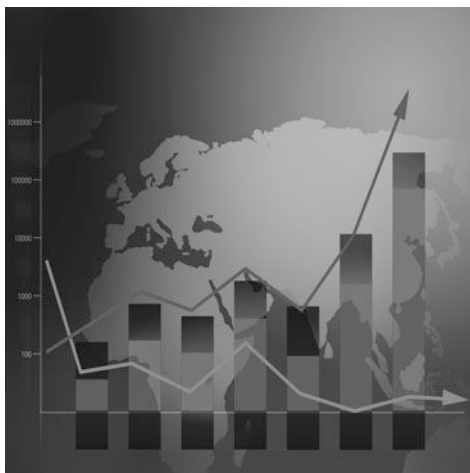




Peter Hackl

Einführung in die Ökonometrie

2., aktualisierte Auflage



Peter Hackl

Einführung in die Ökonometrie

2., aktualisierte Auflage

PEARSON

Higher Education

München • Harlow • Amsterdam • Madrid • Boston
San Francisco • Don Mills • Mexico City • Sydney

a part of Pearson plc worldwide

Die Konsumfunktion $C = \alpha + \beta Y + u$ unterstellt, dass die Ausgaben für Konsum (C) eine Funktion des verfügbaren Einkommens (Y) sind. Das Ergebnis der OLS-Anpassung an österreichische Daten für den Zeitraum 1954 bis 1999 ergibt die Beziehung $\hat{C}_t = 1.948 + 0.897 Y_t$. Die Abbildung 9.1 zeigt den Verlauf der rekursiv geschätzten marginalen Konsumneigung b_t .

Wir sehen, dass am Beginn des Beobachtungszeitraums die marginale Konsumneigung mit etwa 0.84 um einiges geringer war, als es die globale Schätzung ergibt. Im Lauf der Zeit wurde die marginale Konsumneigung ständig, aber nicht ohne Unterbrechungen größer. Zu beachten ist auch die unterschiedliche Breite der Konfidenzintervalle.

9.2 Indikator- oder Dummy-Variable

Einen Regressor, der für eine einzige Beobachtung ($\mathbf{x}'_t, Y_t$) den Wert Eins, für alle anderen Beobachtungen den Wert Null hat, nennen wir eine Indikator-, Schein- oder Dummy-Variable, im Englischen *dummy variable*. Mit einer solchen Dummy-Variable können wir in einem Modell berücksichtigen, dass der datengenerierende Prozess der zu beschreibenden Daten durch qualitative Charakteristika in unterschiedlicher Weise geprägt ist. Wollen wir beispielsweise in einem Modell berücksichtigen, dass ein Teil der Beobachtungen aus einer Wachstumsperiode, andere aus einer Periode der Stagnation stammen, so können wir dem Modell eine Dummy-Variable "Konjunktur-Indikator" hinzufügen, die in Jahren des Wachstums den Wert Eins und sonst den Wert Null hat.

In dieser Weise können wir in einem Modell qualitative Charakteristika berücksichtigen wie

- Konjunktur/Stagnation,
- Zeit vor/nach dem Ölpreisschock,
- Regionen wie Stadt/Land,
- Saisonen des Jahres.

Die datengenerierenden Prozesse in Bereichen, die verschiedenen Ausprägungen dieser Charakteristika entsprechen, auch „Regime“ genannt, unterscheiden sich oft nicht in der Spezifikation des Modells, mit dem diese Prozesse beschrieben werden können, sondern nur in den Parametern, also im Wert des Interzepts oder eines oder mehrerer der anderen Regressionskoeffizienten, also in der Modellstruktur.

Beispiel 9.2

Dummy-Variable für Saisonen

Für die Saisonen des Jahres definieren wir die Dummy-Variablen

$$Q_{it} = \begin{cases} 1 & i\text{-tes Quartal,} \\ 0 & \text{sonst.} \end{cases}$$

Das Frühlings-Dummy Q_{1t} hat für alle t den Wert Eins im ersten Quartal, in den übrigen Quartalen den Wert Null. Analog sind das Sommer-Dummy ($i = 2$), das Herbst- ($i = 3$) und das Winter-Dummy ($i = 4$) definiert. Die Saison-Dummies erfüllen in allen Zeitpunkten $t = 1, \dots, n$ die Bedingung

$$Q_{1t} + Q_{2t} + Q_{3t} + Q_{4t} = 1;$$

jeweils nur eines von ihnen hat den Wert Eins, die anderen haben den Wert Null.

Beispiel 9.3 Modelle für Quartalsdaten

Ausgangspunkt ist das Modell $Y_t = \alpha + \beta X_t + u_t$, das saisonale Effekte unberücksichtigt lässt. Je nachdem, welche Charakteristika wir im Modell berücksichtigen wollen, erweitern wir das Modell um die entsprechenden Parameter. Beispielsweise schreiben wir das Modell mit seasonspezifischem Interzept und Anstieg in folgender Form:

$$Y_t = \alpha_1 + \beta_1 X_t + u_t,$$

$$Y_t = \alpha_2 + \beta_2 X_t + u_t,$$

$$Y_t = \alpha_3 + \beta_3 X_t + u_t,$$

$$Y_t = \alpha_4 + \beta_4 X_t + u_t.$$

Die erste Modellgleichung wird verwendet für Daten aus dem ersten Quartal, die zweite für Daten aus dem zweiten Quartal etc. Aus der Schreibweise der einzelnen Gleichungen geht allerdings nicht hervor, für welches Quartal die jeweilige Modellgleichung verwendet werden soll.

Mit Hilfe der in Beispiel 9.2 definierten Saison-Dummies können wir die vier Modellgleichungen als eine Gleichung schreiben. Sie lautet

$$Y_t = \sum_i \alpha_i Q_{it} + \sum_i \beta_i Q_{it} X_t + u_t$$

oder

$$Y_t = \alpha_1 + \delta_2 Q_{2t} + \delta_3 Q_{3t} + \delta_4 Q_{4t} + \beta_1 X_t + \gamma_2 Q_{2t} X_t + \gamma_3 Q_{3t} X_t + \gamma_4 Q_{4t} X_t + u_t, \quad (9.2.1)$$

mit $\delta_i = \alpha_i - \alpha_1$ und $\gamma_i = \beta_i - \beta_1$, $i = 2, 3, 4$; bei der Umformung haben wir von der Beziehung $Q_{1t} + Q_{2t} + Q_{3t} + Q_{4t} = 1$ Gebrauch gemacht. Der Koeffizient α_1 steht in (9.2.1) für das Interzept. Die Parameter δ_i , $i = 2, 3, 4$ sind die Koeffizienten der Dummies für Sommer, Herbst und Winter; die Dummy-Variable Q_1 für das Frühlingsquartal kommt im Modell nicht vor, da die Summe der vier Quartals-Dummies den Wert Eins hat. Enthielte das Modell vier Quartals-Dummies, so wäre die Summe der entsprechenden vier Spalten aus der Matrix $\mathbf{X}$ gleich der dem Interzept entsprechenden ersten Spalte und der Rang von $\mathbf{X}$ wäre nicht voll!

Ein Modell mit seasonspezifischem Interzept und gemeinsamem Anstieg schreiben wir als $Y_t = \alpha_1 + \delta_2 Q_{2t} + \delta_3 Q_{3t} + \delta_4 Q_{4t} + \beta X_t + u_t$; analog ist $Y_t = \alpha + \beta_1 X_t + \gamma_2 Q_{2t} X_t + \gamma_3 Q_{3t} X_t + \gamma_4 Q_{4t} X_t + u_t$ ein Modell mit gemeinsamem Interzept und seasonspezifischem Anstieg.

Die Anwendung der Dummy-Variablen zum Berücksichtigen von Saisoneffekten kann in analoger Weise auf andere Charakteristika wie Unterschiede zwischen Regionen, Zeiträumen etc. übertragen werden. Insbesondere können wir sie verwenden, um Modelle zu erweitern, so dass Unterschiede in der Modellstruktur berücksichtigt werden. Wir werden im Weiteren sehen, wie wir Dummy-Variablen einsetzen, um einen Strukturbruch, also eine Änderung in der Struktur des datengenerierenden Prozesses, darzustellen und zu identifizieren.

9.3 Analyse von Strukturbrüchen

Von einem Strukturbruch sprechen wir, wenn es Teilbereiche des Beobachtungszeitraums gibt, für die der datengenerierende Prozess durch das gleiche Modell beschrieben werden kann, wobei in den Teilbereichen aber unterschiedliche Werte einiger oder aller Regressionskoeffizienten verwendet werden müssen. Die Teilbereiche oder Regime entsprechen also unterschiedlichen Strukturen. Der Chow-Test erlaubt uns zu entscheiden, ob tatsächlich unterschiedliche Strukturen vermutet werden müssen oder nicht. Der Chow-Test setzt voraus,

- dass Teilbereiche mit konstanter Struktur identifiziert werden können,
- dass uns der Zeitpunkt bekannt ist, zu dem der Übergang zwischen den Regimen stattgefunden hat, und
- dass aus jedem Regime eine ausreichende Anzahl von Beobachtungen zur Verfügung steht, so dass wir das Modell an die Daten jedes einzelnen Regimes anpassen und die Residuen bestimmen können.

Wir werden Dummy-Variablen benutzen, um unterschiedliche Strukturen innerhalb eines Modells zu spezifizieren.

9.3.1 Chow-Test und Strukturbruch

Den Chow-Test verwenden wir, wenn wir die Vermutung überprüfen wollen, dass der datengenerierende Prozess in mehreren Regimen abläuft und das Modell hinsichtlich seiner Koeffizienten regimespezifisch angepasst werden muss. Der Chow-Test ist typischerweise ein Test auf Strukturbruch und beantwortet die Frage nach der Konstanz der Regressionskoeffizienten in der Zeit. Dabei wird die Nullhypothese geprüft, dass die Regressionskoeffizienten in allen Teilbereichen des Beobachtungszeitraums die gleichen sind. Die Alternative ist, dass ab einem bestimmten Zeitpunkt oder zu bestimmten Zeitpunkten das Interzept und einige oder auch alle anderen Regressionskoeffizienten ihren Wert ändern.

Beispiel 9.4 Konsumfunktion, Fortsetzung

Der zeitliche Verlauf der rekursiv geschätzten marginalen Konsumneigung, der in Abbildung 9.1 dargestellt ist, zeigt, dass die Annahme einer konstanten Modellstruktur vermutlich nicht realistisch ist. Nach einer Periode relativer Konstanz beginnt die marginale Konsumneigung ab 1973 zu steigen. Diese Phase dauert bis etwa 1978, wo dieses Wachstum abebbt. Eine zweite Phase steigender marginaler Konsumneigung beginnt etwa mit 1981, eine dritte ab 1992. Vermutungen darüber, mit welchen Ereignissen diese Phasen zusammenhängen, würde vielleicht die Analyse der Preispolitik der Erdölwirtschaft oder der Europäischen Integration bringen. So könnte der Ölpreisschock der frühen 70er Jahre die erste Phase steigender marginaler Konsumneigung ausgelöst haben. Wir definieren vier Regime, nämlich die Perioden 1954 bis 1972, 1973 bis 1980, 1981 bis 1991 und 1992 bis 1994. Mit dem Chow-Test werden wir die Nullhypothese testen, dass die Regressionskoeffizienten für alle vier Regime die gleichen sind. Ende 1972, Ende 1980 und Ende 1992 könnten jene Zeitpunkte sein, in denen Strukturbrüche stattgefunden haben.

Das Vorgehen des Chow-Tests ist am einfachsten zu erklären, wenn wir nur zwei Regime vermuten. Eine Spezifikation des Modells, das die entsprechenden Änderungen der Regressionskoeffizienten berücksichtigt, ist

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} X_1 & 0 \\ 0 & X_2 \end{pmatrix} \begin{pmatrix} \beta_1 \\ \beta_2 \end{pmatrix} + \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}.$$

Die Partitionierung von y , X und u zerlegt in Beobachtungen vor und nach dem Strukturbruch; analog sind β_1 und β_2 die Regressionskoeffizienten vor und nach Strukturbruch. Die Nullhypothese $H_0 : \beta_1 = \beta_2$ bedeutet, dass kein Strukturbruch stattgefunden hat. Zum Testen der Nullhypothese können wir wie schon früher in ähnlichen Situationen den F -Test zum Prüfen der Gültigkeit von linearen Restriktionen verwenden (siehe Abschnitt 7.5). Die Teststatistik nach (7.5.5) lautet

$$F = \frac{S_R - S}{S} \frac{n - 2k}{k}.$$

Dabei ist S die Summe der quadrierten Residuen des Modells mit den partitionierten Matrizen; bei Zutreffen von H_0 vereinfacht sich das Modell zu

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} X_1 \\ X_2 \end{pmatrix} \beta + \begin{pmatrix} u_1 \\ u_2 \end{pmatrix};$$

die Summe der quadrierten Residuen dieses Modells ist S_R . Die Teststatistik F folgt bei Zutreffen der Nullhypothese und normalverteilten Störgrößen einer F -Verteilung mit k und $n - 2k$ Freiheitsgraden. Bei großem Umfang der verfügbaren Daten können wir von der asymptotischen Verteilung des Likelihood-Quotienten Gebrauch machen: Bei Zutreffen der Nullhypothese folgt diese Teststatistik näherungsweise der Chi-Quadrat-Verteilung mit k Freiheitsgraden.

Im allgemeinen Fall spezifizieren wir m Regime, von denen das i -te durch die Parameter β_i charakterisiert ist. Mit dem Chow-Test testen wir die Nullhypothese

$$H_0 : \beta_1 = \dots = \beta_m; \quad (9.3.1)$$

der Chow-Test prüft wiederum die Gültigkeit von linearen Restriktionen. Für die Teststatistik nach (7.5.5) bekommen wir

$$F = \frac{S_R - \sum_i S_i}{\sum_i S_i} \frac{n - mk}{(m - 1)k}, \quad (9.3.2)$$

wobei S_i die Summe der quadrierten Residuen in Regime i ist. Das Modell wird getrennt für die einzelnen Regime und dann für alle Daten gemeinsam geschätzt; die jeweilige Summe S_i und die Summe S_R der quadrierten Residuen bei Zutreffen der Nullhypothese werden in (9.3.2) eingesetzt. Die Teststatistik F folgt bei Zutreffen der Nullhypothese und normalverteilten Störgrößen einer F -Verteilung. Bei großem Umfang der verfügbaren Daten folgt diese Teststatistik unter der Nullhypothese näherungsweise der Chi-Quadrat-Verteilung mit $(m - 1)k$ Freiheitsgraden.

Beispiel 9.5

Konsumfunktion, Fortsetzung

Als Anwendung des Chow-Tests wollen wir die Frage klären, ob sich die marginale Konsumneigung – vielleicht als Folge des ersten Ölpreisschocks – nach dem Jahr 1972 verändert hat. Lassen wir eine Änderung der Regressionskoeffizienten zwischen 1972 und 1973 zu, so erhalten wir für die Summe der quadrierten Residuen den Wert 5381.7; eine gemeinsame Modellstruktur für alle 46 Beobachtungen liefert 6899.7. Einsetzen in die F -Statistik gibt 5.92 und einen p -Wert von 0.0054. Der Likelihood-Quotient hat den Wert 11.43 entsprechend einem p -Wert von 0.0033. Die Nullhypothese der konstanten marginalen Konsumneigung lässt sich also nicht halten. Lassen wir als Alternative drei Strukturbrüche zu (nach 1972, 1980 und 1991), so erhalten wir für die F -Statistik den Wert 17.30 und einen p -Wert von $1.59 \cdot 10^{-9}$. Der optische Eindruck aus Abbildung 9.1 ist ein starker Hinweis auf eine instabile Struktur. Die Anwendung des Chow-Tests bestätigt diesen Eindruck klar.

Der Chow-Test entscheidet die Frage, ob der datengenerierende Prozess tatsächlich Regime-abhängig ist, durch den F -Test, mit dem wir prüfen, ob die Koeffizienten der Dummy-Variablen, die diese Regime charakterisieren, von Null verschieden sind. In analoger Weise können wir das Vorhandensein von saisonalen Effekten in (9.2.1) überprüfen. Dabei lautet die Nullhypothese

$$H_0 : \delta_i = 0, \gamma_i = 0, \quad i = 2, \dots, 4.$$

Die Alternative kann eine der folgenden sein:

- (a) $H_1 : \delta_i \neq 0, \gamma_i = 0,$
- (b) $H_1 : \delta_i = 0, \gamma_i \neq 0,$
- (c) $H_1 : \delta_i \neq 0, \gamma_i \neq 0.$

Als Teststatistik verwenden wir

$$F = \frac{S_R - S}{S} \frac{n - (2 + 3n_s)}{3n_s}$$

mit $n_s = 1$, falls wir gegen (a) und (b), oder $n_s = 2$, falls wir gegen (c) testen; S_R ist die Summe der quadrierten Residuen bei Zutreffen der Nullhypothese.

In ähnlicher Weise verwenden wir den Chow-Test auch, wenn die zu prüfende Instabilität sich nur auf einzelne der Regressionskoeffizienten bezieht.

Beispiel 9.6

Konsumfunktion, Fortsetzung

Wir testen die Stabilität der Konsumfunktion $C = \beta_1 + \beta_2 Y + \beta_3 M1 + u$ gegen die folgenden Alternativen:

$$H_1^{(1)} : C = \beta_1 + \beta_1^s D_{81} + \beta_2 Y + \beta_3 M1 + u;$$

$$H_1^{(2)} : C = \beta_1 + \beta_1^s D_{81} + \beta_2 Y + \beta_2^s Y * D_{81} + \beta_3 M1 + u;$$

$$H_1^{(3)} : C = \beta_1 + \beta_1^s D_{81} + \beta_2 Y + \beta_3 M1 + \beta_3^s M1 * D_{81} + u;$$

$$H_1^{(4)} : C = \beta_1 + \beta_1^s D_{81} + \beta_2 Y + \beta_2^s Y * D_{81} + \beta_3 M1 + \beta_3^s M1 * D_{81} + u;$$

dabei ist M1 eine Geldmenge und D_{81} eine Dummy-Variable, die für alle Beobachtungsperioden ab 1981 den Wert Eins hat. Die Alternative $H_1^{(1)}$ lässt ab dem Jahr 1981 eine Änderung im Interzept, also im autonomen Konsum, zu; die Alternative $H_1^{(2)}$ besagt, dass der autonome Konsum und die marginale Konsumneigung geänderte Werte haben, etc. Der Chow-Test der Nullhypothese gegen $H_1^{(1)}$ läuft auf den t -Test des Koeffizienten β_1^s von D_{81} hinaus; wir finden $t = 2.43$ und einen p -Wert von 0.020. Der Test der Nullhypothese gegen $H_1^{(2)}$ gibt einen Wert der F -Statistik von 20.91, was einem p -Wert von 10^{-6} entspricht. Die p -Werte für die Tests gegen $H_1^{(3)}$ und $H_1^{(4)}$ betragen $2.6 \cdot 10^{-5}$ und $5.0 \cdot 10^{-6}$. Das allgemeinste Modell ist das unter $H_1^{(4)}$ spezifizierte. Allerdings erhalten wir für den p -Wert zum t -Test des Koeffizienten β_3^s in diesem Modell 0.95; die Wechselwirkung $M1 * D_{81}$ sollten wir daher nicht im Modell berücksichtigen. Wir finden damit, dass der autonome Konsum und die marginale Konsumneigung, nicht aber der Koeffizient der Geldmenge mit dem Jahr 1981 ihre Werte geändert haben.

9.3.2 Chow's Prognosetest

Voraussetzung dafür, dass wir den Chow-Test anwenden können, ist es, dass uns aus jedem Regime eine ausreichende Anzahl von Beobachtungen, also mindestens k , zur Verfügung stehen. Wenn wir nun eine Änderung der Struktur beispielsweise gegen Ende des Beobachtungszeitraums vermuten und nach dieser Änderung weniger als k Beobachtungen gemacht wurden, so können wir den Chow-Test nicht anwenden. Wir können aber unser Modell an die Daten vor der vermuteten Änderung anpassen und Prognosen für die Beobachtungen nach der vermuteten Änderung berechnen. Wenn

die Vermutung stimmt, würden wir erwarten, dass die Prognosefehler einen von Null verschiedenen Erwartungswert haben, da wir ein ungültiges Regressionsmodell verwenden; in Abschnitt 8.1 haben wir diesen Fall behandelt. Das ist die Idee, die dem Prognosetest von Chow zugrunde liegt.

Wir gehen davon aus, dass uns n Beobachtungen zur Verfügung stehen; wir vermuten, dass nach $n - p$ Beobachtungen eine Strukturänderung eingetreten ist und dass die letzten p Beobachtungen Ergebnis eines veränderten datengenerierenden Prozesses sind. Ausgangspunkt ist das lineare Regressionsmodell $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$, für das wir durch Anpassen an die $n - p$ Beobachtungen die Schätzer $\mathbf{b}$ erhalten und das uns die (*ex post*) Prognosen $\hat{\mathbf{y}}_f = \mathbf{X}_f \mathbf{b}$ für die letzten p Beobachtungen von Y liefert.

Prognosetests testen die Nullhypothese

$$H_0 : \mathbf{y}_f = \mathbf{X}_f \boldsymbol{\beta} + \mathbf{u}_f \quad \text{mit} \quad E\{\mathbf{u}_f\} = \mathbf{0}, \quad \text{Var}\{\mathbf{u}_f\} = \sigma^2 \mathbf{I}_p, \quad \text{Cov}\{\mathbf{u}_f, \mathbf{u}\} = \mathbf{0} \quad (9.3.3)$$

gegen die Alternative, dass die Nullhypothese nicht zutrifft. Es soll also die Nullhypothese überprüft werden, dass das Modell auch im Prognosezeitraum gültig ist. Natürlich kann der Test nur *ex post* ausgeführt werden, also erst ab dem Zeitpunkt n , zu dem auch die Beobachtungen für $\mathbf{y}_f$ zur Verfügung stehen.

Wenn wir normalverteilte Störgrößen $\mathbf{u}_f$ voraussetzen, gilt für den p -Vektor der Prognosefehler für großes n näherungsweise

$$\mathbf{e}_f \sim N\left(\mathbf{0}, \sigma^2 \left[\mathbf{I}_p + \mathbf{X}_f(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}_f'\right]\right);$$

siehe (8.1.2). Die Verteilung ist die asymptotische Verteilung der Prognosefehler; sie ist die exakte Verteilung, wenn auch die Störgrößen u_t , $t = 1, \dots, n - p$ normalverteilt sind. Wir können auf Basis dieser Verteilung einen Test konstruieren, dessen Teststatistik

$$F = \frac{1}{\sigma^2} \mathbf{e}_f' \left[\mathbf{I}_p + \mathbf{X}_f(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}_f'\right]^{-1} \mathbf{e}_f$$

bei Zutreffen der Nullhypothese (näherungsweise oder exakt) der Chi-Quadrat-Verteilung mit p Freiheitsgraden folgt. In praktischen Situationen wird die Varianz σ^2 nicht bekannt sein und wir ersetzen sie durch den Schätzer $s^2 = \mathbf{e}'\mathbf{e}/(n - k)$. Die Teststatistik

$$F = \frac{\mathbf{e}_f' \left[\mathbf{I}_p + \mathbf{X}_f(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}_f'\right]^{-1} \mathbf{e}_f}{\mathbf{e}'\mathbf{e}} \frac{n - p - k}{p} \quad (9.3.4)$$

ist unter H_0 F -verteilt mit p und $n - p - k$ Freiheitsgraden. Verwenden wir diese Teststatistik zum Test von (9.3.3), so sprechen wir von Chow's Prognosetest. Für große Werte von F bzw. kleine p -Werte werden wir die Nullhypothese (9.3.3) verwerfen und vermuten, dass die an die Daten des Beobachtungszeitraums angepasste Regressionsbeziehung den datengenerierenden Prozess im Prognosezeitraum nicht adäquat beschreibt.

Der Prognosetest von Chow kann verwendet werden, wenn die Anzahl p der Beobachtungen aus dem Regime mit geänderter Struktur geringer als k ist, die Zahl der zu schätzenden Regressionskoeffizienten. Natürlich kann der Prognosetest von Chow auch verwendet werden, wenn $p > k$.

Beispiel 9.7 Konsumfunktion, Fortsetzung

Wir testen die Vermutung, dass die Konsumfunktion, die wir durch Anpassen an die Daten aus dem Zeitraum 1954 bis 1991 erhalten ($n = 38$), auch für den Zeitraum 1992 bis 1999 gültig ist. Die Prognosen ermitteln wir aus der angepassten Regressionsbeziehung $\hat{C} = 14.190 + 0.872Y$. Für die Teststatistik von Chows Prognosetest erhalten wir, wie in Beispiel 9.8 ausgeführt werden wird, $F = 4.63$; der p -Wert dazu beträgt 0.0006. Wir verwerfen also die Nullhypothese. Der Chow-Test ergibt den Wert der Teststatistik von 15.19; der entsprechende p -Wert beträgt $1.1 \cdot 10^{-5}$. Das Beispiel zeigt, dass die p -Werte von Chow-Test und Prognosetest ziemlich unterschiedlich sein können. Sie können auch widersprüchlich sein, indem einer der beiden Tests zum Verwerfen der Nullhypothese führt, der andere aber nicht.

Berechnen der Teststatistik F

Wir haben mehrere Möglichkeiten, die Teststatistik F nach (9.3.4) zu berechnen.

- Wir können von der Beziehung

$$\mathbf{e}_f' [\mathbf{I}_p + \mathbf{X}_f(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}_f']^{-1} \mathbf{e}_f = S_{D+F} - S_D$$

Gebrauch machen, wobei S_D die Summe der $n - p$ quadrierten Residuen ist, die wir beim Anpassen des Modells an die Daten des Beobachtungszeitraums erhalten ($S_D = \mathbf{e}'\mathbf{e}$), und S_{D+F} die Summe der n quadrierten Residuen, wenn wir das Modell an die Daten sowohl des Beobachtungs- als auch des Prognosezeitraums anpassen. Damit können wir die Teststatistik F nach (9.3.4) schreiben als

$$F = \frac{S_{D+F} - S_D}{S_D} \frac{n - p - k}{p}. \quad (9.3.5)$$

Die Teststatistik F nach (9.3.5) ist durch zwei Modellanpassungen recht einfach zu bestimmen. Wir haben dabei folgende Schritte auszuführen:

1. Modellanpassung an die $n - p$ Beobachtungen und Berechnen der Residuen $\mathbf{e}$ und der Summe der quadrierten Residuen $S_D = \mathbf{e}'\mathbf{e}$ zu $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$.
2. Modellanpassung an alle n Beobachtungen und Berechnen der Residuen und der Summe der quadrierten Residuen S_{D+F} zu

$$\begin{pmatrix} \mathbf{y} \\ \mathbf{y}_f \end{pmatrix} = \begin{pmatrix} \mathbf{X} \\ \mathbf{X}_f \end{pmatrix} \boldsymbol{\beta} + \begin{pmatrix} \mathbf{u} \\ \mathbf{u}_f \end{pmatrix}. \quad (9.3.6)$$

3. Einsetzen in (9.3.5) ergibt die Teststatistik F .

- Wir erweitern die Matrix $\mathbf{X}$ um je eine Spalte für jeden der p Prognosezeitpunkte; die entsprechenden Komponenten des Vektors der Regressionskoeffizienten sind die Elemente von $\boldsymbol{\gamma}$:

$$\mathbf{y}^* = \begin{pmatrix} \mathbf{y} \\ \mathbf{y}_f \end{pmatrix} = \begin{pmatrix} \mathbf{X} & \mathbf{0} \\ \mathbf{X}_f & \mathbf{I}_p \end{pmatrix} \begin{pmatrix} \boldsymbol{\beta} \\ \boldsymbol{\gamma} \end{pmatrix} + \begin{pmatrix} \mathbf{u} \\ \mathbf{u}_f \end{pmatrix} = \mathbf{X}^* \boldsymbol{\beta}^* + \mathbf{u}^*. \quad (9.3.7)$$

Die Spalten $k + 1$ bis $k + p$ der Matrix $\mathbf{X}^*$ entsprechen Indikator-Variablen für die Beobachtungen $n - p + 1$ bis n . Wäre $p = 1$, so könnten wir zeigen, dass das Einbeziehen der Indikator-Variablen für die Beobachtung $(\mathbf{x}'_n, Y_n)$ im Modell zur Folge hat, dass (i) die so markierte Beobachtung in die OLS-Schätzung der Regressionskoeffizienten $\boldsymbol{\beta}$ nicht eingeht, und dass (ii) der OLS-Schätzer für γ gleich dem Residuum $e_n = Y_n - \mathbf{x}'_n \mathbf{b}$ ist; dieses Residuum ist der Prognosefehler zu $\hat{Y}_n$. Analog gilt für $p > 1$, dass (i) die Beobachtungen $(\mathbf{x}'_t, Y_t)$ mit $t = n - p + 1, \dots, n$ in die Schätzung von $\boldsymbol{\beta}$ nicht eingehen, und dass (ii) die Komponenten von $\boldsymbol{\gamma}$ die Prognosefehler zu $\mathbf{y}_f$ sind. Das Anpassen des Modells (9.3.7) ergibt demnach als Summe der quadrierten Residuen die Summe S_D . Wenn wir hingegen das Modell (9.3.7) unter der Restriktion $\boldsymbol{\gamma} = \mathbf{0}$ anpassen, so würden wir als Summe der quadrierten Residuen die Größe S_{D+F} erhalten. Wir können die Teststatistik F daher in folgenden Schritten ermitteln:

1. Anpassen des Modells (9.3.7) ergibt die Summe der quadrierten Residuen $S_D = \mathbf{e}'\mathbf{e}$.
2. Anpassen des Modells (9.3.7) unter der Restriktion $\boldsymbol{\gamma} = \mathbf{0}$ liefert die Summe der quadrierten Residuen S_{D+F} .
3. Einsetzen in (9.3.5) ergibt die Teststatistik F .

Der Prognosetest entspricht also dem Test von $H_0 : \boldsymbol{\gamma} = \mathbf{0}$ für das Modell (9.3.7); es handelt sich um den F -Test nach (7.5.5), den wir in Abschnitt 7.5 kennengelernt haben:

$$F = \frac{S_{D+F} - S_D}{S_D} \frac{n - (p + k)}{p}.$$

Schätzen wir das Modell unter der Voraussetzung, dass die Restriktion $\boldsymbol{\gamma} = \mathbf{0}$ gilt, so erhalten wir $S_R = S_{D+F}$; die nichtrestringierte Anpassung liefert $S = S_D$.

Beispiel 9.8

Konsumfunktion, Fortsetzung

Zum Testen der Vermutung, dass die Konsumfunktion, die wir durch Anpassen an die Daten aus dem Zeitraum 1954 bis 1991 erhalten, auch für den Zeitraum 1992 bis 1999 gültig ist, ermitteln wir die Teststatistik F von Chow's Prognosetest nach (9.3.5) und erhalten

$$\begin{aligned} F &= \frac{\mathbf{e}'_{99}\mathbf{e}_{99} - \mathbf{e}'_{91}\mathbf{e}_{91}}{\mathbf{e}'_{91}\mathbf{e}_{91}} \frac{46 - (8 + 2)}{8} \\ &= \frac{6899.69 - 3400.90}{3400.90} \frac{36}{8} = 4.630; \end{aligned}$$

dabei sind $\mathbf{e}_{91}$ die Residuen, die wir erhalten, wenn das Modell an die Daten aus 1954 bis 1991, $\mathbf{e}_{99}$ jene Residuen, die wir erhalten, wenn das Modell an die Daten aus 1954 bis 1999 angepasst wird. Dem Wert der F -Statistik entspricht ein p -Wert von $p = 0.0005$, wie wir ihn auch in Beispiel 9.7 angegeben haben. Die Nullhypothese kann – bei der üblicherweise tolerierten Fehlerwahrscheinlichkeit von 0.05 – nicht gehalten werden.

Achtung! Der Prognosetest nach Chow setzt die Normalverteilung der Störgrößen u_f voraus!

Alternative Teststatistiken

Die Teststatistik F nach (9.3.4) basiert auf den p Prognosefehlern $\mathbf{e}_f$ und ihrer zumindest näherungsweise gültigen Normalverteilung. Auch bei einem großen Umfang der zum Schätzen der Prognosen verwendeten Daten ($n - p$) ist die Normalverteilung von $\mathbf{u}_f$ essenziell für die Verteilungseigenschaften von F und für die Art und Weise, wie der Prognosetest nach Chow ausgeführt wird. Allerdings können wir bei großem n von den folgenden beiden Modifikationen von Chow's Prognosetest Gebrauch machen.

- Der Schätzer $s^2 = \mathbf{e}'\mathbf{e}/(n - p - k)$ ist ein konsistenter Schätzer der Varianz σ^2 der Störgrößen. Wir definieren die Teststatistik

$$pF = \frac{\mathbf{e}'_f \left[\mathbf{I}_p + \mathbf{X}_f (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'_f \right]^{-1} \mathbf{e}_f}{s^2}; \quad (9.3.8)$$

für großes n ist pF unter H_0 näherungsweise Chi-Quadrat-verteilt mit einer Anzahl von Freiheitsgraden, die gleich der Anzahl p der Prognosezeitpunkte ist; vergleiche dazu die Teststatistik (6.5.5) des asymptotischen Tests aus Abschnitt 6.5. Wir können den Prognosetest bei großem n ausführen, indem wir die Teststatistik F nach (9.3.4) durch die Teststatistik pF ersetzen.

- Existiert eine nichtsinguläre Matrix $Q = \lim X'X/n$, trifft also die Annahme A3 zu, dann können wir den Zähler von F nach (9.3.4) schreiben als

$$\mathbf{e}'_f \left[\mathbf{I}_p + \frac{1}{n} \mathbf{X}_f \left(\frac{1}{n} \mathbf{X}'\mathbf{X} \right)^{-1} \mathbf{X}'_f \right]^{-1} \mathbf{e}_f.$$

Bei großem n können wir diesen Ausdruck gleich $\mathbf{e}'_f \mathbf{e}_f$ ersetzen. Wir definieren die Teststatistik

$$T = \frac{\mathbf{e}'_f \mathbf{e}_f}{s^2}; \quad (9.3.9)$$

für großes n ist T unter H_0 näherungsweise Chi-Quadrat-verteilt mit p Freiheitsgraden. Bei großem n können wir den Prognosetest mittels der Teststatistik T ausführen.

Bei beiden Chi-Quadrat-Tests werden wir bei einem großen Wert der Teststatistik bzw. einem kleinen p -Wert die Nullhypothese (9.3.3) verwerfen und wir werden vermuten, dass die an die Daten des Beobachtungszeitraums angepasste Regressionsbeziehung den datengenerierenden Prozess des Prognosezeitraums nicht adäquat beschreibt. Das spezifizierte Modell oder die Modellstruktur, die sich für den Beobachtungszeitraum ergeben hat, ist im Prognosezeitraum nicht gültig.

Copyright

Daten, Texte, Design und Grafiken dieses eBooks, sowie die eventuell angebotenen eBook-Zusatzdaten sind urheberrechtlich geschützt. Dieses eBook stellen wir lediglich als **persönliche Einzelplatz-Lizenz** zur Verfügung!

Jede andere Verwendung dieses eBooks oder zugehöriger Materialien und Informationen, einschließlich

- der Reproduktion,
- der Weitergabe,
- des Weitervertriebs,
- der Platzierung im Internet, in Intranets, in Extranets,
- der Veränderung,
- des Weiterverkaufs und
- der Veröffentlichung

bedarf der **schriftlichen Genehmigung** des Verlags. Insbesondere ist die Entfernung oder Änderung des vom Verlag vergebenen Passwortschutzes ausdrücklich untersagt!

Bei Fragen zu diesem Thema wenden Sie sich bitte an: info@pearson.de

Zusatzdaten

Möglicherweise liegt dem gedruckten Buch eine CD-ROM mit Zusatzdaten bei. Die Zurverfügungstellung dieser Daten auf unseren Websites ist eine freiwillige Leistung des Verlags. **Der Rechtsweg ist ausgeschlossen.**

Hinweis

Dieses und viele weitere eBooks können Sie rund um die Uhr und legal auf unserer Website herunterladen:

<http://ebooks.pearson.de>