

Online-Magazin

Zeitschrift für numerische Simulationsmethoden und angrenzende Gebiete: FEM, CFD, MKS, VR / VIS, PROZESS, SDM

FACHBEITRÄGE

Maschinelles Lernen

- Eine automatisierte Machine Learning Umgebung für industrielle Anwendungen
- Schnellere virtuelle Produktentwicklung und weniger Jobabbrüche: Beschleunigung der CFD-Simulation durch zuverlässige Vorhersagen von Laufzeiten

Keramikentwicklung

- ICME für die Keramikentwicklung

Radialverdichter

- Generierung und Parametrisierung von Radialverdichterkennlinien auf Basis neuronaler Netze

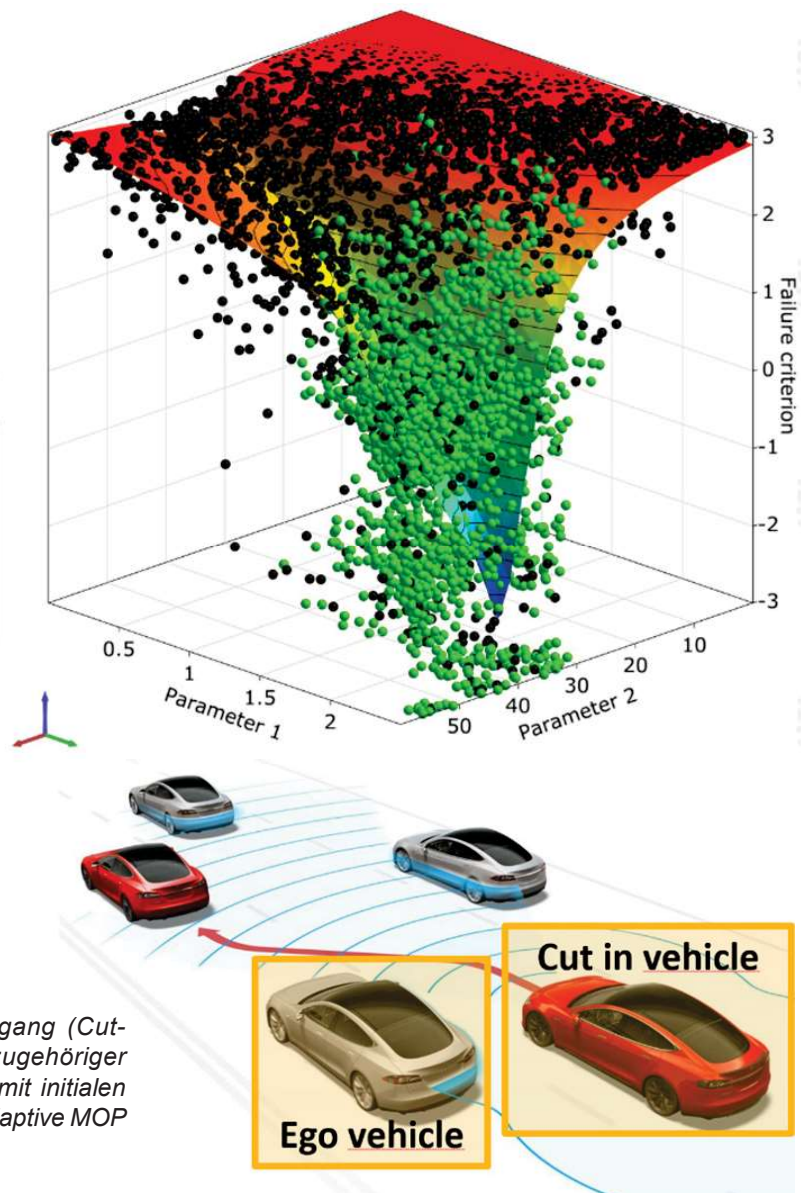


Abb. rechts: Untersuchter Einschervorgang (Cut-In) eines autonomen Fahrzeuges mit zugehöriger Approximation des Versagenskriteriums mit initialen Datenpunkten (schwarz) und durch das Adaptive MOP erzeugte zusätzliche Stützstellen (grün).

Alle bisherigen Ausgaben kostenlos zum Download unter: www.nafems.org/magazin

Sponsoren dieser Ausgabe:

Eine automatisierte Machine Learning Umgebung für industrielle Anwendungen

Thomas Most, Lars Gräning, Johannes Will
Ansys Dynardo GmbH

1 Zusammenfassung

In diesem Beitrag stellen wir Weiterentwicklungen des etablierten Machine Learning Ansatzes, des Metamodells Optimaler Prognosefähigkeit (MOP) [1], [2] vor. In diesem MOP-Rahmen werden verschiedene klassische Metamodelle wie lineare Regression, Gaussprozesse bzw. Kriging direkt an gegebenen Datensätzen trainiert und bewertet. Die Datensätzen können aus beliebigen analytischen oder numerischen Berechnungen und Simulationen oder auch aus Messungen stammen. Die dabei analysierten verschiedenen Approximationsmodelle werden hinsichtlich ihrer Prognosefähigkeit beurteilt, um eine erfolgreiche Weiterverwendung zum Beispiel als Digitaler Zwilling, innerhalb einer Design Optimierung oder auch zur Unsicherheitsquantifizierung im Rahmen einer Robustheitsbewertung zu ermöglichen.

Zusätzlich zu den klassischen Regressionsverfahren können verschiedene Typen von tiefen neuronalen Netzwerken (Deep Learning Networks) in Abhängigkeit der Charakteristik der Daten bewertet werden. Typischerweise benötigen Deep Learning Modelle eine Vielzahl von Datenpunkten, um sogenannte Overfitting-Effekte zu vermeiden. In diesem Beitrag stellen wir einen neuen Regularisierungsansatz vor, der es ermöglicht, diese Deep Learning Modelle hinsichtlich ihrer Topologie optimal zu definieren und auch für wenige Datensätze effizient anzuwenden, wobei eine gute Generalisierung der Modelle ohne Overfitting im Fokus steht. Innerhalb eines automatisierten „Wettbewerbes“ der klassischen und der Deep Learning Modelle werden dann diejenigen Modellansätze gewählt, die für den konkreten Datensatz ausreichend genau jedoch minimal komplex sind. Dieser Ansatz des Metamodells Optimaler Prognosefähigkeit wird durch einen adaptiven Ansatz erweitert, bei dem die Modellqualität zielgerichtet durch neue Datenpunkte in den Bereichen verbessert wird, wo eine Verbesserung der Modellapproximation einen optimalen Nutzen für die konkrete Anwendung besitzt.

2 Maschinelles Lernen

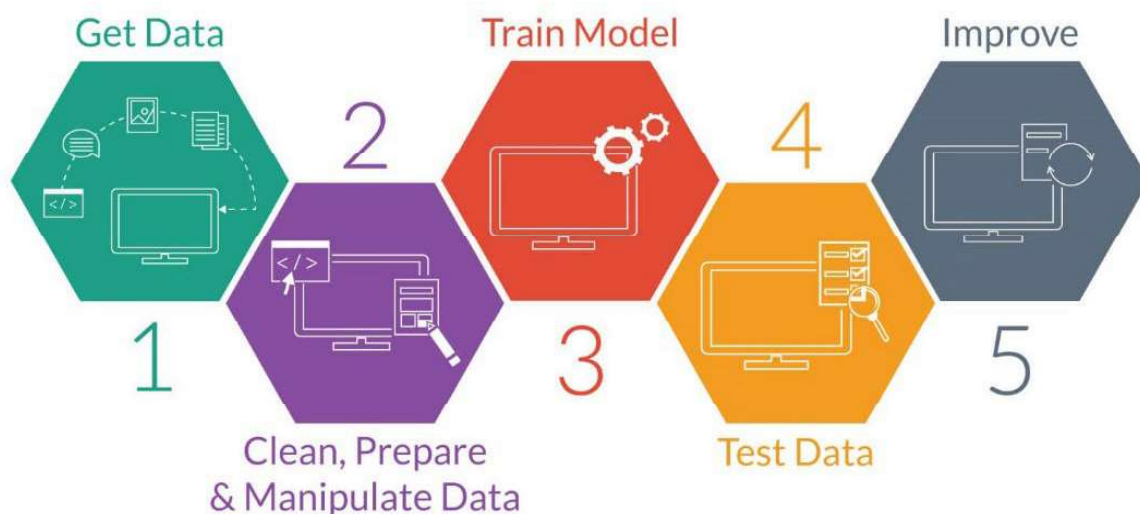


Abb. 1: Typische Schritte eines maschinellen Lernen Ansatzes (www.machinelearning-blog.com)

Im maschinellen Lernen unterscheidet man oft zwischen klassischen Antwortflächenverfahren und Deep Learning Ansätzen. Im klassischen Ansatz werden die Daten entsprechend aufbereitet und dann zum Training der verschiedenen Modellansätze angewendet. Auf Basis eines unabhängigen Testdatensatzes kann die Qualität der einzelnen Modelle bewertet werden und eine adaptive Verbesserung erfolgen (Abbildung 1). Man geht dabei davon aus, dass die relevanten Eingangsgrößen in Regel manuell festgelegt werden, um in mehreren meist manuellen Iterationsschritten ein optimales Ersatzmodell zu erhalten.

In klassischen globalen Regressionsverfahren wird in der Regel ein Satz von Ansatz- oder Basisfunktionen definiert und die zugehörigen, unbekannten Modellkoeffizienten z.B. durch Minimierung von Fehlerquadraten an den Datenpunkten kalibriert [3]. Die Qualität des Approximationsmodells hängt dabei direkt mit der Wahl der Ansatzfunktionen zusammen. Diese globalen Regressionsmodelle sind allerdings oftmals nicht in der Lage, komplexere Zusammenhänge zufriedenstellend genau abzubilden. Eine höhere Flexibilität weisen lokale Regressionsmodelle, wie Moving Least Squares [4] und auch Gaussprozess-Modelle bzw. Kriging [5] auf. Bei diesen lokalen Verfahren dienen die Datenpunkte direkt als Stützstellen einer lokalen Interpolation. Die Ermittlung der optimalen internen Modellparameter wie Einflussradien oder Korrelationslängen ist dabei der kritische Teil des Trainings dieser Modelle, um Overfitting zu vermeiden. Bei Vergrößerung des Datensatzes können diese Modelle eine immer bessere Vorhersage ermöglichen, sind aber oftmals beim Training sehr aufwendig.

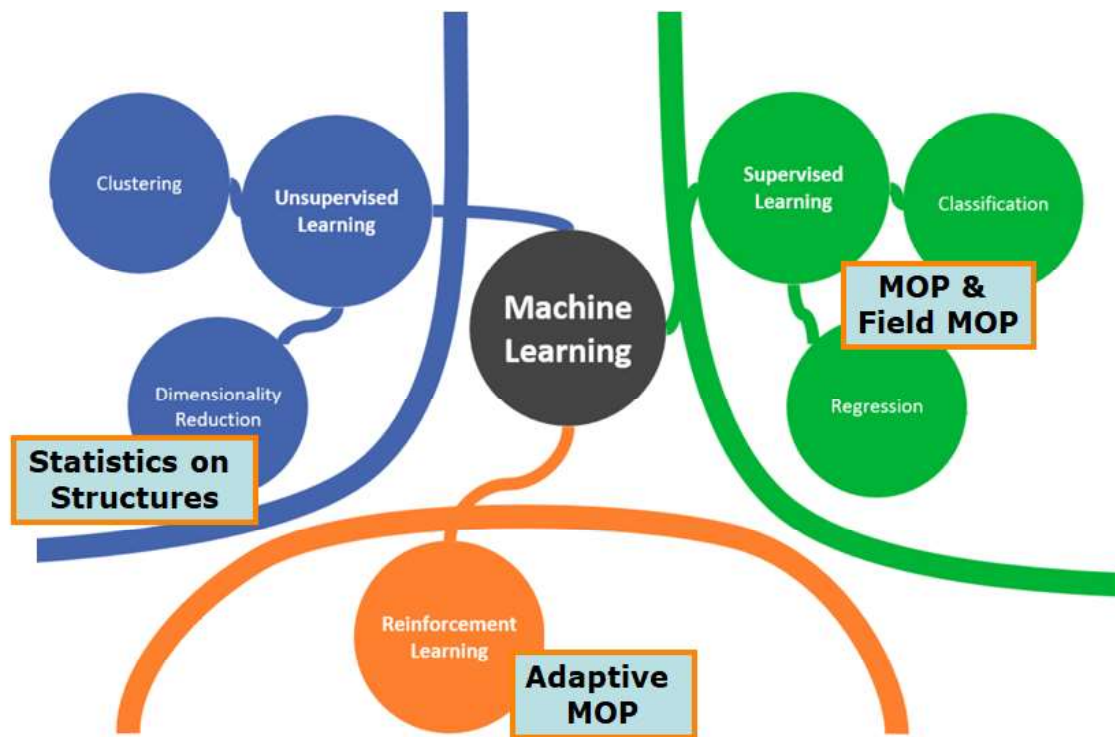


Abb. 2: Überblick über verschiedene Machine Learning Ansätze (www.morethandigital.info)

Der Machine Learning Ansatz kann nun als eine Erweiterung der klassischen Ansätze gesehen werden, bei dem die notwendigen Basistherme und Ansatzfunktionen nicht manuell vorgegeben werden, sondern durch den Trainingsvorgang für die jeweilige Komplexität des Zusammenhangs automatisch gewählt werden. Hierbei lassen sich verschiedene Vorgehensweisen unterscheiden, welche in Abbildung 2 dargestellt sind.

Im Unsupervised Learning (unüberwachtes Lernen) werden Zusammenhänge nur durch die Analyse der Datenpunkte im Raum der Eingangsparameter abgebildet. Dazu zählen zum Beispiel Cluster-Suchverfahren und Dimensionsreduktionsansätze. Da Cluster-Verfahren meist nur den Ort im Eingangsraum zur Unterscheidung nutzen, sind diese oft nur bei wenigen Eingangsvariablen effizient anwendbar. Im Gegensatz dazu sind Dimensionsreduktionsansätze auch bei sehr vielen Eingangsparametern effizient, wenn diese Eingangsgrößen signifikante Zusammenhänge untereinander aufweisen.

Im Supervised Learning, dem überwachten Lernen, werden die Zusammenhänge zwischen Eingangs- und Antwortgrößen aus analytischen und numerischen Berechnungen oder auch Messungen direkt miteinander verknüpft. Meist werden die einzelnen Antwortgrößen durch verschiedene Approximationsmodelle abgebildet. Man unterscheidet dabei zwischen skalaren und Feldgrößen. Unter Anwendung der Dimensionsreduktionsverfahren können auch Feldgrößen wie z.B. Verformungen und Spannungen aus Finite Elemente Berechnungen mit Hilfe weniger Basisvariablen als skalare Subgrößen abgebildet und approximiert werden. Solche Strategien werden unter anderen in Statistics of Structures (SoS) zur Approximation von Feldgrößen angewendet [6]. Generell wird im überwachten Maschinellen Lernen eine automatische Anpassung der Basisfunktionen der Approximationsmodelle angestrebt. Dies geschieht zum Beispiel in der Support Vector Klassifikation und auch Regression durch schrittweises Entfernen einzelner Basistherme innerhalb des Trainingsprozesses. Der klassische Ansatz des Metamodells Optimaler Prognosefähigkeit (MOP) selektiert für den gegebenen Datensatz das am besten geeignete Regressionsmodell für die jeweilige Antwortgrößen indem die verschiedenen trainierten Metamodelle durch eine geeignetes Fehlermaß verglichen werden. Da auch die relevanten Eingangsparameter durch verschiedene Filterschritte automatisch ermittelt werden, kann dieser Ansatz ebenfalls als überwachtes Lernen eingeordnet werden.

Beim Bestärktem Lernen (Reinforcement Learning) von Metamodellen werden die Datenpunkte entsprechend ihrer Güte häufig über Verlustfunktionen bewertet. So können zum Beispiel in einer Designoptimierung die Datenpunkte mit besserer Performance höher gewichtet werden und in einem adaptiven Vorgehen neue Datenpunkte dort konstruiert werden, wo eine maximale Verbesserung erwartet wird. Im Adaptiven Metamodell Optimaler Prognosefähigkeit (AMOP) [7] wird dieser Ansatz dazu verwendet, um einerseits das Metamodell durch neue Datenpunkte dort zu verbessern, wo die Approximationsqualität als nicht ausreichend bewertet wird oder aber um neue Designbereiche mit einer erwarteten Verbesserung von Zielfunktionen zu erkunden.

3 Metamodell Optimaler Prognosefähigkeit (MOP)

Beim Metamodell Optimaler Prognosefähigkeit (MOP) [1,2] wird ein automatisierter Prozess des maschinellen Lernens durchgeführt, bei dem für einen gegebenen Datensatz für jede Antwortgröße das am besten geeignete Approximationsmodell aus einer Bibliothek verfügbarer Modelle ermittelt wird. Um die Approximationsqualität zu verbessern, werden dabei die für die jeweilige Antwortgröße relevanten Eingangsgrößen ermittelt. Da oftmals nur wenige Eingangsgrößen notwendig sind, um die Zusammenhänge ausreichend genau zu beschreiben, kann durch dieses Vorgehen auch eine Anwendbarkeit bei relativ vielen Eingangsparametern und wenigen Datenpunkten realisiert werden.

Ein Kernpunkt im MOP-Ansatz ist der Coefficient of Prognosis (CoP), der als Qualitätsmaß zur Bewertung der Prognosefähigkeit verwendet wird. Im Gegensatz zu klassischen statistischen Maßen, wie dem Bestimmtheitsmaß bzw. dem Coefficient of Determination (CoD) [3], bei dem die quadrierten Abweichungen nur an den Trainingsdaten ermittelt werden, wird beim CoP die Prognosefähigkeit durch Kreuzvalidierung abgeschätzt. Dabei wird der originale Datensatz in mehrere Teildatensätze zerlegt, welche jeweils genau nur einmal als Testdatensatz zur Berechnung der quadrierten Abweichungen verwendet werden. Als Summe der einzelnen Analyseschritte erhält man für jeden Datenpunkt einen Prognosefehler, welcher im CoP zusammengefasst wird. Wichtig für die Robustheit dieses Qualitätsmaßes ist dabei die repräsentative Aufteilung des Gesamtdatensatzes auf die Teildatensätze und das entsprechende Training der Modelle auf den Teildatensätzen. In Abbildung 3 ist das Vorgehen zur CoP Berechnung exemplarisch für zwei Teildatensätze dargestellt. Dabei wird der Datensatz in zwei Teildatensätze zu jeweils 50% in Trainingsdaten und Testpunkte aufgeteilt. In dem dargestellten Interpolationsmodell ist eine perfekte Abbildung der Trainingsdaten möglich, allerdings ist die Vorhersage an den Testdaten weniger genau. Der CoD beurteilt nun die Fitting-Qualität und der CoP die Prognosefähigkeit. Weisen beide Maße eine größere Differenz auf, ist das ein Indiz für ein Overfitting des untersuchten Modells.

Auf Basis des CoPs können nun die verschiedenen Approximationsmodellansätze miteinander verglichen werden. Als Ergebnis des MOP-Trainingsvorgangs wird das Modell mit maximaler Prognosefähigkeit, jedoch mit minimaler Modellkomplexität ermittelt. Dabei werden verschiedene Basistherme, verschiedene Ansätze in den Korrelationsfunktionen sowie auch verschiedene Subräume der Eingangsparameter für den jeweiligen Modellansatz untersucht. Eingangsgrößen, die nur einen sehr kleinen Beitrag zur Prognosefähigkeit haben, werden unter Berücksichtigung eines definierten Schwellwertes aus den Modellen eliminiert.

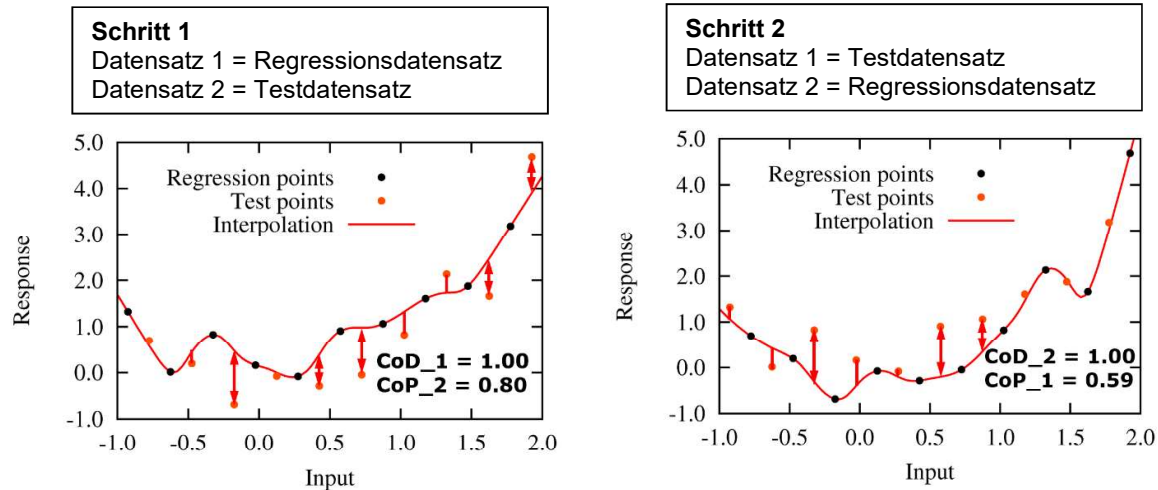


Abb. 3: Ermittlung des Coefficient of Optimal Prognosis (CoP) über Kreuzvalidierung, hier exemplarisch für zwei Teildatensätze dargestellt

4 Künstliche Neuronale Netzwerke

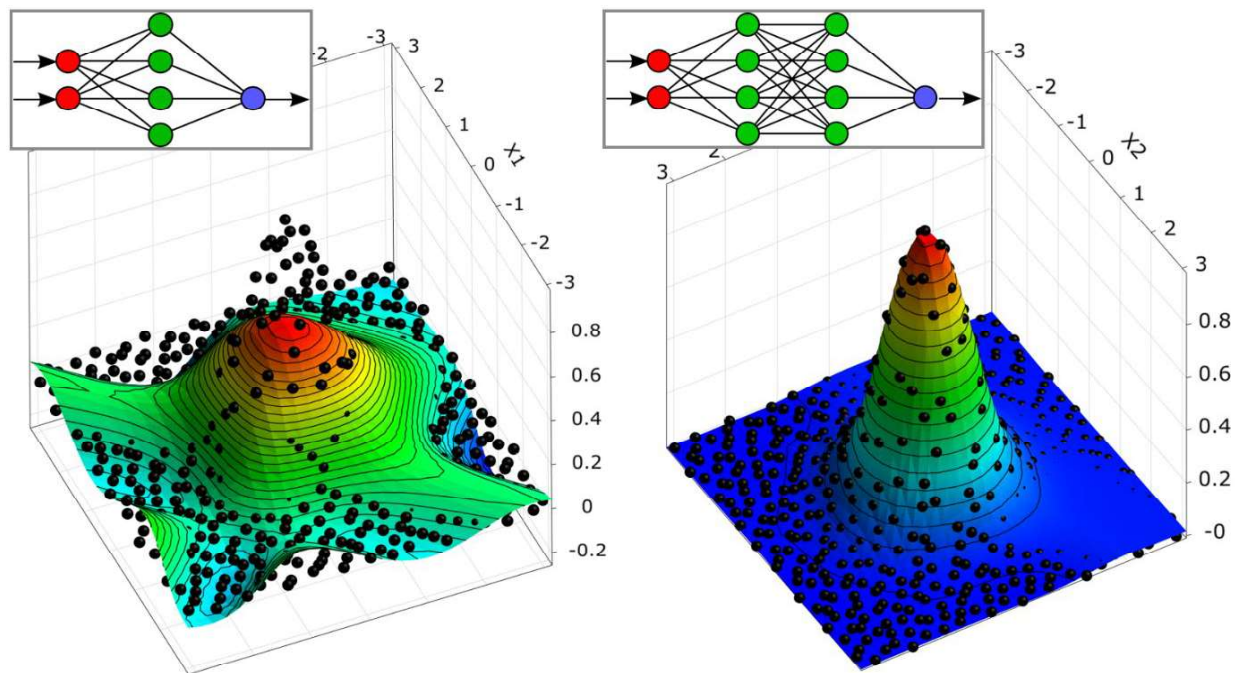


Abb. 4: Approximation einer zweidimensionalen Gauss-Funktion mit einer versteckten Schicht mit vier Neuronen (links) und zwei versteckten Schichten mit jeweils vier Neuronen (rechts)

Im Gegensatz zu klassischen Regressionsmodellen werden bei künstlichen neuronalen Netzwerken keine Basisfunktionen mit entsprechender hoher Komplexität verwendet, um nichtlineare Zusammenhänge in den Daten abzubilden, sondern durch die Verknüpfung einer Vielzahl von einfachen nichtlinearen Operationen kann

eine beliebige Komplexität in der Approximationsfunktion erreicht werden. In sogenannten Neuronen stellen die Aktivierungsfunktionen eine eindimensionale Funktion mit unbekannten Koeffizienten dar. Die Neuronen werden dabei in mehreren Schichten, sogenannten Layern, angeordnet, wobei man bei mehr als einem versteckten Layer von Deep Learning Netzwerken spricht. Wird nur ein versteckter Layer verwendet, können Interaktionen zwischen den Eingangsgrößen nur ungenügend abgebildet werden, wie in Abbildung 4 dargestellt. Bei mehreren versteckten Layern können jedoch sehr komplexe, mehrdimensionale Zusammenhänge sehr flexibel abgebildet werden. Die Kopplung zwischen den Neuronen der einzelnen Layer erfolgt über eine Linearkombination mit zusätzlichen unbekannten Koeffizienten. In dem Trainingsvorgang eines Netzwerkes mit einer vordefinierten Anzahl von Layern und Neuronen werden dann alle unbekannten Koeffizienten z.B. durch eine Minimierung der Fehlerquadrate bzgl. der Trainingsdaten durch einen Optimierungsprozess ermittelt. Die Approximation eines klassischen neuronalen Netzwerkes ist eine globale Funktion im Raum der Eingangsparameter, die einzelne skalare oder sogar vektorielle Antwortgrößen beschreiben kann.

In den vergangenen Jahren sind eine Vielzahl von Deep Learning Bibliotheken veröffentlicht worden, die den Trainingsprozess sehr effizient gestalten und somit die Analyse sehr großer Datenmengen mit komplexen Modellen zulassen. Im Bereich der numerischen Simulation ist allerdings die Anzahl der verfügbaren Datensätze bzw. der Modellberechnung oftmals stark begrenzt. Da in einem komplexeren Deep Learning Modell oftmals mehrere hundert unbekannte Koeffizienten trainiert werden müssen, ist für wenige Datensätze die Anwendbarkeit dieser Modelle durch ein erhebliches Overfitting-Risiko wiederum eingeschränkt. In solchen Fällen ist der MOP-Ansatz, bei dem die irrelevanten Eingangsgrößen aus dem Modell entfernt werden, sehr vielversprechend [8]. Allerdings hat auch die Wahl der Netzwerkarchitektur, also die Anzahl der versteckten Layer und Neuronen sowie die Neuronentypen, Lernrate usw. einen sehr großen Einfluss auf die Prognosefähigkeit der trainierten Modelle.

In diesem Beitrag wurde ein neuer Regularisierungsansatz, das Smart Layout [9,10], angewendet, der es ermöglicht, innerhalb des Trainingsvorganges eine optimale Netzwerkarchitektur automatisch zu identifizieren. In der Verlustfunktion des Trainings werden dabei Regularisierungsterme berücksichtigt, die die unwichtigen Neuronen in den Eingangs- und versteckten Schichten schrittweise eliminieren. Dabei wird eine optimale Prognosefähigkeit für den untersuchten Datensatz innerhalb eines einzigen Trainingslauf durch die modifizierte Verlustfunktion sichergestellt. Der Anwender muss dabei nur eine initiale Netzwerkarchitektur definieren, auf deren Basis die optimale Struktur ermittelt wird. Ein besonderes Augenmerk benötigt dabei die Berechnung des CoP auf Basis der Kreuzvalidierung. Um Overfitting zu vermeiden, darf der jeweilige Testdatensatz nicht beim Training der Teilmodelle berücksichtigt werden. In den numerischen Beispielen dieses Beitrages wird der Effekt der Regularisierungstechnik verdeutlicht.

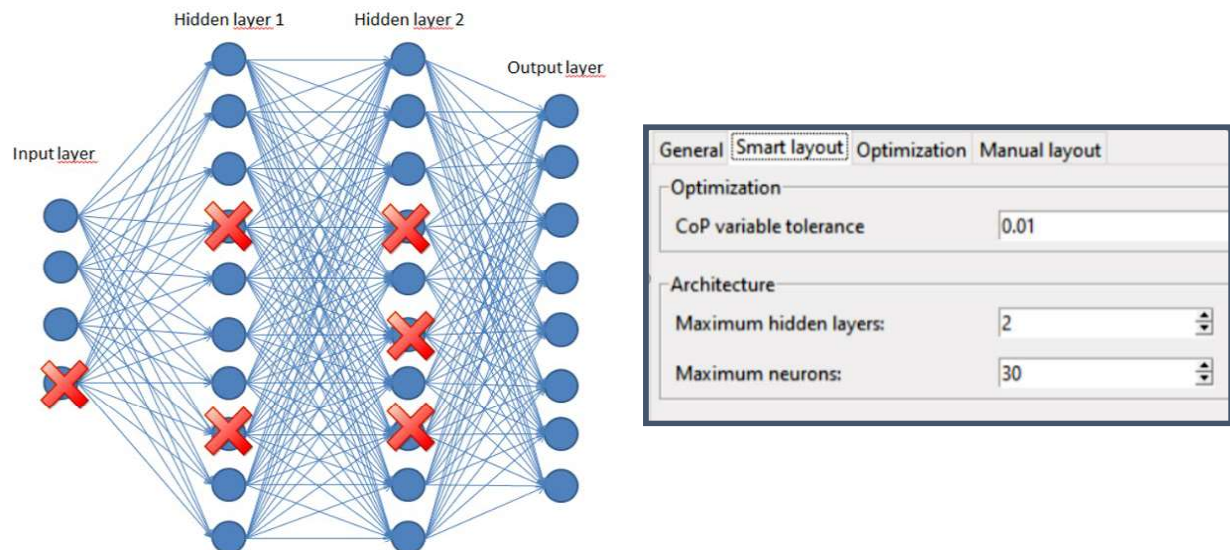


Abb. 5: Automatisierter Regularisierungsansatz zur Identifikation der optimalen neuronalen Netzwerkarchitektur mit minimalen Nutzereinstellungen

5 Adaptives Metamodell Optimaler Prognosefähigkeit (AMOP)

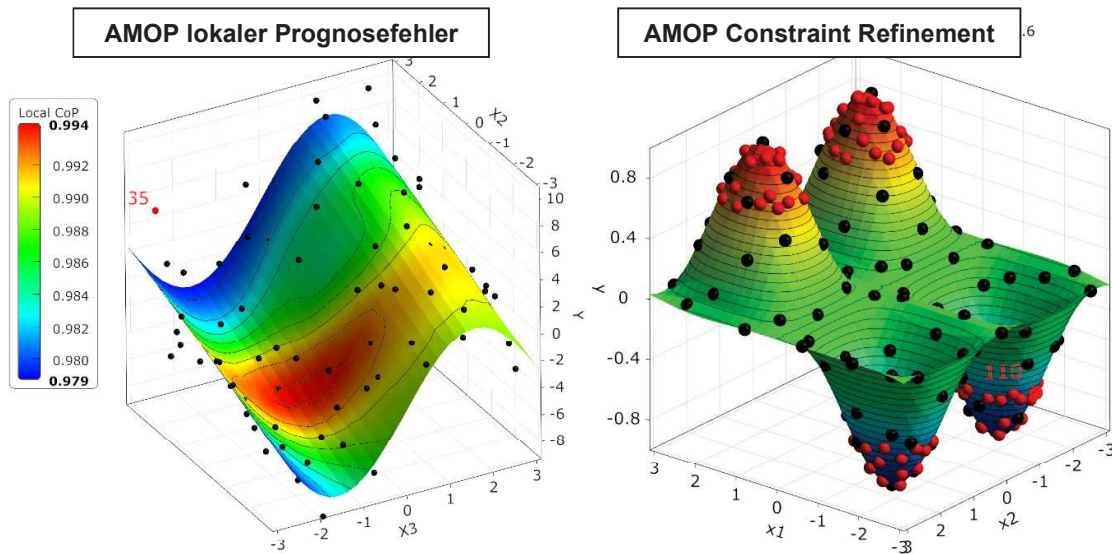


Abb. 6: Adaptives Metamodell Optimaler Prognosefähigkeit: lokale Prognosefehler als Grundlage für ein effizientes Aktualisierungsverfahren durch neue Datenpunkte

Im adaptiven MOP-Ansatz [7] wird das initial trainierte Machine Learning Model iterativ durch neue Datenpunkte in den Bereichen verbessert, wo die benötigte Genauigkeit für die spezifische Anwendung noch nicht erreicht ist. Grundlage dabei stellt die Approximation des Vorhersagefehlers aus der Kreuzvalidierung einer Antwortgröße im Designraum dar, wie in Abbildung 6 links zu sehen. Auf Basis dieses Fehlerschätzers kann einerseits die lokale Qualität des Modells verbessert werden oder aber bestimmte Bereiche, die z.B. für Zielfunktionen interessant sind, durch neue Designs adaptiert werden. Durch die Definition von Nebenbedingungen kann die Verbesserung auch auf bestimmte Bereiche eingeschränkt werden, welche für die jeweilige Anwendung von besonderem Interesse sind. In Abbildung 6 rechts ist ein solcher Fall dargestellt, bei dem nur die Bereiche mit maximalen und minimalen Funktionswerten adaptiert wurden.

6 Beispiele

Im ersten Beispiel wurde die Rotationssteifigkeit eines Turbinenlaufrades analysiert. Als Datensatz standen dafür die Simulationsergebnisse unter Berücksichtigung von 13 Geometrieparametern zur Verfügung. 176 Stichproben eines Latin Hypercube Sampling wurden erzeugt und per multidisziplinärer Simulation durch Finite Elemente, CFD und Fatigue-Berechnungen ausgewertet. Als Ergebnisgrößen wurden die Rotations- und Axialsteifigkeit, die Masse, der isotrope Wirkungsgrad und Massenstrom sowie Lebensdauerabschätzungen berechnet. Anhand dieser Industrieanwendung soll nun die Genauigkeit und der Einfluss des Regularisierungsansatzes der vorgestellten Deep Learning Modelle demonstriert werden.

In Abbildung 7 sind die Approximationsfunktionen der Deep Learning Modelle im Subraum der beiden wichtigsten Eingangsgrößen dargestellt. Dafür wurde das neuronale Netzwerk zunächst im Gesamtraum der 13 Parameter trainiert. Im linken Teil der Abbildung ist ersichtlich, dass das Approximationsmodell den stark nichtlinearen Funktionsverlauf nur ungenügend repräsentieren kann. Durch die hohe Flexibilität des Approximationsmodells ist der CoD annähernd 100% wobei die Vorhersagequalität nur 79% beträgt. Im Trainingsverlauf des Deep Learning Modells ist zu erkennen, dass das Early-Stopping sehr früh den Trainingsprozess abbricht, um zu starkes Overfitting zu vermeiden. Dieses Phänomen ist insbesondere bei wenigen Trainingsdaten zu beobachten. Wird nun die Variablenfilterung im MOP-Ansatz sowie die Regularisierung im Netzwerktraining angewendet, wird automatisch der Subraum der sechs wichtigsten Eingangsparameter als wichtig erkannt, und die

Approximationsqualität des Deep Learning Modells kann bei gleicher Netzwerkarchitektur von 79% auf 99% Vorhersagequalität gesteigert werden. Die gute Übereinstimmung von Fittingqualität (CoD) und Prognosefähigkeit (CoP) indizieren ein sehr geringes Overfitting Risiko.

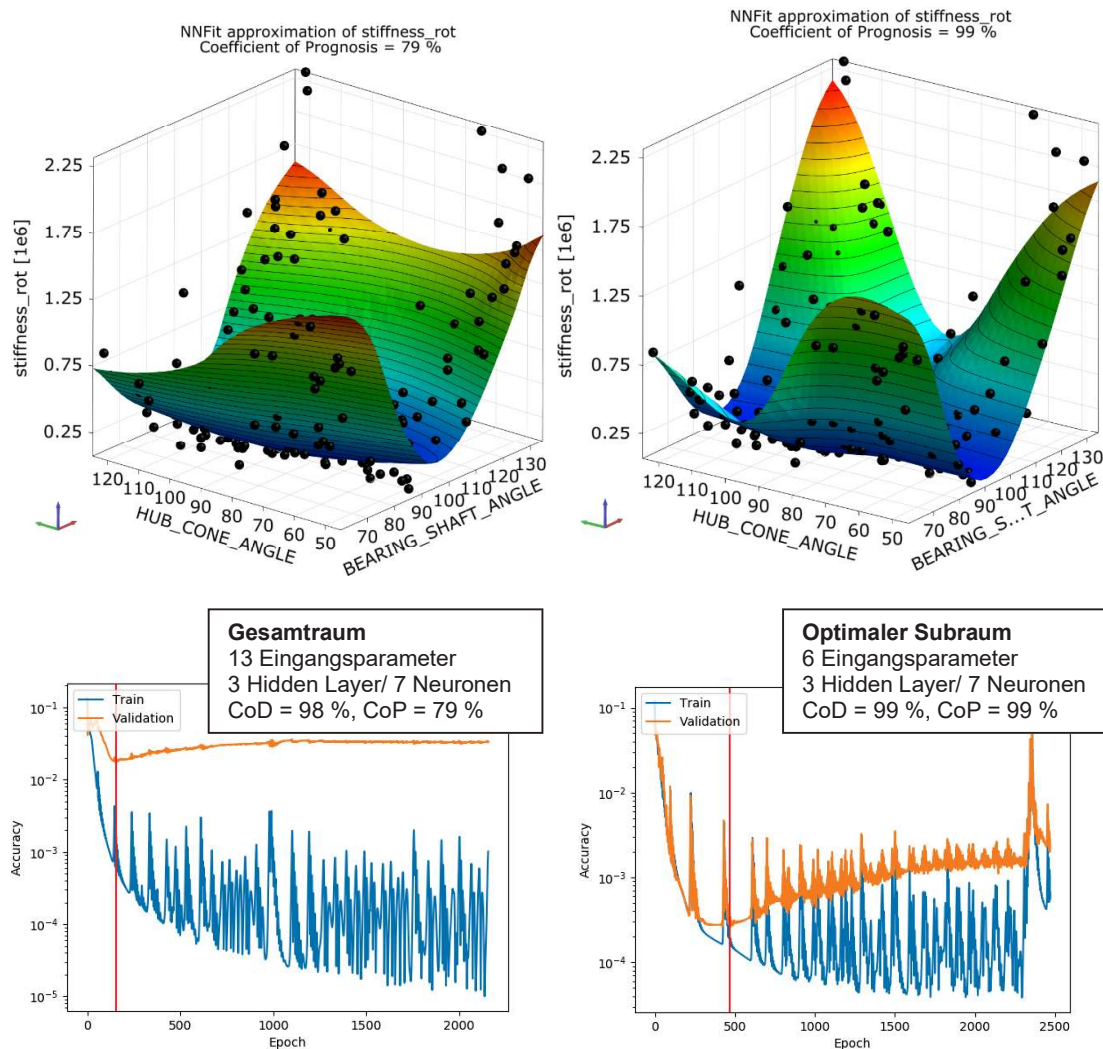


Abb. 7: Approximation der Rotationssteifigkeit eines Turbinenlaufrades durch ein Deep Learning Modell im Gesamttraum (links) und im optimalen Subraum der Eingangsparameter unter Anwendung der automatischen Netzwerk-Regularisierung (rechts)

Wie bereits diskutiert, zeigen Deep Learning Modelle für Anwendungen mit wenigen Datenpunkten (50...200) eine ähnliche Genauigkeit, wie klassische Regressionsansätze. Dies zeigt sich auch in vielen Kundenanwendungen des MOP-Workflows. Im Bereich sehr vieler Datensätze stellen diese Modelle eine echte Alternative dar, da der Trainingsaufwand nur linear oder sublinear mit der Anzahl der Datenpunkte steigt [9]. Im zweiten Beispiel wird daher eine solche Anwendung mit sehr vielen Datenpunkten vorgestellt.

Im zweiten Beispiel wurden die simulierten Daten eines Einschervorgangs (Cut-In) eines autonomen Fahrzeuges, wie in Abbildung 8 dargestellt, analysiert. Dabei wurden 10 Eingangsparameter wie die Geschwindigkeiten des Ego- und Cut-In-Fahrzeuges, der Abstand zum Vorfahrer sowie die Bremsbeschleunigungen und -zeiträume variiert. Ein initialer Datensatz von 2855 Datenpunkten wurde mit Latin Hypercube Sampling gleichförmig verteilt im Designraum erzeugt und anschließend mit einem Verkehrssimulationsprogramm ausgewertet. Dabei wurden Time-Headway (THW), Time-to-Collision (TTC), Kollisionsgeschwindigkeit und andere Antwortgrößen berechnet, und mit einem Deep Learning Modell wurde ein kombiniertes Versagenskriterium approximiert. Im initialen

Sampling wurden hauptsächlich unkritische Parameterkombinationen berechnet und wenige kritische Zustände gefunden, wie im Histogramm in Abbildung 9 dargestellt. Ein Versagenskriterium von kleiner Null indiziert dabei einen unsicheren Zustand, der zum Unfall führt. Um nun den Übergang zwischen einem sicheren und einem unsicheren Zustand besser zu approximieren, wurde das raumfüllende Constraint Refinement des AMOP mit einem Zielbereich für das Versagenskriterium von -1 bis 1 vorgegeben. Innerhalb von neun Iterationen mit insgesamt 2745 weiteren Samples wurde dieser Wertebereich weitaus besser abgedeckt, wie in Abbildung 8 und 9 ersichtlich. Im initialen Sampling befanden sich nur 7.5% der Stützstellen im Wertebereich -1 bis 1. Durch die Adaption konnten 63% der zusätzlichen Samples in diesem Bereich platziert und die Approximation des Deep Learning Modells in diesem Wertebereich erheblich verbessert werden. Das adaptierte Approximationsmodell wurde dann in weiteren Analysen zur Risikobewertung des analysierten Einschervorgangs verwendet. Weitere Details zu diesem Vorgehen sind in [11] zu finden.

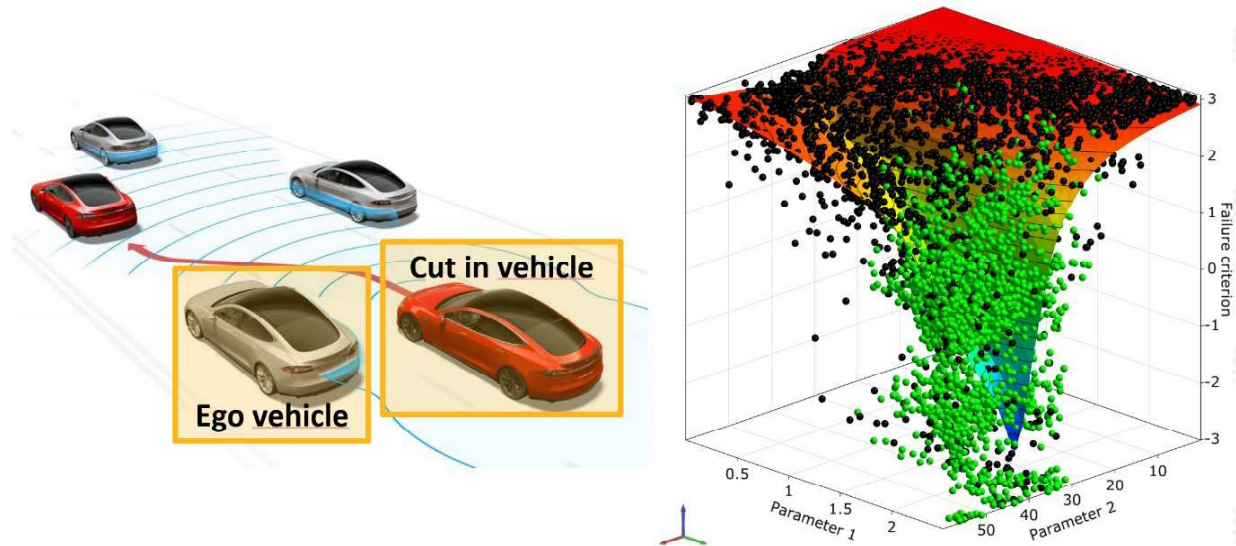


Abb. 8: Untersuchter Einschervorgang (Cut-In) eines autonomen Fahrzeuges mit zugehöriger Approximation des Versagenskriteriums mit initialen Datenpunkten (schwarz) und durch das Adaptive MOP erzeugte zusätzliche Stützstellen (grün)

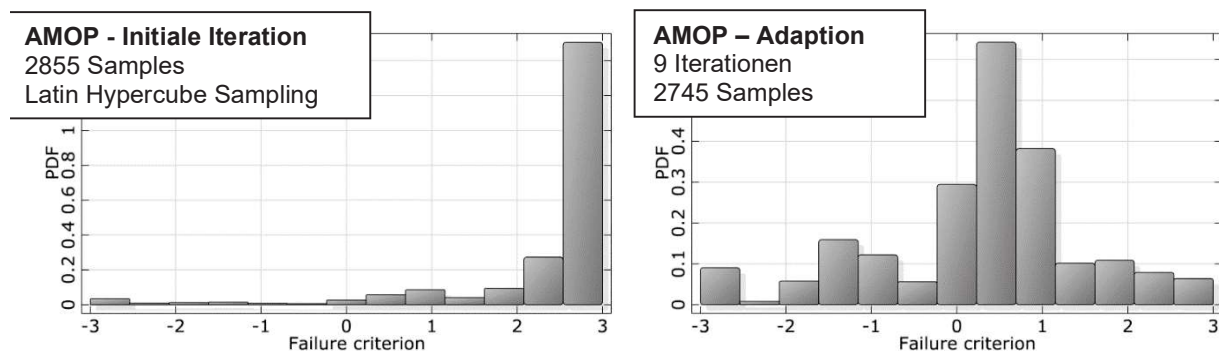


Abb. 9: Histogramme des Versagenskriteriums im initialen Sampling (links) und durch eine schrittweise Adaption im Zielbereich von -1 bis 1

7 Schlussfolgerungen

In diesem Beitrag wurde eine Erweiterung des klassischen Metamodells Optimaler Prognosefähigkeit durch moderne Deep Learning Modelle vorgestellt. Eine neue Regularisierungstechnik erlaubt es, die Deep Learning Modelle während des Trainings hinsichtlich ihrer Architektur automatisch anzupassen und so das Overfitting-Risiko zu minimieren. Für den Endanwender ist das automatisierte Training und die Verwendung dieser Modelle denkbar einfach durch eine robuste Implementation in Ansys optiSLang. Durch die Integration in den MOP-Rahmen und die konsistente Berechnung des Fehlermaßes, dem Coefficient of Prognosis, ist der Vergleich mit klassischen Regressionsansätzen möglich, und das jeweilig am besten geeignete Modell kann automatisiert selektiert werden. Durch die Einbindung in ein adaptives Vorgehen, können die lokalen Modellfehler oder benutzerdefinierte Kenngrößen direkt als Adaption- und Konvergenzkriterium genutzt werden und so die klassischen Regressions- und auch die Deep Learning Modelle optimal an die bestehenden Anforderungen angepasst werden. Eine Erweiterung des MOP Konzeptes für Fieldgrößen und Signale ist bereits verfügbar [6,13] und wird aktuell hinsichtlich der verschiedenen Modellansätze erweitert.

8 Literatur

- [1] Most T., Will J.: "Metamodel of Optimal Prognosis - an automatic approach for variable reduction and optimal metamodel selection",
Proceedings Weimarer Optimization and Stochastic Days 5.0, Weimar, Germany, 2008
- [2] Most T., Will J.: "Sensitivity analysis using the Metamodel of Optimal Prognosis",
Proceedings Weimarer Optimization and Stochastic Days 8.0, Weimar, Germany, 2011
- [3] Montgomery D.C., Runger G.C.: "Applied Statistics and Probability for Engineers",
John Wiley & Sons, third edition, 2003
- [4] Lancaster, P., Salkauskas K.: "Surface generated by moving least squares methods", Mathematics of
Computation, 37:141–158, 1981
- [5] Forrester A., Sobester A., Keane A., "Engineering Design via Surrogate Modelling", Wiley, 2008
- [6] Wolff S.: "Recent Developments of Field Meta Models", Proceedings Weimarer Optimization and
Stochastic Days 18.0, 2021
- [7] Most T., Wolff S.: "Sensitivity analysis using the Metamodel of Optimal Prognosis",
Workshop Weimarer Optimization and Stochastic Days 16.0, Weimar, Germany, 2019
- [8] Most T., Will J., Rotermund J., Gräning L.: "Artificial Intelligence and Machine Learning applied in
Computer Aided Engineering", RDO-Journal, Issue 2/2019, Dynardo GmbH, Weimar
- [9] Most T., Gräning, L., Adam, U.: "Recent Developments in Metamodeling, Optimization, and Uncertainty
Quantification in the optiSLang product line", Proceedings Weimarer Optimization and Stochastic Days
18.0, 2021
- [10] Gräning L., Abdulhkim A., Most T.: "Automated Design of Architectures of Artificial Neural Networks",
ANSYS US-Patent Application 2021-018-DE, 2021
- [11] Rasch M., Ubben P.T., Most T., Bayer V., Niemeier R.: "Safety Assessment and Uncertainty
Quantification of Automated Driver Assistance Systems using Stochastic Analysis Methods", Proceedings
NAFEMS World Congress, Quebec, Canada, 2019
- [12] Niemeier R., Will J., Most T.: "Metamodel of Optimal Prognosis - A Framework for Machine
Learning and Artificial Intelligence applied in CAE", automotive CAE Grand Challenge Conference, 2020
- [13] Bayer V., Kunath S., Niemeier R., Horwege J.: "Signal-Based Metamodels for Predictive Reliability Analysis
and Virtual Testing", Advances in Science, Technology and Engineering Systems Journal, Vol. 3, No. 1,
342-347, 2018