
Maschinelles Lernen und Datenanalyse

In der Werkstoff- und Prüftechnik

Prof. Dr. Stefan Bosse

Universität Koblenz - FB Informatik - Praktische Informatik

Universität Siegen - FB Maschinenbau / LMW

Datenanalyse und Eigenschaftsselektion

Häufig sind die rohen sensorischen Daten(variablen) zu hochdimensional und abhängig voneinander

Reduktion auf wesentliche Merkmale kann ML Qualität deutlich verbessern

Häufig besitzen einzelne Sensorvariablen keine oder nur geringe Aussagekraft (geringe Entscheidbarkeitsqualität)

Datenqualität

- Die Daten $\mathcal{D}$ werden durch fünf wesentliche Eigenschaften beschrieben, die auch mit statistischer Analyse quantifiziert werden können:

Rauschen. Rauschen ist die Verzerrung der Daten. Diese Verzerrung muss entfernt oder Ihre nachteiligen Auswirkungen vermindert werden, bevor ML Algorithmen ausgeführt werden, da die Leistung und Qualität der Algorithmen beeinträchtigen werden kann.



Es gibt eine Vielzahl von Filteralgorithmen um den Einfluss von Rauschen auf das eigentliche Sensorsignal zu vermindern. Aber Rauschminderung erhöht nicht den Informationsgehalt eines Signals und basiert auf einem statistischen Modell (muss bekannt sein).

Ausreißer. Ausreißer sind Instanzen, die sich erheblich von anderen Instanzen im Datensatz unterscheiden.

- Beispiel: Zugversuch.
- Bei einem Zugversuch treten hohe Kräfte auf. Wenn die Einspannung des Werkstücks unzureichend ist könnte sich kurz vor dem bruch das Werkstück etwas lösen und die Dehnungskurve wird verzerrt hin zu größeren Dehnungen. Nicht einfach zu erkennen!



Aber: Ausreißer können in besonderen Fällen nützliche Muster darstellen und die Entscheidung, sie zu entfernen, hängt vom Kontext und der Fragestellung ab. Zum Beispiel vereinzelt vorzeitiges Versagen von Werkstücken durch innere Materialdefekte, die dann zu einer gezielten Untersuchung führen (aber schwierig da selten).

Bias. Systematische Verschiebung von Messdaten (Offset).

- Beispiel: Externe Umwelteinflüsse können die Empfindlichkeit eines Sensors verändern.
- Beispiel: Elektrische Potentialverschiebung verschiebt den Nullpunkt eines Ultraschallsensors
- Aber auch Veränderung des Testobjekts (Lastsituation, Alterung, Materialveränderung)



Systematische Verschiebung kann nur durch zusätzliche Modellfunktionen oder Sensorfusion kompensiert werden!

Fehlende Werte. Fehlende Werte sind Funktionswerte, die in Instanzen fehlen.

- Im Ingenieursbereich sind fehlende Werte der Standardfall. Der Versuchsplan wird immer Lücken aufweisen. Entweder wall nicht alle Parameter der Versuche lückenlos eingestellt werden können, die Anzahl der Parametervariationen nur eine Auswahl darstellt, oder Versuche fehlschlagen.
- Bei der Untersuchung von Schäden kann häufig nur ein Schaden pro Probe erzeugt werden (Einschlagschaden), oder diese sind ungleichmäßig verteilt (Risse).
- Um dieses Problem zu lösen, können wir
 1. Instanzen mit fehlenden Werten entfernen;
 2. Fehlende Werte schätzen (Z. B. durch den gängigsten Wert ersetzen); oder
 3. Fehlende Werte ignorieren, wenn Data Mining Algorithmen ausgeführt werden.

Duplikate. Doppelte Daten treten auf, wenn mehrere Instanzen mit genau denselben Funktionswerten vorhanden sind.

- Duplikate treten im geplanten Versuchsumfeld eher selten auf. Und Duplikat heißt nicht notwendigerweise identisch.
- Je nach Kontext können diese Instanzen entweder entfernt oder beibehalten werden. Wenn Instanzen beispielsweise eindeutig sein müssen, sollten doppelte Instanzen entfernt werden.

Statistische Analyse

- Statistische Analysen von Mess- und Sensordaten können neue Datenvariablen erzeugen und Informationen über die Daten liefern:
 - Eigenschaftsselektion (Feature Selection) für ML und Informationsgewinnung
 - Variablentransformation mit Datenreduktion
- Statistische Analyse liefert eine Reihe von Kennzahlen über Datenvariablen, das können Eigenschaften für die Weiterverarbeitung sein:

$$\begin{aligned} stat(\vec{x}) : \vec{x} &\rightarrow \vec{p}, \\ \vec{p} &= \{mean, \sigma, \dots\} \end{aligned}$$



Welche statistische Größen gibt es? Was können statistische Größen über Daten aussagen?

Statistische Funktionen

Peak amplitude (y_{peak})

$$y_{peak} = \max |y_i|$$

[2:173]

Mean ($\bar{y}$)

$$\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$$

Mean square ($\bar{y}_{sq}$)

$$\bar{y}_{sq} = \frac{1}{n} \sum_{i=1}^n (y_i)^2$$

Root-mean-square (rms)

$$rms = \sqrt{\frac{1}{n} \sum_{i=1}^n y_i^2}$$

Variance (σ^2)^a

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2$$

Statistische Funktionen

Standard deviation (σ)^a

$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2}$$

[2:173]

Skewness (dimensionless) (γ)^a

$$\gamma = \frac{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^3}{\sigma^3}$$

Kurtosis (dimensionless) (κ)^a

$$\kappa = \frac{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^4}{\sigma^4}$$

Crest factor (X_{cf})

$$X_{CF} = y_{peak}/rms$$

K-factor (X_k)

$$X_{CF} = (y_{peak})(rms)$$

WorkBook Live: Statistische_Analyse

CLEAR **LOAD** **+** **-** **R** **dataan1** **dataan2** **dataan3**

Statistische Funktionen in R

Anwendung i.A. auf Vektoren, Matrizen, oder Arrays.

R	Beschreibung
<code>fivenum(x)</code>	Liefert <i>min, q1, median, mean, q3, max</i>
<code>mean(x)</code>	Arithmetischer Mittelwert
<code>min(x), max(x)</code>	Extremwerte
<code>moment(x, order)</code>	Statistisches Moment einer Datenserie
<code>rnorm(n, a, b)</code>	Normalverteilte Zufallszahlen [a,b]
<code>runif(n, a, b)</code>	Uniform verteilte Zufallszahlen [a,b]
<code>sd(x)</code>	Standardabweichung
<code>sum(x)</code>	Arithmetische Summation
<code>var(x)</code>	Varianz (Achtung hier: sd^2)

Korrelation von Datenvariablen

- Variablen **X** sollten möglichst (linear) unabhängig sein,
 - Um eine geeignete, robuste und genaue Modellsynthese (also ML) zu ermöglichen, d.h.
 - Es sollte möglichst keine Zusammenhänge der Form $\exists correlation(X_i, X_j)$ geben!
 - Um den Modellsyntheseprozess zu beschleunigen (also das Training); Rechenzeit reduzieren
 - Um Modelle klein und kompakt zu halten
 - Um Modellspezialisierung zu vermeiden und Varianz zu erhöhen



Abhängige Variablen sollten identifiziert und in unabhängige "transformiert" werden!

Beispiel Prozessanalyse

- Eine Datentabelle D mit experimentellen Messgrößen und Fertigungsparametern (Prozessparameter) von additiv gefertigten Bauteilen hatte zunächst 7 Variablen (numerisch):
 - X_1 : Hatchabstand [mm]
 - X_2 : Scangeschwindigkeit [mm/s]
 - X_3 : Laserleistung [W]
 - X_4 : Schichtstärke [mm]
 - X_5 : Volumenenergiedichte [J/mm^3]
 - X_6 : Bauplatten Position x
 - X_7 : Bauplatten Position y
 - Y_1 : Dichte (%)

- D bestand aus 61 Experimenten mit unterschiedlichen Fertigungsserien
- Mit einer Principle Component Analysis (PCA) konnte die ganze Tabelle auf die Variablen PC_1 und Y_1 reduziert werden!
 - Die Genauigkeit der mit ML synthetisierten Funktion $M(\mathbf{X}): \mathbf{X} \rightarrow \text{Dichte}$ konnte ohne signifikanten Genauigkeitsverlust nur aus PC_1 abgeleitet werden, d.h.
 $M(PC_1): PC_1 \rightarrow \text{Dichte}$
 - Aber: Für die Inferenz (Applikation) von M muss die PCA für die Eingabedaten $\mathbf{X}$ wiederholt werden bzw. die Datentransformation durchgeführt werden!

Analyse Kategorischer Variablen

- Die Analyse von kategorischen Variablen vereint die Konzepte:
 - Mengenlehre
 - Kodierung/Dekodierung
 - Verteilung (Wahrscheinlichkeit des Auftretens)

Attribut	Wertemenge
Aussicht	sonnig, regnerisch, bewölkt
Temperatur	kalt, mild, heiß
Luftfeuchtigkeit	hoch, normal
Windig?	ja, nein

[5]

Abb. 1. Beispiel von rein kategorischen Attributen einer Datentabelle ***D***

Messdaten

Beispiel	Aussicht	Temperatur	Luftfeuchtigk.	Windig?	Klasse
1	sonnig	heiß	hoch	nein	N
2	sonnig	heiß	hoch	ja	N
3	bewölkt	heiß	hoch	nein	P
4	regnerisch	mild	hoch	nein	P
5	regnerisch	kalt	normal	nein	P
6	regnerisch	kalt	normal	ja	N
7	bewölkt	kalt	normal	ja	P
8	sonnig	mild	hoch	nein	N
9	sonnig	kalt	normal	nein	P
10	regnerisch	mild	normal	nein	P
11	sonnig	mild	normal	ja	P
12	bewölkt	mild	hoch	ja	P
13	bewölkt	heiß	normal	nein	P
14	regnerisch	mild	hoch	ja	N

[5:53]

Abb. 2. Beispiel einer rein kategorischen Datentabelle D . Die Zielvariable *Klasse* mit den Werten $\{N,P\}$ ist ebenfalls kategorisch, z.B. Klasse=P $\Rightarrow$ Sportliche Aktivität

Gemischte Variablenklassen

Outlook	Temperature	Humidity	Windy	Play-time
Sunny	85	85	False	5
Sunny	80	90	True	0
Overcast	83	86	False	55
Rainy	70	96	False	40
Rainy	68	80	False	65
Rainy	65	70	True	45
Overcast	64	65	True	60
Sunny	72	95	False	0
Sunny	69	70	False	70
Rainy	75	80	False	45
Sunny	75	70	True	50
Overcast	72	90	True	55
Overcast	81	75	False	75
Rainy	71	91	True	10

[7:47]

Abb. 3. Einige Datenvariablen wurden mit numerischen/metrischen Werten ersetzt (Klasse → Play-time)



Kann aus der vorherigen Datentabelle mit numerischen Variablen noch ein Zusammenhang aus X zu Y hergestellt werden?

Reicht die Anzahl der Experimente im Vergleich zu der rein kategorischen Datentabelle?

Wo liegen die Probleme?

Weiteres Beispiel

Age	Spectacle Prescription	Astigmatism	Tear Production Rate	Recommended Lenses
Young	Myope	No	Reduced	None
Young	Myope	No	Normal	Soft
Young	Myope	Yes	Reduced	None
Young	Myope	Yes	Normal	Hard
Young	Hypermetrope	No	Reduced	None
Young	Hypermetrope	No	Normal	Soft
Young	Hypermetrope	Yes	Reduced	None
Young	Hypermetrope	Yes	Normal	Hard
Prepresbyopic	Myope	No	Reduced	None
Prepresbyopic	Myope	No	Normal	Soft
Prepresbyopic	Myope	Yes	Reduced	None
Prepresbyopic	Myope	Yes	Normal	Hard
Prepresbyopic	Hypermetrope	No	Reduced	None
Prepresbyopic	Hypermetrope	No	Normal	Soft
Prepresbyopic	Hypermetrope	Yes	Reduced	None
Prepresbyopic	Hypermetrope	Yes	Normal	None
Presbyopic	Myope	No	Reduced	None
Presbyopic	Myope	No	Normal	None
Presbyopic	Myope	Yes	Reduced	None
Presbyopic	Myope	Yes	Normal	Hard
Presbyopic	Hypermetrope	No	Reduced	None
Presbyopic	Hypermetrope	No	Normal	Soft
Presbyopic	Hypermetrope	Yes	Reduced	None
Presbyopic	Hypermetrope	Yes	Normal	None

[7:7]

Kodierung

Kategorische Variablen (sowohl Attribute als auch Zielvariablen) können von einer Vielzahl von numerisch basierten ML Verfahren nicht verarbeitet werden (wie neuronale Netze)

- Eine Lösung ist die Abbildung von kategorischen Werten (also Mengen von Symbolen) auf numerische Werte → **Kodierung**
- Kodierte Werte sind aber i.A. weder intervall- noch verhältnisskalierbar!

Kodierungsformate

- *Linear* und nicht akkumulativ (skalar), d.h.
 $\{\alpha, \beta, \gamma, \dots\} \rightarrow \{\delta, 2\delta, 3\delta, \dots\}$
- *Exponentiell* (z.B. zur Basis $B=2$) und akkumulativ (skalar), d.h.
 $\{\alpha, \beta, \gamma, \dots\} \rightarrow \{2^0, 2^1, 2^2, \dots\}$
- *One-hot* und evtl. akkumulativ (vektoriell), d.h.
 $\{\alpha, \beta, \gamma, \dots\} \rightarrow \{[1, 0, 0, \dots], [0, 1, 0, \dots], [0, 0, 1, \dots], \dots\}$



Exponentielle Kodierungen können multiple verschiedene kategorische Werte in einem numerischen Wert darstellen! Z.B. mehrfache kategorische Antworten bei einer Frage einer Umfrage.

Beispiele

```
{sonnig,bewölkt,regnerisch} → {3,2,1}
{ja,nein} → {1,0}
{Schaden A, Schaden B, Schaden C} → { 1,2,3 }
{rot,grün,blau,braun,weiß} → {1,2,4,8,16}
{Sport, Kino, Theater, Musik} → {1,2,4,8}
{heiß,kalt} → {[1,0],[0,1]}
```

- Numerische/metrische Werte können auf kategorische durch Intervallkodierung reduziert werden:

$$\begin{aligned} \text{cat}(x) : x &\rightarrow \{\alpha_1, \alpha_2, \dots, \alpha_n\}, x \in \mathbb{R}/\mathbb{N} \\ \alpha_i &\leftrightarrow x = [x_0 + i\delta, x_0 + (i+1)\delta] \end{aligned}$$

- Z.B. die Kategorisierung von Schadenspositionen in einer mechanischen Struktur durch räumliche Bereiche (Segmente) {S1,S2,S3,...,S9}
- Z.B. Temperaturen durch "gefühlte" Attribute {heiß, warm, moderat, kalt, eiskalt}
- Z.B. Zeitangaben durch Epochen {Steinzeit, Bronzezeit, Kohlezeit, .. }

Dekodierung

Die Kodierung ist umkehrbar mit einer Dekodierungsfunktion (unter Kenntnis der Kodierungsvorschrift)

- Verwendung der switch und factor Funktionen

WorkBook Live: Kodierung_und_Dekodierung

Beispiele in R(+)

```
use math
# Kodierung
c1 = switch('A',A=1,B=2,C=3)
c12 = switch(['A','N'],A=1,B=2,C=3)
c2 = factor(['A','B'],['A','B','C'],[1,2,4])
c3 = switch('heute',heute=0,gestern=-1,morgen=1)
logg(c1,sum(c12),c2,sum(c2))
# Dekodierung
d2 = factor(c2,[1,2,3],['A','B','C'])
logg(d2)
```

Entropie und Informationsgehalt

- Sensorvariablen können unterschiedlichen *Informationsgehalt* besitzen
 - Nur auf den Dateninhalt (Werte) der Variable X_i bezogen (inherenter Informationsgehalt)
 - Oder zusätzlich bezogen auf die Zielvariable Y (abhängiger Informationsgehalt)
- Der *Informationsgehalt* einer Menge X aus Elementen der Menge C wird durch die *Entropie* $E(X)$ gegeben:

$$E(X) = - \sum_{i=1,k} p_i \log_2(p_i), p_i = \frac{\text{count}(c_i, X)}{N}, X = \{c | c \in C\}$$

- Dabei ist k die Anzahl der unterscheidbaren Elemente/Klassen $Val(X) \subseteq C$ in der Datenmenge X (z.B. die Spalte einer Tabelle) und p_i die Häufigkeit des Auftretens eines Elements $c_i \in C$ in X .
- Beispiele:

$X1=\{A, C, B, C, B, C\} \rightarrow Val(X1)=C=\{A, B, C\}, N=6$
 $E(X1)=- (1/6)\log(1/6) - (2/6)\log(2/6) - (3/6)\log(3/6)=1.46$
 $X2=\{A, B, C\} \rightarrow Val(X2)=C=\{A, B, C\}, N=3$
 $E(X2)=- (1/3)\log(1/3) - (1/3)\log(1/3) - (1/3)\log(1/3)=1.58$
 $X3=\{A, A, A, A, B, B\} \rightarrow Val(X3)=\{A, B\} \subset C=\{A, B, C\}, N=6$
 $E(X3)=- (4/6)\log(4/6) - (2/6)\log(2/6) - (0/6)\log(0/6)=0.92$
 $X4=\{A, A, A, A, B, B\} \rightarrow Val(X4)=C=\{A, B\}, N=6$
 $E(X4)=- (4/6)\log(4/6) - (2/6)\log(2/6)=0.92$

- Die Entropie ist Null wenn die Datenmenge X "rein" ist, d.h., nur Elemente einer einzigen Attributklasse $c_1 \in C$ enthält, z.B. $X=\{A,A,A\}$.
- Die Entropie ist $\log_2(|C|)$ wenn alle Werte gleichverteilt vorkommen, wenn nicht dann kleiner (nicht gleichverteilt).
- Die Entropie reicht allein zur Bewertung des Informationsgehaltes nicht aus:

X1	X2	Y
A	C	P
B	C	P
A	D	N
B	D	N

- $E(X1)=1, E(X2)=1$!! Welche Variable X ist für die Entscheidung der Zielvariable Y geeignet?

WorkBook Live: Entropie

CLEAR

LOAD

+

-

R

DataSports

AnalysisSports

DataLenses

AnalysisLenses

Entropie von numerischen Werten

- Die Entropie ist primär bei kategorischen Variablen zu verwenden
- Bei numerischen und vor allem kontinuierlich verteilten Variablen ist eine Intervalldiskretisierung vorzunehmen:

```
# sample from continuous uniform distribution
x1 = runif(10000)
# discretize into 10 categories
y1 = discretize(x1, numBins=10, r=c(0,1))
# compute entropy from counts
entropy(y1) # empirical estimate near theoretical maximum
log(10) # theoretical value for discrete uniform distribution with 10 bins
```


Informationsgewinn (Gain)

- Ansatz: Die Datenmenge Y wird nach den möglichen Werten von X partitioniert, also je eine Partition pro $c_i \in \text{Val}(X)$.

$$G(Y|X) = E(Y) - \sum_{v \in \text{Val}(X)} \frac{|Y_v|}{|Y|} E(Y_v)$$

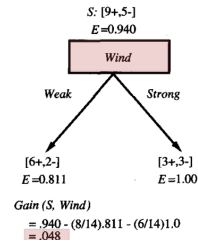
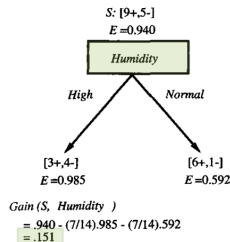
- Die Menge Y_v enthält nur Werte für die $X=v$ ist!
- Ein **Verteilungsvektor** ist dann $\text{Dist}(X)=[|v_1|, |v_2|, \dots]$ und bedeutet wie häufig der bestimmte Wert $v_i \in \text{Val}(X)$ in X auftaucht!

- Ein **Verteilungsvektor** ist dann $Dist(Y_v|X)=[|u_1|,|u_2|,...]$ und bedeutet wie häufig der bestimmte Wert $u \in Val(Y)$ in Y_v auftaucht!

Beispiele

Beispiel	Aussicht	Temperatur	Luftfeuchtigk.	Windig?	Klasse
1	sonnig	heiß	hoch	nein	N
2	sonnig	heiß	hoch	ja	N
3	bewölkt	heiß	hoch	nein	P
4	regnerisch	mild	hoch	nein	P
5	regnerisch	kalt	normal	nein	P
6	regnerisch	kalt	normal	ja	N
7	bewölkt	kalt	normal	ja	P
8	sonnig	mild	hoch	nein	N
9	sonnig	kalt	normal	nein	P
10	regnerisch	mild	normal	nein	P
11	sonnig	mild	normal	ja	P
12	bewölkt	mild	hoch	ja	P
13	bewölkt	heiß	normal	nein	P
14	regnerisch	mild	hoch	ja	N

[12]



WorkBook Live: Gain

CLEAR

LOAD

+

-

R

DataSports

AnalysisSports

DataLenses

AnalysisLenses

Principle Component Analysis

- PCA: Klassische Methode zur unüberwachten linearen Dimensionsreduktion → Analyse der Hauptkomponenten (Eigenwertproblem)



PCA ist noch keine Reduktionsmethode. PCA liefert bei einem n -dimensionalen Vektor $\vec{X}$ genau n Vektoren der Dimensionalität n !

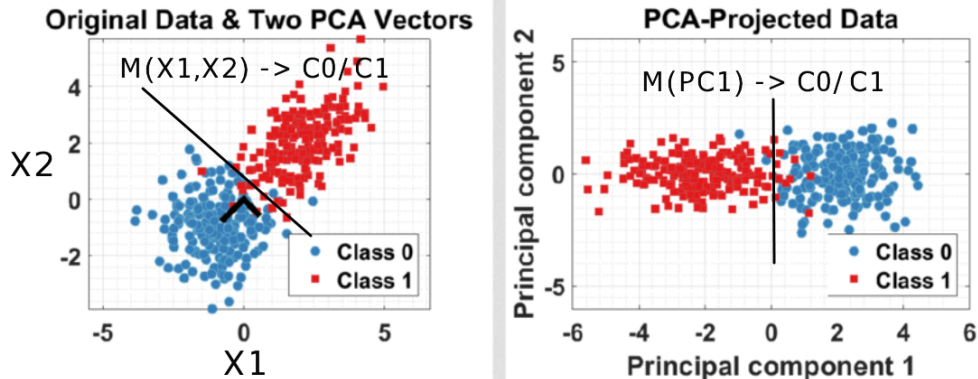
Aber: Reduktion von Redundanz in den Attributen $\mathbf{X}$ ist mit diesen Hauptkomponenten möglich.

- Bessere Trennung bei der Inferenz von kategorischen (und ggfs. auch numerischen) Zielvariablen
- Weitere Verfahren:
 - Lineare Diskriminanzanalyse (LDA)
 - Singuläre Wertzerlegung (SVD)

Beispiel

- $D=[X_1, X_2, Y]$, mit $Val(Y)=\{Class1, Class2\}$
- "Rotation" des zweidimensionalen Attributraums führt zu einer reduzierten Datentabelle $D'=[PC_1, Y]$ (PC_2 kann weg gelassen werden)

[Czarnek, RG]



Spektrale Analyse

Diskrete Fourieranalyse

- (DFT) und schnelle FA (FFT) → (Frequenzspektrum)
 - Spektrum der DFT/FFT hat nur $N/2$ Punkte (Datenreduktion)
 - Aber Problem bei DFT ist: Wenn Anzahl Datenpunkte N klein ist gibt es schlechte Auflösung des Spektrums
 - Vor allem bei "kontinuierlichen" Datenströmen mit fließenden Fensterverfahren ein Problem (Fensterbreite N_{win} ist klein)
- Es gilt: Das aus der DFT erhaltene Frequenzspektrum hat die höchste Frequenz:

$$f_s = 1/\Delta t, \Delta t = t_{i+1} - t_i, f_{\max} = f_s/2$$

Diskrete Fouriertransformation

- Eine zeitaufgelöste Variable X wird in eine frequenz aufgelöste Variable F transformiert:

$$DFT(X) : F(k) = \sum_{n=0}^{N-1} X_n e^{-i2\pi nk/N} \quad (1)$$

für $k = 0, 1, 2, \dots, N-1$

- Dabei ist X eine komplexwertige Datenserie $\{\langle re(x_i), im(x_i) \rangle | i=0, 1, \dots, N-1\}$, und
- Das Ergebnis F ist ebenfalls eine komplexwertige Datenserie $\{\langle re(f_i), im(f_i) \rangle | i=0, 1, \dots, N-1\}$
- Der Imaginärteil von X wird i.A. auf Null gesetzt (Sensoren sind i.A. reellwertig).

Amplitudenspektrum

$$Mag(X) = \frac{\sqrt{[re(DFT(X))]^2 + [im(DFT(X))]^2}}{N} \quad (2)$$

Phasenspektrum

$$Pha(X) = \arctan\left(\frac{im(DFT(X))}{re(DFT(X))}\right) \quad (3)$$

Leistungsspektrum

$$Power(X) = \frac{[re(DFT(X))]^2 + [im(DFT(X))]^2}{N} \quad (4)$$

Beispiel für Spektrale Eigenschaftsselektion

Sinus- und Kosinusschwingungen

WorkBook Live: Spektralanalyse

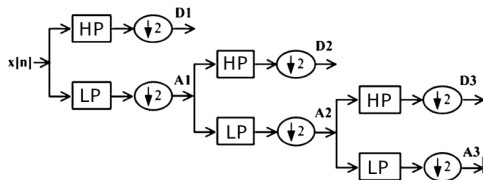
Waveletanalyse

- Diskrete Wavelet Transformation (DWT)
 - DFT liefert nur Informationen im Frequenzraum
 - DWT liefert Informationen aus Zeit- und Frequenzraum
 - Höherer Informationsgehalt

DWT verwendet lange Zeitfenster für niedrige Frequenzen und kurze Zeitfenster für höhere Frequenzen, was zu einer guten zeitfrequenzanalyse führt.

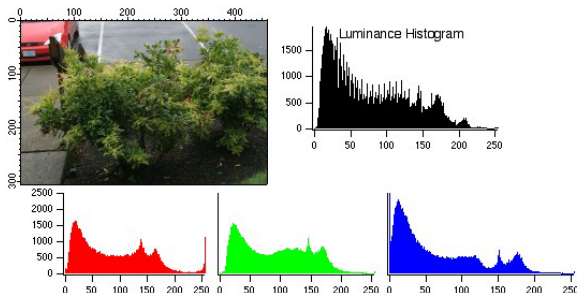
DWT kann mit digitalen Filterkaskaden aufgebaut werden:

- Jede Ebene der Filterkaskade besteht aus einem Hoch- und einem Tiefpassfilter
- Der Hochpassfilter liefert die Details, der Tiefpassfilter die Approximation der DWT auf der n -ten Ebene
- Die Approximation der Ebene n ist das Eingangssignal für die Ebene $n+1$
- In jeder Ebene wird der Eingangsvektor $|\mathbf{x}_i|=N$ auf $N/2$ reduziert, d.h. $|\mathbf{x}_{i+1}|=N/2$



Histogrammanalyse

- Ein Histogramm gibt die (intervallbasierte) Verteilung von unterscheidbaren Elementen in einer Datenmenge X an
 - Beispiel sind Histogramme von Bildern die die Verteilung der Farb-/Grauwerte mit den Werten $\{0,1,...,255\}$ wiedergeben.



[wavemetrics]

Histogrammanalyse

WorkBook Live: Histogrammanalyse

Analyse von Datenserien



Welcher Sensoren können bei der Bauteilprüfung und Schadensüberwachung Zeit- oder Datenserien erzeugen..

Messdaten

[C. R. Farrar and K. Worden, Structural Health Monitoring: A Machine Learning Perspective. Wiley-Interscience, 2013, pp. 164]

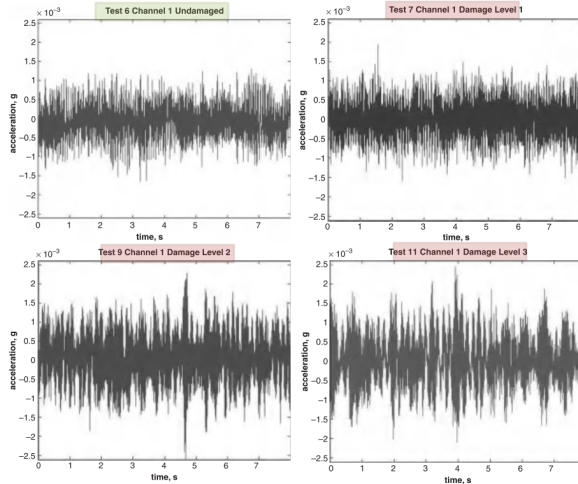


Abb. 4. Zeitaufgelöste Sensordaten $s(t)$ eines Beschleunigungssensors einer Maschine ohne und mit Schäden

Korrelationsanalyse

[C. R. Farrar and K. Worden, Structural Health Monitoring: A Machine Learning Perspective. Wiley-Interscience, 2013, pp. 164]

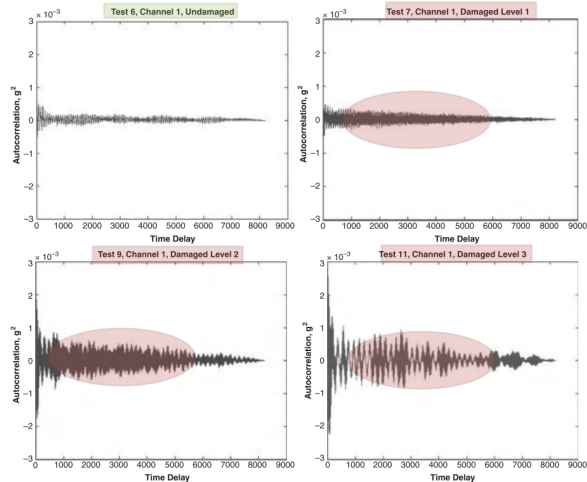


Abb. 5. Autokorrelation der zeitaufgelösten Sensordaten $s(t)$ eines Beschleunigungssensors einer Maschine ohne und mit Schäden

Spektralanalyse

[C. R. Farrar and K. Worden, Structural Health Monitoring: A Machine Learning Perspective. Wiley-Interscience, 2013, pp. 167]

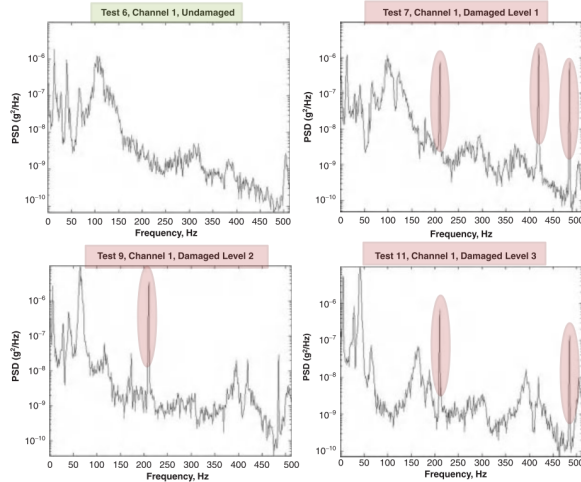


Abb. 6. Spektralanalyse der zeitaufgelösten Sensordaten $s(t)$ eines Beschleunigungssensors einer Maschine ohne und mit Schäden

Merkmalsselektion

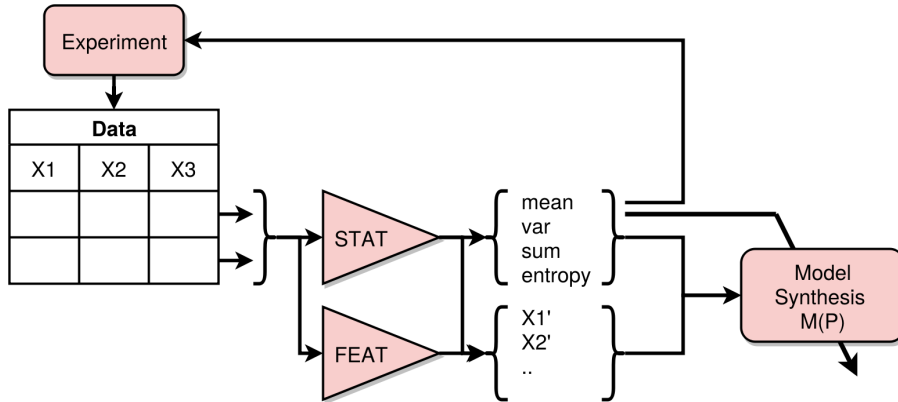


Abb. 7. Die statistische und weitere Analysen können die Eingabe für ML liefern, aber auch die Modellsynthese parametrisieren bzw. beeinflussen

Zusammenfassung

- Die statistische Analyse von Datentabellen liefert wichtige Informationen über die Qualität der Daten
- Die Merkmalsselektion transformiert die Rohdaten auf neue möglichst linear unabhängige Attribute
 - Datenreduktion → Dimensionalität
 - Datenreduktion → Datengröße
 - Datenqualitätserhöhung
- Es werden Verfahren für kategorische und numerische Datenvariablen unterschieden