

Klassifikation von Oberbaufehlern am Beispiel Weichen

Auf Daten des DB-Projekts „Continuous Track Monitoring“ (CTM) basierende Fehlerprognosen

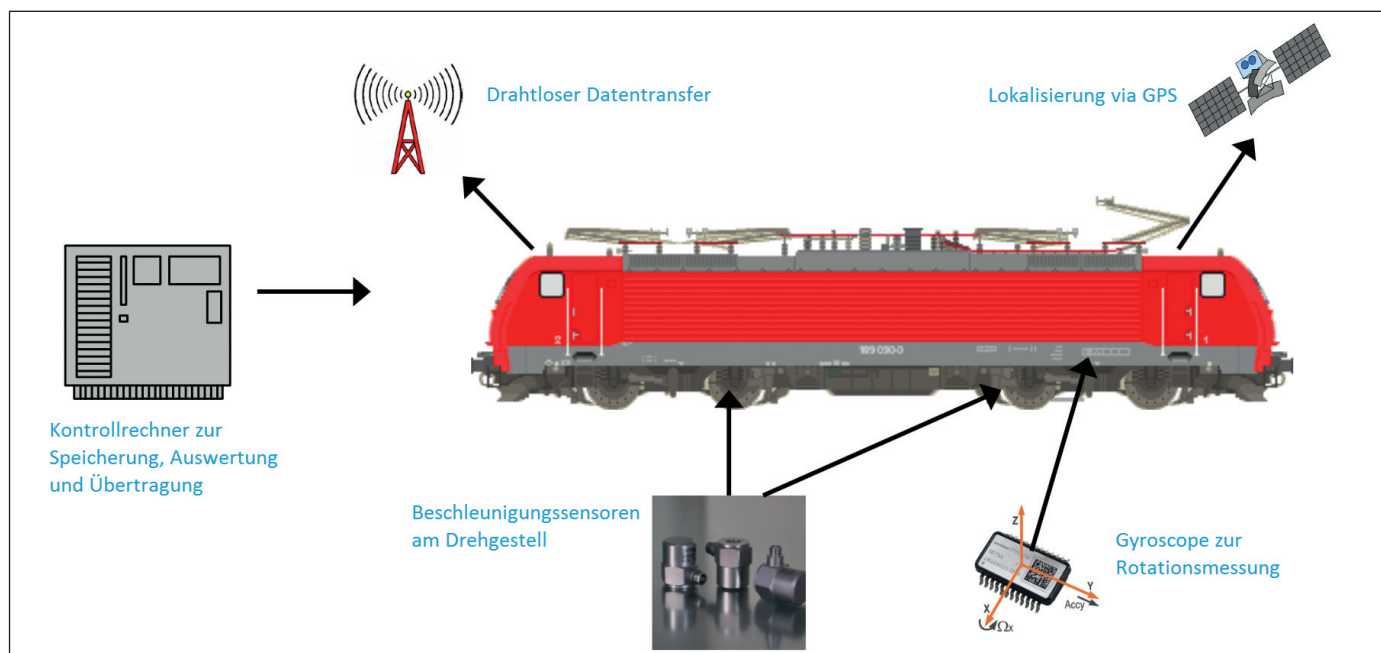


Abb. 1: Schematische Darstellung des Messaufbaus auf einer im Regelbetrieb befindlichen Güterlok zu „in service“-Messkampagnen im Rahmen von AutoMain
Grafik: [3]

Christian Linder
Alexander Oehler

Neben langwelligen Gleisgeometriefehlern, kommt es bei der Instandhaltung des Oberbaus immer wieder zu punktuellen Einzelfehlern der vertikalen Längshöhe mit Ausdehnungen von unter 3 m wie z. B. Hohllagen oder Übergangszonen. Dabei wird für die Beurteilung der Kritikalität dieser Fehler gemäß Konzernrichtlinie (KoRil) 821 [1] allein deren maximale vertikale Ausdehnung, also die Amplitude des Messsignals, betrachtet. Die KoRil 821 stellt darüber hinaus auch die Beurteilungsmaßstäbe für die Bewertung von Einzelfehlern bereit. Hierbei gibt es ein dreistufiges Beurteilungssystem. Ist die höchste Stufe (SR_{LIM}) erreicht, der Fehler also besonders weit fortgeschritten, muss durch den Anlagenverantwortlichen möglichst sofort gehandelt werden. Es kommt vor, dass durch die teils mehrmonatigen Inspektionszyklen, Fehler erst dann entdeckt werden, wenn sie bereits den SR_{LIM} -Grenzwert überschritten haben. In diesen Fällen ist der Betreiber zu

wirtschaftlich ungünstigen Sofortmaßnahmen gezwungen. Eine kontinuierlichere Überwachung dieser Oberbaufehler und eine Prognoseabschätzung sind daher von besonderem Interesse.

Im Rahmen des EU-Projektes AutoMain ist genau dies adressiert worden. Dabei ist eine Methode entwickelt worden, die Einzelfehler der Gleisgeometrie identifiziert, im zeitlichen Verlauf beobachtet und anhand der Degradationsgeschwindigkeit eine Prognose ableitet. Die Datenbasis dafür entstand in Verbindung mit dem DB-Projekt Continuous-Track-Monitoring (CTM) sowie einer eigens umgerüsteten Güterlok für „in service“-Messfahrten (Abb. 1).

Die genannten Prognosen des AutoMain-Projektes basieren auf der Erstellung eines linearen Regressionsmodells für die Fehleramplituden. Hierbei wird jedoch – wie auch in der KoRil 821 – weder die Fehlerform, noch die Art des Fehlers in diese Betrachtung einbezogen. Diese Information könnte sich jedoch als nützlich erweisen, da bei der Prognoseabschätzung für verschiedene Fehlerarten auch unterschiedliches Degradationsverhalten angenommen wird.

So entwickeln sich im zeitlichen Verlauf Fehler durch Hohllagen anders als Fehler an Übergangszonen. Wäre die Art eines Fehlers durch Datenanalyse klassifizierbar, könnten Algorithmen jedem neu entdeckten Fehler, eine Klasse zuweisen und ein spezifisches Prognosemodell, statt des bisherigen linearen Ansatzes, anwenden. Die Prognose wäre genauer.

Neben der verbesserten Prognose, existiert eine weitere interessante Anwendungsmöglichkeit. Werden anhand einer kontinuierlichen Gleisüberwachung (z. B. durch CTM) Fehler entdeckt und auch sofort klassifiziert, könnte diese Information dem Instandhaltungspersonal vorab mitgeteilt werden. Es wäre also vor dem ersten Besuch der Schadstelle möglich zu wissen, dass vor Ort beispielsweise eine Hohllage zu erwarten ist. Benötigte Ressourcen, Personal oder Maschinen könnten vorab organisiert werden. Dies eröffnet wirtschaftliches Optimierungspotential.

Daher wird als konsequente Weiterentwicklung im Folgenden gezeigt, wie sich Techniken der Computational-Intelligence bzw. des Data-Minings anwenden lassen,

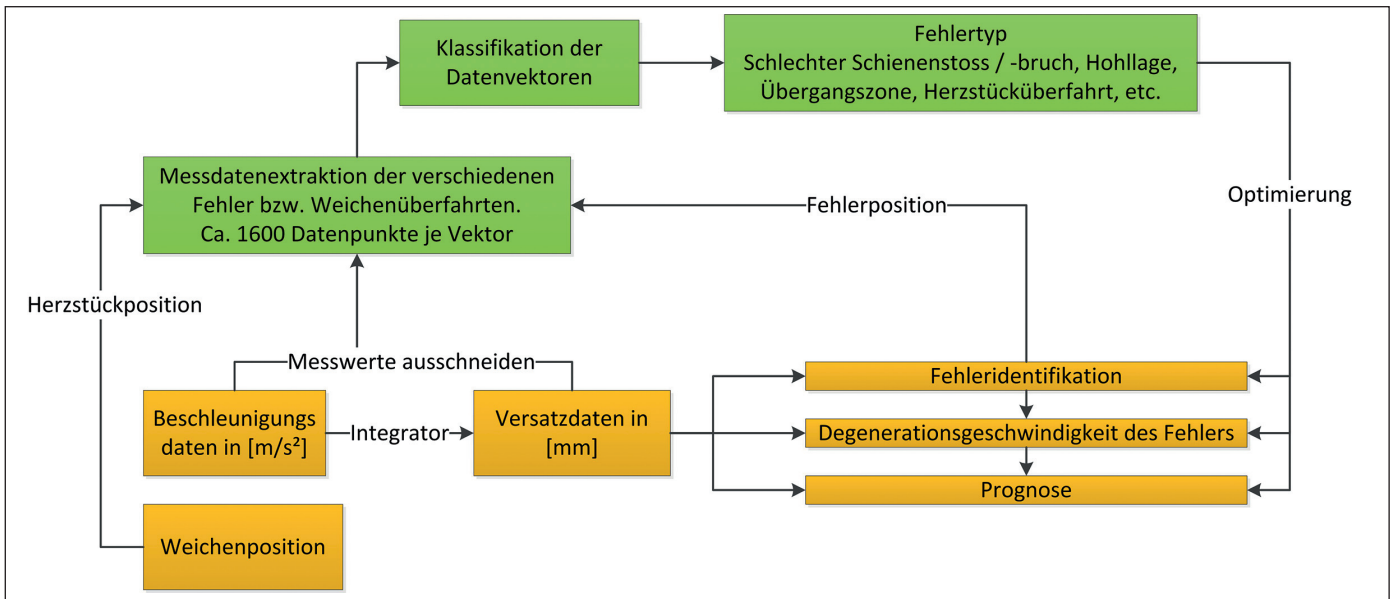


Abb. 2: Schematische Darstellung des Ansatzes aus AutoMain (Gelb) und der Weiterentwicklung (Grün)

um weiterführende Aussagen über die Art des Fehlers zu treffen. Abb. 2 visualisiert das Vorgehen und zeigt die Zusammenhänge zwischen den AutoMain-Ergebnissen (gelb hinterlegt) und dem neuen Ansatz (grün

hinterlegt). Über die gelb hinterlegten Bestandteile der Methode wurde bereits in [2] berichtet. Um dies zu erreichen, ist es zunächst erforderlich, dass Messwerte von Fehlern als

Referenz vorliegen, deren Fehlertyp bekannt ist. Dies ist in den vorliegenden Daten jedoch nicht der Fall. Stattdessen sind von der mittels CTM vermessenen Strecke die Positionen der Weichen bekannt. Diese Information wird ausgenutzt, um anhand der Messdatencharakteristik einer Weichenüberfahrt zu zeigen, dass Unregelmäßigkeiten wie Einzelfehler oder Herzstücke, die sich entlang eines Streckenabschnitts befinden, klassifiziert werden können. Zum besseren Verständnis zeigt Abb. 3 zwei in AutoMain identifizierte Vertikalausschläge. In beiden Bildern sind jeweils 13 Messdatenreihen des gleichen Fehlers zu sehen, die an verschiedenen Tagen aufgenommen wurden. Das linke Bild zeigt die Vertikalausschläge einer Weichenüberfahrt, das Rechte einen bislang unklassifizierten Einzelfehler. Bereits optisch wird klar, dass sich die Form der Fehler unterscheidet.

Die folgenden Abschnitte beschreiben das Vorgehen zur Klassifikation der Datenvektoren anhand der bekannten Charakteristik der Weichenüberfahrt. Die Aufgabe wurde am Institut für Verkehrssystemtechnik des Deutschen Zentrums für Luft- und Raumfahrt (DLR) in Braunschweig im Rahmen einer Masterarbeit [4] detailliert bearbeitet. Für das Validieren der Methode ist es dabei unerheblich, dass eine Weichenüberfahrt zwar fehlerartige Messwerte produziert, jedoch im eigentlichen Sinne kein Einzelfehler ist. Vielmehr soll gezeigt werden, dass sich CTM-Daten eignen, um bestimmte punktuelle Phänomene in der Gleisgeometrie zu klassifizieren und dadurch voneinander zu unterscheiden.

Datenextraktion und Vorverarbeitung

Das Unterscheiden von Weichenüberfahrten gegenüber anderen Fehlerarten kann

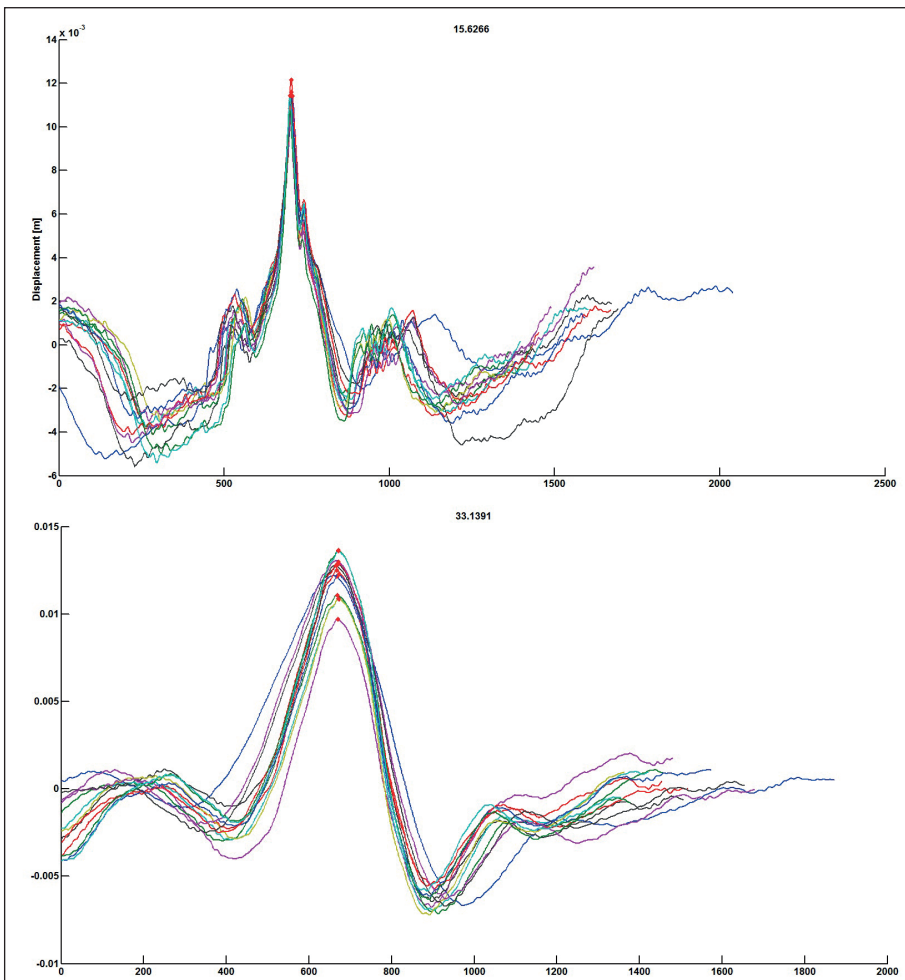


Abb. 3: Herzstücküberfahrten bei km 15,6 (links), Einzelfehler bei km 33,1 (rechts)

zunächst als 2-Klassen-Problem definiert werden. Eine Klasse enthält dabei die bereits bekannten Weichenüberfahrten, die andere Klasse alle übrigen Arten von Einzelfehlern wie in Abb. 2 exemplarisch dargestellt ist, sowie weitere Nicht-Weichen-Datensätze, die zufällig ausgewählt wurden. Die einzelnen Datenvektoren haben dabei eine variierende Länge von jeweils ca. 1600 Datenpunkten. Diese Variabilität in der Anzahl der Datenpunkte ist in der Geschwindigkeit des Messfahrzeuges begründet. Führt der Zug langsamer über die fragliche Stelle, generiert das Messsystem mehr Datenpunkte pro Meter, da dessen Abtastrate konstant bleibt. Ein weiterer Aspekt, der bei der Datenvorverarbeitung zur späteren Klassifikation beachtet werden muss, ist die Verschiebung der Datenpunkte entlang der X-Achse. Führt das Messfahrzeug an verschiedenen Tagen über den gleichen Fehler, entsteht wegen der Ungenauigkeiten des GPS-Systems eine geringfügige Abweichung in der Position der Datenvektoren. Mittels Korrelation kann diese Abweichung jedoch behoben und die Messreihen somit exakt übereinander geschoben werden. Um der variablen Länge der Datenvektoren zu begegnen, wurden mehrere Ansätze verfolgt. Zunächst können die Vektoren künstlich, durch Hinzufügen fehlender Punkte mittels Interpolation, bzw. Entfernen von Datenpunkten und Beschneidung der Vektorenränder auf eine einheitliche Länge gebracht werden. Ferner besteht die

Möglichkeit, die Vektoren durch verschiedene Transformationen auf Koeffizienten zu reduzieren und die Klassifikation anhand dieser Koeffizienten durchzuführen. Zum Beispiel lässt sich mittels der Fourier-Transformation (FT) das Frequenzspektrum des zu betrachtenden Datenvektors berechnen. Eine weitere Variante stellen die Wavelet-Transformationen (WT) dar, dabei können neben den Frequenzen auch Aussagen über die Zeit bzw. den Ort getroffen werden. Die so vereinheitlichten Datenvektoren werden im Folgenden mit zwei verschiedenen Verfahren klassifiziert:

Durchführung der Klassifikation mittels SVM und GMLVQ

Selbstlernende Klassifizier-Algorithmen bieten mittels überwachten Lernens die Möglichkeit, Muster in einem Datensatz zu finden. Dabei wird die Datenmenge in zwei Teilmengen aufgeteilt, einen Trainings- und einen Testdatensatz. Antrainiert wird der Algorithmus mit den vorverarbeiteten Weichen und Nicht-Weichen-Datenvektoren, bzw. den daraus berechneten Fourier- oder Waveletkoeffizienten. Für die Lösung des 2-Klassen-Problems eignet sich unter anderem die Support-Vector-Machine (SVM) [5], die mit klassenrand-typischen Prototypen, den Support-Vektoren, arbeitet. Die SVM berechnet in der Trainingsphase eine optimale Hyperebene für die Trennung der beiden Datensätze. Da in der Praxis die bei-

den Klassen meist nicht vollständig linear trennbar sind, wird die Separation in einem höher-dimensionalen Raum berechnet. Diese kann anschließend wieder im ursprünglichen Datenraum angewandt werden.

Im Gegensatz zur SVM arbeiten die Varianten der Linear-Vector-Quantization (LVQ) mit klassen-typischen Prototypen. Ziel ist es dabei die Prototypen durch eine entsprechende Adaption so zu platzieren, dass die jeweilige Klasse möglichst gut repräsentiert wird. Unter Verwendung einer gegebenen Lernrate richten sich dabei die Prototypen im Datenraum aus und werden so platziert, dass sie möglichst gut im Klassenzentrum liegen und somit als Repräsentant der jeweiligen Klasse genutzt werden können. Betrachtet werden dabei in jedem Schritt des Algorithmus sogenannte Gewinner- und Verliererprototypen, die mittels eines Ähnlichkeitsmaßes und der entsprechenden Klassenzugehörigkeit bestimmt werden.

Eine Erweiterung des LVQ stellt der Generalized Matrix Linear Vector Quantization (GMLVQ) Algorithmus dar, der in [6] beschrieben wird und in dieser Arbeit verwendet wurde. Hier wird zusätzlich eine Energiefunktion eingeführt. Die Energiefunktion wird mittels stochastischen Gradientenabstiegs minimiert und approximiert die Anzahl der Fehlklassifikationen. In jedem Schritt wird ein Datenvektor zufällig ausgewählt und der Gewinnerprototyp, der Prototyp mit dem geringsten Abstand und

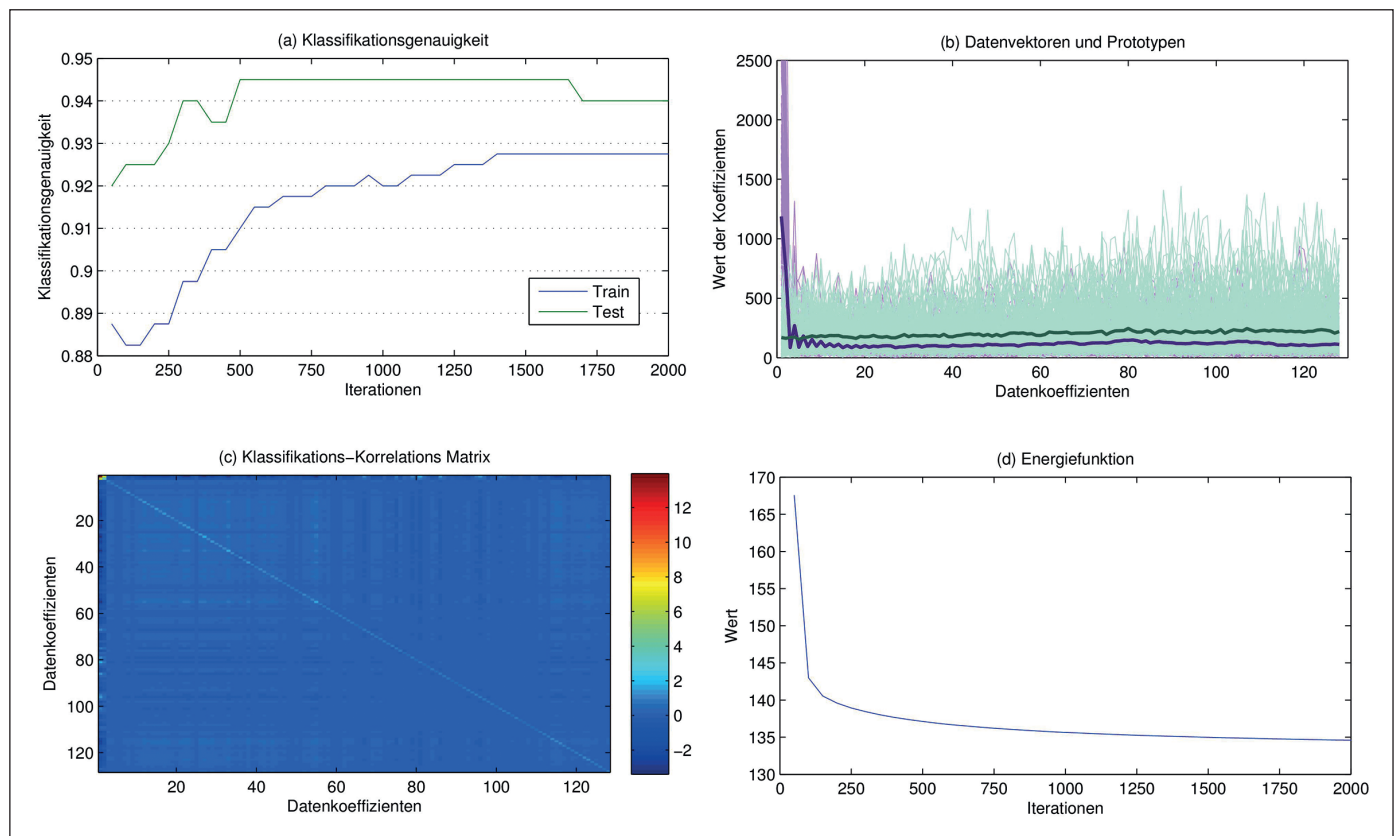


Abb. 4: Grafische Repräsentation des GMLVQ-Klassifikationsergebnisses

dem gleichen Klassenlabel, herangezogen. Für die Distanz bzw. den Abstand wird hierfür ein Ähnlichkeitsmaß verwendet. Der Prototyp mit kleinster Distanz, der jedoch einer anderen Klasse zugeordnet ist, wird weggestoßen. Weiterhin liefert die durch den GMLVQ adaptierte Klassifikations-Korrelations-Matrix Informationen über die für die Klassifikation relevanten Koeffizienten. Die LVQ-Algorithmen funktionieren auch mit mehr als zwei Klassen und können somit für unterschiedliche Fehlertypen eingesetzt werden. Die Multiclass SVM bietet ebenfalls diese Möglichkeit, wurde im Zusammenhang mit den vorliegenden Daten jedoch nicht weiter betrachtet. In der anschließenden Testphase des LVQ werden die bisher unbekannten Datenvektoren der Testmenge jeweils der Klasse zugeordnet, zu dessen Prototyp sie den geringsten Abstand haben.

Sowohl die SVM als auch der GMLVQ Algorithmus wurden im Rahmen dieser Arbeit für das Zwei-Klassenproblem verwendet, sodass im nächsten Abschnitt erklärt werden kann, welche Ergebnisse erreicht werden konnten.

Klassifikationsergebnisse

In Abb. 4 sind die Ergebnisse des GMLVQ dargestellt, welche aus der Verwendung der Fourier-Koeffizienten hervorgingen. Diese Frequenzspektren sind in (b) zu sehen, einschließlich der dazu antrainierten Prototypen (die beiden hervorgehobenen Verläufe). Aus mehreren Versuchen ging hervor, dass die Anzahl der Prototypen bei den vorliegenden Daten keinen Einfluss auf das Ergebnis hat, weshalb für jede Klasse nur ein Prototyp initialisiert wurde. Die Klassifikationsgenauigkeit nach 2000 Schritten ist in (a) zu sehen. Die Genauigkeit bei den Trainingsdaten erreicht einen Wert von ca. 93 %. Bei den Testdaten, die mittels Kreuzvalidierung generiert wurden, gab es ein Maximum von 94,5 % Klassifikationsgenauigkeit. Zu erkennen ist außerdem, dass bereits nach 500 Iterationsschritten ein akzeptables Resultat mit Werten über 90 % erzielt werden konnte. Der Verlauf der Energiefunktion, die mit einem stochastischen Gradientenabstieg minimiert wird, ist in (d) dargestellt. Durch diesen Gradientenabstieg ist die Konvergenz der Energiefunktion sichergestellt. Da der Lernprozess nach etlich vielen Schritten abgebrochen werden muss (in diesem Fall 2000), kann jedoch keine Aussage darüber getroffen werden, ob möglicherweise nur ein lokales Minimum der Energiefunktion erreicht wurde. Die Klassifikations-Korrelations-Matrix Ω , wie in (c) zu sehen, enthält auf der Hauptdiagonale Information über die Relevanz der einzelnen Datenkoeffizienten bezüglich der Klassifikationseigenschaft. Bei den vorliegenden Daten ist jedoch zu erkennen, dass kaum Unterschiede existieren

Länge [m]	GMLVQ		SVM	
	BD	VD	BD	VD
8	65,5%	55,5%	51%	52%
16	64%	77%	50,5%	68%
32	71,5%	80%	50%	64,5%
64	94,5%	79%	50%	67%
128	60,5%	80,5%	50,5%	62%

Tab. 1: Klassifikationsergebnisse des Frequenzspektrums der Beschleunigungsdaten (BD) und Versatzdaten (VD) in Abhängigkeit der Intervalllänge

und somit alle Koeffizienten relevant sind. Getestet wurden die aufgezeichneten Beschleunigungs- und Versatzdaten mit zwei selbstlernenden Klassifizierern. Weiterhin wurden mit der FT und der WT zwei Methoden der Signalverarbeitung angewandt, um dadurch Informationen über Frequenzen zu extrahieren. Die besten Ergebnisse lieferte der GMLVQ unter Verwendung der Fourier-Koeffizienten. Hierbei konnte eine Klassifikationsgenauigkeit der Testdatensmenge von über 90 % erreicht werden. Die Resultate der SVM waren meist schlechter als die des GMLVQ. Außerdem waren oft sehr viele Datenvektoren Support-Vektoren, was auf ein Auswendiglernen hindeutet. Die in den Messdaten enthaltenen Frequenzen, die mittels FT und WT berechnet werden können, sind für die Klassifikation deutlich besser geeignet als die reinen Messpunkte. In Tab. 1 sind die Ergebnisse der Klassifikation der Frequenzspektren der FT in Abhängigkeit der Intervalllänge dargestellt. Das beste Resultat wurde bei den 64 m langen Überfahrten erreicht. Die Klassifikationsergebnisse der wavelettransformierten Daten waren etwas schlechter, aber immer noch deutlich besser als bei den komplett unverarbeiteten Beschleunigungs- bzw. Versatzdaten. Die Schwierigkeit bei der WT ist, das passende Wavelet zu finden. Getestet wurden einige Daubechies-Wavelets und Symlets, mit teilweise stark abweichenden Klassifikationsresultaten. Möglicherweise erzielen andere Basis-Wavelets bessere Ergebnisse.

Fazit und Zusammenfassung

Mit der vorliegenden Methode wurde gezeigt, dass es möglich ist, punktuelle und kurzzeitige Unregelmäßigkeiten der Längshöhe in CTM-Daten zu finden und auch anhand der Signalfrequenz über die Beurteilung der Amplitude hinaus zu betrachten. Es wurde gezeigt, dass bei bekannter Referenzform eines Fehlers, dieser von anderen unterschieden werden kann. Als bekannte Referenz wurde hier auf den charakteristischen Messdatenverlauf einer Weichenüberfahrt zurückgegriffen, da zwar Messdaten von

Fehlern, aber nicht deren genauer Fehlertyp bekannt waren. Die Erhebung dieser Referenzen und das Migrieren und Validieren des vorgestellten Verfahrens auf andere Fehlertypen wird das Ziel weiterer Untersuchungen sein. Darauf aufbauend sollen je nach Fehlerart verbesserte Prognoseverfahren abgeleitet werden.

LITERATUR

- [1] Deutsche Bahn AG: Konzernrichtlinie 821 – Oberbau inspizieren
- [2] Linder, C.; Schenkendorf, R.; Lackhove, C.: Prognoseverfahren zur Gleislageabweichung bei Einzelfehlern, in: EI – DER EISENBAHNINGENIEUR, Heft 2|2014, Seiten 17-20
- [3] Baumann, G.; Linder, C.: Inspection of track from in service freight trains, AutoMain Abschlussbericht zu Arbeitspaket 3.1, Online (inkl. Vollversion): http://www.automain.eu/IMG/pdf/task_3.1_inspection_of_track_from_in-service_freight_trains.pdf
- [4] Oehler, A.: Entwicklung einer Methode zur Einzelfehlererkennung und -klassifikation im Gleisoberbau, Masterarbeit der Hochschule Mittweida (2014), Online (inkl. Vollversion): <http://elib.dlr.de/89186/>
- [5] Cortes, C.; Vapnik, V.: Support-vector networks, Machine Learning 20 (3): 273, 1995
- [6] Schneider, P.; Biehl, M.; Hammer, B.: „Adaptive relevance matrices in Learning Vector Quantization“ (2009), in: Neural Computation, Ausgabe 21, S. 3532-3561



Dipl.-Geolnf. Christian Linder

Deutsches Zentrum für
Luft- und Raumfahrt e.V.
Institut für Verkehrssystemtechnik
Braunschweig
christian.linder@dlr.de



Alexander Oehler M.Sc.

Absolvent der
Hochschule Mittweida
alexander.oehler@gmx.de

Summary

Superstructure fault classification by the example of turnouts

By extending a recently published method for forecasting of single failures in track, based on continuous track monitoring data within the AutoMain project, the German Aerospace Center (DLR) presents a method, which proves, that different kinds of single failures can be separated from each other. While state of the art assessment of track in Germany is based on a limit value consideration about the failures amplitude, the failures type is neglected. This work shows that with a given reference dataset of a known type of failure, a classification can be achieved, which aims to optimise the previously presented forecasting algorithms by developing type specific prognostic models for each failure type. This method can contribute to improved maintenance efficiency in terms of better prognostics and optimised preparation for repairs by anticipating the failure type before seeing the worksite.

HOCHLEISTUNG | PRÄZISION | ZUVERLÄSSIGKEIT

Plasser & Theurer



Der Standard der Bettungsreinigung

Zeit und Ressourcen sparen. Hohe Reinigungsqualität sichern. Und darüber hinaus auch noch Flexibilität im Einsatz beweisen. Das ist die Philosophie von Plasser & Theurer für eine wirtschaftliche Bettungsreinigung. Die wichtigste Voraussetzung für die spätere Gleislage ist ein einwandfreies, gerades Planum. Dafür sind die Gleis- und Weichenreinigungsmaschinen von Plasser & Theurer mit einer Aushubkette mit querliegender Kettenführung ausgestattet - präzise auf die gewünschte Räumtiefe und Planumsneigung einstellbar.

