

LOGISTISCHE REGRESSION

Joern Ploennigs · AI4SC

MODELL GRUNDLAGEN

Wie bei der linearen Regression beginnt die logistische Regression mit einer linearen Kombination der Eingabevariablen (Prädiktoren). Diese lineare Kombination kann als z_i dargestellt werden:

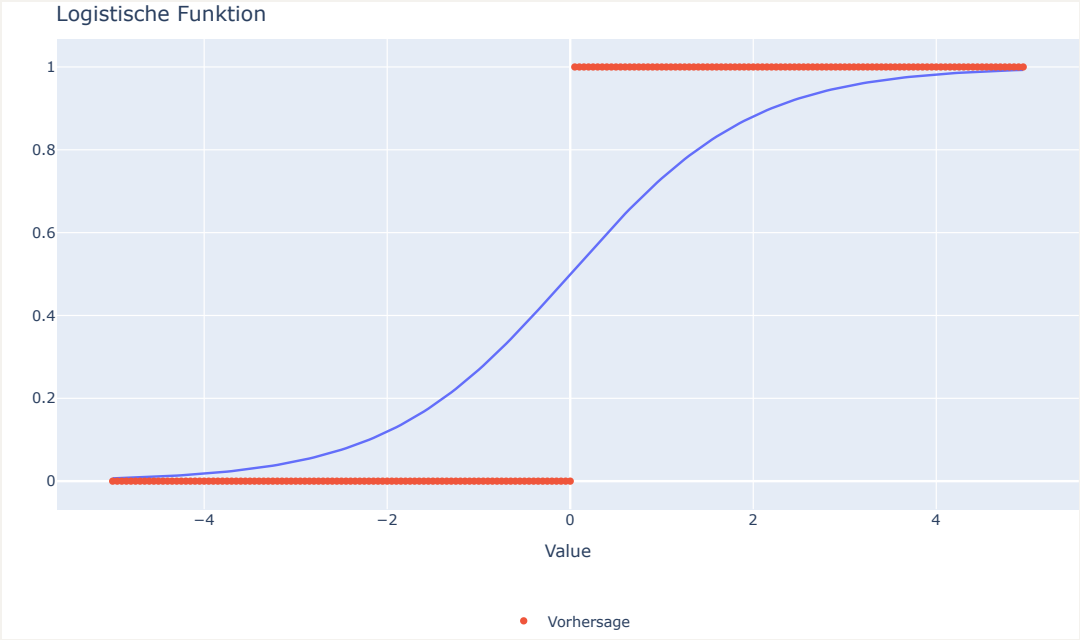
$$z_i = \beta_0 + \beta_1 \cdot x_{1,i} + \beta_2 \cdot x_{2,i} + \dots + \beta_n \cdot x_{n,i}$$

Statt die lineare Kombination direkt als Ergebnis zu nehmen, wird sie durch die logistische Funktion (auch Sigmoid-Funktion genannt) transformiert. Die logistische Funktion ist definiert als:

$$P(y_i = 1 \mid \vec{x}_i) = \sigma_i = \frac{1}{1 + e^{-z_i}} = \frac{1}{1 + e^{-\beta_0 - \beta_1 \cdot x_{1,i} - \beta_2 \cdot x_{2,i} - \dots - \beta_n \cdot x_{n,i}}}$$

Diese Funktion nimmt beliebige reelle Zahlen und transformiert sie in den Bereich (0, 1), was als Wahrscheinlichkeit interpretiert werden kann. Die Ausgabe der logistischen Funktion σ_i gibt die Wahrscheinlichkeit an, dass das Ergebnis 1 ist (z. B. Kunde kauft das Produkt) gegeben die Eingabewerte. Die Wahrscheinlichkeit, dass das Ergebnis 0 ist, ist einfach $P(y_i = 0 \mid \bar{x}_i) = 1 - \sigma_i$.

Die lineare Kombination (z.B. aus Temperatur, Alter und Druck) wird durch die Sigmoid-Funktion in eine Wahrscheinlichkeit zwischen 0 und 1 umgewandelt: z.B. 0.84 (84 %) Wahrscheinlichkeit für Ereignis 'Ja'.



Die Regressionskoeffizienten $\beta_0, \beta_1, \dots, \beta_n$ geben an, wie sich eine Einheitsänderung der jeweiligen unabhängigen Variablen $x_1, \dots, x_n$ auf die Logarithmus der Chancen (Odds Ratio) des Ereignisses $y=1$ auswirkt.

- Ein positiver Regressionskoeffizient β_i bedeutet, dass eine Zunahme der Variablen x_i um eine Einheit die Chancen des Ereignisses $y_i = 1$ erhöht.
- Ein negativer Regressionskoeffizient β_i bedeutet, dass eine Zunahme der Variablen x_i um eine Einheit die Chancen des Ereignisses $y_i = 1$ verringert.
- Ein Wert von 0 des Regressionskoeffizienten β_i bedeutet, dass die Variable x_i keinen Einfluss auf die Wahrscheinlichkeit des Ereignisses $y_i = 1$ hat.

Die Interpretation der Regressionskoeffizienten in Bezug auf die Wahrscheinlichkeit des Ereignisses $y_i = 1$ erfordert die Berechnung der Chancen und deren anschließende Interpretation.

PARAMETERANPASSUNG

Die Anpassung des logistischen Regressionsmodells an Daten erfolgt typischerweise durch *Maximum-Likelihood-Schätzung (MLE)*. Dies bedeutet, dass die Koeffizienten $\bar{\beta}$ so gewählt werden, dass die Wahrscheinlichkeit der beobachteten Daten maximiert wird.

Die Likelihood-Funktion für eine Menge von n Trainingsbeispielen x_i, y_i ist:

$$L(\bar{\beta}) = \prod_{i=1}^n P(y_i \mid \bar{x}_i) = \prod_{i=1}^n \sigma_i^{y_i} (1 - \sigma_i)^{1-y_i} = \prod_{i=1}^n \begin{cases} \sigma_i & y_i = 1, \\ 1 - \sigma_i & y_i = 0. \end{cases}$$

Da das Produkt schwerer zu berechnen ist, nutzt man lieber die Log-Likelihood. Dafür berechnen wir den Logarithmus der Likelihood-Funktion:

$$\log L(\bar{\beta}) = \log \left(\prod_{i=1}^n P(y_i | \tilde{x}_i) \right) = \sum_{i=1}^n \left(y_i \log(\sigma_i) + (1 - y_i) \log(1 - \sigma_i) \right).$$

Um die Koeffizienten $\bar{\beta}$ zu bestimmen, berechnet man wieder die Nullstellen der partiellen Ableitung nach $\beta_0, \beta_1, \dots, \beta_n$.

$$0 = \frac{\partial \ell}{\partial \beta_0} = \sum_{k=1}^n (y_k - \sigma_k)$$

$$0 = \frac{\partial \ell}{\partial \beta_i} = \sum_{k=1}^n (y_k - \sigma_k) x_{k,i}$$

und löst das resultierende Gleichungssystem mit numerischen Methoden. Das ist notwendig, da die Log-Likelihood-Funktion nicht linear ist (sie enthält die e -Funktion). Deshalb wird zur Maximierung oft ein iteratives Verfahren wie das Newtonverfahren (Newton-Raphson-Algorithmus) verwendet. Dieser Algorithmus bestimmt die Parameter $\bar{\beta}$ iterativ:

$$\bar{\beta}_{t+1} = \bar{\beta}_t - H(\bar{\beta}_t)^{-1} \Delta \log L(\bar{\beta}_t),$$

bis eine hinreichende Genauigkeit erzielt wird. Hier ist $\Delta \log L(\bar{\beta}_l)$ der Gradient (Vektor der partiellen Ableitungen) und $H(\bar{\beta}_l)$ die Hessian-Matrix (Matrix der zweiten Ableitungen) der Log-Likelihood-Funktion in Bezug auf die Parameter $\bar{\beta}$ mit

$$\Delta \log L(\bar{\beta}_l) = \sum_{k=1}^n (y_k - \sigma_k) \bar{x}_k$$

$$H(\bar{\beta}_l) = - \sum_{k=1}^n \sigma_k (1 - \sigma_k) \bar{x}_k \bar{x}_k^T$$

LOGISTIC REGRESSION IN PYTHON

Wie bei der linearen Regression können wir sowohl die Bibliotheken *statsmodels* als auch *sklearn* auch für die logistische Regression verwenden. Die Verwendung ist ähnlich wie bei der linearen Regression, aber beide Bibliotheken haben ihre eigenen Eigenheiten.

Als Beispiel nutzen wir wieder den Energiedatensatz der Universität Rostock. Wir wollen hier vorhersagen, wann die Anlagen in “HT 740” aktiviert sind, also Energie verbrauchen.

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

egywth = pd.read_csv("../data/UR05/Energy1D_weather_clean.csv", parse_dates=[0])
egywth["Weekday"] = egywth["Date"].dt.day_name()
egywth["EV_HT_740_IS_ON"] = egywth.EV_HT_740==0
egywth.EV_HT_740_IS_ON.value_counts()
```

```
EV_HT_740_IS_ON
False      553
True       545
Name: count, dtype: int64
```

LOGISTISCHE REGRESSION IN PYTHON MIT DER STATSMODELS FORMEL API

Bei der Verwendung von *statsmodels* müssen wir eine ähnliche Formel wie bei der linearen Regression angeben. Bevor wir das jedoch tun können, müssen wir eine Eigenheit in statsmodels ansprechen - nämlich, dass es nicht erlaubt ist, dass das Ergebnis eine logische oder Zeichenfolgenvariable ist. Es *muss* eine 0/1 (numerische) Variable sein, ansonsten tritt ein Fehler auf:

ValueError: wthegy hat sich zu einem Array mit mehreren Spalten ausgewertet, das die Form (2675, 2) hat. Dies tritt auf, wenn die Variable, die in wthegy umgewandelt wurde, nicht numerisch ist (z. B. bool oder str).

Um also ein Modell zu erstellen, bei dem wir die Wahrscheinlichkeit der Behandlung schätzen, müssen wir zuerst eine neue Variable erstellen, die nur die numerische Version von `EV_HT_740_IS_ON` ist. Hierfür wandeln wir sie in ein Integer.

```
egywth["EV_HT_740_ON"] = egywth.EV_HT_740_IS_ON.astype(int)
egywth.EV_HT_740_ON.value_counts()
```

```
EV_HT_740_ON
0      553
1      545
Name: count, dtype: int64
```

Wir sehen, dass beide Werte relativ gleich verteilt sind. Mit dieser numerischen Variable können wir jetzt ein logistische Regressionsmodell erstellen. Hierfür geben wir wie beim linearen Modell eine Formel der Eingangsparameter an, die kontinuierlich sein können als auch kategorisch wie der Wochentag

```
import statsmodels.formula.api as smf
smm0 = smf.logit("EV_HT_740_ON ~ TMK + SDK + NM + VPM + Weekday", data=egywth)
smm = smm0.fit()
```

Optimization terminated successfully.
Current function value: 0.594367
Iterations 8

Dieser Code ähnelt sehr der linearen Regression, der einzige Unterschied besteht darin, dass `smf.logit` anstelle von `smf.ols` verwendet wird. Genau wie bei der linearen Regression verwendet die logistische Regression R-Style-Formeln. Daher ist das spezifische Modell, das wir verwenden
$$P(y_i = 1 \mid \mathbf{X}) = \frac{1}{1 + e^{-\mathbf{z}}} \text{ wobei}$$
$$\mathbf{z} = \beta_0 + \beta_{\text{T.Monday}} \mathbf{x}_{\text{T.Monday}} + \beta_{\text{T.Saturday}} \mathbf{x}_{\text{T.Saturday}} + \dots + \beta_{\text{T.Wednesday}} \mathbf{x}_{\text{T.Wednesday}} + \beta_{\text{TMK}} \mathbf{x}_{\text{TMK}} + \dots + \beta_{\text{VPM}} \mathbf{x}_{\text{VPM}}.$$

Wir können die geschätzten Werte mit anzeigen.

smm.params		
Intercept	0.123141	
Weekday[T.Monday]	0.054098	
Weekday[T.Saturday]	0.087200	
Weekday[T.Sunday]	4.196551	
Weekday[T.Thursday]	-0.015990	
Weekday[T.Tuesday]	-0.085271	
Weekday[T.Wednesday]	-0.058213	
TMK	-0.000715	
SDK	-0.043613	
NM	-0.001606	
VPM	-0.014591	
dtype: float64		

Wir können auch eine vollständige Zusammenfassung der Ausgabe erhalten.

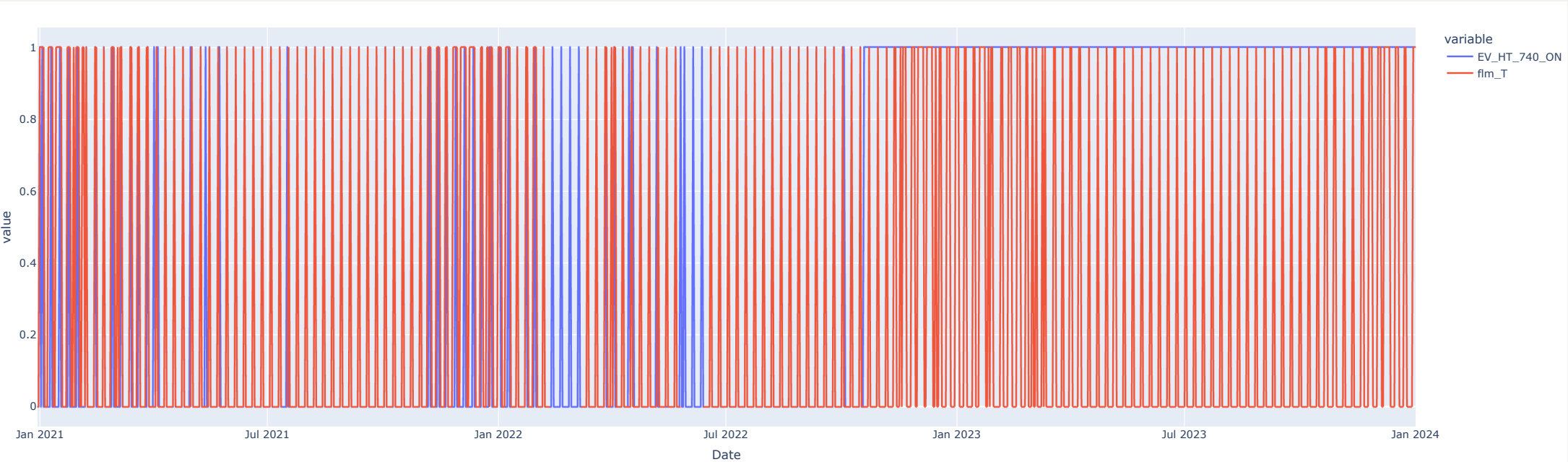
```
smm.summary()
```

Logit Regression Results

Dep. Variable:	EV_HT_740_ON		No. Observations:		1055	
Model:	Logit		Df Residuals:		1044	
Method:	MLE		Df Model:		10	
Date:	Sat, 06 Jun 2026		Pseudo R-squ.:		0.1423	
Time:	11:23:43		Log-Likelihood:		-627.06	
converged:	True		LL-Null:		-731.06	
Covariance Type:	nonrobust		LLR p-value:		3.439e-39	
	coef	std err	z	P> z	[0.025	0.975]
Intercept	0.1231	0.528	0.233	0.816	-0.913	1.159
Weekday[T.Monday]	0.0541	0.234	0.231	0.817	-0.405	0.513
Weekday[T.Saturday]	0.0872	0.234	0.373	0.709	-0.371	0.546
Weekday[T.Sunday]	4.1966	0.607	6.913	0.000	3.007	5.386
Weekday[T.Thursday]	-0.0160	0.234	-0.068	0.946	-0.475	0.443
Weekday[T.Tuesday]	-0.0853	0.235	-0.363	0.716	-0.545	0.375
Weekday[T.Wednesday]	-0.0582	0.235	-0.248	0.804	-0.518	0.402
TMK	-0.0007	0.035	-0.020	0.984	-0.070	0.069
SDK	-0.0436	0.031	-1.385	0.166	-0.105	0.018
NM	-0.0016	0.059	-0.027	0.978	-0.117	0.113

Die Ausgabe ähnelt der linearen Regression, weist jedoch bestimmte Kennzahlen für die Modellqualität von logistischer Regression.

```
egywth["flm_P"] = smm.predict(egywth)
egywth["flm_T"] = egywth.flm_P.map(lambda x: 0 if x < 0.5 else 1, na_action='ignore')
px.line(egywth, x="Date", y=["EV_HT_740_ON", "flm_T"])
```



In dem Linien-Diagramm sieht man die Übereinstimmung nicht wirklich gut. Für binäre Werte eignen sich eher Heatmaps.

```
px.imshow(egywth[["EV_HT_740_ON", "flm_T"]].T, x=egywth["Date"], color_continuous_scale='Viridis')
```



Hier zeigt sich eine gute Übereinstimmung bis Mitte 2022 aber dann falsche Vorhersagen für den Zeitraum danach.

LOGISTISCHE REGRESSION MIT SCIKIT-LEARN

Die Durchführung einer logistischen Analyse mit *sklearn* ähnelt in vieler Hinsicht der Durchführung einer linearen Regression in *sklearn*. Die Tatsache, dass wir den gleichen Ansatz mit logistischer Regression verwenden können wie im Fall der linearen Regression, ist ein großer Vorteil von *sklearn*: Der gleiche Ansatz gilt auch für andere Modelle, daher ist es sehr einfach, mit verschiedenen Modellen zu experimentieren. Sehr wenig zusätzlicher Code und keine zusätzliche Datenverarbeitung sind erforderlich.

Zuerst importieren und richten wir das Modell ein:

```
# we need to import the model
from sklearn.linear_model import LogisticRegression
# create the model
sklm = LogisticRegression()
```

Wir können verschiedene Optionen für `LogisticRegression` angeben, aber im Moment sind wir mit den Standardwerten zufrieden.

Als nächstes erstellen wir die Designmatrix, wir müssen nur die “age”-Variable extrahieren und sicherstellen, dass das Ergebnis ein Datenrahmen (oder eine Matrix) ist:

```
egywthNoNA = egywth[["SDK", "NM", "VPM", "TMK", "EV_HT_740_ON"]].dropna() # Drop NA rows

# Modell trainieren
X = egywthNoNA.iloc[:, :-1].values # alles außer letzte Spalte
Y = egywthNoNA.EV_HT_740_ON.values
sklm.fit(X, Y)
```

▼ LogisticRegression ⓘ ?

▸ Parameters

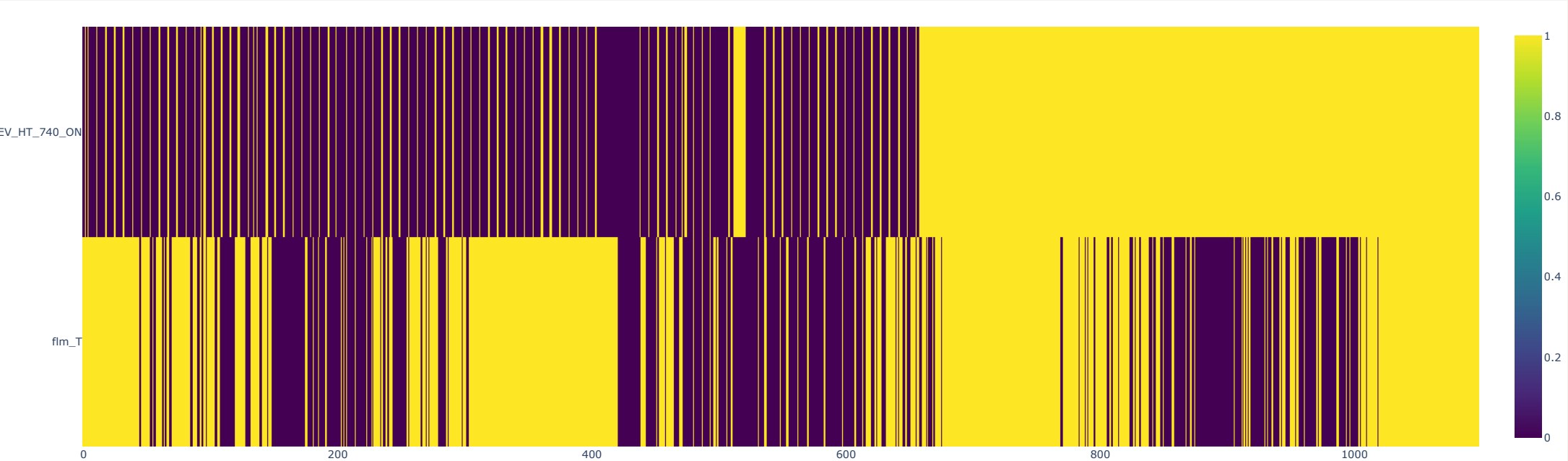
Wir können die gelernten Parameter wieder anschauen mit

```
sklm.intercept_, sklm.coef_  
  
(array([0.26382612]),  
 array([[ -0.0285415 ,  0.01756992, -0.01759732,  0.00113189]]))
```

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

```
egywthNoNA["flm_P"] = sklm.predict_proba(X)[: , 1]
egywthNoNA["flm_T"] = egywthNoNA.flm_P.map(lambda x: 0 if x < 0.5 else 1, na_action='ignore')

px.imshow(egywthNoNA[["EV_HT_740_ON", "flm_T"]].T, x=egywthNoNA.index, color_continuous_scale='Viridis')
```



Hier zeigt sich, dass das Modell im hinteren Bereich nach 2022 besser ist, aber infolgedessen auch mehr Fehler im vorderen Bereich macht.

MODELLQUALITÄT

GENAUIGKEIT

Eine der einfachsten Kenngrößen für die Modellqualität logistischer Regressionsmodelle, bzw. bei der Vorhersage kategorischer oder binärer Variablen ist die Genauigkeit der Vorhersagen, also in wie vielen Fällen, die dem Zielwert überstimmen. Der Wert wird oft relativ über der Gesamtzahl angegeben und kann einfach berechnet werden:

<pre>(egywth.flm_T==egywth.EV_HT_740_ON).sum() / egywth.EV_HT_740_ON.count()</pre>
<pre>np.float64(0.6056466302367942)</pre>
<pre>(egywthNoNA.flm_T==egywthNoNA.EV_HT_740_ON).sum() / egywthNoNA.EV_HT_740_ON.count()</pre>
<pre>np.float64(0.5440758293838862)</pre>

Im Falle der logistischen Regression gibt die `score` -Methode die Genauigkeit (Prozentsatz der korrekten Vorhersagen) anstatt von R^2 aus.

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

```
0.5440758293838862
```

Dabei zeigt sich, dass das Statsmodell hier eine höhere Genauigkeit erreicht. Dieser Vergleich ist allerdings nur eingeschränkt aussagekräftig, da beide Modelle unterschiedliche Eingangsvariablen verwenden. Jetzt klingt die Genauigkeit von 54% Prozent erstmal “ok”. Vergleichen wir die Genauigkeit allerdings mit einer reinen Zufallsreihe. Da die Originalwerte ja recht gleich verteilt waren, können wir eine binomiale Gleichverteilung nutzen. Hier ergibt sich eine Genauigkeit von ca. 50%. Das heißt das SKlearn-Modell ist eigentlich nur in 4% der Fälle besser als eine reine Zufallswahl.

```
(np.random.randint(2, size=egywth.EV_HT_740_ON.count())==egywth.EV_HT_740_ON).sum() / egywth.EV_HT_740_ON.count()

np.float64(0.5136612021857924)
```

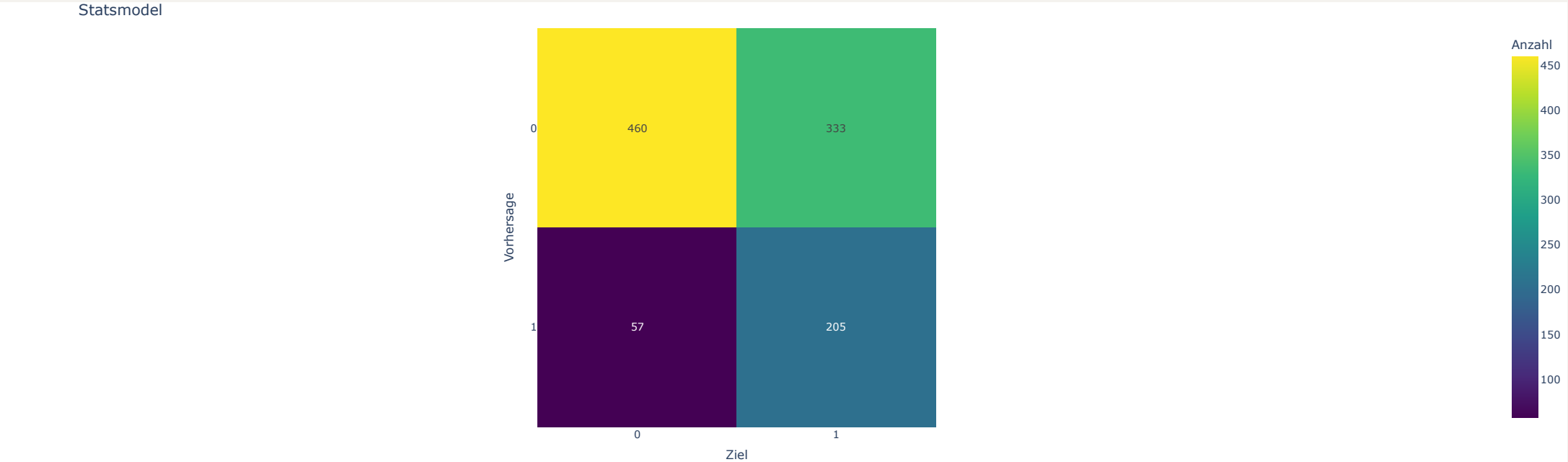
Konfusionsmatrix

Eine Konfusionsmatrix ist ein wichtiges Werkzeug zur Bewertung der Leistung von Modellen im Bereich des maschinellen Lernens zur Vorhersage von kategorischen Variablen. Sie vergleichen die vorhergesagten Werte gegen die originalen Werte.

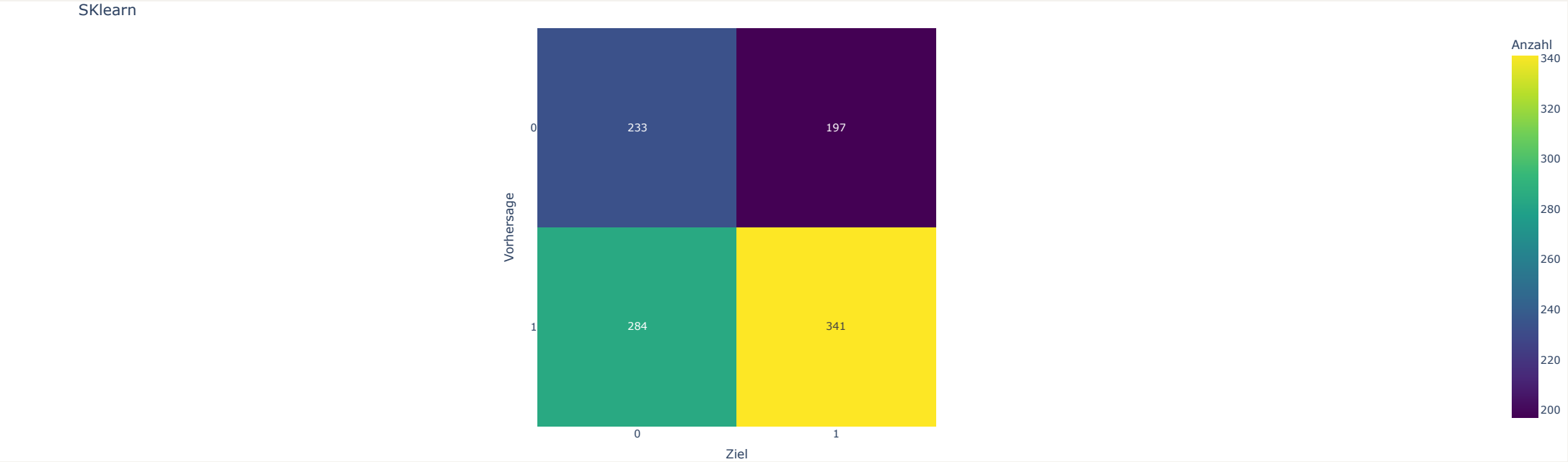
```
confusion_matrix=egywth[['EV_HT_740_ON', 'flm_T']].value_counts().reset_index().sort_values(['EV_HT_740_ON','flm_T'])
confusion_matrix
```

	EV_HT_740_ON	flm_T	count
0	0	0.0	460
3	0	1.0	57
1	1	0.0	333
2	1	1.0	205

```
px.imshow(confusion_matrix["count"].values.reshape(2, 2).T, x=["0","1"], y=["0","1"], color_continuous_scale='Viridis', text_auto=True, labels=dict(x="Ziel", y="Vorhersage", color="Anzahl"), title="Statsmodel")
```



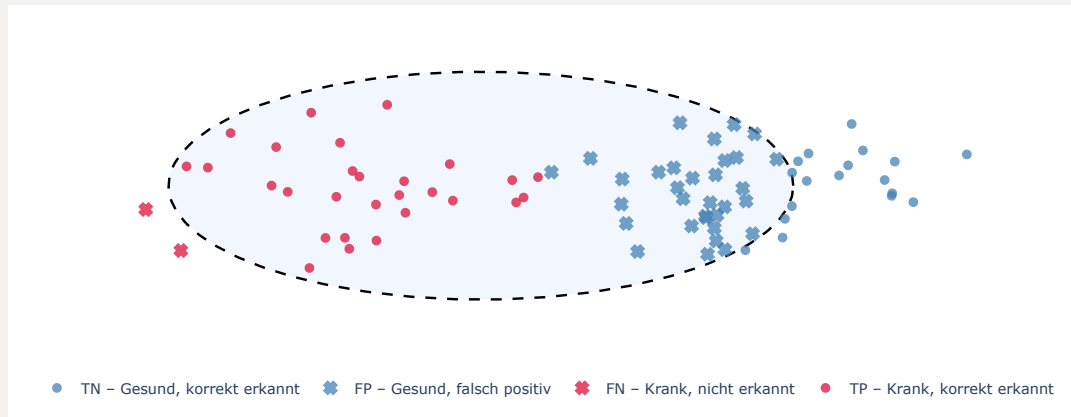
```
from sklearn.metrics import confusion_matrix
px.imshow(confusion_matrix(egywthNoNA.EV_HT_740_ON, egywthNoNA.flm_T).T, x=["0","1"], y=["0","1"], color_continuous_scale='Viridis', text_auto=True, labels=dict(x="Ziel", y="Vorhersage", color="Anzahl"), title="SKlearn")
```



WAHRE UND FALSCH VORHERSAGEN

Die Diskussion der Konfusionsmatrix hat bereits gezeigt, dass die Genauigkeit nicht hinreichend ist, um die Modellqualität gut zu bewerten, da nicht nur wichtig ist ob etwas häufig richtig ist, sondern ob das Modell systematische Fehler macht, also zum Beispiel nur eine häufige Kategorie korrekt vorhersagt, aber bei seltenen Kategorien meist falsch liegt und somit zur Vorhersage dieser Kategorien ungeeignet ist. Deshalb unterscheidet man bei binären Problemen die Vorhersagen detailliert in:

- *True Positiv*: Korrekt als Ja klassifiziert,
- *True Negativ*: Korrekt als Nein klassifiziert,
- *False Positiv*: Falsch als Ja klassifiziert und
- *False Negativ*: Falsch als Nein klassifiziert.



GENAUIGKEIT UND PRÄZISION

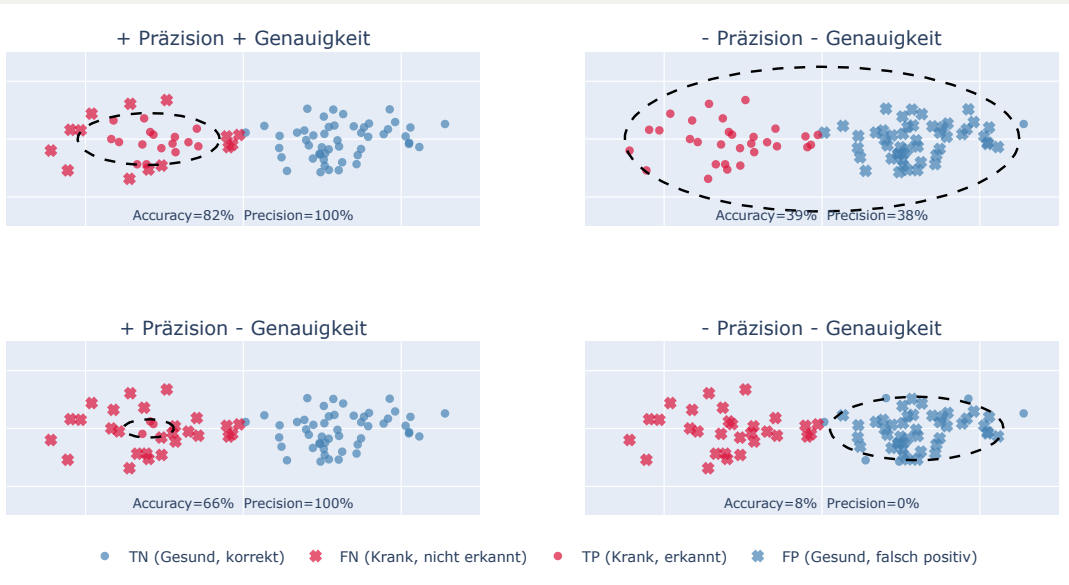
Hieraus berechnet man verschiedene Kenngrößen:

- *Genauigkeit (Accuracy)*: Die Genauigkeit ist das Verhältnis der korrekt vorhergesagten Instanzen zur Gesamtanzahl der Instanzen:

$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$

- *Präzision (Precision)*: Die Präzision gibt an, wie viele der als positiv vorhergesagten Instanzen tatsächlich positiv sind:

$Precision = \frac{TP}{TP+FP}$



SENSITIVITÄT, SPEZIFITÄT UND FALSCHALARMRATE

- *Sensitivität (Recall) oder Trefferquote (True Positive Rate, TPR)*: Die Sensitivität misst den Anteil der tatsächlich positiven Instanzen, die korrekt als positiv vorhergesagt wurden:

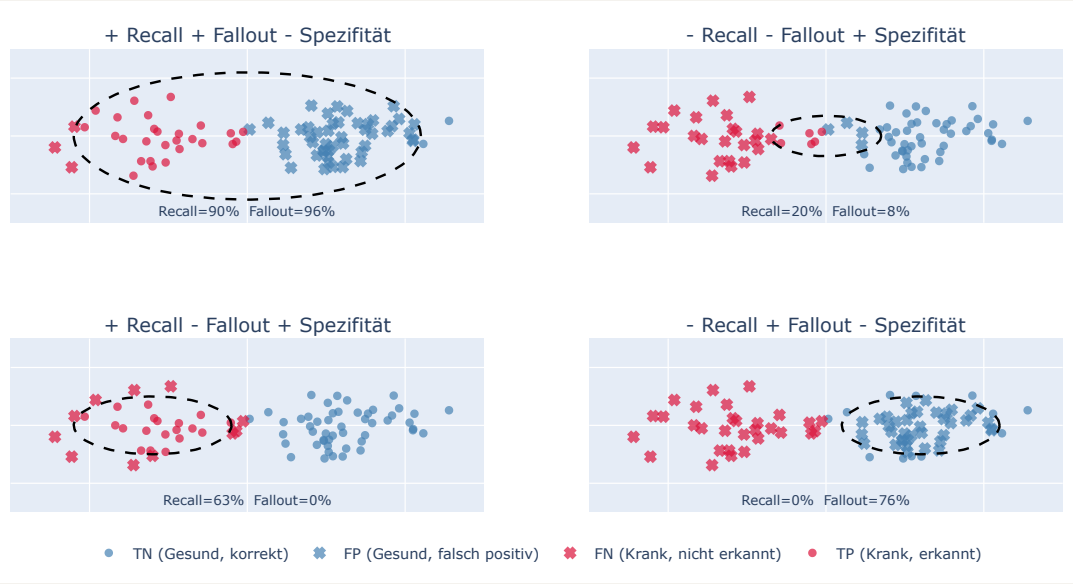
$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}$$

- *Spezifität (Specificity) oder True Negative Rate (TNR)*: Die Spezifität misst den Anteil der tatsächlich negativen Instanzen, die korrekt als negativ vorhergesagt wurden:

$$\text{Specificity} = \frac{\text{TN}}{\text{TN} + \text{FP}}$$

- *Falschalarmrate (Fallout) oder False Positive Rate (FPR)*: Die Falschalarmrate bestimmt den Anteil der falschen als Positiv vorhergesagten Werte:

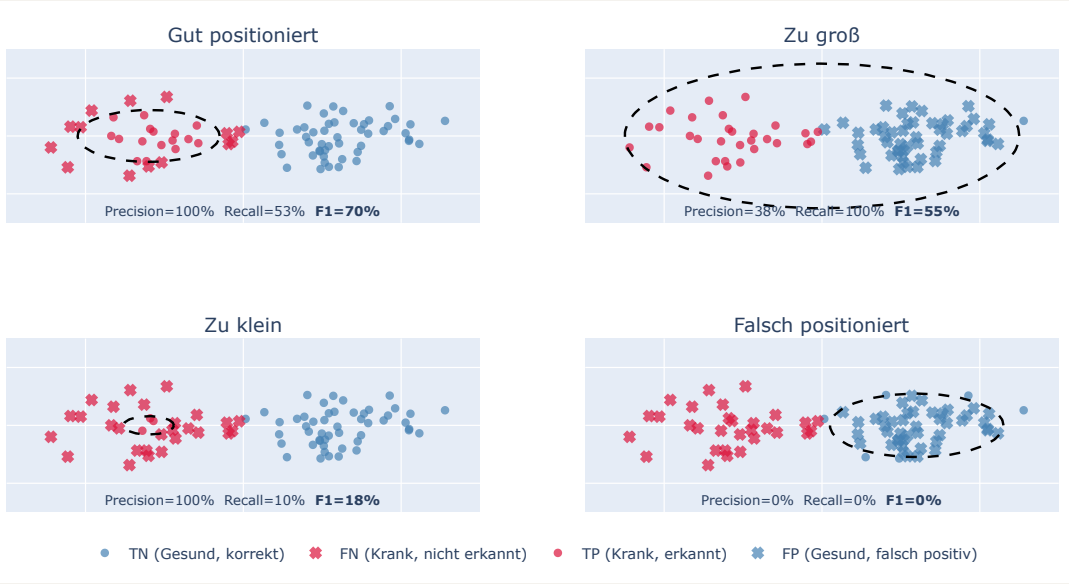
$$\text{Fallout} = \frac{\text{FP}}{\text{TN} + \text{FP}}$$



F1-SCORE

- *F1-Score*: Der F1-Score ist das harmonische Mittel von Präzision und Sensitivität und gibt eine ausgewogene Messung der Modellleistung:

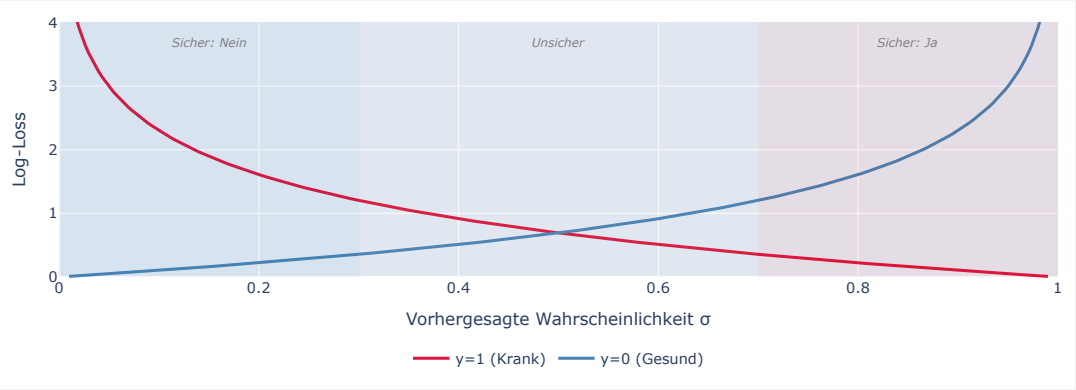
$$F1 = 2 * \frac{\text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall}} = \frac{2 \text{ TP}}{2 \text{ TP} + \text{FP} + \text{FN}}$$



LOG-LOSS ODER BINÄRE KREUZENTROPIE

- *Log-Loss oder Binäre Kreuzentropie*: Log-Loss misst die Leistung eines Klassifikationsmodells, dessen Vorhersagen Wahrscheinlichkeitswerte zwischen 0 und 1 sind. Es bestraft falsche Klassifikationen, insbesondere wenn die Wahrscheinlichkeiten stark danebenliegen:

$$\text{LogLoss} = -\frac{1}{n} \sum_{i=1}^n \left(y_i \log(\sigma_i) + (1-y_i) \log(1-\sigma_i) \right).$$



Wir können die Kenngrößen mit SKLearn bestimmen:

```
from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score

print("Statsmodels")
sm = egywth[["EV_HT_740_ON", "flm_T"]].dropna()
print('Accuracy: {:.2f}'.format(accuracy_score(sm.EV_HT_740_ON, sm.flm_T)))
print('Precision: {:.2f}'.format(precision_score(sm.EV_HT_740_ON, sm.flm_T)))
print('Recall: {:.2f}'.format(recall_score(sm.EV_HT_740_ON, sm.flm_T)))
print('F1-Score: {:.2f}'.format(f1_score(sm.EV_HT_740_ON, sm.flm_T)))

Statsmodels
Accuracy: 0.63
Precision: 0.78
Recall: 0.38
F1-Score: 0.51

print("SKlearn")
print('Accuracy: {:.2f}'.format(accuracy_score(egywthNoNA.EV_HT_740_ON, egywthNoNA.flm_T)))
print('Precision: {:.2f}'.format(precision_score(egywthNoNA.EV_HT_740_ON, egywthNoNA.flm_T)))
print('Recall: {:.2f}'.format(recall_score(egywthNoNA.EV_HT_740_ON, egywthNoNA.flm_T)))
print('F1-Score: {:.2f}'.format(f1_score(egywthNoNA.EV_HT_740_ON, egywthNoNA.flm_T)))

SKlearn
Accuracy: 0.54
Precision: 0.55
Recall: 0.63
F1-Score: 0.59
```

Hier sieht man, dass das sklearn-Modell zwar eine gewisse Präzision erreicht, aber Recall und F1-Score insgesamt noch nicht überzeugend sind.

FINALES MODEL

Um das Problem mit unserem Modell zu beheben, können wir das Jahr mit in unser Logistisches Modell integrieren.

```
egywth["Year"] = egywth["Date"].dt.year

smmY = smf.logit("EV_HT_740_ON ~ TMK + SDK + NM + VPM + Weekday + Year", data=egywth)
smmY = smmY.fit()
```

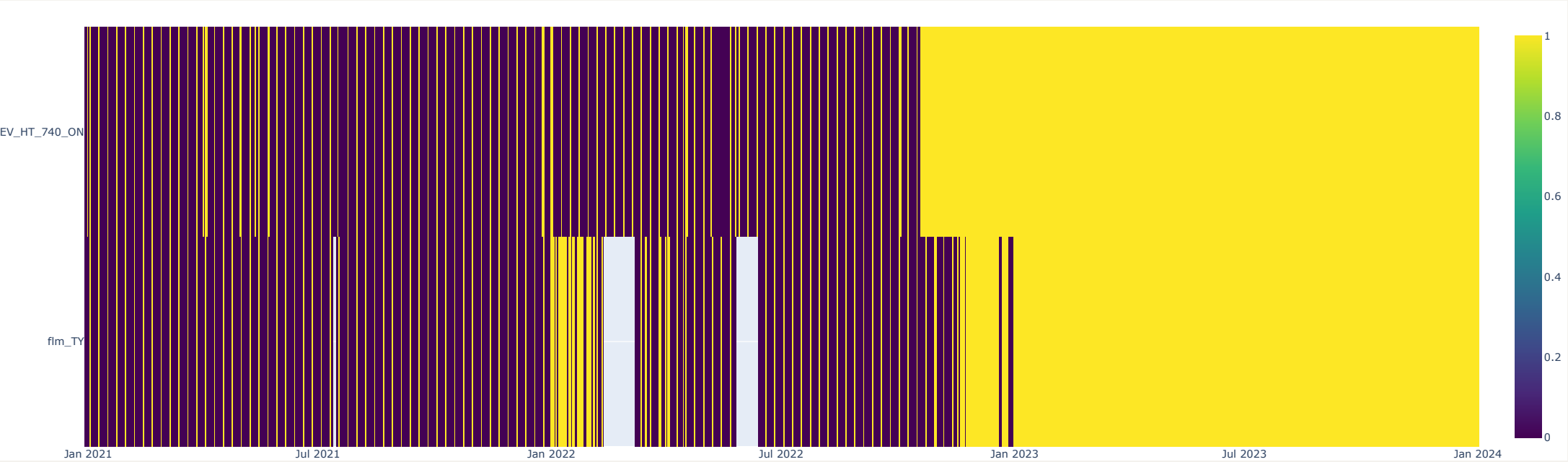
```
Optimization terminated successfully.
      Current function value: 0.187412
      Iterations 15
```

smmY.summary()

Logit Regression Results

Dep. Variable:	EV_HT_740_ON			No. Observations:	1055	
Model:	Logit			Df Residuals:	1043	
Method:	MLE			Df Model:	11	
Date:	Sat, 06 Jun 2026			Pseudo R-squ.:	0.7295	
Time:	11:23:43			Log-Likelihood:	-197.72	
converged:	True			LL-Null:	-731.06	
Covariance Type:	nonrobust			LLR p-value:	8.490e-222	
	coef	std err	z	P> z	[0.025	0.975]
Intercept	-1.033e+04	779.169	-13.257	0.000	-1.19e+04	-8802.049
Weekday[T.Monday]	0.1781	0.453	0.393	0.694	-0.709	1.065
Weekday[T.Saturday]	0.3216	0.443	0.726	0.468	-0.546	1.190
Weekday[T.Sunday]	9.1750	0.849	10.804	0.000	7.511	10.839
Weekday[T.Thursday]	0.0110	0.455	0.024	0.981	-0.881	0.903
Weekday[T.Tuesday]	-0.3040	0.459	-0.662	0.508	-1.203	0.595
Weekday[T.Wednesday]	-0.2047	0.457	-0.448	0.654	-1.100	0.690
TMK	-0.1358	0.069	-1.972	0.049	-0.271	-0.001
SDK	-0.1333	0.062	-2.132	0.033	-0.256	-0.011
NM	0.0118	0.108	0.109	0.914	-0.200	0.224
VPM	0.0924	0.096	0.967	0.334	-0.095	0.280

```
egywth["flm_PY"] = smmY.predict(egywth)
egywth["flm_TY"] = egywth.flm_PY.map(lambda x: 0 if x < 0.5 else 1, na_action='ignore')
px.imshow(egywth[["EV_HT_740_ON", "flm_TY"]].T, x=egywth["Date"], color_continuous_scale='Viridis')
```



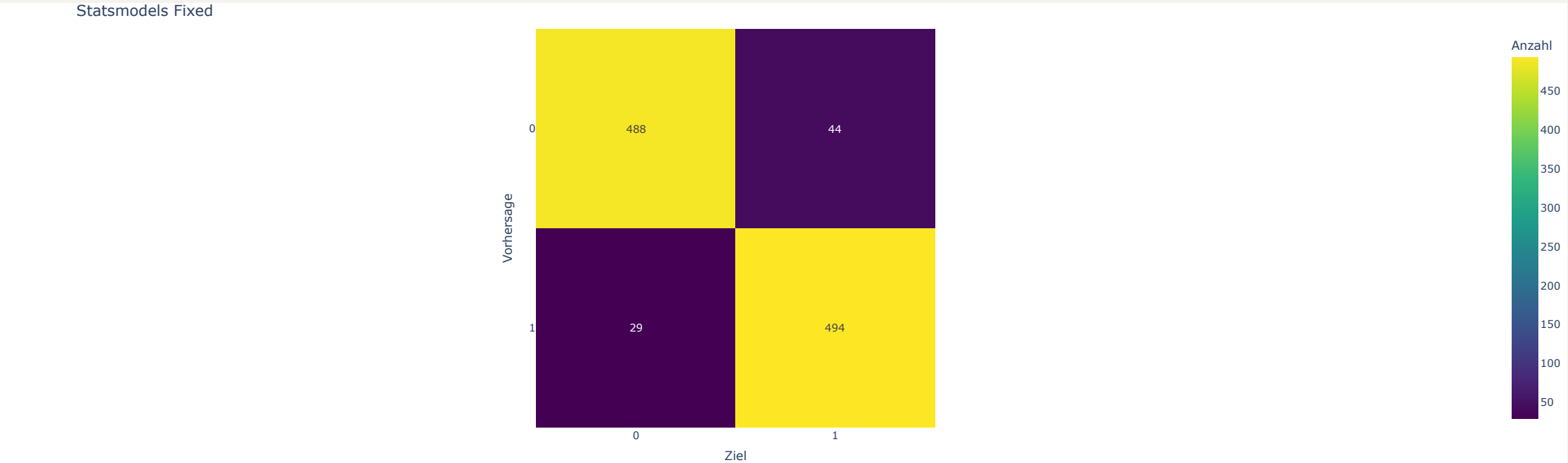
Mit deutlich besserer Genauigkeit, Präzision und F1-Score.

```
print("Statsmodels Fixed")
sm = egwth[["EV_HT_740_ON", "flm_TY"]].dropna()
print('Accuracy: {:.2f}'.format(accuracy_score(sm.EV_HT_740_ON, sm.flm_TY)))
print('Precision: {:.2f}'.format(precision_score(sm.EV_HT_740_ON, sm.flm_TY)))
print('Recall: {:.2f}'.format(recall_score(sm.EV_HT_740_ON, sm.flm_TY)))
print('F1-Score: {:.2f}'.format(f1_score(sm.EV_HT_740_ON, sm.flm_TY)))
```

```
Statsmodels Fixed
Accuracy: 0.93
Precision: 0.94
Recall: 0.92
F1-Score: 0.93
```

und besserer Konfusionsmatrix.

```
px.imshow(confusion_matrix(sm.EV_HT_740_ON, sm.flm_TY).T, x=["0","1"], y=["0","1"], color_continuous_scale='Viridis', text_auto=True, labels=dict(x="Ziel", y="Vorhersage", color="Anzahl"), title="Statsmodels Fixed")
```



QUESTIONS