

Lineare Regression

Joern Ploennigs · AI4SC



Das Ergebnis habe ich schon, jetzt brauche ich nur noch den Weg, der zu ihm führt.

— Carl Friedrich Gauß

Regressionsmodelle sind eine Klasse von statistischen Modellen zum überwachten Lernen, die verwendet werden, um die Beziehung zwischen einer abhängigen Variable und einer oder mehreren unabhängigen Variablen zu modellieren. Sie werden häufig eingesetzt, um Vorhersagen zu treffen, Muster in Daten zu erkennen oder als Basismodell zum Vergleich. Zum Beispiel zur Vorhersage der Betonfestigkeit in Abhängigkeit von Mischungsverhältnis und Aushärtungszeit, zur Schätzung der Setzung eines Fundaments basierend auf Bodenkennwerten oder zur Prognose von Energieverbrauch eines Gebäudes.

Wir wollen verschiedene univariate und multivariate Modelle behandeln. Hierbei unterscheidet sich die Anzahl der Eingangsvariablen. Modelle mit einer Eingangsvariable werden *univariat* genannt und Modelle mit mehreren Eingangsvariablen heißen *multivariat*. Dies ist wichtig, da in der Praxis viele Ergebnisse nicht nur von einem Faktor abhängen. Zum Beispiel, die Aushärtung von Beton hängt von Temperatur, Luftfeuchtigkeit und Mischung ab – ein multivariates Modell kann das berücksichtigen.

Univariate Lineare Regression

Eines der einfachsten Machine Learning Modelle ist die *Lineare Regression*. Es basiert auf der Idee einer linearen Funktion aus der linearen Algebra. Dort ist eine lineare Funktion definiert als:

$$\hat{y}_i = \hat{f}(x_i) = \beta_0 + \beta_1 \cdot x_i.$$

Die Funktion modelliert Beziehung zwischen der Zielvariable $\hat{y}_i$ und einer unabhängigen Variable x_i in Form einer einfachen Linie definiert durch die Funktion $\hat{f}(x_i)$. Die Parameter sind die y-Achsenabschnitt (auch: Schnittpunkt mit der y-Achse) (en. intercept) β_0 und der Anstieg β_1 .

Wenn man z.B. weiß, dass bei einem Beton mit einer 0.45 Wasser-Zement-Wert eine Druckfestigkeit von 45 MPa erreicht wird, und bei einem 0.55 Wasser-Zement-Wert nur 35 MPa, kann ein lineares Modell helfen, diese Beziehung zu approximieren.

Bestimmen der Parameter für zwei Punkte

Aus der linearen Algebra wissen wir, dass wir die Parameter β_0 und β_1 bestimmen können (en. *fitting*), wenn wir zwei Punkte $[y_0, x_0]$ und $[y_1, x_1]$ kennen und das Gleichungssystem für diese aufstellen und lösen

Wir stellen y_0 nach β_0 um

$$\beta_0 = y_0 - \beta_1 \cdot x_0$$

und setzen es in y_1 ein und erhalten

Damit können wir die Parameter für einen Datensatz mit zwei Punkten bestimmen

```
import numpy as np # Import von NumPy
import pandas as pd # Import von Pandas
import plotly.express as px # Import von Plotly

xy = pd.DataFrame({"x": [20, 80], "y": [2, 3]})
px.scatter(xy, x="x", y="y")
```

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

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

```
beta_1 = (xy.y[1] - xy.y[0]) / (xy.x[1] - xy.x[0])
beta_0 = xy.y[0] - beta_1 * xy.x[0]
beta_0, beta_1
```

```
(np.float64(1.6666666666666667), np.float64(0.016666666666666666))
```

Jetzt haben wir ein einfaches Modell bestimmt und wir können damit beliebige Punkte auf dieser Linie berechnen. Sind wir innerhalb des Bereiches $20 \leq x \leq 80$ so nennt man das *Interpolation*. Sind wir außerhalb $x < 20$ oder $x > 80$ so nennt man es *Extrapolation*.

```
xy = pd.DataFrame({"x": range(100)})
xy["y"] = beta_0 + beta_1 * xy.x
xy["type"] = np.where((20 <= xy.x) & (xy.x <= 80), "Interpolation",
```

```
"Extrapolation")
```

```
px.line(xy, x="x", y="y", color="type", markers=True)
```

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

Bestimmen der Parameter für mehrere Punkte

Diese Lösung funktioniert nicht, wenn wir eine lineare Funktion für mehrere Punkte bestimmen wollen, da diese dann überbestimmt ist, sofern nicht alle Punkte auf einer Linie liegen.

Um die Parameter β_0 und β_1 zu schätzen, verwenden wir die Methode der kleinsten Quadrate (en. Ordinary Least Squares - OLS). Sie wurde von Carl Friedrich Gauß entwickelt zum Annähern von Planetenbahnen. OLS minimiert die Summe der quadrierten Anpassungsfehler. Der *Anpassungsfehler* ε_i berechnet sich aus der Differenz des Zielwertes y_i und der Vorhersage des linearen Modells $\hat{y}_i$

$$\varepsilon_i = \hat{y}_i - y_i = \beta_0 + \beta_1 x_i - y_i.$$

Dieser Anpassungsfehler wird auch als *Residuum* bezeichnet. Man schreibt deshalb das lineare Modell auch oft als

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i.$$

Zu beachten ist, dass y_i der Zielwert ist und nicht die Vorhersage $\hat{y}_i$ von zuvor.

Hierbei wird angenommen, dass der Fehler ε_i *normalverteilt* ist. Unter dieser Annahme lassen sich die Koeffizienten β_0 und β_1 mit der kleinsten Summe der Fehlerquadrate aus folgendem Optimierungsproblem bestimmen

$$\min_{\beta_0, \beta_1} \sum_{i=1}^n \varepsilon_i^2.$$

Aus der Mathematik ist bekannt, dass eine Funktion ein lokales Extremum erreicht, wenn ihre erste Ableitung zu Null wird. Wir können also das Minimum bestimmen, indem wir die erste Ableitung gleich Null setzen. Da β_0 und β_1 unsere variablen Parameter sind, bilden wir hierfür die partielle Ableitung nach β_0 und β_1 und bestimmen davon die Nullstellen. Dabei hilft uns das Fehlerquadrat, dass es dadurch eine eindeutige Lösung gibt, welche gegeben ist durch das lineare Gleichungssystem

mit der Lösung für den Anstieg

$$\beta_1 = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

mit der Summe der Abweichungsprodukte im Zähler und der Summe der Abweichungsquadrate im Nenner. Für die Abweichung ergibt sich

EVDHE_NF_AB_Lab_DAS_TMKONS-DATUM_N_3 ..TMKUPMTXKTNKTGK eorTemHeiQNSNF_4										Kuehl- ra- Ta- tur- ge Klas- se	
2	2021001-008010221290.0	2021407-00210001.0	2.0	90.0	3.0	1.1	0.3	eor	Cold	Heiz	Rad-5
	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	2021407-00210002.0	3.3	95.0	4.0	2.6	2.0	eor	Cold	Heiz	Rad-5
	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	2021407-00210003.0	3.5	81.0	4.6	2.4	0.8	eor	Cold	Heiz	Rad-5
	00:00:00+00:00	00:00:00+00:00								g- al- rad-le tag Pa- ra- me- ter kor- ri- giert	

Wir interessieren uns als erstes den Energieverbrauch ES_Lab und die mittlere Außentemperatur TMK. Wir betrachten uns das Streudiagramm. Darin lässt sich ein leichter linearer Trend erkennen, da der Energieverbrauch bei hohen Temperaturen über 10 Grad steigt, was auf eine aktive Kühlung hinweist.

```
fig = px.scatter(egywth, x="TMK", y="ES_Lab", opacity=.5)
fig
```

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

Mit den Formeln oben können wir uns jetzt ein lineares Modell berechnen. Hierbei ist ES_Lab unsere Zielvariable y und TMK unsere Eingangsvariable x .

```
beta_1 = ((egywth.TMK - egywth.TMK.mean())*(egywth.ES_Lab -
egywth.ES_Lab.mean())).sum()
beta_1 /= (egywth.TMK - egywth.TMK.mean()).pow(2).sum()

beta_0 = egywth.ES_Lab.mean() - beta_1 * egywth.TMK.mean()

beta_0, beta_1
```

```
(np.float64(7.95973736889691), np.float64(3.1496156006907596))
```

Um das Lineare Modell mit dem Originalwerten zu vergleichen, können wir die Vorhersagen entsprechend der ersten Formel bestimmen

```
# Wir nehmen noch einmal die TMK Werte und sortieren sie
pred_x = np.sort(egywth.TMK.values)

pred_y = beta_0 + beta_1 * pred_x
```

```
fig.add_scatter(x=pred_x, y=pred_y, name="manual")
fig
```

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

Hier zeigt sich die im Streudiagramm leicht sichtbare Korrelation nun sehr deutlich. Das lineare Regressionsmodell erlaubt uns also insbesondere lineare Zusammenhänge zu extrahieren, die sonst in der Streuung untergehen.

Das können wir uns sogar im Streudiagramm mit Plotly direkt berechnen lassen, indem wir die Trendlinie ols angeben. Dieses Lineare Regressionsmodell wird von Plotly mit der Bibliothek statsmodels berechnet.

```
px.scatter(egywth, x="TMK", y="ES_Lab", opacity=.5, trendline="ols")
```

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

ML-Frameworks für Lineare Regression

Plotly zeigt bereits, dass es nicht notwendig ist, dass man sich das lineare Modell selbst berechnet. Es gibt verschiedene Pakete die ML-Modelle implementieren und genutzt werden können. Wir diskutieren zwei beliebte Bibliotheken für die lineare Regression in Python. Das ist zum einen das von Plotly genutzte `statsmodels.formula.api`. Sie ist nützlich, wenn wir die Modellkomponenten selbst bestimmen wollen und die ideale Kombination an Variablen identifizieren wollen, da es einen einfachen Formelsyntax bietet und bietet gut gestaltete Zusammenfassungstabellen. Die andere Bibliothek, `sklearn`, ist eher zum Vergleich verschiedener Vorhersagemodelle geeignet, da die Bibliothek ein standardisiertes Vorgehen und Methoden bietet. Allerdings gibt es kein Formelsyntax, sondern erwartet die Eingaben als Numpy Matrizen und bietet auch keine schön formatierten Zusammenfassungstabellen.

SciKit-Learn

Die erste Bibliothek ist SciKit-Learn. Sie enthält viele ML-Modelle inklusive der Linearen Regression.

```
from sklearn.linear_model import LinearRegression

sklm = LinearRegression() # Das legt nur eine Instanz der Modellklasse an
```

Dieser Aufruf erstellt nur eine Modellinstanz aber enthält noch keine Daten. Diese müssen wir trainieren, was wir mit der Methode `fit()` machen können, der wir die Eingabedaten und die Zieldaten übergeben. SciKit-Learn arbeitet allerdings nicht direkt auf Pandas DataFrames sondern mit Numpy Arrays, weshalb wir die Pandas Serien in ein Numpy Arrays umwandeln müssen, wofür wir das `values` Attribut und die Methode `reshape` nutzen (um das Array in eine Matrix umzuwandeln).

```
X = egywth.TMK.values.reshape(-1, 1) # Wir erstellen eine Matrix der
Eingangsdaten
Y = egywth.ES_Lab.values # Wir erstellen einen Vektor der
Zielaten
sklm.fit(X, Y)
```

ValueError: Input y contains NaN.

```
[31m-----[?]
[39m
[31mValueError[39m                                Traceback (most recent
call last)
[36mCell[39m[36m [39m[32mIn[11]:[39m[32m, line 3[39m
[32m      1[39m X = egywth.TMK.values.reshape(-[32m1[39m,
[32m[39m) [38;5;66;03m# Wir erstellen eine Matrix der
Eingangsdaten[39;00m
[32m      2[39m Y = egywth.ES_Lab.values [38;5;66;03m# Wir
```

```

erstellen einen Vektor der Zieldaten[39;00m
[32m---> [39m[32m3[39m sklm.fit(X, Y)

[36mFile [39m[32m~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/
base.py:1336[39m, in
[36m_fit_context.<locals>.decorator.<locals>.wrapper[39m[34m(estimator,
*args, **kwargs)[39m
[32m    1329[39m     estimator._validate_params()
[32m    1331[39m [38;5;28;01mwith[39;00m config_context(
[32m    1332[39m     skip_parameter_validation=(
[32m    1333[39m         prefer_skip_nested_validation
[38;5;129;01mor[39;00m global_skip_validation
[32m    1334[39m     )
[32m    1335[39m ):
[32m-> [39m[32m1336[39m     [38;5;28;01mreturn[39;00m
[43mfit_method[49m[43m([49m[43mestimator[49m[43m,[49m[43m
[49m[43m*[49m[43margs[49m[43m,[49m[43m
[49m[43m*[49m[43m*[49m[43m*[49m[43mkwarg[49m[43m)[49m

[36mFile [39m[32m~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/
linear_model/_base.py:630[39m, in
[36mLinearRegression.fit[39m[34m(self, X, y, sample_weight)[39m
[32m    626[39m n_jobs_ = [38;5;28;01mself[39m.n_jobs
[32m    628[39m accept_sparse = [38;5;28;01mFalse[39;00m
[38;5;28;01mif[39;00m [38;5;28;01mself[39m.positive
[38;5;28;01melse[39;00m [33m"[39m[33mcsr[39m[33m"[39m,
[33m"[39m[33mcs[39m[33m"[39m,
[33m"[39m[33mcoo[39m[33m"[39m]
[32m--> [39m[32m630[39m X, y = [43mvalidate_data[49m[43m([49m
[32m    631[39m [43m     [49m[38;5;28;43;01mself[39;49m[43m,[49m
[32m    632[39m [43m     [49m[43mX[49m[43m,[49m
[32m    633[39m [43m     [49m[43my[49m[43m,[49m
[32m    634[39m [43m
[49m[43maccept_sparse[49m[43m=[49m[43maccept_sparse[49m[43m,[49m
[49m
[32m    635[39m [43m
[49m[43my_numeric[49m[43m=[49m[38;5;28;43;01mTrue[39;49;00m[43m,[49m
[49m
[32m    636[39m [43m
[49m[43mmulti_output[49m[43m=[49m[38;5;28;43;01mTrue[39;49;00m[43m,[49m
[49m

```

```

637 force_writeable = True
638 if not force_writeable:
640     has_sw = sample_weight is not None
641     if has_sw:
        File ~/Documents/code/Lehre/IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/utils/validation.py:2919: in validate_data
        2917         y = check_array(y,
        input_name='y', **check_y_params)
        2918         else:
        2919         X, y =
        2920         check_X_y(X, y, accept_sparse=accept_sparse,
        2921                    dtype=dtype, order=order, copy=copy, force_writeable=force_writeable,
        2922                    ensure_2d=ensure_2d, allow_nd=allow_nd, multi_output=multi_output,
        2923                    ensure_min_samples=ensure_min_samples, ensure_min_features=ensure_min_features,
        2924                    y_numeric=y_numeric, estimator=estimator)
        1310         raise ValueError(
        1311             f'{estimator_name} requires y to be passed, but the target y is
        1312             None')
        1314     X = check_array(
        1315         X,
        1316         accept_sparse=accept_sparse,
        1317         ...
        1318         input_name='X',

```

```

1329 )
1330 y =
1331 check_array(y, dtype=None, order='C',
1332              multi_output=False,
1333              estimator=None,
1334              ensure_2d=True,
1335              ensure_min_samples=1,
1336              ensure_min_features=1,
1337              accept_sparse='csr',
1338              max_depth=None,
1339              max_features=None,
1340              max_leaf_nodes=None,
1341              min_features_split=0.5,
1342              min_samples_leaf=1,
1343              min_samples_split=1,
1344              min_weight_fraction=0.0,
1345              min_weight_fraction_leaf=0.0,
1346              verbose=0,
1347              warm_start=False,
1348              y_numeric=False)
1349 estimator_name = _check_estimator_name(estimator)
1350
1351 File ~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/utils/
validation.py:1074, in check_array(y, multi_output,
y_numeric, estimator)
1352 """Isolated part of check_X_y
dedicated to y validation"""
1353 if multi_output:
1354     y = check_array(y, dtype=None, order='C',
1355                     multi_output=False,
1356                     estimator=None,
1357                     ensure_2d=True,
1358                     ensure_min_samples=1,
1359                     ensure_min_features=1,
1360                     accept_sparse='csr',
1361                     max_depth=None,
1362                     max_features=None,
1363                     max_leaf_nodes=None,
1364                     min_features_split=0.5,
1365                     min_samples_leaf=1,
1366                     min_samples_split=1,
1367                     min_weight_fraction=0.0,
1368                     min_weight_fraction_leaf=0.0,
1369                     verbose=0,
1370                     warm_start=False,
1371                     y_numeric=False)
1372 else:
1373     y = check_array(y, dtype=None, order='C',
1374                     multi_output=False,
1375                     estimator=None,
1376                     ensure_2d=True,
1377                     ensure_min_samples=1,
1378                     ensure_min_features=1,
1379                     accept_sparse='csr',
1380                     max_depth=None,
1381                     max_features=None,
1382                     max_leaf_nodes=None,
1383                     min_features_split=0.5,
1384                     min_samples_leaf=1,
1385                     min_samples_split=1,
1386                     min_weight_fraction=0.0,
1387                     min_weight_fraction_leaf=0.0,
1388                     verbose=0,
1389                     warm_start=False,
1390                     y_numeric=True)
1391
1392 File ~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/utils/
validation.py:1074, in check_array(y, multi_output,
y_numeric, estimator)

```

```

accept_sparse, accept_large_sparse, dtype, order, copy, force_writeable,
ensure_all_finite, ensure_non_negative, ensure_2d, allow_nd,
ensure_min_samples, ensure_min_features, estimator, input_name)
    1068 raise ValueError(
    1069         Found array with
dim {array.ndim}
    1070         while dim <= 2 is
required{
    1071     )
    1073     if ensure_all_finite:
    1074         assert_all_finite(
    1075             array,
    1076             input_name=
    1077             estimator_name=
    1078             allow_nan=
    1079             ensure_all_finite=
    1080             ==
    1081             if
    1082             if
    1083             # only make a copy if `array` and
`array_orig` may share memory`

File ~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/utils/
validation.py:133, in assert_all_finite(X, allow_nan,
msg_dtype, estimator_name, input_name)
    130     if
    131         return
    132 -->
    133     assert_all_finite_element_wise(
    134         X,
    135         X,
    136         X,
    137         X,
    138         X,
    139         X,
    140         X,
    141         X,
    142         X,
    143         X,
    144         X,
    145         X,
    146         X,
    147         X,
    148         X,
    149         X,
    150         X,
    151         X,
    152         X,
    153         X,
    154         X,
    155         X,
    156         X,
    157         X,
    158         X,
    159         X,
    160         X,
    161         X,
    162         X,
    163         X,
    164         X,
    165         X,
    166         X,
    167         X,
    168         X,
    169         X,
    170         X,
    171         X,
    172         X,
    173         X,
    174         X,
    175         X,
    176         X,
    177         X,
    178         X,
    179         X,
    180         X,
    181         X,
    182         X,
    183         X,
    184         X,
    185         X,
    186         X,
    187         X,
    188         X,
    189         X,
    190         X,
    191         X,
    192         X,
    193         X,
    194         X,
    195         X,
    196         X,
    197         X,
    198         X,
    199         X,
    200         X,
    201         X,
    202         X,
    203         X,
    204         X,
    205         X,
    206         X,
    207         X,
    208         X,
    209         X,
    210         X,
    211         X,
    212         X,
    213         X,
    214         X,
    215         X,
    216         X,
    217         X,
    218         X,
    219         X,
    220         X,
    221         X,
    222         X,
    223         X,
    224         X,
    225         X,
    226         X,
    227         X,
    228         X,
    229         X,
    230         X,
    231         X,
    232         X,
    233         X,
    234         X,
    235         X,
    236         X,
    237         X,
    238         X,
    239         X,
    240         X,
    241         X,
    242         X,
    243         X,
    244         X,
    245         X,
    246         X,
    247         X,
    248         X,
    249         X,
    250         X,
    251         X,
    252         X,
    253         X,
    254         X,
    255         X,
    256         X,
    257         X,
    258         X,
    259         X,
    260         X,
    261         X,
    262         X,
    263         X,
    264         X,
    265         X,
    266         X,
    267         X,
    268         X,
    269         X,
    270         X,
    271         X,
    272         X,
    273         X,
    274         X,
    275         X,
    276         X,
    277         X,
    278         X,
    279         X,
    280         X,
    281         X,
    282         X,
    283         X,
    284         X,
    285         X,
    286         X,
    287         X,
    288         X,
    289         X,
    290         X,
    291         X,
    292         X,
    293         X,
    294         X,
    295         X,
    296         X,
    297         X,
    298         X,
    299         X,
    300         X,
    301         X,
    302         X,
    303         X,
    304         X,
    305         X,
    306         X,
    307         X,
    308         X,
    309         X,
    310         X,
    311         X,
    312         X,
    313         X,
    314         X,
    315         X,
    316         X,
    317         X,
    318         X,
    319         X,
    320         X,
    321         X,
    322         X,
    323         X,
    324         X,
    325         X,
    326         X,
    327         X,
    328         X,
    329         X,
    330         X,
    331         X,
    332         X,
    333         X,
    334         X,
    335         X,
    336         X,
    337         X,
    338         X,
    339         X,
    340         X,
    341         X,
    342         X,
    343         X,
    344         X,
    345         X,
    346         X,
    347         X,
    348         X,
    349         X,
    350         X,
    351         X,
    352         X,
    353         X,
    354         X,
    355         X,
    356         X,
    357         X,
    358         X,
    359         X,
    360         X,
    361         X,
    362         X,
    363         X,
    364         X,
    365         X,
    366         X,
    367         X,
    368         X,
    369         X,
    370         X,
    371         X,
    372         X,
    373         X,
    374         X,
    375         X,
    376         X,
    377         X,
    378         X,
    379         X,
    380         X,
    381         X,
    382         X,
    383         X,
    384         X,
    385         X,
    386         X,
    387         X,
    388         X,
    389         X,
    390         X,
    391         X,
    392         X,
    393         X,
    394         X,
    395         X,
    396         X,
    397         X,
    398         X,
    399         X,
    400         X,
    401         X,
    402         X,
    403         X,
    404         X,
    405         X,
    406         X,
    407         X,
    408         X,
    409         X,
    410         X,
    411         X,
    412         X,
    413         X,
    414         X,
    415         X,
    416         X,
    417         X,
    418         X,
    419         X,
    420         X,
    421         X,
    422         X,
    423         X,
    424         X,
    425         X,
    426         X,
    427         X,
    428         X,
    429         X,
    430         X,
    431         X,
    432         X,
    433         X,
    434         X,
    435         X,
    436         X,
    437         X,
    438         X,
    439         X,
    440         X,
    441         X,
    442         X,
    443         X,
    444         X,
    445         X,
    446         X,
    447         X,
    448         X,
    449         X,
    450         X,
    451         X,
    452         X,
    453         X,
    454         X,
    455         X,
    456         X,
    457         X,
    458         X,
    459         X,
    460         X,
    461         X,
    462         X,
    463         X,
    464         X,
    465         X,
    466         X,
    467         X,
    468         X,
    469         X,
    470         X,
    471         X,
    472         X,
    473         X,
    474         X,
    475         X,
    476         X,
    477         X,
    478         X,
    479         X,
    480         X,
    481         X,
    482         X,
    483         X,
    484         X,
    485         X,
    486         X,
    487         X,
    488         X,
    489         X,
    490         X,
    491         X,
    492         X,
    493         X,
    494         X,
    495         X,
    496         X,
    497         X,
    498         X,
    499         X,
    500         X,
    501         X,
    502         X,
    503         X,
    504         X,
    505         X,
    506         X,
    507         X,
    508         X,
    509         X,
    510         X,
    511         X,
    512         X,
    513         X,
    514         X,
    515         X,
    516         X,
    517         X,
    518         X,
    519         X,
    520         X,
    521         X,
    522         X,
    523         X,
    524         X,
    525         X,
    526         X,
    527         X,
    528         X,
    529         X,
    530         X,
    531         X,
    532         X,
    533         X,
    534         X,
    535         X,
    536         X,
    537         X,
    538         X,
    539         X,
    540         X,
    541         X,
    542         X,
    543         X,
    544         X,
    545         X,
    546         X,
    547         X,
    548         X,
    549         X,
    550         X,
    551         X,
    552         X,
    553         X,
    554         X,
    555         X,
    556         X,
    557         X,
    558         X,
    559         X,
    560         X,
    561         X,
    562         X,
    563         X,
    564         X,
    565         X,
    566         X,
    567         X,
    568         X,
    569         X,
    570         X,
    571         X,
    572         X,
    573         X,
    574         X,
    575         X,
    576         X,
    577         X,
    578         X,
    579         X,
    580         X,
    581         X,
    582         X,
    583         X,
    584         X,
    585         X,
    586         X,
    587         X,
    588         X,
    589         X,
    590         X,
    591         X,
    592         X,
    593         X,
    594         X,
    595         X,
    596         X,
    597         X,
    598         X,
    599         X,
    600         X,
    601         X,
    602         X,
    603         X,
    604         X,
    605         X,
    606         X,
    607         X,
    608         X,
    609         X,
    610         X,
    611         X,
    612         X,
    613         X,
    614         X,
    615         X,
    616         X,
    617         X,
    618         X,
    619         X,
    620         X,
    621         X,
    622         X,
    623         X,
    624         X,
    625         X,
    626         X,
    627         X,
    628         X,
    629         X,
    630         X,
    631         X,
    632         X,
    633         X,
    634         X,
    635         X,
    636         X,
    637         X,
    638         X,
    639         X,
    640         X,
    641         X,
    642         X,
    643         X,
    644         X,
    645         X,
    646         X,
    647         X,
    648         X,
    649         X,
    650         X,
    651         X,
    652         X,
    653         X,
    654         X,
    655         X,
    656         X,
    657         X,
    658         X,
    659         X,
    660         X,
    661         X,
    662         X,
    663         X,
    664         X,
    665         X,
    666         X,
    667         X,
    668         X,
    669         X,
    670         X,
    671         X,
    672         X,
    673         X,
    674         X,
    675         X,
    676         X,
    677         X,
    678         X,
    679         X,
    680         X,
    681         X,
    682         X,
    683         X,
    684         X,
    685         X,
    686         X,
    687         X,
    688         X,
    689         X,
    690         X,
    691         X,
    692         X,
    693         X,
    694         X,
    695         X,
    696         X,
    697         X,
    698         X,
    699         X,
    700         X,
    701         X,
    702         X,
    703         X,
    704         X,
    705         X,
    706         X,
    707         X,
    708         X,
    709         X,
    710         X,
    711         X,
    712         X,
    713         X,
    714         X,
    715         X,
    716         X,
    717         X,
    718         X,
    719         X,
    720         X,
    721         X,
    722         X,
    723         X,
    724         X,
    725         X,
    726         X,
    727         X,
    728         X,
    729         X,
    730         X,
    731         X,
    732         X,
    733         X,
    734         X,
    735         X,
    736         X,
    737         X,
    738         X,
    739         X,
    740         X,
    741         X,
    742         X,
    743         X,
    744         X,
    745         X,
    746         X,
    747         X,
    748         X,
    749         X,
    750         X,
    751         X,
    752         X,
    753         X,
    754         X,
    755         X,
    756         X,
    757         X,
    758         X,
    759         X,
    760         X,
    761         X,
    762         X,
    763         X,
    764         X,
    765         X,
    766         X,
    767         X,
    768         X,
    769         X,
    770         X,
    771         X,
    772         X,
    773         X,
    774         X,
    775         X,
    776         X,
    777         X,
    778         X,
    779         X,
    780         X,
    781         X,
    782         X,
    783         X,
    784         X,
    785         X,
    786         X,
    787         X,
    788         X,
    789         X,
    790         X,
    791         X,
    792         X,
    793         X,
    794         X,
    795         X,
    796         X,
    797         X,
    798         X,
    799         X,
    800         X,
    801         X,
    802         X,
    803         X,
    804         X,
    805         X,
    806         X,
    807         X,
    808         X,
    809         X,
    810         X,
    811         X,
    812         X,
    813         X,
    814         X,
    815         X,
    816         X,
    817         X,
    818         X,
    819         X,
    820         X,
    821         X,
    822         X,
    823         X,
    824         X,
    825         X,
    826         X,
    827         X,
    828         X,
    829         X,
    830         X,
    831         X,
    832         X,
    833         X,
    834         X,
    835         X,
    836         X,
    837         X,
    838         X,
    839         X,
    840         X,
    841         X,
    842         X,
    843         X,
    844         X,
    845         X,
    846         X,
    847         X,
    848         X,
    849         X,
    850         X,
    851         X,
    852         X,
    853         X,
    854         X,
    855         X,
    856         X,
    857         X,
    858         X,
    859         X,
    860         X,
    861         X,
    862         X,
    863         X,
    864         X,
    865         X,
    866         X,
    867         X,
    868         X,
    869         X,
    870         X,
    871         X,
    872         X,
    873         X,
    874         X,
    875         X,
    876         X,
    877         X,
    878         X,
    879         X,
    880         X,
    881         X,
    882         X,
    883         X,
    884         X,
    885         X,
    886         X,
    887         X,
    888         X,
    889         X,
    890         X,
    891         X,
    892         X,
    893         X,
    894         X,
    895         X,
    896         X,
    897         X,
    898         X,
    899         X,
    900         X,
    901         X,
    902         X,
    903         X,
    904         X,
    905         X,
    906         X,
    907         X,
    908         X,
    909         X,
    910         X,
    911         X,
    912         X,
    913         X,
    914         X,
    915         X,
    916         X,
    917         X,
    918         X,
    919         X,
    920         X,
    921         X,
    922         X,
    923         X,
    924         X,
    925         X,
    926         X,
    927         X,
    928         X,
    929         X,
    930         X,
    931         X,
    932         X,
    933         X,
    934         X,
    935         X,
    936         X,
    937         X,
    938         X,
    939         X,
    940         X,
    941         X,
    942         X,
    943         X,
    944         X,
    945         X,
    946         X,
    947         X,
    948         X,
    949         X,
    950         X,
    951         X,
    952         X,
    953         X,
    954         X,
    955         X,
    956         X,
    957         X,
    958         X,
    959         X,
    960         X,
    961         X,
    962         X,
    963         X,
    964         X,
    965         X,
    966         X,
    967         X,
    968         X,
    969         X,
    970         X,
    971         X,
    972         X,
    973         X,
    974         X,
    975         X,
    976         X,
    977         X,
    978         X,
    979         X,
    980         X,
    981         X,
    982         X,
    983         X,
    984         X,
    985         X,
    986         X,
    987         X,
    988         X,
    989         X,
    990         X,
    991         X,
    992         X,
    993         X,
    994         X,
    995         X,
    996         X,
    997         X,
    998         X,
    999         X,
    1000        X

```

```

[36mFile ~/Documents/code/Lehre/
IntroductionMachineLearning/.venv/lib/python3.12/site-packages/sklearn/utils/
validation.py:182[39m, in
[36m_assert_all_finite_element_wise[39m[34m(X, xp, allow_nan, msg_dtype,
estimator_name, input_name)[39m
[32m    165[39m [38;5;28;01mif[39;00m estimator_name
[38;5;129;01mand[39;00m input_name ==
[33m"[39m[33mX[39m[33m"[39m [38;5;129;01mand[39;00m
has_nan_error:
[32m    166[39m [38;5;66;03m# Improve the error message on how to
handle missing values in[39;00m
[32m    167[39m [38;5;66;03m# scikit-learn.[39;00m
[32m    168[39m     msg_err += (
[32m    169[39m
[33mf[39m[33m"[39m[38;5;130;01m\n[39;00m[38;5;132;01m{[39;00mest
[39;00m[33m does not accept missing values[39m[33m"[39m
[32m    170[39m [33m"[39m[33m encoded as NaN natively. For
supervised learning, you might want[39m[33m"[39m
[32m    (...) [39m[32m    180[39m [33m"[39m[33m#estimators-
that-handle-nan-values[39m[33m"[39m
[32m    181[39m )
[32m--> [39m[32m182[39m [38;5;28;01mraise[39;00m
[38;5;167;01mValueError[39;00m(msg_err)

```

Hier gibt es allerdings einen Fehler. SciKit beschwert sich, dass der Datensatz NaN enthält. Zum Glück haben wir beim Data Cleaning gelernt, wie wir diese mit `.dropna()` entfernen können. Das müssen wir allerdings für den Eingangsvektor `X` und `Y` gleichzeitig machen, damit sie sich nicht verschieben oder unterschiedlich lang sind. Deshalb wenden wir `.dropna()` auf die Originaldaten an und splitten die Daten danach in `X` und `Y`.

12

```
# Modell trainieren
X = egywthNoNA.TMK.values.reshape(-1, 1)
Y = egywthNoNA.ES_Lab.values
sklm.fit(X, Y)
```

	<code>fit_intercept</code> <code>fit_intercept</code>: bool, default=True Whether to calculate the intercept for this model. If set to False, no intercept will be used in calculations (i.e. data is expected to be centered).	True
	<code>copy_X</code> <code>copy_X</code>: bool, default=True If True, X will be copied; else, it may be overwritten.	True
	<code>tol</code> <code>tol</code>: float, default=1e-6 The precision of the solution (<code>'coef_'</code>) is determined by <code>'tol'</code> which specifies a different convergence criterion for the <code>'lsqr'</code> solver. <code>'tol'</code> is set as <code>'atol'</code> and <code>'btol'</code> of <code>:func:'scipy.sparse.linalg.lsqr'</code> when fitting on sparse training data. This parameter has no effect when fitting on dense data. .. versionadded:: 1.7	1e-06
	<code>n_jobs</code> <code>n_jobs</code>: int, default=None The number of jobs to use for the computation. This will	None

	only provide speedup in case of sufficiently large problems, that is if firstly <code>`n_targets > 1`</code> and secondly <code>`X`</code> is sparse or if <code>`positive`</code> is set to <code>`True`</code>. <code>`None`</code> means 1 unless in a <code>:obj:`joblib.parallel_backend`</code> context. <code>`-1`</code> means using all processors. See :term:`Glossary` for more details.	
	positive positive: bool, default=False When set to <code>`True`</code>, forces the coefficients to be positive. This option is only supported for dense arrays. For a comparison between a linear regression model with positive constraints on the regression coefficients and a linear regression without such constraints, see :ref:`sphx_glr_auto_examples_linear_model_plot_nnls.py`. .. versionadded:: 0.24	False

Wir können die gelernten Parameter uns anschauen mit

```
sklm.intercept_, sklm.coef_
```

```
(np.float64(7.778394959065665), array([3.18532247]))
```

Eine Vorhersage können wir bestimmen mit der Methode `predict` welcher wir die vorherzusagenden X-Matrix übergeben.

```
egywthNoNA["pred"] = sklm.predict(X)
```

```
fig.add_scatter(x=egywthNoNA.TMK, y=egywthNoNA.pred, name="sklearn")  
fig
```

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

Wir erkennen eine Abweichung zwischen der manuell berechneten Funktion und der Funktion von SciKit. Das lässt sich mit den entfernten Werten erklären.

Statsmodels

Alternativ zu SciKit-Learn kann man Statsmodels nutzen. Es fokussiert stärker auf statistisch Modelle, einschließlich *ols*, *ordinary least squares*, d.h. lineare Regression. Die grundlegende Syntax von *ols* ist

```
smf.ols(equation, data)
```

Die Besonderheit ist, dass der Zusammenhang hier mit *equation* als Formel in der Form "*y* ~ *x*" angebar ist und *data* direkt ein Dataframe ist, der die Variablen enthält. Es ist also nicht erforderlich, eine explizite Designmatrix zu erstellen, wie bei *sklearn*. Diese Formeln stammen aus ursprünglich aus der Programmiersprache *R* für maschinelles Lernen, wo sie ähnlich verwendet werden.

In der Formel wird das, was vor der Tilde ~ steht, als abhängige Variable betrachtet (hier *y*), und was danach steht, sind unabhängige Variablen (hier nur *x*). Sofern nicht ausdrücklich anders gefordert, beinhaltet das Modell auch die y-Achsenabschnitt (auch: Schnittpunkt mit der y-Achse) β_0 .

Zum Beispiel kann die Regression der Energiedaten wie folgt durchgeführt werden:

```
import statsmodels.formula.api as smf  
  
smfols = smf.ols("ES_Lab ~ TMK", egywth)
```

Diese Zeile konfiguriert auch wieder nur das lineare Regressionsmodell, lernt aber noch nicht die Parameter. Hierfür muss auch die Methode `.fit()` aufgerufen werden, allerdings ohne Daten und Formel, da die ja schon definiert sind.

```
smflm = smfols.fit() # Der Aufruf gibt das eigentliche Modell zurück
```

Die Parameter lassen sich hier am Attribut `params` auslesen oder mit der Methode `summary()` inklusiver Fehlerstatistiken strukturiert anzeigen. Auf diese gehen wir später noch ein.

```
smflm.params
```

```
Intercept    7.778395  
TMK          3.185322  
dtype: float64
```

```
smflm.summary()
```

Dep. Variable:	ES_Lab	R-squared:	0.381
Model:	OLS	Adj. R-squared:	0.380
Method:	Least Squares	F-statistic:	663.7
Date:	Sat, 06 Jun 2026	Prob (F-statistic):	1.79e-114
Time:	11:23:33	Log-Likelihood:	-5107.2
No. Observations:	1082	AIC:	1.022e+04
Df Residuals:	1080	BIC:	1.023e+04
Df Model:	1		
Covariance Type:	nonrobust		

	coef	std err	t	P> t	[0.025	0.975]
Intercept	7.7784	1.540	5.050	0.000	4.756	10.801
TMK	3.1853	0.124	25.762	0.000	2.943	3.428

Omnibus:	61.376	Durbin-Watson:	0.603
Prob(Omnibus):	0.000	Jarque-Bera (JB):	68.207
Skew:	0.596	Prob(JB):	1.55e-15
Kurtosis:	2.696	Cond. No.	23.3

Eine Vorhersage machen wir hier wieder mit der Methode `predict()`. Hier ist zu beachten, dass die Methode jetzt einen DataFrame in der gleichen Gestalt, wie beim Training erwartet, also einen mit der Spalte „TMK“.

```
pred_x_df = pd.DataFrame({"TMK":pred_x})  
pred_y_smf = smflm.predict(pred_x_df)
```

Wenn wir das Ergebnis darstellen, sehen wir, dass das Modell identisch ist mit dem von SciKit-Learn.

```
fig.add_scatter(x=pred_x, y=pred_y_smf, name="statsmodel")
fig
```

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

Multivariate Lineare Regression

Interessant wird die lineare Regression, wenn man nicht nur Modelle trainiert, die von einer Variable abhängen, sondern von mehreren Eingangsvariablen. Methodisch ist das nur eine Verallgemeinerung des bereits bekannten Modells in dem man die Beziehung zwischen einer abhängigen Zielvariable y und mehreren unabhängigen Eingangsvariablen $x_1, x_2, \dots, x_k$ modelliert. Sie wird häufig in maschinellen Lernanwendungen eingesetzt, um Vorhersagen zu treffen oder Muster in Daten zu erkennen.

Hierfür erweitern wir die ursprüngliche univariate Gleichung der linearen Regression einfach um weitere Anstiege $\beta_1, \dots, \beta_k$ für jede unabhängige Eingangsvariable $x_1, x_2, \dots, x_k$:

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k + \epsilon = \mathbf{X}\beta + \epsilon$$

Um die Parameter $\beta_1, \beta_2, \dots, \beta_k$ zu schätzen, verwenden wir wieder die Methode der kleinsten Quadrate (OLS) und minimiert die Summe der quadrierten Fehler:

$$Q = \sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_{i,1} + \beta_2 x_{i,2} + \dots + \beta_k x_{i,k}))^2$$

wobei n die Anzahl der Datenpunkte ist und $y_i, x_{i,1}, x_{i,2}, \dots, x_{i,k}$ die Werte der Variablen für den i -ten Datenpunkt sind.

Um die Parameter $\beta_0, \dots, \beta_k$ zu schätzen bestimmen wir wieder die Nullstellen der partiellen Ableitungen $\frac{\partial Q}{\partial \beta_0} = 0$ bis $\frac{\partial Q}{\partial \beta_i} = 0$ für $i = 1, 2, \dots, k$.

Die Lösung des Minimierungsproblems ergibt die folgenden Gleichungen für die Regressionskoeffizienten:

$$\beta_0 = \frac{\sum_{i=1}^n (y_i - \bar{y}) \bar{x}_i}{\sum_{i=1}^n \bar{x}_i^2}$$
$$\beta_j = \frac{\sum_{i=1}^n (y_i - \bar{y}) (x_{i,j} - \bar{x}_j)}{\sum_{i=1}^n (x_{i,j} - \bar{x}_j)^2}$$

wobei:

- $\bar{y}$ der Mittelwert der abhängigen Variable ist
- $\bar{x}_j$ der Mittelwert der j -ten unabhängigen Variable ist
- $x_{i,j}$ der Wert der j -ten unabhängigen Variable für den i -ten Datensatzpunkt ist

Das lässt sich auch numerisch durch Matrizeninversion lösen

$$\beta = (X^T \cdot X)^{-1} \cdot (X^T \cdot y)$$

SciKit-Learn

Wir können in SciKit-Learn sehr einfach ein multivariates Modell erstellen, indem wir eine Matrix mit allen Eingangspalten übergeben.

```
sklm_mv = LinearRegression()

egywthNoNA_MV = egywth[["SDK", "NM", "VPM", "TMK",
"ES_Lab"]].dropna().sort_values("TMK") # Drop NA rows

# Modell trainieren
X_MV = egywthNoNA_MV.iloc[:, :-1].values # alles außer letzte Spalte
Y_MV = egywthNoNA_MV.ES_Lab.values
sklm_mv.fit(X_MV, Y_MV)
```

	fit_intercept fit_intercept: bool, default=True Whether to calculate the intercept for this model. If set to False, no intercept will be used in calculations (i.e. data is expected to be centered).	True
	copy_X copy_X: bool, default=True If True, X will be copied; else, it may be overwritten.	True
	tol tol: float, default=1e-6 The precision of the solution (`coef_`) is determined by `tol` which specifies a different convergence criterion for the `lsqr` solver. `tol` is set as `atol` and `btol` of :func:`scipy.sparse.linalg.lsqr` when fitting on sparse training data. This parameter has no ef-	1e-06

	fect when fitting on dense data. .. versionadded:: 1.7	
	<code>n_jobs</code> <code>n_jobs</code>: int, default=None The number of jobs to use for the computation. This will only provide speedup in case of sufficiently large problems, that is if firstly <code>`n_targets > 1`</code> and secondly <code>`X`</code> is sparse or if <code>`positive`</code> is set to <code>`True`</code>. <code>`None`</code> means 1 unless in a <code>:obj: `joblib.parallel_backend`</code> context. <code>`-1`</code> means using all processors. See :term:`Glossary` for more details.	None
	<code>positive</code> <code>positive</code>: bool, default=False When set to <code>`True`</code>, forces the coefficients to be positive. This option is only supported for dense arrays. For a comparison between a linear regression model with positive constraints on the regression coefficients and a linear regression without such constraints, see :ref:`sphx_glr_auto_examples_linear_model_plot_nnls.py`. .. versionadded:: 0.24	False

```
egywthNoNA_MV["pred"] = sklm_mv.predict(X_MV)
fig.add_scatter(x=egywthNoNA_MV.TMK, y=egywthNoNA_MV.pred, name="sklearn MV")
fig
```

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

Wir sehen hier, dass die neue multivariate Vorhersage zumindest in dieser Darstellung des Streudiagramms schlechter scheint. Wir schauen später, ob diese Beobachtung stimmt.

Statsmodels

Auch in Statsmodels können wir einfach ein multivariates Modell erzeugen, indem wir einfach mehr Variablen in der Formel aufnehmen. Im Falle von mehr unabhängigen Variablen können diese mit einem Pluszeichen hinzugefügt werden, zum Beispiel " $y \sim x_1 + x_2 + x_3$ ". Dies entspricht dem Regressionsmodell.

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \epsilon$$

Daher bedeutet das Pluszeichen in der Formel nicht „addieren“, sondern „in das Modell aufnehmen“. Die Formel sollte also so gelesen werden: „Verwende y als abhängige Variable und nimm x_1 , x_2 und x_3 als erklärende Variablen auf.“

Zum Beispiel können wir den Energieverbrauch unter zu Hilfenahme der Sonnenscheindauer (SDK), dem mittleren Bedeckungsgrade (NM) und des Dampfdruckes (VPM) bestimmen:

```
smflm_mv = smf.ols("ES_Lab ~ TMK + SDK + NM + VPM", data=egywth).fit()
```

```
smflm_mv.params
```

```
Intercept    -20.611237
TMK           1.997753
SDK           6.936602
NM            3.690851
VPM          -1.571116
dtype: float64
```

```
smflm_mv.summary()
```

Dep. Variable:	ES_Lab	R-squared:	0.872
Model:	OLS	Adj. R-squared:	0.871
Method:	Least Squares	F-statistic:	1754.
Date:	Sat, 06 Jun 2026	Prob (F-statistic):	0.00

Time:	11:23:33	Log-Likelihood:	-4088.9
No. Observations:	1039	AIC:	8188.
Df Residuals:	1034	BIC:	8212.
Df Model:	4		
Covariance Type:	nonrobust		

	coef	std err	t	P> t	[0.025	0.975]
Intercept	-20.6112	2.932	-7.030	0.000	-26.364	-14.858
TMK	1.9978	0.202	9.884	0.000	1.601	2.394
SDK	6.9366	0.180	38.489	0.000	6.583	7.290
NM	3.6909	0.336	10.986	0.000	3.032	4.350
VPM	-1.5711	0.295	-5.333	0.000	-2.149	-0.993

Omnibus:	72.688	Durbin-Watson:	1.079
Prob(Omnibus):	0.000	Jarque-Bera (JB):	184.271
Skew:	-0.378	Prob(JB):	9.68e-41
Kurtosis:	4.920	Cond. No.	140.

```
egywthNoNA_MV["pred2"] = smflm_mv.predict(egywthNoNA_MV)
fig.add_scatter(x=egywthNoNA_MV.TMK, y=egywthNoNA_MV.pred2, name="statsmodel
MV")
fig
```

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

Umgang mit kategorischen Eingangsvariablen und Kombinationen

Neben der bloßen Spezifizierung der Variablen bieten die Formeln viele weitere Optionen. So können wir auch kategorische Eingangsvariablen einfach zum Modell hinzufügen. Erstens ist zu beachten, dass wenn die Variable kategorisch ist (d.h. ein String), wird sie automatisch in entsprechende Dummy-Variablen umgewandelt. Fügen wir zu unserem Datensatz den Wochentag als kategorische Variable hinzu

```
egywth["Weekday"] = egywth["Date"].dt.day_name()
egywth["Weekday"]
```

0	Wednesday
1	Thursday

```

2      Friday
3      Saturday
4      Sunday
...
1093   Thursday
1094   Friday
1095   Saturday
1096   Sunday
1097   Monday
Name: Weekday, Length: 1098, dtype: str

```

Trainieren wir darauf ein Modell können wir die kategorische Variable direkt mit angeben.

```

smflm_mv_wd = smf.ols("ES_Lab ~ TMK + SDK + NM + VPM + Weekday",
data=egywth).fit()

```

```

smflm_mv_wd.summary()

```

Dep. Variable:	ES_Lab	R-squared:	0.872
Model:	OLS	Adj. R-squared:	0.871
Method:	Least Squares	F-statistic:	699.5
Date:	Sat, 06 Jun 2026	Prob (F-statistic):	0.00
Time:	11:23:33	Log-Likelihood:	-4087.5
No. Observations:	1039	AIC:	8197.
Df Residuals:	1028	BIC:	8251.
Df Model:	10		
Covariance Type:	nonrobust		

	coef	std err	t	P> t	[0.025	0.975]
Intercept	-20.8895	3.053	-6.842	0.000	-26.881	-14.898
Weekday[T.Monday]	0.0295	1.446	0.020	0.984	-2.809	2.868
Weekday[T.Saturday]	-0.7237	1.443	-0.502	0.616	-3.555	2.107
Weekday[T.Sunday]	0.7229	1.449	0.499	0.618	-2.120	3.566
Weekday[T.Thursday]	-0.2152	1.442	-0.149	0.881	-3.045	2.615
Weekday[T.Tuesday]	1.3861	1.442	0.961	0.337	-1.444	4.216
Weekday[T.Wednesday]	0.4161	1.445	0.288	0.773	-2.419	3.251
TMK	2.0007	0.203	9.874	0.000	1.603	2.398

SDK	6.9364	0.181	38.341	0.000	6.581	7.291
NM	3.6994	0.337	10.962	0.000	3.037	4.362
VPM	-1.5740	0.295	-5.328	0.000	-2.154	-0.994

Omnibus:	70.515	Durbin-Watson:	1.075
Prob(Omnibus):	0.000	Jarque-Bera (JB):	176.111
Skew:	-0.370	Prob(JB):	5.73e-39
Kurtosis:	4.877	Cond. No.	156.

Die Zeile `Weekday[T.Monday]` zeigt den Effekt der Dummy-Variable mit dem Wert *Monday*. Darauf folgt *Saturday* da die Kategorien alphabetisch geordnet werden. Daraus folgt ein lineares Modell in der Form

$$y = \beta_0 + \beta_{T.Monday} x_{T.Monday} + \beta_{T.Saturday} x_{T.Saturday} + \dots + \beta_{T.Wednesday} x_{T.Wednesday} + \beta_{TMK} x_{TMK} + \dots + \beta_{VPM} x_{VPM} + \varepsilon$$

mit den binären Dummy-Variablen $x_{T.Monday}$ für den Wochentag. Sie ist nur an Montagen 1 und an allen anderen Wochentagen 0. Das bedeutet, dass für jeden Wochentag ein spezifischer Intersept $\beta_{T.Monday}, \dots, \beta_{T.Sunday}$ im Modell aktiviert wird.

In manchen Fällen sind die Kategorien keine Zeichenketten, sondern haben numerische Labels. Hier kann das Modell diese numerischen Kategorien nicht von numerischen Werten unterscheiden und würde einen linearen Koeffizienten hinzufügen, anstatt eines binären. Zum Beispiel können wir die Wochentage auch numerisch codieren, für das Modell ist aber nicht die Reihenfolge relevant, sondern welcher Wochentag es ist. Dies können wir erzwingen, indem wir die Transformationsfunktion „C(Variable)“ mit angeben.

```
egywth["WeekdayN"] = egywth["Date"].dt.dayofweek
egywth["WeekdayN"]
```

```
0      2
1      3
2      4
3      5
4      6
. . .
1093   3
1094   4
1095   5
1096   6
1097   0
Name: WeekdayN, Length: 1098, dtype: int32
```

```
smflm_mv_wd = smf.ols("ES_Lab ~ TMK + SDK + NM + VPM + C(WeekdayN)",
data=egywth).fit()
```

```
smflm_mv_wd.summary()
```

Dep. Variable:	ES_Lab	R-squared:	0.872
Model:	OLS	Adj. R-squared:	0.871
Method:	Least Squares	F-statistic:	699.5
Date:	Sat, 06 Jun 2026	Prob (F-statistic):	0.00
Time:	11:23:33	Log-Likelihood:	-4087.5
No. Observations:	1039	AIC:	8197.
Df Residuals:	1028	BIC:	8251.
Df Model:	10		
Covariance Type:	nonrobust		

	coef	std err	t	P> t	[0.025	0.975]
Intercept	-20.8600	3.111	-6.705	0.000	-26.965	-14.755
C(WeekdayN)[T.1]	1.3566	1.445	0.939	0.348	-1.479	4.192
C(WeekdayN)[T.2]	0.3866	1.447	0.267	0.789	-2.452	3.225
C(WeekdayN)[T.3]	-0.2447	1.443	-0.170	0.865	-3.077	2.588
C(WeekdayN)[T.4]	-0.0295	1.446	-0.020	0.984	-2.868	2.809
C(WeekdayN)[T.5]	-0.7533	1.444	-0.522	0.602	-3.586	2.080
C(WeekdayN)[T.6]	0.6934	1.448	0.479	0.632	-2.148	3.535
TMK	2.0007	0.203	9.874	0.000	1.603	2.398
SDK	6.9364	0.181	38.341	0.000	6.581	7.291
NM	3.6994	0.337	10.962	0.000	3.037	4.362
VPM	-1.5740	0.295	-5.328	0.000	-2.154	-0.994

Omnibus:	70.515	Durbin-Watson:	1.075
Prob(Omnibus):	0.000	Jarque-Bera (JB):	176.111
Skew:	-0.370	Prob(JB):	5.73e-39
Kurtosis:	4.877	Cond. No.	161.

Hierbei lassen sich auch Kombinationen betrachten. Zum Beispiel mag der Energieverbrauch nicht nur von den Wochentagen abhängen, sondern nutzt auch an Heiztagen und Kühltagen unterschiedliche Betriebsstrategien. Um das auszudrücken kann man kategorische Variablen mit * kombinieren, was in diesem Fall keine Multiplikationsoperator ist, sondern für die Kombinationen steht.

```
smflm_mv_wd2 = smf.ols("ES_Lab ~ TMK + SDK + NM + VPM +  
C(WeekdayN)*HeizKuehlTage", data=egywth).fit()
```

```
smflm_mv_wd2.summary()
```

Dep. Variable:	ES_Lab	R-squared:	0.876
Model:	OLS	Adj. R-squared:	0.873
Method:	Least Squares	F-statistic:	298.8
Date:	Sat, 06 Jun 2026	Prob (F-statistic):	0.00
Time:	11:23:33	Log-Likelihood:	-4070.0
No. Observations:	1039	AIC:	8190.
Df Residuals:	1014	BIC:	8314.
Df Model:	24		
Covariance Type:	nonrobust		

	coef	std err	t	P> t	[0.025	0.975]
Intercept	-18.3083	3.315	-5.523	0.000	-24.813	-11.803
C(WeekdayN)[T.1]	2.1577	1.754	1.230	0.219	-1.284	5.600
C(WeekdayN)[T.2]	1.0729	1.769	0.606	0.544	-2.398	4.544
C(WeekdayN)[T.3]	0.3539	1.752	0.202	0.840	-3.084	3.791
C(WeekdayN)[T.4]	0.2426	1.761	0.138	0.890	-3.213	3.698
C(WeekdayN)[T.5]	-0.3289	1.747	-0.188	0.851	-3.758	3.100
C(WeekdayN)[T.6]	-0.3372	1.755	-0.192	0.848	-3.781	3.106
HeizKuehlTage[T.Kühlgradtag]	3.9698	5.915	0.671	0.502	-7.637	15.577
HeizKuehlTage[T.Normaltag]	5.2019	2.512	2.071	0.039	0.272	10.132
C(WeekdayN)[T.1]:HeizKuehlTa- ge[T.Kühlgradtag]	-3.4876	9.169	-0.380	0.704	-21.480	14.505
C(WeekdayN)[T.2]:HeizKuehlTa- ge[T.Kühlgradtag]	-18.8997	9.195	-2.055	0.040	-36.943	-0.856

C(WeekdayN)[T.3]:HeizKuehlTa- ge[T.Kühlgradtag]	-5.9965	9.175	-0.654	0.514	-24.001	12.008
C(WeekdayN)[T.4]:HeizKuehlTa- ge[T.Kühlgradtag]	-15.5864	8.474	-1.839	0.066	-32.214	1.042
C(WeekdayN)[T.5]:HeizKuehlTa- ge[T.Kühlgradtag]	-4.4747	7.990	-0.560	0.576	-20.154	11.204
C(WeekdayN)[T.6]:HeizKuehlTa- ge[T.Kühlgradtag]	-0.5844	7.663	-0.076	0.939	-15.621	14.452
C(WeekdayN)[T.1]:HeizKuehlTa- ge[T.Normaltag]	-2.1827	3.128	-0.698	0.486	-8.322	3.956
C(WeekdayN)[T.2]:HeizKuehlTa- ge[T.Normaltag]	-1.1488	3.096	-0.371	0.711	-7.225	4.927
C(WeekdayN)[T.3]:HeizKuehlTa- ge[T.Normaltag]	-1.3698	3.121	-0.439	0.661	-7.494	4.755
C(WeekdayN)[T.4]:HeizKuehlTa- ge[T.Normaltag]	0.6019	3.118	0.193	0.847	-5.516	6.720
C(WeekdayN)[T.5]:HeizKuehlTa- ge[T.Normaltag]	-0.4470	3.157	-0.142	0.887	-6.643	5.749
C(WeekdayN)[T.6]:HeizKuehlTa- ge[T.Normaltag]	4.1547	3.175	1.309	0.191	-2.075	10.385
TMK	1.9744	0.204	9.677	0.000	1.574	2.375
SDK	6.8849	0.182	37.827	0.000	6.528	7.242
NM	3.6449	0.336	10.851	0.000	2.986	4.304
VPM	-1.9031	0.309	-6.158	0.000	-2.510	-1.297

Omnibus:	64.597	Durbin-Watson:	1.101
Prob(Omnibus):	0.000	Jarque-Bera (JB):	184.425
Skew:	-0.282	Prob(JB):	8.97e-41
Kurtosis:	4.985	Cond. No.	777.

Modellqualität

Residuen

Ein wichtiger Aspekt im Maschinellen Lernen ist die Bewertung der Qualität von Modellen, um auf dieser Basis zu entscheiden, ob die Fehler im Rahmen sind und welche Modelle geeignet sind. Die Grundlage für diese Bewertung sind meist die *Residuen* oder der Vorhersagefehler, die wir bereits zuvor kennen gelernt haben.

Betrachten wir einmal die Residuen für unser Beispiel von dem univariaten und multivariaten Linearen Modellen. Wir berechnen es, indem wir die Vorhersage vom Zielwert subtrahieren.

```
egywth["Pred_UV"]=smflm.predict(egywth)
egywth["Residum_UV"]= egywth.ES_Lab - egywth.Pred_UV
```

Wie bereits diskutiert, basiert die Annahme des Linearen Modells darauf, dass dieser Fehler Normalverteilt ist. Mit einem Histogramm können wir uns die Verteilung anschauen.

```
px.histogram(egywth.Residum_UV)
```

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

Die Verteilung des Residuums ist nicht Normalverteilt, sondern linkslastig. Das ist bereits ein Indiz, dass das Model vielleicht nicht optimal ist.

Schauen wir uns statt dem Streudiagramm, die Vorhersage in einem Liniendiagramm an.

```
px.line(egywth, x="Date", y=["ES_Lab", "Pred_UV"])
```

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

Jetzt sehen wir, dass das ursprüngliche lineare Modell im Streudiagramm recht gut aussah, weil es den mittleren Trend gut erfasst hat. In der Zeitreihe offenbart sich allerdings eine recht klare Abweichung vom Original.

Vergleichen wir das mit dem multivariaten Modell

```
egywth["Pred_MV"]=smflm_mv_wd2.predict(egywth)
egywth["Residum_MV"]= egywth.ES_Lab - egywth.Pred_MV
```

```
px.histogram(egywth.Residum_MV)
```

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

```
px.line(egywth, x="Date", y=["ES_Lab", "Pred_MV"])
```

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

Hier sehen wir, dass die Verteilung des Residuums eher Normalverteilt ist und die Grundform der Zeitreihe gut erfasst wird. Dennoch sind deutliche Abweichungen sichtbar.

Mean Square Error (MSE)

Um diesen Vergleich von Modellen professionell verwendet man unterschiedliche Kennwerte. Der erste Kennwert ist der Mean Square Error, auf Deutsch auch die mittlere quadratische Abweichung genannt, ist eine wichtige Kenngröße hierfür.

Der MSE berechnet im Wesentlichen den durchschnittlichen quadratischen Fehler zwischen den tatsächlichen Werten y_i und den von einem Modell vorhergesagten Werten $\hat{y}_i$.

$$\text{MSE} = \frac{1}{N} \sum_i e_i^2$$

Ein kleinerer MSE-Wert deutet auf ein besseres Modell hin. Vergleichen wir die Werte für unsere beiden Modelle

```
egywth.Residum_UV.pow(2).mean()
```

```
np.float64(736.8486474358979)
```

```
egywth.Residum_MV.pow(2).mean()
```

```
np.float64(147.90321019931775)
```

und wir sehen, dass der MSE für das zweite Modell besser ist.

Der MSE ist in der Einheit des quadrierten Zielwertes. Daher kann die Interpretation der absoluten Größe des MSE manchmal schwierig sein. Der MSE ist empfindlich gegenüber Ausreißern in den Daten, die den Wert verfälschen können.

Root Mean Squared Error (RMSE)

Die RMSE ist die Quadratwurzel des MSE und hat somit die Einheit wie der Zielwert, was die Interpretation einfacher macht. Gleichzeitig können bestehende Interpretationen des MSE übernommen werden.

$$\text{RMSE} = \sqrt{\text{MSE}}$$

```
np.sqrt(egywth.Residum_UV.pow(2).mean())
```

```
np.float64(27.144956206188617)
```

```
np.sqrt(egywth.Residum_MV.pow(2).mean())
```

```
np.float64(12.16154637368611)
```

Mean Absolute Error (MAE)

Der MAE misst die durchschnittliche absolute Abweichung zwischen den vorhergesagten und den tatsächlichen Werten.

$$\text{MAE} = \frac{1}{N} \sum_i |e_i|$$

Der MAE ist weniger empfindlich gegenüber Ausreißern als der MSE. Die Einheit des MAE ist die gleiche wie die Einheit des Zielwertes, was die Interpretation einfacher macht.

```
egywth.Residum_UV.abs().mean()
```

```
np.float64(21.892016248549634)
```

```
egywth.Residum_MV.abs().mean()
```

```
np.float64(9.110498926081684)
```

Mean Absolute Percentage Error (MAPE)

Der MAPE misst den durchschnittlichen prozentualen absoluten Fehler. Er ist zu interpretieren, ohne den originalen Wertebereich zu kennen. Ein niedriger MAPE-Wert deutet auf eine gute Übereinstimmung zwischen den vorhergesagten und den tatsächlichen Werten hin.

$$\text{MAPE} = \frac{100}{N} \sum_i \left| \frac{e_i}{y_i} \right|$$

```
egywth.Residum_UV.abs().mean()*100
```

```
np.float64(2189.2016248549635)
```

```
egywth.Residum_MV.abs().mean()*100
```

```
np.float64(911.0498926081684)
```

Bestimmungskoeffizient R^2

Der Bestimmungskoeffizient R^2 gibt den Anteil der durch das Regressionsmodell erklärten Variabilität an der gesamten Variabilität in den Daten wieder. Bei gut passenden Modellen liegt er oft nahe bei 1. Ein Wert von 0 bedeutet, dass das Modell auf diesen Daten nicht besser ist als die Vorhersage mit dem Mittelwert. In manchen Situationen kann R^2 auch negativ werden, zum Beispiel auf Testdaten oder bei schlecht passenden Modellen.

Er bestimmt sich aus der Sum of Squared Error (SSE) und der Total Sum of Squares (TSS) zu

$$R^2 = 1 - \frac{\text{SSE}}{\text{TSS}}$$

Die Gesamtsumme der Quadrate (TSS) misst die Summe der quadrierten Abweichungen der tatsächlichen Werte (y_i) vom Mittelwert ($\bar{y}$) aller Datenpunkte. Sie repräsentiert die Gesamtvariabilität in den Daten. Er wird berechnet durch

$$\text{TSS} = \sum_{i=1}^N (y_i - \bar{y})^2.$$

Die Summe der quadratischen Abweichungen (SSE) misst die Summe der quadrierten Abweichungen zwischen den vorhergesagten Werten $\hat{y}_i$ und den tatsächlichen Werten y_i . Sie dient zur Beurteilung der Genauigkeit des Modells durch:

$$\text{SSE} = \sum_{i=1}^N \varepsilon_i^2 = \sum_{i=1}^N (y_i - \hat{y}_i)^2.$$

Wir können das manuell berechnen

```
TSS = (egywth.ES_Lab - egywth.ES_Lab.mean()).pow(2).sum()
SSE = egywth.Residum_UV.pow(2).sum()
R2 = 1 - SSE/TSS
R2
```

```
np.float64(0.38061649401642605)
```

```
TSS = (egywth.ES_Lab - egywth.ES_Lab.mean()).pow(2).sum()
SSE = egywth.Residum_MV.pow(2).sum()
R2 = 1 - SSE/TSS
R2
```

```
np.float64(0.8806156958265965)
```

Oder direkt das von statsmodel bereitgestellte rsquared-Attribut nutzen

```
smflm.rsquared
```

```
np.float64(0.38061649401642617)
```

```
smflm_mv_wd2.rsquared
```

```
np.float64(0.876122441552206)
```

Hier ist wichtig, dass das Modell besser ist, wenn der R^2 Wert größer ist und nicht kleiner, wie bei den anderen Fehlerwerten.

Bei *sklearn* können wir dafür die Methode `score` nutzen.

```
sklm.score(X, Y)
```

```
0.38061649401642617
```

```
sklm_mv.score(X_MV, Y_MV)
```

```
0.8715365639979806
```

Erneut erhalten wir genau die gleiche Zahl wie bei der Verwendung von *statsmodels* oben, sofern wir jeweils dasselbe Modell vergleichen. Beachten Sie, dass die Methode `.score` zwei Argumente erwartet: die Designmatrix X und das Ziel y , in dieser Reihenfolge. (Es ist die gleiche Reihenfolge wie für die Methode `.fit`.) Intern berechnet sie die Werte für jede Zeile von X , berechnet Fehler zwischen den Vorhersagen und y und transformiert dies in R^2 . Aber indem wir X und y der Methode `.score` übergeben, können wir R^2 auf anderen Daten berechnen - auf Validierungsdaten anstelle von Trainingsdaten (siehe Abschnitt [@ref\(overfitting-validation\)](#) für weitere Informationen).

Aber *sklearn* hat keine Funktion, um eine ähnliche Zusammenfassungstabelle wie `.summary` in *Statsmodels* auszudrucken. Wir können Standardfehler, p -Werte oder statistische Signifikanz nicht so einfach sehen. *sklearn* ist für die Vorhersagemodellierung konzipiert, wo solche Tabellen von geringem Wert sind. In komplexen Fällen, zum Beispiel in der Bildverarbeitung, möchten wir möglicherweise Werte von Millionen von Pixeln in das Modell einbeziehen, und hier ist die individuelle Spezifikation jedes Pixels klarerweise nicht machbar. Wir bevorzugen den Ansatz der Designmatrix. Auf ähnliche Weise sind wir nicht daran interessiert, Werte von Millionen von Koeffizienten zu interpretieren. *sklearn* ist genau für solche Aufgaben konzipiert.

Bibliographie