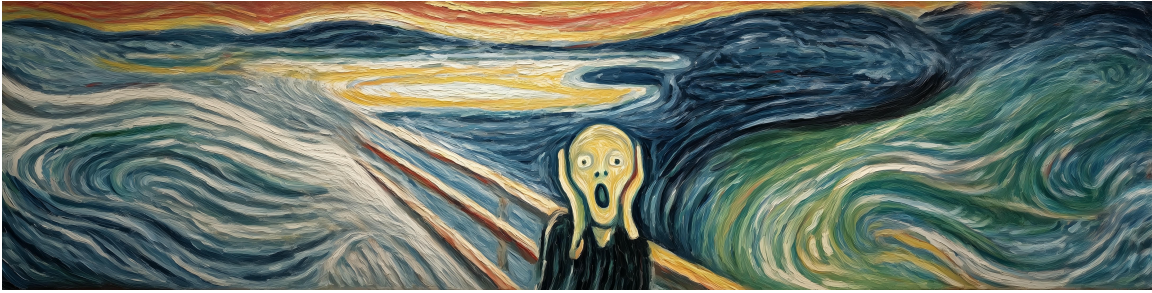


Dimensionsreduktion

Joern Ploennigs · AI4SC



To tell you the truth, I've never met anybody who can envision more than three dimensions.

— Brian Greene

Methoden zur *Dimensionsreduktion* werden im Machine Learning immer dann genutzt, wenn Datensätze vereinfacht werden sollen, um sie besser zu verstehen oder schneller zu erlernen. Hierbei teilt man die Verfahren in zwei Anwendungsfälle:

- Reduktion der Anzahl an Variablen in sehr großen Datensätzen mittels der **Hauptkomponentenanalyse**
- Reduktion der Anzahl an Beobachtungen, um typische Muster abzuleiten mittels **Matrixfaktorisierung**

Dimensionsreduktion wird häufig in der Vorverarbeitung von ML-Modellen eingesetzt, um etwa *Overfitting* zu vermeiden, die Visualisierung zu erleichtern oder Trainingszeiten zu verkürzen. So können die Verfahren zur Reduktion von Sensor-Messwerten in Brücken oder Gebäuden benutzt werden oder zur Verdichtung komplexer Materialprüfdaten auf wenige Hauptkomponenten.

Zum Vergleich: PCA sucht lineare Richtungen maximaler Varianz und eignet sich besonders zur Verdichtung stark korrelierter Variablen. NMF zerlegt einen Datensatz dagegen in additive, nichtnegative Bausteine und liefert dadurch oft besser interpretierbare Muster oder Prototypen. PCA ist also meist das Mittel der Wahl für kompakte Projektionen, NMF eher für interpretierbare Komponenten in nichtnegativen Daten.

Hauptkomponentenanalyse (PCA)

Methode

Die Hauptkomponentenanalyse (en. *Principal Component Analysis*, PCA) wird eingesetzt, um Datensätze mit vielen Variablen zu vereinfachen, damit wir sie besser interpretieren und nutzen

können und damit sich kompaktere Modelle trainieren lassen. Das ist insbesondere bei Datensätzen mit vielen, stark korrelierten Variablen sinnvoll, da sowohl Nutzer die vielen Variablen nur schwer verstehen als auch viele lineare Modelle bei stark korrelierten Variablen instabil werden oder zusätzliche Regularisierung benötigen.

Stellen wir uns vor, wir überwachen eine Brücke mit 10 verschiedenen Sensoren. Viele der Messreihen beschreiben ähnliche Effekte und sind stark korreliert. PCA hilft dabei, die wichtigsten Muster, z.B. eine Vibration in X/Y-Zugrichtung, aus diesen Daten herauszufiltern und auf 2 bis 3 Hauptachsen zu reduzieren.

Karl Pearson beschrieb die Methode bereits 1901 [1]. Ziel ist es, eine Matrix mit vielen Spalten (Dimensionen) auf einige wenige lineare Komponenten (Hauptkomponenten) zu reduzieren, die den Datensatz möglichst gut beschreiben. Dafür suchen wir neue, synthetische Achsen, entlang derer die Varianz im Datensatz möglichst hoch ist. Gleichzeitig werden die Daten dekorreliert. Anders als bei linearen Regressionsmodellen ist es hierbei nicht das Ziel, eine einzelne Ausgangsvariable zu modellieren, sondern den vollständigen Datensatz.

Wir betrachten den bekannten Energie- und Wetter Datensatz.

```
import numpy as np # Import von NumPy
import pandas as pd # Import von Pandas

egywth = pd.read_csv("../data/UR05/EnergyID_weather_clean.csv",
parse_dates=[0])
egywth.head()
```

EVIDENCE_ID_AIR_SAB_IDEILAMKONSDATUN_3 ..TMKUPMTXKTNKTGK eorTemHeizNSQNF_4															Kuehl- ra- Ta- tur- ge Klas- se					
0	2020-12-30	Na	Na	Na	Na	Na	2020-12-30	2020-12-30	10.3	0.0	4.0	81.0	4.6	2.9	2.2	eor	Cold	Heiz	5	
	00:00:00	+00:00					00:00:00	+00:00										g- al- rad-le tag Pa- ra- me- ter kor- ri- giert		
1	2020-12-30	Na	Na	Na	125629	1.5	0.0	2020-12-30	2020-12-30	10.3	0.0	3.4	83.0	4.4	2.3	0.9	eor	Cold	Heiz	5
	00:00:00	+00:00						00:00:00	+00:00									g- al-		

EVL_HV_Nr_7AB_Sab_DASILAMKONSDATUN_3 ..TMKUPMTXKTNKTGK eorTemHeiQNSNF_4
 Kuehl-
 ra- Ta-
 tur- ge
 Klas-
 se

rad-le
 tag Pa-
 ra-
 me-
 ter
 kor-
 ri-
 giert

2

2021001-008010221290.0 2021427-00210001.0 2.0 90.0 3.0 1.1 0.3 eor ColdHeiznich5
 00:00:00+00:00 00:00:00+00:00

g- al-
 rad-le
 tag Pa-
 ra-
 me-
 ter
 kor-
 ri-
 giert

3

2021001-008010221290.0 2021427-00210002.0 3.3 95.0 4.0 2.6 2.0 eor ColdHeiznich5
 00:00:00+00:00 00:00:00+00:00

g- al-
 rad-le
 tag Pa-
 ra-
 me-
 ter
 kor-
 ri-
 giert

4

2021001-008010221290.0 2021427-00210003.0 3.5 81.0 4.6 2.4 0.8 eor ColdHeiznich5
 00:00:00+00:00 00:00:00+00:00

g- al-
 rad-le
 tag Pa-
 ra-
 me-

EV_HT_740 EV_NT_740 EV_AV_Lab EV_SV_Lab FX FM RSK RSKF SDK SHK_TAG NM VPM PMT MKUP MTXK TNKTGK eorTemH eorTNS eorNF_4

Kuehl-
ra- Ta-
tur- ge
Klas-
se

ter
kor-
ri-
giert

Wir analysieren vom Datensatz allerdings nur die numerischen, nicht-NA Spalten, weshalb wir sie selektieren.

```
egywthNumber = egywth.select_dtypes(include='number').dropna()
egywthNumber = egywthNumber[['EV_HT_740', 'EV_NT_740', 'E_AV_Lab', 'E_SV_Lab',
'E_S_Lab', 'FX', 'FM', 'RSK', 'RSKF', 'SDK', 'SHK_TAG', 'NM', 'VPM', 'PM',
'TMK', 'UPM', 'TXK', 'TNK', 'TGK']]
```

Schritt 1 - Datenzentrierung: Die Daten werden zentriert, indem der Mittelwert jeder Dimension subtrahiert wird. Dadurch liegen alle Daten um den Ursprung herum.

$$\mathbf{X}_{\text{cent}} = \mathbf{X} - \boldsymbol{\mu}_X.$$

```
egywthCentered = egywthNumber - egywthNumber.mean()
egywthCentered.head()
```

EV_HT_740 EV_NT_740 EV_AV_Lab EV_SV_Lab FX FM RSK RSKF SDK SHK_TAG NM VPM PMT MKUP MTXK TNKTGK

2	-1245.258071	1716.925364	293.668162	226.920174	563.97305	197.948534	108.537838	106.090025	88.470497	4891
3	-75.840614	1578.925364	293.668162	20.820783	633.97305	197.948525	100.537831	109.082588	4703.974891	
4	-1249.281081	1716.925364	293.668162	221.420174	563.97305	197.948534	108.537838	106.090025	88.470497	4891
5	1994.159189	1752.925364	293.668162	783.800983	633.97305	197.948534	107.537831	106.090025	88.470497	4891
6	2224.159019	1752.925364	293.668162	783.800983	633.97305	197.948528	107.537831	106.090025	88.470497	4891

Schritt 2 - Berechnung der Kovarianzmatrix: Die Kovarianzmatrix C wird berechnet, um die Streuung und Korrelation zwischen den verschiedenen Dimensionen zu erfassen.

$$\mathbf{C} = \frac{\mathbf{X}_{\text{cent}}^T \mathbf{X}_{\text{cent}}}{n - 1}.$$

```
cov_matrix = np.cov(egywthCentered, rowvar=False)
```

Schritt 3 - Eigenwertzerlegung der Kovarianzmatrix: Die Kovarianzmatrix wird einer Eigenwertzerlegung unterzogen, um die Eigenwerte und Eigenvektoren zu berechnen. Die Eigenvektoren repräsentieren die Richtungen der größten Varianz, und die Eigenwerte repräsentieren die Größe der Varianz in diesen Richtungen. Der Eigenvektor einer quadratischen Matrix ist ein Vektor, der durch die Matrix nur um einen skalaren Faktor, den Eigenwert λ , gestreckt oder gestaucht wird. Mathematisch ausgedrückt:

$$Cv = \lambda v$$

dabei ist λ der skalare Eigenwert und v der zugehörige Eigenvektor.

```
# Berechne die Eigenwerte und Eigenvektoren
eigenvalues, eigenvectors = np.linalg.eig(cov_matrix)
```

Schritt 4 - Auswahl der Hauptkomponenten: Die Hauptkomponenten (*Principal Components*) sind die Eigenvektoren, die den größten Eigenwerten entsprechen. Diese Eigenvektoren definieren die neuen Achsen des transformierten Datenraums. Hierfür wählen wir die besten n Komponenten in der Matrix W .

```
# Sortiere die Eigenvektoren nach absteigenden Eigenwerten
idx = np.argsort(eigenvalues)[::-1]
eigenvalues = eigenvalues[idx]
eigenvectors = eigenvectors[:, idx]

# Wähle die n Hauptkomponenten
n_components = 2
eigenvalues_sub = eigenvalues[:n_components]
W = eigenvectors[:, :n_components]
```

Schritt 5 - Transformation der Daten: Die ursprünglichen Daten werden auf die neuen Hauptkomponenten projiziert, um die transformierten Daten zu erhalten.

$$X_{pca} = X_{cent} \cdot W$$

```
# Transformiere die Daten
egywth_pca = np.dot(egywthCentered, W)
egywth_pca
```

```
array([[1764.35300097, -144.63239642],
       [-70.67495979,  177.86671771],
       [1547.04272956,   99.44054371],
       ...,
       [ 920.66526878,   888.17908265],
       [ 914.15695    ,   895.99134945],
       [ 803.76433239, 1030.65117329]], shape=(916, 2))
```

PCA mit SciKit-Learn

Wichtig in der Praxis ist, dass PCA empfindlich gegenüber der Skalierung der Variablen ist. Variablen mit deutlich unterschiedlichen Skalen oder Größenordnungen können die ersten Hauptkomponenten sonst dominieren. Deshalb standardisiert man die Daten häufig vor einer PCA, wenn die Variablen unterschiedliche Einheiten oder Größenordnungen haben.

Für die Praxis gibt es dafür auch eine einfache Implementierung in sklearn, die wir wie gewohnt einsetzen. Wir können zum Training `fit` benutzen. Wollen wir allerdings gleich die Komponenten erhalten, bietet sich alternativ `fit_transform` an:

```
from sklearn.decomposition import PCA

pcam = PCA(n_components=2)
components = pcam.fit_transform(egywthNumber)
```

Wir können uns mit dem Attribut `explained_variance_` die Varianz der einzelnen Hauptkomponenten ausgeben lassen. Diese Werte entsprechen den Eigenwerten der Kovarianzmatrix. Das Attribut `singular_values_` enthält dagegen die Singulärwerte der zentrierten Datenmatrix aus der SVD, die sklearn intern zur Berechnung der PCA verwendet.

```
pcam.singular_values_
```

```
array([61197.4192215 , 11416.98508508])
```

Wir können uns auch mit `explained_variance_ratio_` die erklärte Varianz ausgeben lassen.

```
print(pcam.explained_variance_ratio_)
```

```
[0.96270602 0.0335066 ]
```

Hier sehen wir, dass die erste Komponente mit 96.7% den Großteil der Varianz erklärt, die zweite mit 3% deutlich weniger.

Als nächstes wollen wir die Komponenten visualisieren. Das ist ein weiterer wichtiger Anwendungsbereich von PCA. Da sich Datensätze mit sehr vielen Variablen schlecht visualisieren lassen, kann man sie mit PCA auf eine 2D- oder 3D-Darstellung reduzieren. Dies wird z.B. häufig beim Clustering verwendet, um höherdimensionale Cluster zu visualisieren.

```
import plotly.express as px
# Visualisierung der transformierten Daten
fig = px.scatter(components, x=0, y=1, width=500, height=500)
fig.update_layout(xaxis_title="Hauptkomponente 1",
yaxis_title="Hauptkomponente 2")
# Visualisierung der Eigenvektoren
```

```

loadings = pcam.components_.T * np.sqrt(pcam.explained_variance_)
for i, feature in enumerate(egywthCentered.columns):
    fig.add_annotation(ax=0, ay=0, axref="x", ayref="y", x=loadings[i, 0],
y=loadings[i, 1], showarrow=True, arrowsize=2, arrowhead=2, xanchor="right",
yanchor="top")
    fig.add_annotation( x=loadings[i, 0], y=loadings[i, 1], ax=0, ay=0,
xanchor="center", yanchor="bottom", text=feature, yshift=5 )
fig.show()

```

Unable to display output for mime type(s): text/html

Unable to display output for mime type(s): text/html

Welche Merkmale laden besonders stark auf die erste Hauptkomponente, und was sagt das über deren physikalische Bedeutung aus? Im Diagramm erkennen wir drei Gruppen um "EV_NT_740", "EV_HT_740" und "E_AV_Lab". Um letztere gruppieren sich auch die gesamten Wetterwerte. Das deutet darauf hin, dass diese deutlich stärker mit "E_AV_Lab" korreliert sind als mit den anderen Werten.

Nichtnegative Matrixfaktorisierung (NMF)

Methode

Matrixfaktorisierung wird genutzt, um die Anzahl der Beobachtungen (die Anzahl der Variablen bleibt gleich) zu reduzieren. Dabei zerteilt man einen Datensatz in zwei niedrigdimensionale Matrizen. Das ist sinnvoll, wenn man die Daten komprimieren oder latente Merkmale extrahieren möchte. Es wird z.B. bei der Bilderkennung benutzt, aber auch als Alternative zum Clustering, wenn keine klaren Cluster erkennbar sind.

Die *nichtnegative Matrixfaktorisierung* ist ein typischer Ansatz, bei dem alle Elemente der Matrizen nichtnegativ sind (sein müssen). Dies ist besonders nützlich für Anwendungen, bei denen die Daten nichtnegativ sind, wie z.B. Bilder, Texte und Empfehlungssysteme. Gegeben ist die nichtnegative Matrix $\mathbf{X}$ mit Dimension $m \times n$, die wir in die zwei Matrizen $\mathbf{W}$ und $\mathbf{H}$ mit den Dimensionen $m \times k$ und $k \times n$ zerlegen wollen, so dass

$$\mathbf{X} \approx \mathbf{W}\mathbf{H}.$$

wobei k die Anzahl der latenten Merkmale ist.

Wir wollen $\mathbf{W}$ und $\mathbf{H}$ finden, so dass die Differenz zwischen $\mathbf{X}$ und $\mathbf{W}\mathbf{H}$ minimiert wird. Die folgenden multiplikativen Aktualisierungsregeln minimieren den Rekonstruktionsfehler auf Basis der Frobenius-Norm [2]:

$$\min_{\mathbf{W}, \mathbf{H}} \|\mathbf{X} - \mathbf{W}\mathbf{H}\|_F^2$$

Eine gebräuchliche Methode dafür ist der folgende iterative Algorithmus.

Schritt 1 - Initialisiere: Initialisiere W und H mit zufälligen nichtnegativen Werten.

Schritt 2 - Iterative Aktualisierung: Aktualisiere W und H iterativ unter Verwendung der multiplikativen Aktualisierungsregeln:

$$H \leftarrow H \odot \frac{W^T X}{W^T W H}$$
$$W \leftarrow W \odot \frac{X H^T}{W H H^T}$$

wobei $\odot$ das elementweise Produkt und $\frac{A}{B}$ die elementweise Division bezeichnen. Neben dieser Variante auf Basis der Frobenius-Norm gibt es auch andere Zielfunktionen, z.B. Divergenzmaße für Zählraten.

Schritt 3 - Konvergenzkriterium: Überprüfe die Konvergenz, z.B. ob die Änderung des Rekonstruktionsfehlers unter einen bestimmten Schwellenwert fällt oder eine maximale Anzahl von Iterationen erreicht ist.

Schritt 4 - Rekonstruktion: Die rekonstruierte Matrix $\hat{X}$ ist das Produkt der Faktormatrizen W und H :

$$\hat{X} = W H.$$

NMF in SciKit-Learn

Auch für die nichtnegative Matrixfaktorisierung gibt es eine Implementierung in `sklearn`. Zu beachten ist, dass sie allerdings nur nichtnegative Werte verarbeiten kann. Deshalb müssen wir unsere Daten zuerst auf einen positiven Wertebereich *normalisieren*. Dafür können wir den `MinMaxScaler` verwenden, der unsere Daten auf den Wertebereich 0, ..., 1 transformiert.

$$X_{\text{mx}} = \frac{X - \min_X}{\max_X - \min_X}.$$

mit dem Vektor der minimalen Werte $\min_X$ und maximalen Werte $\max_X$ aller Variablen in X .

```
from sklearn.preprocessing import MinMaxScaler

mms = MinMaxScaler()
egywthMinMax = mms.fit_transform(egywthNumber)
mms.min_, mms.scale_
```

```
(array([ 0.          , -0.14705882, -0.17349138, -0.19689119,
         0.          , -0.14126394, -0.08982036,  0.          ,
         0.          ,  0.          ,  0.          ,  0.          ,
        -0.13333333, -14.3641576 ,  0.1627907 , -0.55099483,
         0.09004739,  0.28315412,  0.415625  ]),
 array([0.00022624, 0.00022624, 0.00053879, 0.00259067, 0.00847458,
        0.03717472, 0.05988024, 0.02427184, 0.125          , 0.06198475,
```

```
0.08333333, 0.125      , 0.05128205, 0.01470156, 0.02906977,  
0.01566661, 0.02369668, 0.03584229, 0.03125    ]))
```

Die Anwendung der NMF ist nun einfach.

```
from sklearn.decomposition import NMF  
  
nmfm = NMF(n_components=2)  
W = nmfm.fit_transform(egywthMinMax)  
H = nmfm.components_
```

Hier gibt nun die Matrix W an, aus welchen Komponenten sich jede Beobachtung in unserem Datensatz zusammensetzt. Zum Beispiel sehen wir, dass sich die ersten zwei Zeilen nur aus der Komponente 1 zusammensetzen.

```
pd.DataFrame(W, columns=[f'Merkmal {i+1}' for i in range(n_components)])
```

	Merkmal 1	Merkmal 2
0	0.237477	0.000000
1	0.241596	0.000000
2	0.223237	0.016886
3	0.222422	0.028341
4	0.252235	0.000574
...	...	...
911	0.202307	0.057166
912	0.186173	0.083753
913	0.200195	0.037590
914	0.209245	0.025975
915	0.205023	0.028051

Interessanter ist die Matrix H , weil sie die extrahierten Komponenten darstellt. Hierfür transformieren wir sie zurück in den originalen Datenbereich, indem wir den MinMaxScaler mit `inverse_transform` anwenden.

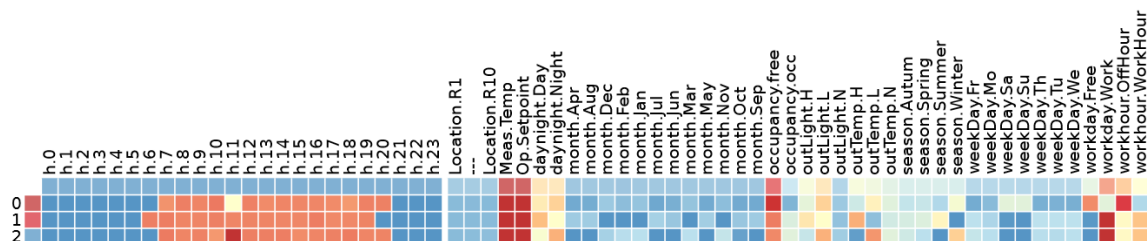
```
H_rescaled=mms.inverse_transform(H)  
pd.DataFrame(H_rescaled, columns=egywthNumber.columns)
```

EV_HEV_NE_WEISV_Lab FX FM RSKRSKSKTAG NMVPM PMTMKUPMTXKTNTKGK

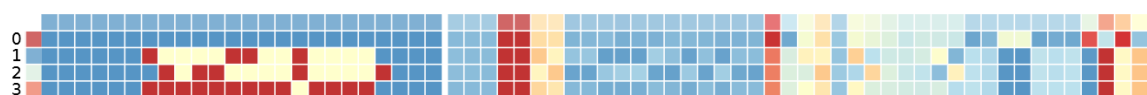
0 5016.197110826193.823007453438531252161344000033849511333011692363024452358957289504538804
1 2445.9525833622846195.9402348920776000002650599040107132429763.235893831504407605735254881

Diese zwei Komponenten sind nun zwei charakteristische Kombinationen, mit denen der Datensatz in dieser niedrigdimensionalen Darstellung approximiert wird. So lassen sich aus großen Datensätzen wichtige Muster ableiten. Man kann sie mit Vorsicht ähnlich wie typische Muster oder Prototypen auffassen.

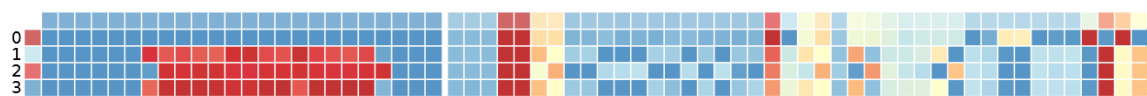
Die folgende Abbildung aus der Publikation „Understanding building operation from semantic context“ [3] zeigt das an einem Beispiel. In der Publikation werden die Steuerungsstrategien von Gebäuden aus Verbrauchsdaten abgeleitet. Die Abbildung vergleicht die Muster, die durch Clustering (k -Means), SiVM-Matrixfaktorisierung (vergleichbar mit NMF, die insbesondere charakteristische Muster identifiziert) und SiVM-Clustering (eine Kombination beider) identifiziert werden. Hierbei zeigt sich, dass die Muster, die durch SiVM-Matrixfaktorisierung identifiziert werden, deutlich ausgeprägter sind. Das liegt an der besonderen Methode SiVM zur Identifikation der Komponenten. Daraus wurden anschließend über *Assoziationsanalyse* Betriebsregeln abgeleitet, die Entscheidungsbäumen ähneln.



(b) Daily pattern extracted by k -means clustering



(c) Daily pattern extracted by SiVM matrix factorisation



(d) Daily pattern extracted by SiVM k -means clustering

Referenzen

Bibliographie

- [1] K. Pearson, „On Lines and Planes of Closest Fit to Systems of Points in Space“, *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science*, Bd. 2, Nr. 11, S. 559–572, 1901, doi: 10.1080/14786440109462720.
- [2] D. D. Lee und H. S. Seung, „Algorithms for Non-negative Matrix Factorization“, in *Advances in Neural Information Processing Systems 13*, 2001, S. 556–562.
- [3] J. Ploennigs, B. Chen, und A. Schumann, „Understanding building operation from semantic context“, in *IECON - An. Conf. IEEE Ind. Elec. Soc.*, 2015.