Der Modellentwurfsprozess wurde durch systematische Experimente mittels SPO zielgerichtet gesteuert. Da SPO selbst verschiedene Modelle und Designs zur Verfügung stellt, ist eine Analyse der Abhängigkeit der Ergebnisse der Meta-Modellierung von dem verwendeten Meta-Modell und Design von großem Interesse. Die Meta-Modelle aus der Klasse der Kriging-Modelle (also `regpoly0`, `regpoly1`, `regpoly2`) zeigten ein ähnliches Verhalten. Die besten Ergebnisse wurden mit kleinen Initialdesigns erzielt, siehe Bild 7. Für das baumbasierte Modell ist die Situation nicht so eindeutig: Ein Trend ist hier nicht zu erkennen. Die folgende Aussage trifft wiederum auf alle Meta-Modelltypen (Kriging und Regressionebäume) zu: Die von SPO verwendete sequentielle Vorgehensweise ist von Vorteil für die Optimierung. Für eine initiale Designgröße von $k = 500^9$ wurde das schlechteste Ergebnis erzielt.

7 Zusammenfassung

Zeitreihen mit intermittierender Aktivität treten in zahlreichen Modellierungs- und Prognoseanwendungen der Automatisierungstechnik auf und sind schwierig zu prognostizieren. Wir betrachteten exemplarisch einen Anwendungsfall aus der Wasserwirtschaft. Es liegt nicht an ungünstig gewählten Parametern, dass Modelle wie FIR oder NARX eine schlechtere Performance zeigen. Mit hoher Wahrscheinlichkeit sind FIR und NARX für diese Modellierungsaufgabe nicht gut geeignet. Die bereits recht guten manuellen Ergebnisse für ODE und INT (Faktor 2 bis 6 besserer RMSE als FIR bzw. NARX) konnten durch SPO nochmals um 30-40% gesteigert werden (s. Tabelle 2). Durch die SPO erschließen sich oftmals überraschende Einsichten über die Modellierung sowie Ansätze, wie sich möglicherweise ein verbessertes Modell bauen läßt.

Die SPO selbst zeigte sich als robustes Optimierungsverfahren, das mit einer kleinen Samplezahl gestartet werden sollte. Anschließend sollten weitere Punkte hinzugenommen werden. Das Meta-Modell ermöglicht eine signifikante Verbesserung der Optimierung. Dabei ist die Wahl des Meta-Modells als nicht sehr kritisch zu bewerten. Das baumbasierte Modell zeigte größere Schwankungen in der Optimierungsqualität, ist aber ca. doppelt so schnell zu berechnen wie die Kriging Modelle. Es kann zusätzlich mit kategorischen Daten umgehen – dies ist für viele Modelle ein entscheidender Vorteil. Insgesamt kann der SPO-basierte Findungsprozess als erfolgreich bezeichnet werden. SPO konnte die manuell gefundenen Modelldesigns signifikant verbessern. Aus den Tab. 2 und 3 ist

⁹Zur Erinnerung: Insgesamt stand ein Budget von 500 Modellberechnungen zur Verfügung. Dieses war somit bereits ohne Meta-Modellbildung ausgeschöpft.

aber auch ersichtlich, dass das vom Entscheider (Experten, Anwender) bevorzugte Modell (INT) mit dem von SPO gefundenen Modell überstimmt. Lediglich wurde von SPO eine bessere Einstellung der Modellfaktoren gefunden.

Danksagung: Für die Bereitstellung von Daten und Fotos aus dem KANNST-Projekt danken wir Prof. Dr. Michael Bongards und Dipl.-Ing. Tanja Hilmer. – Teile dieser Arbeit wurden durch die FH Köln im Rahmen des Forschungsschwerpunktes COSA gefördert.

Literatur

- [1] Beyth-Marom, R.; Fischhoff, B.; Quadrel, M.; Furby, L.: Teaching decision making to adolescents: A critical review. In: *Teaching Desision Making to Adolescents* (Baron, J.; Brown, R., Hg.), S. 19–60. Mahwah, NY: Erlbaum. 1991.
- [2] Bartz-Beielstein, T.; Parsopoulos, K. E.; Vrahatis, M. N.: Design and analysis of optimization algorithms using computational statistics. *Applied Numerical Analysis and Computational Mathematics (ANACM)* 1 (2004) 2, S. 413–433.
- [3] Bartz-Beielstein, T.; Lasarczyk, C.; Preuß, M.: Sequential Parameter Optimization. In: *Proceedings 2005 Congress on Evolutionary Computation (CEC'05), Edinburgh, Scotland* (McKay, B.; et al., Hg.), Bd. 1, S. 773–780. Piscataway NJ: IEEE Press. 2005.
- [4] Bartz-Beielstein, T.: *Experimental Research in Evolutionary Computation—The New Experimentalism*. Natural Computing Series. Berlin, Heidelberg, New York: Springer. 2006.
- [5] Konen, W.; Zimmer, T.; Bartz-Beielstein, T.: Optimierte Modellierung von Füllständen in Regenüberlaufbecken mittels CI-basierter Parameterselektion. *at – Automatisierungstechnik* (2008). Eingereicht.
- [6] Bongards, M.: Online-Konzentrationsmessung in Kanalnetzen – Technik und Betriebsergebnisse. Techn. Ber., Cologne University of Applied Sciences. 2007.
- [7] Grüning, H.: Abflusssteuerung - quo vadis? *KA Korrespondenz Abwasser, Abfall* 55 (2008), S. 358–364.
- [8] Hilmer, T.; Bongards, M.: Integrierte Steuer- und Regelstrategien für Kanalnetz und Kläranlage. Techn. Ber., Cologne University of Applied Sciences. 2008.
- [9] Bartz-Beielstein, T.; Bongards, M.; Claes, C.; Konen, W.; Westenberger, H.: Datenanalyse und Prozessoptimierung für Kanalnetze und Kläranlagen mit CI-Methoden. In: *Proc. 17th Workshop Computational Intelligence* (Mikut, R.; Reischl, M., Hg.), S. 132–138. Universitätsverlag, Karlsruhe. 2007.
- [10] U.S. Environmental Protection Agency: Storm Water Management Model. www.epa.gov/ednnrmrl/models/swmm/index.htm, Online; Stand 16.08.08. 2008.
- [11] Siegelmann, H. T.; Horne, B. G.; Giles, C. L.: Computational capabilities of recurrent NARX neural networks. Techn. Ber. UMIACS-TR-95-12 and CS-TR-3408, University of Maryland. URL citeseeer.ist.psu.edu/siegelmann97computational.html. 1995.
- [12] Jaeger, H.: The echo state approach to analysing and training recurrent neural networks. Techn. Ber. 148, Fraunhofer Institute for Autonomous Intelligent Systems (AIS), Sankt Augustin. 2001.
- [13] Bellmann, R. E.: *Adaptive Control Processes*. Princeton University Press. 1961.
- [14] Kleijnen, J. P. C.: *Statistical Tools for Simulation Practitioners*. New York NY: Marcel Dekker. 1987.
- [15] Sacks, J.; Welch, W. J.; Mitchell, T. J.; Wynn, H. P.: Design and analysis of computer experiments. *Statistical Science* 4 (1989) 4, S. 409–435.
- [16] Santner, T. J.; Williams, B. J.; Notz, W. I.: *The Design and Analysis of Computer Experiments*. Berlin, Heidelberg, New York: Springer. 2003.
- [17] Bartz-Beielstein, T.; Preuss, M.: Moderne Methoden zur experimentellen Analyse evolutionärer Verfahren. In: *Proc. 16th Workshop Computational Intelligence* (Mikut, R.; Reischl, M., Hg.), S. 25–32. Universitätsverlag, Karlsruhe. 2006.

- [18] Konen, W.; Bartz-Beielstein, T.; Westenberger, H.: Computational Intelligence und Data Mining – Datenanalyse und Prozessoptimierung am Beispiel Kläranlagen. Techn. Ber., Cologne University of Applied Sciences. 2007.
- [19] Kleijnen, J. P. C.: *Design and analysis of simulation experiments*. New York NY: Springer. 2008.
- [20] Lophaven, S.; Nielsen, H.; Søndergaard, J.: Aspects of the Matlab Toolbox DACE. Techn. Ber. IMM-REP-2002-13, Informatics and Mathematical Modelling, Technical University of Denmark, Copenhagen, Denmark. 2002.
- [21] Bartz-Beielstein, T.: Review: Design and Analysis of Simulation Experiments by Jack P. C. Kleijnen. *ICS News* (2008). <http://computing.society.informs.org/newsletter.php>.