

# Datenvorverarbeitung

Joern Ploennigs · AI4SC



Eine Person, die weiß, dass sie nichts weiß, weiß mehr als eine Person, die nicht weiß, dass sie nichts weiß.

— Socrates

Dieser Abschnitt erklärt und demonstriert bestimmte Datenbereinigungs- und Vorbereitungsaufgaben unter Verwendung von Pandas. Die Aufgabe hier besteht hauptsächlich darin, Sie mit verschiedenen nützlichen Funktionen vertraut zu machen und zu zeigen, wie häufige Aufgaben gelöst werden. Hierbei werden nicht grundlegende Datenverarbeitungsprobleme gelöst.

In der Baupraxis stammen viele Daten aus Sensorik, Messungen oder manuellen Eingaben. Diese sind häufig unvollständig oder fehlerhaft, was zu Problemen in der Auswertung und Modellbildung führen kann. So können, fehlende oder falsch codierte Werte ML-Modelle zum Absturz bringen oder das Training verzerren. Eine saubere Vorbereitung der Daten ist daher unverzichtbar.

## Fehlende Beobachtungen

In vielen Datensätzen aus der Praxis stoßen wir auf fehlende Werte. Sie entstehen weil (i) Werte nicht regelmäßig erfasst wurden, (ii) Werte historisch nicht gemessen wurden, (iii) Messgeräte ausfallen oder (iv) Werte entfernt worden sind, weil man dachte, dass sie nicht benötigt werden.

Solche fehlenden Werte, können sehr schnell Berechnungen verfälschen oder unmöglich machen. Zum Beispiel, meldet ein Temperatursensor auf einer Brücke im Fehlerfall -9999. Wird dieser Wert bei der Mittelwertbildung nicht erkannt, könnte er die statistische Analyse verfälschen. Deshalb ist es wichtig, fehlende Werte zu identifizieren und diese richtig zu behandeln.

### Wie werden fehlende Daten codiert?

Grundsätzlich ist es wichtig das man fehlende Werte in den Daten codiert und abspeichert, denn nur so kann man unterscheiden, ob sie *bewusst* oder *zufällig* fehlen.

| Sprache / Tool          | Fehlender Wert       |
|-------------------------|----------------------|
| Python (NumPy)          | <code>np.nan</code>  |
| Python (Pandas ab v1.0) | <code>pd.NA</code>   |
| R                       | NA                   |
| Julia                   | <code>missing</code> |
| Excel                   | Leeres Feld          |
| Geräte                  | -9999, 99999 o.ä.    |

Dabei unterscheidet sich auch die Bedeutung der Codierung. In Numpy steht `np.nan` eigentlich für „Not a Number“. In Pandas steht `pd.NA` für „Not Available“ und ist eindeutiger als `np.nan`. Diese Vielfalt macht es notwendig, sich immer die Datendokumentation anzusehen.

Zum Beispiel schauen Sie sich eine Lieferliste in Excel an. Bei einigen Lieferungen ist die Zelle der gelieferten Teile leer. Jetzt ist die Beutung unklar: (i) wurde nichts geliefert; (ii) wurde etwas geliefert, aber nicht erfasst; (iii) wurde der Eintrag gelöscht oder (iv) ist der Eintrag selbst falsch und es gab keine Lieferung. Deshalb ist die Aussage das ein Wert nicht da ist stets zu unterscheiden von einer leeren Menge.

Zur Codierung dieser fehlenden Werte bieten einige Sprachen für Machine Learning spezielle Werte an. So gibt es in *R* den Wert NA („not available“) und in *Julia* oder *Matlab* den Wert `missing` (fehlend). Zusätzlich gibt es den numerischen Wert NaN (not a number) für ungültige mathematische Ergebnisse (wie den Logarithmus einer negativen Zahl) zu kennzeichnen. Der entscheidende Unterschied ist, dass für NaN mathematische Operationen definiert sind, die meist NaN als Ergebnis haben. Ein numerischer Algorithmus, kann also ohne Fehler weiter rechnen. Das ist bei NA / `missing` nicht der Fall, wo Berechnungen scheitern.

Numpy bietet für die Codierung nicht vorhandener **nur** den speziellen numerischen Datentyp `np.nan` an. Dadurch wird gewährleistet, dass numerische Berechnungen ausführbar sind. Es ist zu beachten, dass `np.nan` numerisch ist, wodurch eine Serie mit fehlenden Werten entweder numerisch sein sollte oder bei der Verwendung mit nicht-numerischen Werten (*str*, *bool*) zum generischen Typ *object* wechselt. Wenn beispielsweise ein fehlender Wert `np.nan` zu einer logischen Serie hinzugefügt wird, ist das Ergebnis vom Typ `_object`. In manchen Fällen ist es dennoch wichtig zu unterscheiden, ob der Wert das Ergebnis einer ungültigen Berechnung ist (z.B.  $\log(-1)$ ) oder ein fehlender Wert. Deshalb hat Pandas 2.0.3 der spezielle Wert `pd.NA` eingeführt, bei deren Verwendung sich der Datentyp einer Serie nicht verändert. Alternativ kann auch der Python Datentyp `None` zur Codierung von fehlenden Werten verwendet werden, er wird automatisch in `pd.NA` konvertiert. Es ist empfehlenswert `pd.NA` anstatt von `np.nan` zu benutzen, da es eindeutig ist und die Datentypen nicht verändert.

Aber nicht alle fehlenden Werte werden mit `np.nan`/`pd.NA` codiert. Zum Beispiel ist es in den Sozialwissenschaften üblich, fehlende Werte als negative Zahlen zu kennzeichnen. Die World Value Survey codiert alle fehlenden Antworten als -1: weiß nicht; -2: keine Antwort; -3: nicht zutreffend; -4 nicht in der Umfrage gestellt; -5: fehlend, unbekannt. Fehlende Zeichenfolgen

können als leere Zeichenfolge, als „N/A“ oder auf unzählige andere Arten markiert werden. Andere Umfragen können positive Werte außerhalb des Bereichs verwenden, z.B. kann fehlendes Einkommen als 90000 oder 99999 oder ähnliches codiert sein. Auch Messgeräte nutzen oft numerische Werte, die außerhalb des Messbereiches liegen um fehlende Werte zu codieren, z.B. -9999 bei einer Temperatur. Der erste Schritt, den man hier unternehmen sollte, ist, die Datendokumentation zu konsultieren. Einige Datensätze verfügen über recht umfangreiche Dokumentationen, in denen die Codierung von Werten sehr gut erklärt wird. Aber viele Datensätze haben sehr knappe Dokumentationen, wenn überhaupt, und selbst gut dokumentierte hochwertige Datensätze können Fehler enthalten. Deshalb sollte man immer auf solche Werte achten, da sie bei großen Datensätzen schnell untergehen können und Ergebnisse stark verzerren.

Im Folgenden gehen wir die wichtigsten Tools in Pandas und Numpy durch, die dabei helfen, fehlende Werte zu identifizieren, zu entfernen oder zu ersetzen. Da die speziellen Funktionen nur mit np.nan-Codes funktionieren, geben wir auch Beispiele dafür, wie man benutzerdefinierte Codes und Eingabefehler behandeln kann.

### Umgang mit fehlenden Werten

Um den Umgang mit fehlenden Werten zu lernen, schauen wir uns den Wetter Datensatz des DWD für die Wetterstation in Warnemünde vor 1960 an.

```
import numpy as np # Import von NumPy
import pandas as pd # Import von Pandas
np.random.seed(3) # Fixierung des Zufallszahlengenerators zur besseren Erklärung

wetter = pd.read_csv("../data/Wetter/warnemuende_1960.csv", sep=';')
wetter.shape
```

```
(7209, 19)
```

```
from myst_nb import glue
glue(„wetter_size1“, wetter.shape[0]) glue(„wetter_size2“, wetter.shape[1])
```

```
wetter.loc[2:4, „TMK“]=pd.NA wetter.loc[ 7, „TMK“]=pd.NA
```

Die Originaldaten enthalten {glue}wetter\_size1 Zeilen und {glue}wetter\_size2 Spalten.

Als erstes schauen wir uns die Daten mal an. Hierfür bieten sich die \*.head() und \*.tail() Funktionen, da sie uns den Anfang und das Ende des Datensatzes zeigen.

```
wetter.head()
```

| STATION | MESSDATUM | Q_N_3        | FX | FM | Q_N_4  | RSK         | RSK | SK | SK | TAG     | NMVPM | PMT  | TKUP | PMTXK | TNKTGK | eor |
|---------|-----------|--------------|----|----|--------|-------------|-----|----|----|---------|-------|------|------|-------|--------|-----|
| 0       | 427119470 | 1999-999-999 | 3  |    | -999.0 | 999-999-999 | 999 | 41 |    | 1019.32 | 178.0 | -1.5 | -3.3 | -5.8  | eor    |     |
| 1       | 427119470 | 1999-999-999 | 3  |    | -999.0 | 999-999-999 | 999 | 36 |    | 1032.44 | 80.0  | -2.4 | -6.2 | -4.5  | eor    |     |

```
STATION_ID DATE TIME MON_3 FX FMON_4 RSKRSKISSK TAG NMVPM PMTMKUPMTXKTNKTGK eor
2 4271 19470-999-999.0 -999.0999-999.0999-999.19 1031.912.979.0 -9.4 -14.0-16.0 eor
3 4271 19470-999-999.0 -999.0999-999.0999-999.19 1023.412.982.0 -9.8 -15.3-19.0 eor
4 4271 19470-999-999.0 -999.0999-999.0999-999.24 1017.59.9 84.0 -8.7 -13.7-19.0 eor
```

Als erstes schauen wir uns die Daten mal an. Hierfür bieten sich die `head` und `tail` Funktionen, da sie uns den Anfang und das Ende des Datensatzes zeigen.

```
wetter.tail()
```

```
STATION_ID DATE TIME MON_3 FX FMON_4 RSKRSKISSK TAG NMVPM PMTMKUPMTXKTNKTGK eor
7204 4271 19666926 -9996.6 5 0.3 1 2.2 0 6.6 12.8 1014.03.1 86.0 15.1 9.5 7.9 eor
7205 4271 19666927 -99910.0 5 0.0 1 6.1 0 5.8 10.1 1014.82.0 72.0 13.3 10.9 10.4 eor
7206 4271 19666928 -9994.5 5 0.0 0 9.9 0 2.6 9.1 1017.5.7 74.0 12.5 7.7 8.7 eor
7207 4271 19666929 -9993.3 5 0.0 0 7.8 0 4.8 9.2 1015.8.7 78.0 14.9 5.1 3.1 eor
7208 4271 19666930 -9993.5 5 0.0 0 7.8 0 2.2 10.9 1007.91.1 83.0 16.1 5.8 4.8 eor
```

Es gibt also am Anfang des Datensatzes sehr viele Werte die -999 oder -999.0 sind. Am Ende allerdings nur noch in der Spalte FX. Wir können also davon ausgehen, dass in dem Datensatz fehlende Werte mit -999 codiert worden sind. Die fehlenden Werte sind auch erklärbar, da 1947 Messungen noch händisch durchgeführt wurden und nur einfache Größen wie Temperatur erfasst wurden.

Diese fehlenden Werte müssen wir mit `np.nan` oder `pd.NA` ersetzen. Hierfür können wir die Dataframe-Methode `replace` an.

```
wetter = wetter.replace(-999, pd.NA)
wetter.head()
```

```
STATION_ID DATE TIME MON_3 FX FMON_4 RSKRSKISSK TAG NMVPM PMTMKUPMTXKTNKTGK eor
0 4271 19470-NA*NA*NA5 <NA*NA*NA*NA*NA4.1 1019.32.1 78.0 -1.5 -3.3 -5.8 eor
1 4271 19470-NA*NA*NA3 <NA*NA*NA*NA*NA3.6 1032.44.8 80.0 -2.4 -6.2 -4.5 eor
2 4271 19470-NA*NA*NA5 <NA*NA*NA*NA*NA1.9 1031.912.979.0 -9.4 -14.0-16.0 eor
3 4271 19470-NA*NA*NA5 <NA*NA*NA*NA*NA1.9 1023.412.982.0 -9.8 -15.3-19.0 eor
4 4271 19470-NA*NA*NA5 <NA*NA*NA*NA*NA2.4 1017.59.9 84.0 -8.7 -13.7-19.0 eor
```

Alternativ können wir die Daten direkt aus der CSV mit NA Werten laden, indem wir den Parameter `na_values` angeben.

```
wetter = pd.read_csv("../data/Wetter/warnemuende_1960.csv", sep=';',
na_values=['-999', 'NA'])
```

Jetzt haben wir einen Datensatz, in dem die fehlenden Werte richtig codiert sind. Jetzt ist es sinnvoll zu evaluieren, wie viele Daten fehlen. Wir können mit der Methode `count` beginnen: Diese Methode gibt die Anzahl der *gültigen* Beobachtungen für ein Dataframe oder eine einzelne Variable an. Und „gültig“ bedeutet hier alles außer `np.nan`. Die Anzahl der gültigen Beobachtungen für jede Variable ist:

```
wetter.count()
```

```
STATIONS_ID    7209
MESS_DATUM     7209
QN_3           4656
FX              0
FM            4656
QN_4           7209
RSK            5752
RSKF           5752
SDK            5752
SHK_TAG        5752
NM             5752
VPM            7209
PM             7209
TMK            7209
UPM            7209
TXK            7209
TNK            7209
TGK            7209
eor            7209
dtype: int64
```

Wir können sehen, dass einige Variablen, z.B. *TNK* und *TGK*, gültige Werte für alle `{glue}wetter_size1` Fälle haben, während *FX* keinen einzigen gültigen Werte hat.

Alternativ können wir die Anzahl der fehlenden Werte direkt zählen. Hierfür können wir zwei Methoden kombinieren: `isna` (oder `isnull`) gibt für jede Beobachtung an, ob sie fehlt (wahr oder falsch), und `sum` addiert diese Wahrheitswerte und Falschwerte, wobei Wahrheitswerte in Einsen und Falschwerte in Nullen umgewandelt werden. Auf diese Weise kann dieser Ansatz wie `count` die Anzahl der fehlenden Werte für jede Variable in einem Dataframe oder nur für eine einzelne Variable in einer Serie produzieren. So können wir die Anzahl der fehlenden Werte wie folgt erhalten:

```
wetter.isna().sum()
```

```

STATIONS_ID      0
MESS_DATUM       0
QN_3             2553
FX               7209
FM              2553
QN_4             0
RSK             1457
RSKF            1457
SDK             1457
SHK_TAG          1457
NM              1457
VPM             0
PM              0
TMK             0
UPM             0
TXK             0
TNK             0
TGK             0
eor             0
dtype: int64

```

Offensichtlich können wir das gleiche Ergebnis erzielen, indem wir die Anzahl gültigen Werte von der Anzahl der Zeilen abziehen

```
wetter.shape[0] - wetter.count()
```

Diese beiden Beispiele zeigen `count` und `isna`, die auf Dataframes angewendet werden. Wenn wir nur die Anzahl der fehlenden Werte in einer einzelnen Variablen zählen wollen, können wir einfach diese Variable auswählen. Zum Beispiel auf die Spalte `QN_3`:

```
wetter.QN_3.count()
```

```
np.int64(4656)
```

Ein guter Weg, um fehlende Beobachtungen zu entfernen, ist die Verwendung der Methode `dropna`. Standardmäßig werden *alle Zeilen entfernt, die einen fehlenden Wert enthalten*.

```
t = wetter.dropna()
t.shape
```

```
(0, 19)
```

Wir haben keine Beobachtungen mehr! Es gibt keine einzige Zeile, die gültige Werte in allen Zellen hat. Während dieser Fehler hier leicht zu erkennen ist—da nachfolgende Analysen ohne

Werte scheitern werden—können solche Probleme insbesondere bei großen Datensätzen und in automatisierten Datenanalyseverfahren zu Fehlern und falschen Ergebnissen führen, die nicht augenscheinlich falsch sind. Stellen Sie sich den Fall vor, dass das Entfernen von fehlenden Werten unseren Datensatz um 90 % verkleinert, weil eine irrelevante Variable größtenteils fehlt. Es sei denn, wir überprüfen die resultierende Anzahl von Zeilen, könnten wir unsere Analyse auf einem völlig falschen Subset abschließen.

Wie die obigen Tabellen nahelegen, ist der Hauptverursacher die Spalte *FX* die keine einzigen gültigen Wert enthält. In solchen Fällen ist es sinnvoll solche Spalten komplett zu entfernen. Dies können wir auch mit der Methode `dropna` machen, indem wir sie nur auf Spalten anwenden (`axis=1`) und diese entfernen, wenn alle Werte darin fehlen (`how="all"`):

```
t = wetter.dropna(axis=1, how="all")
t.shape
```

```
(7209, 18)
```

Wenn unsere Analyse nur bestimmte Variablen erfordert, sollten wir nicht alle solchen Zeilen entfernen, sondern nur diejenigen Zeilen entfernen, in denen unsere wichtigen Variablen fehlen. Wollen wir z.B. nur die Temperatur (*TMK*) und Luftfeuchtigkeit (*UPM*) analysieren, können wir die fehlenden Zeilen nur in diesen beiden Variablen entfernen, indem wir das Argument `subset` verwenden:

```
w = t.dropna(subset=["TMK", "UPM"])
w.shape
```

```
(7209, 18)
```

```
w.head()
```

| STATION | DATE | TIME     | QON_3 | FMQON_4 | RSKRSKF | SHK_TAG | NMVPM | PMTMKUPM | TXK | TNK | TGK | eor        |      |      |       |       |     |
|---------|------|----------|-------|---------|---------|---------|-------|----------|-----|-----|-----|------------|------|------|-------|-------|-----|
| 0       | 4271 | 19470101 | NaN   | NaN     | 5       | NaN     | NaN   | NaN      | NaN | NaN | 4.1 | 1019.32.1  | 78.0 | -1.5 | -3.3  | -5.8  | eor |
| 1       | 4271 | 19470101 | NaN   | NaN     | 3       | NaN     | NaN   | NaN      | NaN | NaN | 3.6 | 1032.44.8  | 80.0 | -2.4 | -6.2  | -4.5  | eor |
| 2       | 4271 | 19470101 | NaN   | NaN     | 5       | NaN     | NaN   | NaN      | NaN | NaN | 1.9 | 1031.912.9 | 79.0 | -9.4 | -14.0 | -16.0 | eor |
| 3       | 4271 | 19470101 | NaN   | NaN     | 5       | NaN     | NaN   | NaN      | NaN | NaN | 1.9 | 1023.412.9 | 82.0 | -9.8 | -15.3 | -19.0 | eor |
| 4       | 4271 | 19470101 | NaN   | NaN     | 5       | NaN     | NaN   | NaN      | NaN | NaN | 2.4 | 1017.59.9  | 84.0 | -8.7 | -13.7 | -19.0 | eor |

Wenn unsere Analyse nur bestimmte Variablen erfordert, sollten wir nicht alle solchen Zeilen entfernen, sondern nur diejenigen Zeilen entfernen, in denen unsere wichtigen Variablen fehlen. Stellen Sie sich vor, wir möchten hier nur die Variablen *TMK* und *UPM* analysieren. Wir können

die fehlenden Zeilen nur in diesen beiden Variablen entfernen, indem wir das Argument `subset` verwenden:

## Falsche Beobachtungen

Bisher haben wir diskutiert, wie wir mit fehlenden Werten umgehen. Es gibt darüber hinaus weitere Gründe für fehlerhafte und ungewöhnliche Daten, wie Messfehler, Eingabefehler, Einheitenfehler, Skalierungsfehler, Berechnungsfehler, etc. Diese Werte können Ergebnisse sehr stark verzerren. Deshalb ist es wichtig sie zu identifizieren und zu behandeln. Man unterscheidet die folgenden Arten:

- **Fehlende Werte:** Sind Werte bei denen bekannt ist, dass sie fehlen.
- **Außerhalb des Wertebereichs** (Out-of-range): Sind Werte die logisch ausschließbar sind. Zum Beispiel weil sie außerhalb des definierten Wertebereichs liegen (z.B. Temperatur  $< -273.15^{\circ}\text{C}$ ), außerhalb des Messbereichs liegen (z.B. Raumtemperatursensor  $> 60^{\circ}\text{C}$ ), oder offensichtlich falsch sind (z.B. 200 Jahre alter Mensch).
- **Ausreißer:** Sind Werte, die statistisch deutlich von den restlichen Daten abweichen. Sie können entweder viel größer oder viel kleiner sein als die meisten anderen Werte. Sie sind insbesondere dann ein Problem, wenn sie häufiger auftreten als statistisch erklärbar Werte.
- **Anomalien:** Sind Werte, die Kausal erklärbar sind, aber vom erwarteten Verhalten abweichen, z.B. Wetteranomalien.

Diese fehlerhaften oder ungewöhnlichen Werte werden meist unterschiedlich behandelt:

- Werte außerhalb des Wertebereichs werden meist als fehlende Werte codiert und genauso behandelt, weil sie ja nicht erklärbar auftreten können.
- Die Behandlung von Ausreißer ist schwieriger, da sie statistisch vorkommen können und somit erklärbar wären. Man behandelt sie meist nur dann, wenn sie statistisch ungewöhnlich häufig auftreten (ein Indiz für Fehler) oder aber z.B. beim Berechnen von stabilen ML-Modellen stören, da sie sehr starke Varianzen erzeugen.
- Anomalien sind ungewöhnlich, aber erklärbar und häufig ist es ein Ziel von ML-Modellen diese zu identifizieren (Anomalieerkennung). Deshalb entfernt man sie nur selten, wenn man anomaliefreie Datensätze braucht, um ML-Modelle zu trainieren die das Normalverhalten vorhersagen (z.B. zur Anomalieerkennung).

Sie können die Ergebnisse statistischer Analysen verzerren und die Modellierung erschweren. Daher ist es wichtig, Ausreißer zu identifizieren und den richtigen Umgang mit ihnen zu wählen.

Der erste Ansatz ist die Wertebereiche der Daten zu prüfen. Hier kann man deskriptive Methoden benutzen, um eventuelle verdächtige Werte zu erkennen. Bei numerischen Werten ist der offensichtlichste Ausgangspunkt, den Wertebereich zu überprüfen. Zum Beispiel - was ist der kleinste und größte Wert in unserem Wetter Datensatz:

Prüfen wir zum Beispiel die Mittlere Tagestemperatur (TMK):

```
wetter.TMK.min(), wetter.TMK.max()
```

```
(np.float64(-14.8), np.float64(27.8))
```

Die Werte liegen durchaus im realistischen Bereich und erfordern keine Änderung notwendig ist. Schauen wir auf die relative Luftfeuchtigkeit (UPM):

```
wetter.UPM.min(), wetter.UPM.max()
```

```
(np.float64(42.0), np.float64(1000.0))
```

Dann stellen wir fest, dass diese zwischen 42% und 1000% liegt. Letzteres ist nicht möglich und könnte ein Tippfehler sein. Bei solchen fehlerhaften Werten wissen wir allerdings nicht, was der originale Wert war, deshalb ist es sinnvoll ihn mit `pd.NA` ersetzen. Hierfür können wir die Funktion `mask` nutzen:

```
wetter["UPM"] = wetter.UPM.mask(cond=wetter.UPM > 100, other=pd.NA)
```

Wenn wir jetzt die Min und Max bestimmen, sehen wir, dass die Werte sich jetzt im Definitionsbereich liegen:

```
wetter.UPM.min(), wetter.UPM.max()
```

```
(np.float64(42.0), np.float64(100.0))
```

Wie wir mit Ausreißern und Anomalien umgehen, werden wir später behandeln.

### Fehlende Werte Ersetzen

Manchmal möchten man fehlenden Werte durch neuen Wert ersetzen. Hier unterscheidet man unterschiedliche Fälle:

- *Fehlende Werte mit sinnvollen Standardwerten:* Es gibt Datensätze, wo man durchaus sinnvolle Standardwerte annehmen kann für fehlende Werte. Zum Beispiel, zählen wir die Unfälle auf Baustellen pro Tag. Da diese nur registriert werden, wenn welche auftreten, ist für die anderen Tage anzunehmen (aber nicht garantiert), dass keine aufgetreten sind.

Man könnte zum Beispiel annehmen, dass die Schneetiefe in unserem Wetterdatensatz 0 ist (Das wäre allerdings falsch), wenn sie nicht angegeben ist. In diesem Fall können wir die fehlenden Werte `pd.NA` mit `fillna` ersetzen:

```
wetter.SHK_TAG.fillna(value=0)
```

```
0      0.0
1      0.0
2      0.0
```

```

3      0.0
4      0.0
...
7204    0.0
7205    0.0
7206    0.0
7207    0.0
7208    0.0
Name: SHK_TAG, Length: 7209, dtype: float64

```

### Wenige Fehlende Werte Ersetzen

Bei Zeitreihen, wie den Wetterdaten, kann man fehlende Werte zwischendurch ersetzen, wenn man davon ausgehen kann, dass sie sich nicht schnell ändern. Dazu nutzt man Methoden der Interpolation. Pandas unterstützt

- `ffill` - Ersetzt NA/NaN-Werte, indem Sie die letzte gültige Beobachtung verwenden, um die Lücke zu füllen.
- `bfill` - Ersetzt NA/NaN-Werte, indem Sie die nächste gültige Beobachtung nutzt, um die Lücke zu füllen.
- `interpolate` - Ersetzt NA/NaN-Werte mit einer Interpolationsmethode, die mit dem Parameter `method` angegeben wird:
  - `linear` - Interpoliert linear zwischen dem letzten und dem nächsten Wert.
  - `nearest` - Wählt entweder den letzten und dem nächsten Wert, je nachdem welcher zeitlich näher liegt.
  - `polynomial` - Interpoliert zwischen dem letzten und dem nächsten Wert mit einem Polynom der angegebenen Ordnung (z.B. `order=5`).
  - `cubic` - Interpoliert zwischen dem letzten und dem nächsten Wert mit einem cubischen Polynom (das Überspringen vermeidet).

Bei allen diesen Methoden ist empfohlen den Parameter `limit` zu nutzen. Er bestimmt wie viele nacheinander folgende Werte ersetzt werden dürfen (andere bleiben NA). So mag es akzeptabel sein, ein einzelner Tag zu interpolieren, allerdings nicht eine ganze Woche, da der anzunehmende Interpolationsfehler zu groß wird.

Schauen wir uns die mittlere Temperatur an:

```
wetter.TMK[:10]
```

```

0    -2.1
1    -4.8
2   -12.9
3   -12.9
4    -9.9
5    -9.1
6   -10.8

```

```
7    -9.3
8    -1.1
9     2.6
Name: TMK, dtype: float64
```

Hier fehlen mehrere Werte. Wir können diese interpolieren mit:

```
wetter["TMK"]=wetter.TMK.interpolate(method="cubic", limit=1)
wetter.TMK[:10]
```

```
0    -2.1
1    -4.8
2   -12.9
3   -12.9
4    -9.9
5    -9.1
6   -10.8
7    -9.3
8    -1.1
9     2.6
Name: TMK, dtype: float64
```

### Fehlende Werte Am Anfang und Ende

Fehlende Werte am Anfang oder Ende einer Zeitreihe kann diese nicht interpolieren, da Information über Vorgängerwerte oder Nachfolgerwerte fehlen. In solchen Fällen spricht man von Extrapolation. Prinzipiell kann man hier die Funktionen `ffill` und `bfill` verwenden. Da halt allerdings Randwertinformationen fehlen ist der anzunehmende Extrapolationsfehler meist groß, weshalb unbedingt `limit` zu verwenden ist.

```
wetter["RSK"]=wetter.RSK.ffill(limit=1)
```

### Fehlende Datumswerte hinzufügen

Gerade bei der Analyse von Zeitreihen ist es oft wichtig dass Messungen in festem Abstand (Equidistant) aufgenommen werden. Um das zu prüfen müssen wir als erstes die Datumsangabe in unserm Wetterdatensatz in einen nutzbaren Datentyp für Zeiten umwandeln. Dies können wir mit der Funktion `pd.to_datetime` umwandeln. Dafür müssen wir die originale Datumsspalte `MESS_DATUM` in einen String umwandeln, der von `to_datetime` verwendet werden kann und ein String der das format repräsentiert (siehe Umgang mit Zeitstempeln):

```
wetter["DATUM_DT"] = pd.to_datetime(wetter.MESS_DATUM.astype(str),
format="%Y%m%d", utc=True)
```

Als nächstes empfiehlt es sich den Datensatz nach der Zeit zu sortieren:

```
wetter.sort_values(by="DATUM_DT", inplace=True) # Sortier nach der Spalte
"MESS_DATUM" und ersetzt den Dataframe direkt (inplace)
```

Dann können wir prüfen, mit welcher Häufigkeit, welche zeitlichen Differenzen auftreten. Wir können dafür die `diff` Funktion nutzen, um die Zeitdifferenz zu bestimmen. Diese verketteten wir mit `value_counts`, um die Häufigkeit der Differenzen zu berechnen.

```
wetter.DATUM_DT.diff().value_counts()
```

```
DATUM_DT
1 days    7205
2 days      2
3 days      1
Name: count, dtype: int64
```

Wir sehen, dass in unserem Datensatz mehrere Messwerte fehlen, da Differenzen von mehr als 1 Tag vorkommen.

Um diese Daten wiederherzustellen können wir die Funktion `resample` benutzen. Dafür müssen wir zuerst die Datums-Spalte zum Index des Dataframes machen, so dass eine richtige mehrdimensionale Zeitreihe daraus wird.

```
wetter.set_index("DATUM_DT", inplace=True)
```

Jetzt können wir mit `resample` den Dataframe in eine regelmäßige tägliche Auflösung bringen. Dabei müssen wir eine Funktion zur Interpolation der fehlenden Werte mit angeben.

```
wetter1D=wetter.resample('1d').ffill(limit=1)
```

```
/var/folders/0c/805v0049471f04d_w7xm50rc0000gn/T/
ipykernel_86485/727961184.py:1: Pandas4Warning: 'd' is deprecated and will be
removed in a future version, please use 'D' instead.
  wetter1D=wetter.resample('1d').ffill(limit=1)
```

Nun können wir den Index wieder zu einer Spalte machen

```
wetter1D.reset_index(inplace=True)
```

Der neue Dataframe hat jetzt einen lückenlosen täglichen Zeitindex. Es können aber weiterhin fehlende Werte in einzelnen Variablen vorhanden sein, wenn diese durch `ffill(limit=1)` nicht ersetzt wurden.

```
wetter1D.DATUM_DT.diff().value_counts()
```

```
DATUM_DT
1 days    7212
Name: count, dtype: int64
```

## Fehlende Werte und mathematische Operationen

Die Methoden `.sum` und `.mean` von Pandas ignorieren fehlende Werte. Berechnen wir den Mittelwert der täglich Sonnenscheindauer (SDK):

```
wetter1D.SDK.mean()
```

```
np.float64(4.904363699582754)
```

Sie liegt bei 4.9 Stunden.

Wenn wir jedoch den Mittelwert entsprechend der Formel für den arithmetischen Mittelwert

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i.$$

selbst berechnen

```
wetter1D.SDK.sum()/len(wetter1D.SDK)
```

```
np.float64(3.9109801746845974)
```

Dann ergibt sich ein komplett anderes Ergebnis. Dies liegt daran, dass `len` fehlende Werte nicht ignoriert und daher die Summe durch die Gesamtzahl der Fälle teilen, inklusive der fehlenden Fälle.

In dem Fall wäre die Berechnung mit `count` richtig, da dann NA Werte ignoriert werden.

```
wetter1D.SDK.sum()/wetter1D.SDK.count()
```

```
np.float64(4.904363699582754)
```

Wenn wir jedoch einen fehlenden Wert mit einer Zahl vergleichen, ergibt sich `False` und nicht ein fehlender Wert `pd.NA`, was formal richtig wäre, da für die fehlenden Werte ja keine Aussage getroffen werden kann. Das kann sehr schnell zu falschen Ergebnissen führen. Wollen wir zum Beispiel die Anzahl der Sonnentage mit mehr als 8 Stunden Sonnenschein bestimmen, so ist es intuitiv zu berechnen:

```
sonnentage = wetter1D.SDK > 8
sonnentage.value_counts(dropna=False)
```

```
SDK
True      4180
False     3033
Name: count, dtype: int64
```

Das ist allerdings Falsch. Wir haben die Informationen über die fehlenden Werte im Prozess verloren und Tage mit fehlenden Werten automatisch zu keinem Sonnentag gemacht. Die richtige Berechnung kann man mit der Funktion `map` bestimmen, da mit dem Parameter `na_action='ignore'` enthaltene NA Werte belassen werden:

```
sonnentage = wetter1D.SDK.map(lambda x: x < 8, na_action='ignore')
sonnentage.value_counts(dropna=False)
```

```
SDK
True      4180
False     1572
NaN       1461
Name: count, dtype: int64
```

Dies unterstreicht die Bedeutung der Behandlung von fehlenden Werten in Daten und das man sich mit dem Umgang der Standardfunktionalität von Pandas mit NA vertraut machen sollte.

## Konvertierung von Variablen

### Konvertierung von kategorischen Variablen

Manchmal stoßen wir auf Situationen, in denen die in den Daten bereitgestellten Kategorien nicht gut zu den Anforderungen unserer Analyse passen. Zwei häufige Fälle sind, wenn die Originaldaten zu granularen Informationen enthalten (zum Beispiel, wenn wir an ein paar breiten Wirtschaftssektoren interessiert sind, aber die Daten 27 Branchen auflisten) oder Kategorien zu ausführlich sind (zum Beispiel kann der Beruf als „Handwerker, Vorarbeiter und ähnliche“ aufgeführt sein, während wir lieber den kürzeren Begriff „qualifiziert“ verwenden würden). In diesen Fällen müssen wir die ursprünglichen Kategorien in neue umwandeln und möglicherweise mehrere ursprüngliche Kategorien in eine einzige neue Kategorie zusammenführen.

In diesem Fall gibt es mehrere Möglichkeiten, wie man vorgehen kann. Wir zeigen die Optionen anhand des Falls, in dem die Daten numerisch codiert sind, aber nicht alle Codes echte numerische Werte darstellen.

### Verwenden von `pd.cut`

`pd.cut` ist eine dedizierte Funktion von Pandas, um kontinuierliche numerische Daten in vordefinierte Intervalle zu schneiden. Es wird wie folgt aufgerufen:

```
pd.cut(x, bins, right=True, labels=None)
```

wo  $x$  die Variable ist, die geschnitten werden soll, bins sind die Intervallgrenzen, labels sind die Intervallnamen und right gibt an, ob der rechte Grenzwert im Intervall enthalten ist (right = True bedeutet, dass die Grenzen rechts offen sind). Es ist hauptsächlich für das Aufteilen von kontinuierlichen Werten gedacht, kann aber auch leicht für diskrete Werte verwendet werden. Lassen Sie uns diese Daten in gewünschte Intervalle aufteilen:

```
wetter1D["TemperaturKlasse"] = pd.cut(wetter1D.TMK,
    bins = [-np.inf, 13, 18, 24, 30, np.inf],
    labels = ["Cold", "Cool", "Mild", "Warm", "Hot"],
    right=False)
```

```
wetter1D["TemperaturKlasse"].value_counts(dropna=False)
```

```
TemperaturKlasse
Cold      4861
Cool      1836
Mild       503
Warm        12
NaN         1
Hot         0
Name: count, dtype: int64
```

Wir setzen die Grenzen bei 13, 18, 24 und 30 °C und halten die Ränder offen, indem wir die erste und letzte Grenze auf Unendlich setzen. Wir geben den Intervallen explizite Labels, aufgrund der begrenzten Funktionalität von np.nan (es gibt keine fehlenden Werte vom Typ String), fügen wir explizit den fehlenden Wert „NA“ in den String ein. Schließlich sagen wir cut, dass der rechte Grenzwert aus dem Intervall ausgeschlossen werden soll. Andernfalls würde der Wert 0 in die Kategorie „NA“, der Wert 1 in „0“ usw. fallen.

### Verwendung von np.where

Eine weitere Option ist die Verwendung von np.where, einer vektorisierten Version einer if-else-Anweisung. np.where nimmt drei Argumente entgegen: eine vektorisierte logische Bedingung, den Wert, wenn die Bedingung wahr ist, und den Wert, wenn sie falsch ist. Wir können die Funktion verschachteln, um alternative Bedingungen zu prüfen. Zum Beispiel wollen wir die Heizgradtage bestimmen (Mittlere Temperatur < 15°C) und die Kühlgradtage (Mittlere Temperatur > 22°C).

```
wetter1D["HeizKuehlTage"] = np.where(wetter1D.TMK > 22, "Kühlgradtag",
    np.where(wetter1D.TMK < 15, "Heizgradtag", "Normaltag"))
wetter1D["HeizKuehlTage"].value_counts(dropna=False)
```

```
HeizKuehlTage
Heizgradtag    5611
```

```
Normaltag      1559
Kühlgradtag    43
Name: count, dtype: int64
```

Es ist zu beachten, dass die logische Bedingung in `np.where` wieder NA Werte auf False mapt. Deshalb sollte man das wieder gezielt behandeln:

```
wetter1D["HeizKuehlTage"] = np.where(wetter1D.TMK.isna(), pd.NA,
np.where(wetter1D.TMK > 22, "Kühlgradtag", np.where(wetter1D.TMK < 15,
"Heizgradtag", "Normaltag"))))
wetter1D["HeizKuehlTage"].value_counts(dropna=False)
```

```
HeizKuehlTage
Heizgradtag    5611
Normaltag      1558
Kühlgradtag    43
NaN            1
Name: count, dtype: int64
```

Mit der `map` Funktion lässt sich das auch realisieren:

```
wetter1D["HeizKuehlTage"] = wetter1D.TMK.map(lambda x: "Heizgradtag" if x < 15
else "Kühlgradtag" if x > 22 else "Normaltag", na_action='ignore')
```

### Ersetzen ausgewählter Werte im Dataframe

Häufig ist es notwendig bestimmte Werte gezielt zu ersetzen, auf Basis einer logischen Bedingung zu modifizieren. Dies kann man indem man den `.loc`-Index eines Dataframes nutzt. Das kann man auch nutzen, um neue Spalten zu berechnen. Wir erstellen eine neue leere Spalte und passen sie anschließend nach unseren Bedürfnissen an. Beachten Sie, dass dies einen gemischten Indexierungsansatz erfordert. Wir werden die Zeilen anhand eines logischen Vektors (logische Bedingung) und die Spalten anhand ihres Namens indizieren. Die Spalte `QN_3` im Datensatz zum Beispiel codiert die Datenqualität, mit den folgenden kategorischen Werten.

```
wetter1D["QNS_4"] = pd.NA
wetter1D.loc[wetter1D.QN_4 == 1, "QNS_4"] = "nur formale Prüfung"
wetter1D.loc[wetter1D.QN_4 == 2, "QNS_4"] = "nach individuellen Kriterien
geprüft"
wetter1D.loc[wetter1D.QN_4 == 3, "QNS_4"] = "automatische Prüfung und
Korrektur"
wetter1D.loc[wetter1D.QN_4 == 5, "QNS_4"] = "historische, subjektive
Verfahren"
wetter1D.loc[wetter1D.QN_4 == 7, "QNS_4"] = "geprüft, gepflegt, nicht
korrigiert"
wetter1D.loc[wetter1D.QN_4 == 8, "QNS_4"] = "Qualitätsicherung ausserhalb
```

```
ROUTINE"
wetter1D.loc[wetter1D.QN_4 == 9, "QNS_4"] = "nicht alle Parameter korrigiert"
wetter1D.loc[wetter1D.QN_4 == 10, "QNS_4"] = "Qualitätsprüfung und Korrektur
beendet"

wetter1D["QNS_4"].value_counts(dropna=False)
```

```
QNS_4
historische, subjektive Verfahren      7202
automatische Prüfung und Korrektur      3
nur formale Prüfung                    2
nach individuellen Kriterien geprüft    2
Qualitätsicherung ausserhalb ROUTINE    2
<NA>                                    1
nicht alle Parameter korrigiert         1
Name: count, dtype: int64
```

Für solche komplexeren Mappings kann man auch die map Funktion verwenden und in dieser eine speziell definierte Mappingfunktion.

```
def QN_Mapping(x):
    if x == 1: return "nur formale Prüfung"
    if x == 2: return "nach individuellen Kriterien geprüft"
    if x == 3: return "automatische Prüfung und Korrektur"
    if x == 5: return "historische, subjektive Verfahren"
    if x == 7: return "geprüft, gepflegt, nicht korrigiert"
    if x == 8: return "Qualitätsicherung ausserhalb ROUTINE"
    if x == 9: return "nicht alle Parameter korrigiert"
    if x == 10: return "Qualitätsprüfung und Korrektur beendet"
    return pd.NA

wetter1D["QNS_4"] = wetter1D.QN_4.map(QN_Mapping, na_action='ignore')
```

## Konvertierung kategorischer Variablen in Dummy-Variablen

Im vorherigen Abschnitt haben wir uns damit beschäftigt, Variablen in Kategorien umzuwandeln. In vielen Fällen müssen Kategorien für Analysen oder Modelle in eine maschinenlesbare Darstellung überführt werden.

Ein üblicher Ansatz für nominale Kategorien sind sogenannte *Dummies* oder One-Hot-Variablen. Diese kann man mit `pd.get_dummies` erzeugen. Dabei wird pro Kategorie eine eigene Spalte erstellt. Das ist oft die sicherste Darstellung für ML-Modelle, weil keine künstliche Reihenfolge zwischen den Kategorien eingeführt wird. Beachten Sie auch, dass der originale Index erhalten bleibt. Man kann bei Bedarf einen Namenspräfix für die Spalten angeben:

```
QN_Dummies = pd.get_dummies(wetter1D.QNS_4, prefix="QN")
QN_Dummies
```

|      | QN_Qualitäts-<br>sicherung<br>ausserhalb<br>ROUTINE | QN_automatische<br>Prüfung<br>und Korrektur | QN_historische,<br>subjektive<br>Verfahren | QN_nach<br>individuel-<br>len Kriteri-<br>en geprüft | QN_nicht<br>alle Para-<br>meter korri-<br>giert | QN_nur<br>formale<br>Prüfung |
|------|-----------------------------------------------------|---------------------------------------------|--------------------------------------------|------------------------------------------------------|-------------------------------------------------|------------------------------|
| 0    | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 1    | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 2    | False                                               | True                                        | False                                      | False                                                | False                                           | False                        |
| 3    | False                                               | True                                        | False                                      | False                                                | False                                           | False                        |
| 4    | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| ...  | ...                                                 | ...                                         | ...                                        | ...                                                  | ...                                             | ...                          |
| 7208 | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 7209 | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 7210 | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 7211 | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |
| 7212 | False                                               | False                                       | True                                       | False                                                | False                                           | False                        |

Für bestimmte technische Zwecke gibt es auch die Funktion `pd.factorize`. Sie ordnet jeder Kategorie eine ganze Zahl zu und gibt zusätzlich das zugehörige Mapping zurück. Diese Darstellung ist praktisch, wenn man Kategorien kompakt speichern oder später wieder zurückübersetzen möchte. Für nominale Kategorien sollte man sie jedoch nicht unkritisch als metrische Modellvariable interpretieren, weil die Zahlen eine künstliche Reihenfolge erzeugen können.

```
QN_Factors_Values, QN_Factors_Mapping = pd.factorize(wetter1D.QNS_4)
QN_Factors_Values, QN_Factors_Mapping
```

```
(array([0, 0, 1, ..., 0, 0, 0], shape=(7213,)),
Index(['historische, subjektive Verfahren',
      'automatische Prüfung und Korrektur', 'nur formale Prüfung',
      'nach individuellen Kriterien geprüft',
      'Qualitätsicherung ausserhalb ROUTINE',
      'nicht alle Parameter korrigiert'],
      dtype='str'))
```

Zu beachten ist, dass die Funktion ein Tupel mit den numerischen Werten und dem Mapping zurück gibt. Wir können die Werte als neue Spalte dem Dataframe hinzufügen, sollten aber im





| DATUM | SS       | DT       | UN_3   | FX  | FM  | UN_4 | RS  | SK  | FSDK | ..  | ON      | Heiz | ON           | hist         | ON      | ON      | ON     | nur            |                     |     |     |
|-------|----------|----------|--------|-----|-----|------|-----|-----|------|-----|---------|------|--------------|--------------|---------|---------|--------|----------------|---------------------|-----|-----|
|       |          |          |        |     |     |      |     |     |      |     | Kuehl-  |      | aus-Prü-sub- | in-          | al-     | for-    |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     | ra- Ta- |      | serfung jek- | di-          | le      | ma-     |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     | tur- ge |      | halb und ti- | vi-          | Pa-     | le      |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     | Klas-   |      | ROU-Kor-     | ve           | du-     | ra-Prü- |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     | se      |      | TI-rek-Ver-  | el-          | mefung  |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      | NE tur fah-  | len          | ter     |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              | ren Kri-kor- |         |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              | te- ri- |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              | rigiert |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              | en      |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              | ge-     |         |        |                |                     |     |     |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              | prüft   |         |        |                |                     |     |     |
| 3     | 1947     | 407      | 109    | 47  | Na  | Na   | Na  | Na  | 3.0  | Na  | Na      | Na   | Na           | ...          | Cold    | Heiz    | au- 1  | Fal- True      | Fal- Fal- Fal- Fal- |     |     |
|       | 00:00:00 | 00:00:00 | +00:00 |     |     |      |     |     |      |     |         |      |              |              |         | g- to-  | se     | se             | se                  | se  | se  |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              |         | rad-ma- |        |                |                     |     |     |
| 4     | 1947     | 407      | 109    | 47  | Na  | Na   | Na  | Na  | 5.0  | Na  | Na      | Na   | Na           | ...          | Cold    | Heiz    | his- 0 | Fal- Fal- True | Fal- Fal- Fal-      |     |     |
|       | 00:00:00 | 00:00:00 | +00:00 |     |     |      |     |     |      |     |         |      |              |              |         | g- to-  | se     | se             | se                  | se  | se  |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              |         | rad-ri- |        |                |                     |     |     |
| 7208  | 1966     | 407      | 120    | 66  | 59  | 02   | 60  | Na  | 6.6  | 5.0 | 0.3     | 1.0  | 2.2          | ...          | Cool    | Heiz    | his- 0 | Fal- Fal- True | Fal- Fal- Fal-      |     |     |
|       | 00:00:00 | 00:00:00 | +00:00 |     |     |      |     |     |      |     |         |      |              |              |         | g- to-  | se     | se             | se                  | se  | se  |
|       |          |          |        |     |     |      |     |     |      |     |         |      |              |              |         | rad-ri- |        |                |                     |     |     |
| ...   | ...      | ...      | ...    | ... | ... | ...  | ... | ... | ...  | ... | ...     | ...  | ...          | ...          | ...     | ...     | ...    | ...            | ...                 | ... | ... |

|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              |         |         |      |    |     |
|-------|----|----|------|----|----|------|----|----|----|----|----|----|---------|----|--------------|--------------|---------|---------|------|----|-----|
| DATUM | SS | DT | UN_3 | FX | FM | UN_4 | RS | SK | FS | DK | .. | ON | Heiz    | ON | Heiz         | ON           | Heiz    | ON      | Heiz | ON | nur |
|       |    |    |      |    |    |      |    |    |    |    |    |    | Kuehl-  |    | aus-Prü-sub- | in-          | al-     | for-    |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    | ra- Ta- |    | serfung jek- | di-          | le      | ma-     |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    | tur- ge |    | halb und ti- | vi-          | Pa-     | le      |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    | Klas-   |    | ROU-Kor-     | ve           | du-     | ra-Prü- |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    | se      |    | TI-rek-Ver-  | el-          | mefung  |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    | NE tur fah-  | len          | ter     |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              | ren Kri-kor- |         |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              | te- ri- |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              | rigiert |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              | en      |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              | ge-     |         |      |    |     |
|       |    |    |      |    |    |      |    |    |    |    |    |    |         |    |              |              | prüft   |         |      |    |     |

---

|  |  |  |  |  |  |  |  |  |  |  |  |  |      |  |  |  |  |  |  |  |  |
|--|--|--|--|--|--|--|--|--|--|--|--|--|------|--|--|--|--|--|--|--|--|
|  |  |  |  |  |  |  |  |  |  |  |  |  | sub- |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | jek- |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | ti-  |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | ve   |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | Ver- |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | fah- |  |  |  |  |  |  |  |  |
|  |  |  |  |  |  |  |  |  |  |  |  |  | ren  |  |  |  |  |  |  |  |  |

|            |           |    |      |     |     |     |     |     |      |      |        |      |      |      |      |      |      |
|------------|-----------|----|------|-----|-----|-----|-----|-----|------|------|--------|------|------|------|------|------|------|
| 1966407120 | 196659270 | Na | 10.0 | 5.0 | 0.0 | 1.0 | 6.1 | ... | Cold | Heiz | his- 0 | Fal- | Fal- | True | Fal- | Fal- | Fal- |
| 00:00:00   | +00:00    |    |      |     |     |     |     |     |      |      |        | g-   | to-  |      | se   | se   | se   |

|      |  |  |  |  |  |  |  |  |  |  |  |         |       |  |  |  |  |
|------|--|--|--|--|--|--|--|--|--|--|--|---------|-------|--|--|--|--|
|      |  |  |  |  |  |  |  |  |  |  |  | rad-ri- |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | tag     | sche, |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  |         | sub-  |  |  |  |  |
| 7209 |  |  |  |  |  |  |  |  |  |  |  | jek-    |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | ti-     |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | ve      |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | Ver-    |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | fah-    |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | ren     |       |  |  |  |  |

|            |           |    |     |     |     |     |     |     |      |      |        |      |      |      |      |      |      |
|------------|-----------|----|-----|-----|-----|-----|-----|-----|------|------|--------|------|------|------|------|------|------|
| 1966407120 | 196659280 | Na | 4.5 | 5.0 | 0.0 | 0.0 | 9.9 | ... | Cold | Heiz | his- 0 | Fal- | Fal- | True | Fal- | Fal- | Fal- |
| 00:00:00   | +00:00    |    |     |     |     |     |     |     |      |      |        | g-   | to-  |      | se   | se   | se   |

|      |  |  |  |  |  |  |  |  |  |  |  |         |       |  |  |  |  |
|------|--|--|--|--|--|--|--|--|--|--|--|---------|-------|--|--|--|--|
|      |  |  |  |  |  |  |  |  |  |  |  | rad-ri- |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | tag     | sche, |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  |         | sub-  |  |  |  |  |
| 7210 |  |  |  |  |  |  |  |  |  |  |  | jek-    |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | ti-     |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | ve      |       |  |  |  |  |
|      |  |  |  |  |  |  |  |  |  |  |  | Ver-    |       |  |  |  |  |





DAIUNEDSDATUN\_3 FX FQN\_4RSRKFSDK ..TMKUPMTXKTNKTGK eorTemHqNSQNF\_4

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

1947407-10947NaNNaNNaN5.0 NaNNaNNaN... -2.178.00-1.5-3.3-5.8e or ColdHeizhis- 0.0  
00:00:00+00:00 g- to-

0

1947407+004711NHNANNa5.0 NaNNaNNaN... -2.178.00-1.5-3.3-5.8e or ColdHeizhis- 0.0  
00:00:00+00:00 g- to-

1

1947407409474030NaNNaNNa3.0 NaNNaNNa... -4.880.00-2.4-6.2-4.5e or ColdHeizau- 1.0  
00:00:00+00:00 g- to-

2



Obwohl beide Dataframes eine unterschiedliche Anzahl an Spalten haben (Es fehlen bei `wetter_ab_1960` die Spalten, die wir hinzugefügt haben), kann `concat` beide Dataframes kombinieren. Hierbei werden fehlende Spalten mit `np.NaN` aufgefüllt. Deshalb ist es sinnvoll erst zusammenzufügen, bevor man abgeleitete Werte berechnet. Dies können wir korrigieren, indem wir die fehlenden Werte neu berechnen.

```
wetter_all["TemperaturKlasse"]=pd.cut(wetter_all.TMK, bins = [-np.inf, 13, 18,
24, 30, np.inf], labels = ["Cold", "Cool", "Mild", "Warm", "Hot"],
right=False)
wetter_all["HeizKuehlTage"] = wetter_all.TMK.map(lambda x: "Heizgradtag" if x
< 15 else "Kühlgradtag" if x > 22 else "Normaltag", na_action='ignore')
wetter_all["QNS_4"] = wetter_all.QN_4.map(QN_Mapping, na_action='ignore')
wetter_all["QNF_4"], _ = pd.factorize(wetter_all.QNS_4)
```

### Dataframe kombinieren

Meist will man nicht nur Daten kombinieren, die nicht gleich groß sind oder den gleichen Index haben. Zum Beispiel wollen wir die Wetterdaten mit den Energiedaten der Universität Rostock kombinieren, um diese später zu analysieren.

```
uros_egy=pd.read_csv("../data/UR05/EnergyID_kW.csv", parse_dates=["Date"])
uros_egy.head()
```

|   | Date                         | EV_HT_740 | EV_NT_740 | E_AV_Lab | E_SV_Lab | ES_Lab |
|---|------------------------------|-----------|-----------|----------|----------|--------|
| 0 | 2020-12-30<br>00:00:00+00:00 | NaN       | NaN       | NaN      | NaN      | NaN    |
| 1 | 2020-12-31<br>00:00:00+00:00 | NaN       | NaN       | 1256.0   | 291.0    | 5.0    |
| 2 | 2021-01-01<br>00:00:00+00:00 | 0.0       | 4080.0    | 1221.0   | 290.0    | 1.0    |
| 3 | 2021-01-02<br>00:00:00+00:00 | 1170.0    | 2630.0    | 1243.0   | 284.0    | 2.0    |
| 4 | 2021-01-03<br>00:00:00+00:00 | 0.0       | 3750.0    | 1222.0   | 283.0    | 2.0    |

Die Daten können wir mit der Methode `pd.merge()` kombinieren, die wie ein SQL Join funktioniert. Hierbei können wir die Spalten die gematcht werden sollen mit `left_on` und `right_on` spezifizieren (oder `on` wenn die Spalten gleich sind). Mit `how="left"` spezifizieren wir, dass nur die Zeilen genommen werden sollen, die im Energiedatensatz vorhanden sind.

```
uros_egy_weater=pd.merge(uros_egy, wetter_all, right_on="DATUM_DT",
left_on="Date", how="left")
uros_egy_weater.head()
```

ValueError: You are trying to merge on datetime64[us, UTC] and object columns for key 'Date'. If you wish to proceed you should use pd.concat

```
[31m-----[?]
[39m
[31mValueError[39m                                Traceback (most recent
call last)
[36mCell[39m[36m [39m[32mIn[49][39m[32m, line 1[39m
[32m----> [39m[32m1[39m uros_egy_weater=pd.merge(uros_egy, wetter_all,
right_on=[33m"DATUM_DT"[39m, left_on=[33m"Date"[39m,
how=[33m"left"[39m)
[32m      2[39m uros_egy_weater.head()

[36mFile [39m[32m~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/pandas/core/
reshape/merge.py:385[39m, in [36mmerge[39m[34m(left, right, how, on,
left_on, right_on, left_index, right_index, sort, suffixes, copy, indicator,
validate)[39m
[32m    371[39m     [38;5;28;01mreturn[39;00m _cross_merge(
[32m    372[39m         left_df,
[32m    373[39m         right_df,
[32m    (...) [39m[32m    382[39m         validate=validate,
[32m    383[39m     )
[32m    384[39m     [38;5;28;01melse[39;00m:
[32m--> [39m[32m385[39m         op =
[43m_MergeOperation[49m[43m([49m
[32m    386[39m     [43m         [49m[43mleft_df[49m[43m, [49m
[32m    387[39m     [43m         [49m[43mright_df[49m[43m, [49m
[32m    388[39m     [43m
[49m[43mhow[49m[43m=[49m[43mhow[49m[43m, [49m
[32m    389[39m     [43m
[49m[43mon[49m[43m=[49m[43mon[49m[43m, [49m
[32m    390[39m     [43m
[49m[43mleft_on[49m[43m=[49m[43mleft_on[49m[43m, [49m
[32m    391[39m     [43m
[49m[43mright_on[49m[43m=[49m[43mright_on[49m[43m, [49m
[32m    392[39m     [43m
[49m[43mleft_index[49m[43m=[49m[43mleft_index[49m[43m, [49m
[32m    393[39m     [43m
[49m[43mright_index[49m[43m=[49m[43mright_index[49m[43m, [49m
```

```

    394    395    396    397    398    399
    400    401    402    403    404    405
    406    407    408    409    410    411
    412    413    414    415    416    417
    418    419    420    421    422    423
    424    425    426    427    428    429
    430    431    432    433    434    435
    436    437    438    439    440    441
    442    443    444    445    446    447
    448    449    450    451    452    453
    454    455    456    457    458    459
    460    461    462    463    464    465
    466    467    468    469    470    471
    472    473    474    475    476    477
    478    479    480    481    482    483
    484    485    486    487    488    489
    490    491    492    493    494    495
    496    497    498    499    500    501
    502    503    504    505    506    507
    508    509    510    511    512    513
    514    515    516    517    518    519
    520    521    522    523    524    525
    526    527    528    529    530    531
    532    533    534    535    536    537
    538    539    540    541    542    543
    544    545    546    547    548    549
    550    551    552    553    554    555
    556    557    558    559    560    561
    562    563    564    565    566    567
    568    569    570    571    572    573
    574    575    576    577    578    579
    580    581    582    583    584    585
    586    587    588    589    590    591
    592    593    594    595    596    597
    598    599    600    601    602    603
    604    605    606    607    608    609
    610    611    612    613    614    615
    616    617    618    619    620    621
    622    623    624    625    626    627
    628    629    630    631    632    633
    634    635    636    637    638    639
    640    641    642    643    644    645
    646    647    648    649    650    651
    652    653    654    655    656    657
    658    659    660    661    662    663
    664    665    666    667    668    669
    670    671    672    673    674    675
    676    677    678    679    680    681
    682    683    684    685    686    687
    688    689    690    691    692    693
    694    695    696    697    698    699
    700    701    702    703    704    705
    706    707    708    709    710    711
    712    713    714    715    716    717
    718    719    720    721    722    723
    724    725    726    727    728    729
    730    731    732    733    734    735
    736    737    738    739    740    741
    742    743    744    745    746    747
    748    749    750    751    752    753
    754    755    756    757    758    759
    760    761    762    763    764    765
    766    767    768    769    770    771
    772    773    774    775    776    777
    778    779    780    781    782    783
    784    785    786    787    788    789
    790    791    792    793    794    795
    796    797    798    799    800    801
    802    803    804    805    806    807
    808    809    810    811    812    813
    814    815    816    817    818    819
    820    821    822    823    824    825
    826    827    828    829    830    831
    832    833    834    835    836    837
    838    839    840    841    842    843
    844    845    846    847    848    849
    850    851    852    853    854    855
    856    857    858    859    860    861
    862    863    864    865    866    867
    868    869    870    871    872    873
    874    875    876    877    878    879
    880    881    882    883    884    885
    886    887    888    889    890    891
    892    893    894    895    896    897
    898    899    900    901    902    903
    904    905    906    907    908    909
    910    911    912    913    914    915
    916    917    918    919    920    921
    922    923    924    925    926    927
    928    929    930    931    932    933
    934    935    936    937    938    939
    940    941    942    943    944    945
    946    947    948    949    950    951
    952    953    954    955    956    957
    958    959    960    961    962    963
    964    965    966    967    968    969
    970    971    972    973    974    975
    976    977    978    979    980    981
    982    983    984    985    986    987
    988    989    990    991    992    993
    994    995    996    997    998    999
    1000

```

```
needs_i8_conversion(rk.dtype):
[32m 1841[39m [38;5;28;01mraise[39;00m
[38;5;167;01mValueError[39;00m(msg)

[31mValueError[39m: You are trying to merge on datetime64[us, UTC] and
object columns for key 'Date'. If you wish to proceed you should use pd.concat
```

Hierbei kann es bei Spalten vom Typ Datetime schnell zu Problemen kommen. Was daran liegt dass es bei dem Datentypen schnell zu Inkompatibilitäten gibt, die ggf. nicht direkt ersichtlich sind. Wenn das auftritt hilft es meist, beide Spalten nochmal neu zu konvertieren.

```
uros_egy["Date"]=pd.to_datetime(uros_egy['Date'], utc=True)
wetter_all["DATUM_DT"]=pd.to_datetime(wetter_all['DATUM_DT'], utc=True)

uros_egy_weater=pd.merge(uros_egy, wetter_all, right_on="DATUM_DT",
left_on="Date", how="left")
uros_egy_weater.head()
```

| VDW_Nr | AD       | AE     | Sub_ID | AS | AM     | AN  | ASD | AT   | UN       | N_3    | ..TMKUPMTXKTNKTGK | eor  | Tem                                         | Heiz | QNF_4 |     |     |      |      |                                                                              |       |
|--------|----------|--------|--------|----|--------|-----|-----|------|----------|--------|-------------------|------|---------------------------------------------|------|-------|-----|-----|------|------|------------------------------------------------------------------------------|-------|
|        |          |        |        |    |        |     |     |      |          |        |                   |      | Kuehl-<br>ra- Ta-<br>tur- ge<br>Klas-<br>se |      |       |     |     |      |      |                                                                              |       |
| 0      | 2020     | Na     | Na     | Na | Na     | Na  | Na  | Na   | 2020     | 427.30 | 2010              | 30.0 | 4.0                                         | 81.0 | 4.6   | 2.9 | 2.2 | eor  | Cold | Heiz                                                                         | nich5 |
|        | 00:00:00 | +00:00 |        |    |        |     |     |      | 00:00:00 | +00:00 |                   |      |                                             |      |       |     |     |      |      | g- al-<br>rad-<br>le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |       |
| 1      | 2020     | Na     | Na     | Na | 125629 | 1.6 | 0   | 2020 | 427.30   | 2010   | 31.0              | 3.4  | 83.0                                        | 4.4  | 2.3   | 0.9 | eor | Cold | Heiz | nich5                                                                        |       |
|        | 00:00:00 | +00:00 |        |    |        |     |     |      | 00:00:00 | +00:00 |                   |      |                                             |      |       |     |     |      |      | g- al-<br>rad-<br>le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-                 |       |

| EVDHE_NF_AB_Sab_DASUMMONSDATUN_3 ..TMKUPMTXKTNKTGK eorTemHeQNSQNF_4 |                          |                    |     |      |     |     |     |     |          | Kuehl-<br>ra- Ta-<br>tur- ge<br>Klas-<br>se |                                                                          |
|---------------------------------------------------------------------|--------------------------|--------------------|-----|------|-----|-----|-----|-----|----------|---------------------------------------------|--------------------------------------------------------------------------|
| 2                                                                   |                          |                    |     |      |     |     |     |     |          |                                             | ri-<br>giert                                                             |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210001.0 | 2.0 | 90.0 | 3.0 | 1.1 | 0.3 | eor | ColdHeiz | naich5                                      |                                                                          |
| 3                                                                   | 00:00:00+00:00           | 00:00:00+00:00     |     |      |     |     |     |     |          |                                             | g- al-<br>rad-le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210002.0 | 3.3 | 95.0 | 4.0 | 2.6 | 2.0 | eor | ColdHeiz | naich5                                      |                                                                          |
| 4                                                                   | 00:00:00+00:00           | 00:00:00+00:00     |     |      |     |     |     |     |          |                                             | g- al-<br>rad-le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210003.0 | 3.5 | 81.0 | 4.6 | 2.4 | 0.8 | eor | ColdHeiz | naich5                                      |                                                                          |
|                                                                     | 00:00:00+00:00           | 00:00:00+00:00     |     |      |     |     |     |     |          |                                             | g- al-<br>rad-le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210004.0 | 3.5 | 81.0 | 4.6 | 2.4 | 0.8 | eor | ColdHeiz | naich5                                      |                                                                          |
|                                                                     | 00:00:00+00:00           | 00:00:00+00:00     |     |      |     |     |     |     |          |                                             | g- al-<br>rad-le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210005.0 | 3.5 | 81.0 | 4.6 | 2.4 | 0.8 | eor | ColdHeiz | naich5                                      |                                                                          |
|                                                                     | 00:00:00+00:00           | 00:00:00+00:00     |     |      |     |     |     |     |          |                                             | g- al-<br>rad-le<br>tag Pa-<br>ra-<br>me-<br>ter<br>kor-<br>ri-<br>giert |
|                                                                     | 2021001-008010221290.0.0 | 2021407-00210006.0 | 3.5 | 81.0 | 4.6 | 2.4 | 0.8 | eor | ColdHeiz | naich5                                      |                                                                          |

## Duplikate entfernen

Nachdem wir beide Datensätze zusammengefügt haben, ergibt sich das neue Problem, duplizierter Daten. Das kann in Datensätzen vorkommen, so dass es immer sinnvoll ist, auf Duplikate zu prüfen. Hierfür gibt es die Funktion `uplicated`. Sie gibt eine boolesche Serie zurück, die Wahr ist für Zeilen, wo alle Werte doppelt vorkommen.

```
sum(wetter_all.duplicated())
```

```
2465
```

Das ist nur sinnvoll, wenn sich Zeilen vollständig wiederholen. In unserm Fall ist das Datum, der eigentliche Index und es ist empfohlen `duplicated` auf diesen anzuwenden. Wenn wir unseren zusammengefügte Dataframe prüfen, so sehen wir, dass wir einige Duplikate haben, d.h. Daten für die mehr als eine Messung existiert.

```
sum(wetter_all.duplicated(subset=['DATUM_DT']))
```

```
2465
```

Um diese herauszufiltern können wir die Funktion `drop_duplicates` nutzen, die ähnliche Parameter verwendet.

```
wetter_all.drop_duplicates(subset=['DATUM_DT'], inplace=True)  
wetter_all.shape
```

```
(28192, 24)
```

Damit haben wir jetzt einen gereinigten Dataframe, den wir uns für die weitere Verwendung abspeichern.

```
wetter_all.to_csv("../data/Wetter/warnemuende_clean.csv", sep=';',  
index=False)  
uros_egy_weater.to_csv("../data/UR05/EnergyID_weather_clean.csv", index=False)
```

## Dataframes Transformieren

Bisher haben wir mit einem Datensatz in Spaltendarstellung gearbeitet. Das ist für uns Menschen eine gewohnte tabellarische Darstellung. Für viele Berechnungen und Visualisierungen, ist das allerdings nicht immer das beste Format. Da braucht man häufig die Daten in einer Zeilendarstellung.

| Index Variables to Keep |            |                  | Column names to transform to „variable“ |      |     |         |     |      |        |      |      |      |       |       |
|-------------------------|------------|------------------|-----------------------------------------|------|-----|---------|-----|------|--------|------|------|------|-------|-------|
|                         | DATUM_DT   | TemperaturKlasse | RSK                                     | RSKF | SDK | SHK_TAG | NM  | VPM  | PM     | UPM  | TXK  | TMK  | TNK   | TGK   |
| 1461                    | 1951-01-01 | Cold             | 0.8                                     | 7.0  | 0.5 | 2.0     | 8.0 | 4.0  | 997.2  | 89.0 | -1.0 | -3.6 | -10.0 | -11.6 |
| 1462                    | 1951-01-02 | Cold             | 0.0                                     | 0.0  | 0.0 | 3.0     | 7.2 | 5.2  | 995.6  | 94.0 | 0.2  | -1.1 | -4.3  | -8.1  |
| 1463                    | 1951-01-03 | Cold             | 0.0                                     | 8.0  | 0.0 | 2.0     | 8.0 | 5.7  | 1001.1 | 96.0 | 0.4  | -0.1 | -1.5  | -2.2  |
| 1464                    | 1951-01-04 | Cold             | 0.4                                     | 7.0  | 0.0 | 2.0     | 7.8 | 5.9  | 1011.2 | 87.0 | 2.0  | 1.3  | -0.4  | -1.7  |
| 1465                    | 1951-01-05 | Cold             | 3.6                                     | 1.0  | 0.0 | 1.0     | 8.0 | 7.1  | 1006.3 | 96.0 | 4.2  | 3.0  | 0.0   | -1.2  |
| ...                     | ...        | ...              | ...                                     | ...  | ... | ...     | ... | ...  | ...    | ...  | ...  | ...  | ...   | ...   |
| 7208                    | 1966-09-26 | Cool             | 0.3                                     | 1.0  | 2.2 | 0.0     | 6.6 | 12.8 | 1014.0 | 86.0 | 15.1 | 13.1 | 9.5   | 7.9   |
| 7209                    | 1966-09-27 | Cold             | 0.0                                     | 1.0  | 6.1 | 0.0     | 5.8 | 10.1 | 1014.3 | 72.0 | 13.3 | 12.0 | 10.9  | 10.4  |
| 7210                    | 1966-09-28 | Cold             | 0.0                                     | 0.0  | 9.9 | 0.0     | 2.6 | 9.1  | 1017.5 | 74.0 | 12.5 | 9.7  | 7.7   | 8.7   |
| 7211                    | 1966-09-29 | Cold             | 0.0                                     | 0.0  | 7.8 | 0.0     | 4.8 | 9.2  | 1015.8 | 78.0 | 14.9 | 9.7  | 5.1   | 3.1   |
| 7212                    | 1966-09-30 | Cold             | 0.0                                     | 0.0  | 7.8 | 0.0     | 2.2 | 10.9 | 1007.9 | 83.0 | 16.1 | 11.1 | 5.8   | 4.8   |
| 5752 rows × 14 columns  |            |                  | Data to transformed to „value“          |      |     |         |     |      |        |      |      |      |       |       |

> melt >

< pivot <

| Index Variables to Keep |            |                  | Column names |       |
|-------------------------|------------|------------------|--------------|-------|
|                         | DATUM_DT   | TemperaturKlasse | variable     | value |
| 0                       | 1951-01-01 | Cold             | RSK          | 0.8   |
| 1                       | 1951-01-02 | Cold             | RSK          | 0.0   |
| 2                       | 1951-01-03 | Cold             | RSK          | 0.0   |
| 3                       | 1951-01-04 | Cold             | RSK          | 0.4   |
| 4                       | 1951-01-05 | Cold             | RSK          | 3.6   |
| ...                     | ...        | ...              | ...          | ...   |
| 69019                   | 1966-09-26 | Cool             | TGK          | 7.9   |
| 69020                   | 1966-09-27 | Cold             | TGK          | 10.4  |
| 69021                   | 1966-09-28 | Cold             | TGK          | 8.7   |
| 69022                   | 1966-09-29 | Cold             | TGK          | 3.1   |
| 69023                   | 1966-09-30 | Cold             | TGK          | 4.8   |
| 69024 rows × 4 columns  |            |                  | Data         |       |

> melt >

< pivot <

| Index Variables to Keep |            | Column names     |                |
|-------------------------|------------|------------------|----------------|
|                         | DATUM_DT   | TemperaturKlasse | variable value |
| 0                       | 1951-01-01 | Cold             | RSK 0.8        |
| 1                       | 1951-01-02 | Cold             | RSK 0.0        |
| 2                       | 1951-01-03 | Cold             | RSK 0.0        |
| 3                       | 1951-01-04 | Cold             | RSK 0.4        |
| 4                       | 1951-01-05 | Cold             | RSK 3.6        |
| ...                     | ...        | ...              | ...            |
| 69019                   | 1966-09-26 | Cool             | TGK 7.9        |
| 69020                   | 1966-09-27 | Cold             | TGK 10.4       |
| 69021                   | 1966-09-28 | Cold             | TGK 8.7        |
| 69022                   | 1966-09-29 | Cold             | TGK 3.1        |
| 69023                   | 1966-09-30 | Cold             | TGK 4.8        |

69024 rows × 4 columns

Data

Abbildung 1: Dataframe Format in Spalten und Zeilendarstellung

Nehmen wir als Beispiel die Temperaturwerte und Klassen aus unserem Datensatz. Wir wollen jetzt den Zusammenhang zwischen der Temperaturklasse und der unteren TNK, mittleren TMK und oberen Temperatur TXK bestimmen. Wir könnten das jetzt für jede Spalte einzeln machen, aber es wäre praktischer, wenn wir statt dessen nur eine Spalte mit Werten hätten. Das könnte man erreichen indem man jede Spalte raus löst und dann mit `pd.concat` in einem neuen Dataframe zusammenführt. Oder man spart sich den Aufwand und nutzt die Methode `melt`. Dieser übergeben wir zwei Parameter. Der Parameter `id_vars` enthält eine Liste der Indexwerte, die wir behalten wollen und `value_vars` spezifiziert alle Spalten die wir in eine Reihendarstellung überführen wollen.

```
wetter_melted = wetter_all.melt(id_vars=['DATUM_DT', "TemperaturKlasse"],
value_vars=["TXK", "TMK", "TNK"])
wetter_melted
```

|       | DATUM_DT                  | TemperaturKlasse | variable | value |
|-------|---------------------------|------------------|----------|-------|
| 0     | 1947-01-01 00:00:00+00:00 | Cold             | TXK      | -1.5  |
| 1     | 1947-01-02 00:00:00+00:00 | Cold             | TXK      | -1.5  |
| 2     | 1947-01-03 00:00:00+00:00 | Cold             | TXK      | -2.4  |
| 3     | 1947-01-04 00:00:00+00:00 | Cold             | TXK      | -2.4  |
| 4     | 1947-01-05 00:00:00+00:00 | Cold             | TXK      | -9.4  |
| ...   | ...                       | ...              | ...      | ...   |
| 84571 | 2024-03-04 00:00:00+00:00 | Cold             | TNK      | 3.1   |
| 84572 | 2024-03-05 00:00:00+00:00 | Cold             | TNK      | 2.1   |
| 84573 | 2024-03-06 00:00:00+00:00 | Cold             | TNK      | 2.8   |
| 84574 | 2024-03-07 00:00:00+00:00 | Cold             | TNK      | 0.6   |
| 84575 | 2024-03-08 00:00:00+00:00 | Cold             | TNK      | -0.4  |

Im neuen Dataframe sehen wir jetzt, dass die ursprünglichen Spalten TNK,TMK,TxK mit Werten fehlen und stattdessen es eine neue Spalte `variable` und `value` gibt. Wenn man sich die

Werte in der Spalte `variable` ansieht versteht man, was passiert ist: Hier stehen jetzt die alten Spaltennamen `TNK`, `TMK`, `TxK` und in der neuen Spalte `value` die neuen Werte. Damit haben wir den Spaltennamen in eine beschreibende kategorische Variable umgewandelt. Der neue Dataframe ist dreimal so lang wie der ursprüngliche Dataframe, weshalb man auch von einer Zeilendarstellung redet.

Aus der Zeilendarstellung könnten wir jetzt ein einzelnes ML-Modell bauen, was den Zusammenhang zwischen der Temperaturklasse und den Werten in `value` lernt und `variable` nur als Unterscheidungskriterium nutzt. Alternativ können wir damit auch die Zeitreihen einfacher darstellen. So lässt sich damit ein sehr übersichtliches Diagramm zum Vergleich aller Zeitreihen erstellen. Im nächsten Abschnitt Visuelle Datenanalyse gehen wir detaillierte darauf ein.

```
import plotly.express as px

px.line(wetter_melted, x="DATUM_DT", y="value", facet_row="variable")
```

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

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

Diese Zeilendarstellung wird auch oft bei komplexen Dataframes mit vielen kategorischen Variablen benutzt. So gibt es viele Datensätze des Statistischen Bundesamtes auf dem GENESIS-Portal im CSV-Flat format, was deutlich einfacher zu verarbeiten ist. Allerdings muss dann der Datensatz aus der Reihendarstellung in eine Spaltendarstellung überführt werden.

Hierfür bietet Pandas die Methode `pivot` an. Jeder, der sich gewundert hat, wofür man die Pivot-Darstellung in Excel nutzt, hat jetzt eine Antwort. Ähnlich zur `melt`-Methode, müssen wir bei der `pivot`-Methode die Spalten identifizieren, welche die Werte `values` enthält, die Spalte(n) `columns` welche in neue Spaltennamen umzuwandeln sind und die beizubehaltenden Indexe `index`.

```
wetter_hardened = wetter_melted.pivot(columns='variable', index=['DATUM_DT',
"TemperaturKlasse"], values='value')
wetter_hardened
```

|                           |                  | variable | TMK   | TNK   | TXK  |
|---------------------------|------------------|----------|-------|-------|------|
| DATUM_DT                  | TemperaturKlasse |          |       |       |      |
| 1947-01-01 00:00:00+00:00 | Cold             |          | -2.1  | -3.3  | -1.5 |
| 1947-01-02 00:00:00+00:00 | Cold             |          | -2.1  | -3.3  | -1.5 |
| 1947-01-03 00:00:00+00:00 | Cold             |          | -4.8  | -6.2  | -2.4 |
| 1947-01-04 00:00:00+00:00 | Cold             |          | -4.8  | -6.2  | -2.4 |
| 1947-01-05 00:00:00+00:00 | Cold             |          | -12.9 | -14.0 | -9.4 |

|     |                           | variable         | TMK | TNK  | TXK  |
|-----|---------------------------|------------------|-----|------|------|
|     | DATUM_DT                  | TemperaturKlasse |     |      |      |
| ... | ...                       | ...              | ... | ...  | ...  |
|     | 2024-03-04 00:00:00+00:00 | Cold             | 7.1 | 3.1  | 12.0 |
|     | 2024-03-05 00:00:00+00:00 | Cold             | 3.7 | 2.1  | 6.0  |
|     | 2024-03-06 00:00:00+00:00 | Cold             | 4.2 | 2.8  | 6.2  |
|     | 2024-03-07 00:00:00+00:00 | Cold             | 4.4 | 0.6  | 9.5  |
|     | 2024-03-08 00:00:00+00:00 | Cold             | 2.5 | -0.4 | 8.9  |

Mit der bekannten Methode `reset_index()` wandeln wir den Index in Spalten um:

```
wetter_hardened.reset_index()
```

|       | variable | DATUM_DT                  | TemperaturKlasse | TMK   | TNK   | TXK  |
|-------|----------|---------------------------|------------------|-------|-------|------|
| 0     |          | 1947-01-01 00:00:00+00:00 | Cold             | -2.1  | -3.3  | -1.5 |
| 1     |          | 1947-01-02 00:00:00+00:00 | Cold             | -2.1  | -3.3  | -1.5 |
| 2     |          | 1947-01-03 00:00:00+00:00 | Cold             | -4.8  | -6.2  | -2.4 |
| 3     |          | 1947-01-04 00:00:00+00:00 | Cold             | -4.8  | -6.2  | -2.4 |
| 4     |          | 1947-01-05 00:00:00+00:00 | Cold             | -12.9 | -14.0 | -9.4 |
| ...   | ...      | ...                       | ...              | ...   | ...   | ...  |
| 28187 |          | 2024-03-04 00:00:00+00:00 | Cold             | 7.1   | 3.1   | 12.0 |
| 28188 |          | 2024-03-05 00:00:00+00:00 | Cold             | 3.7   | 2.1   | 6.0  |
| 28189 |          | 2024-03-06 00:00:00+00:00 | Cold             | 4.2   | 2.8   | 6.2  |
| 28190 |          | 2024-03-07 00:00:00+00:00 | Cold             | 4.4   | 0.6   | 9.5  |
| 28191 |          | 2024-03-08 00:00:00+00:00 | Cold             | 2.5   | -0.4  | 8.9  |

## Bibliographie