



BCEAO

BANQUE CENTRALE DES ETATS
DE L'AFRIQUE DE L'OUEST

LES PRECIS DU COFEB

N°003 / 2013

ECONOMETRIE : GUIDE PRATIQUE

TOME 1

LE MODELE LINEAIRE GENERAL

Dr Fodiyé Bakary DOUCOURE

Octobre 2013



BCEAO
BANQUE CENTRALE DES ÉTATS
DE L'AFRIQUE DE L'OUEST

Direction Générale des Ressources Humaines et de la Formation
Centre Ouest Africain de Formation et d'Études Bancaires (COFEB)

LES PRECIS DU COFEB

N°003 / 2013

ECONOMETRIE : GUIDE PRATIQUE

TOME 1

LE MODELE LINEAIRE GENERAL

Dr Fodiyé Bakary DOUCOURE

Octobre 2013

Les avis exprimés engagent
la responsabilité des seuls auteurs.

TABLE DES MATIERES

PREAMBULE	5
Introduction	7
1. Exposé du problème	9
2. Notation matricielle du modèle linéaire général	10
3. Estimation et propriétés des estimateurs	11
3.1 Hypothèses d'application de la méthode des moindres carrés ordinaires	11
3.1.1 Hypothèses structurelles	11
3.1.2 Hypothèses stochastiques	12
3.2 Estimation des paramètres par la méthode des moindres carrés ordinaires	12
4. Interprétation économique des paramètres	13
4.1 Modèle sans logarithme	13
4.2 Modèle log-linéaire	13
4.3 Modèle semi logarithmique	14
4.4 Modèle réciproque	16
5. Théorème de Gauss et Markov	16
6. Equation d'analyse de la variance et qualité d'un ajustement	17
6.1 Equation d'analyse de la variance	17
6.2 Qualité d'un ajustement	18
6.2.1 Coefficient de détermination	18
6.2.2 Coefficient de détermination corrigé	18
6.2.3 Interprétation du coefficient de détermination	19
6.2.4 Utilisation du R^2 pour sélectionner une théorie	19
7. Tests économétriques	20
7.1 Test d'hypothèses	20
7.1.1 Définitions	20
7.1.2 Types d'erreurs	20
7.1.3 Niveau de signification d'un test	20
7.2 Test du coefficient de corrélation linéaire	21
7.3 Test de normalité de Jarque-Bera	24
7.3.1 Moment centré d'ordre r	24

7.3.2 Test de Jarque-Bera	25
7.4 test de student.....	27
7.4.1 comparaison d'un paramètre à une valeur fixée a	27
7.4.2 Comparaison d'un paramètre à la valeur 0	28
7.4.3 Test de l'hypothèse.....	28
7.4.4 Intervalle de confiance pour les paramètres	29
7.5 test de signification globale d'une régression	30
7.6 Tests d'autocorrélation des erreurs	38
7.6.1 Définition et causes de l'autocorrélation des erreurs	38
7.6.2 Estimation en présence d'autocorrélation des erreurs	38
7.6.3 Les tests d'autocorrélation des erreurs	38
7.6.3.1 Test de Durbin et Watson	38
7.6.3.2 Test de Breusch et Godfrey	40
7.6.3.3 Procédures d'estimation en cas d'autocorrélation des erreurs	41
7.7 Tests d'hétéroscédasticité des erreurs	45
7.7.1 Définition	45
7.7.2 Tests de détection de l'hétéroscédasticité des erreurs	45
7.7.2.1 Test de White	46
7.7.2.2 Test ARCH	47
7.7.3 Procédure d'estimation en présence de l'hétéroscédasticité des erreurs	48
7.8 Test de spécification de Ramsey	50
7.9 tests de stabilité	52
7.9.1 test de chow	52
7.9.2 Tests Cusum	54
8. La prévision à l'aide du modèle linéaire général	58
Conclusion	61
Annexe : Instructions et codes des logiciels Eviews et Stata	
pour la résolution des applications proposées	63

PREAMBULE

Le présent opusculé constitue le troisième numéro de la série « Les Précis du COFEB ». Il porte sur la théorie économétrique et ses applications.

Ecrit dans un langage relativement simple, par un universitaire, ce document expose avec clarté, précision et méthode les principes de base des méthodes économétriques. Cet ouvrage s'efforce, en tout cas, d'exprimer par des phrases ce que beaucoup d'autres laissent le soin à leur lecteur d'imaginer. Toutes les démonstrations sont omises, parce qu'elles n'apparaissent pas nécessaires à la compréhension de la démarche.

La bonne compréhension du contenu oblige toutefois à avoir assimilé un certain nombre d'éléments de statistique descriptive et de statistique mathématique.

Ce document est le fruit de divers enseignements d'économétrie dispensés par l'auteur au Centre Ouest Africain de Formation et d'Etudes Bancaires (COFEB) de la Banque Centrale des Etats de l'Afrique de l'Ouest et à la Faculté des Sciences Economiques et de Gestion de l'Université Cheikh Anta Diop de Dakar. Il s'adresse ainsi principalement aux agents de la BCEAO et aux auditeurs du COFEB. Il sera également utile aux professionnels qui utilisent les techniques économétriques dans la mesure où chaque chapitre est systématiquement illustré d'applications empiriques.

Le premier tome traite l'économétrie du modèle linéaire général illustrée par des applications entièrement traitées à l'aide des logiciels économétriques, principalement EvIEWS et Stata. On y présente tout d'abord les méthodes d'estimation et les tests applicables au modèle linéaire général. Des versions moins restrictives de ce modèle sont ensuite présentées avec des perturbations corrélées entre elles ou de variances non constantes.

INTRODUCTION

Le présent document sur le thème « Econométrie : guide pratique », s'inscrit dans le cadre de la politique de vulgarisation des concepts économétriques à l'intention des agents de la Banque Centrale des Etats de l'Afrique de l'Ouest. Il repose sur une approche essentiellement descriptive des méthodes économétriques et se propose de présenter de façon simple des concepts à priori complexes.

L'économétrie est la mesure des phénomènes économiques. Elle constitue une branche de la science économique qui fait appel conjointement à la théorie économique, la statistique, les mathématiques et l'informatique.

L'économétrie, en tant que discipline, est née en 1930 lors de la création de la Société d'économétrie par Ragnar Frisch, Charles Roos et Irving Fisher.

L'économétrie est un ensemble de méthodes statistiques appliquées à l'économie. Elle a deux fonctions essentielles :

- i) tester les théories économiques ou, plus modestement certaines assertions de la théorie économique;
- ii) évaluer les paramètres en jeu dans les relations économiques.

L'économétrie est une branche de la science économique consistant à établir des lois ou à vérifier des hypothèses à partir de données chiffrées tirées de la réalité.

Il est à peu près impossible de faire de la recherche en sciences économiques sans se trouver devant la nécessité de lire ou de réaliser des travaux d'économétrie à un moment ou à un autre.

C'est la raison pour laquelle, dans tous les pays, la formation des économistes suppose l'acquisition des techniques économétriques. A titre d'illustration sur les derniers Prix Nobel d'Economie, trois furent attribués à des économètres : Heckman et McFadden en 2000 (économétrie des variables qualitatives), Engle et Granger en 2003 (économétrie des séries temporelles et économétrie financière) et Sargent et Sims en 2011 (recherches sur la cause et l'effet en macroéconomie, modèles VAR, ...).

De fait donc, l'économiste ne doit-il pas être aussi économètre ? John Maynard Keynes, dans les années 1930, écrivait :

« L'économiste doit être mathématicien, historien, philosophe, homme d'Etat, ... »

S'il faut transférer la pensée de Keynes aujourd'hui, n'aurait-il pas lui-même ajouté l'économiste doit être économètre ?

L'économétrie est donc une discipline dont le contenu opérationnel est très important. A titre d'exemple, l'économétrie peut permettre de quantifier un phénomène, d'établir une relation entre plusieurs variables, de valider ou d'infirmer empiriquement une théorie, d'évaluer les effets d'une mesure de politique économique, etc.

L'économétrie peut permettre de répondre à des interrogations comme :

- Les termes de l'échange sont-ils un instrument de la valeur des taux de change ? D'autres variables économiques ont-elles plus d'impact ?
- La théorie de la parité des pouvoirs d'achat est-elle vérifiée empiriquement ?
- Une hausse de l'inflation permet-elle de réduire le chômage ?
- Y a-t-il une convergence des PIB par tête au niveau international ?
- Quel est l'impact d'une hausse de 10% du revenu sur la consommation d'un ménage ? Cet impact diffère-t-il selon le pays de résidence du ménage ?
- etc.

De manière générale, pour répondre à ces interrogations, l'économetre devra construire un modèle visant à mettre en relation les diverses variables d'intérêt.

L'économétrie peut en effet être appliquée à toutes les branches de l'économie : la macroéconomie, la finance, l'économie du travail, l'économie industrielle, l'économie publique, etc. Les techniques économétriques peuvent aussi être mobilisées dans d'autres disciplines telles que la gestion (notamment en marketing), la biologie, l'agronomie, la sociologie, la psychologie, la science politique, etc.

Pour terminer, il est important de relever que les techniques économétriques sont régulièrement utilisées dans les banques, les ministères, les agences gouvernementales, les grandes entreprises, les institutions internationales, la recherche agronomique, la finance, etc.

CHAPITRE I

LE MODELE LINEAIRE GENERAL

On parle de modèle de régression simple si une seule variable explicative est considérée. La fonction de consommation keynésienne

$$C_t = a_1 + a_2 R_t + \varepsilon_t, \quad t = 1, \dots, n$$

où C est la consommation et R le revenu est un modèle de régression simple.

En pratique, il est cependant fréquent qu'une variable (expliquée) dépende de plusieurs variables explicatives. Prenons quelques exemples visant à illustrer les questions auxquelles on peut répondre dans le chapitre. Les dépenses de consommation d'une famille dépendent-elles plus de son revenu, de sa taille ou des prix ? Le salaire d'un chef de ménage dépend-t-il plus du niveau de formation ou de l'expérience professionnelle ? La production d'une entreprise dépend-t-elle plus du facteur capital ou du facteur travail ? Le PIB d'un pays dépend-t-il plus de ses exportations, de l'investissement, du capital humain, de l'encours de la dette extérieure, du taux d'inflation, de la bonne gouvernance ou du taux de change ?

Dans ces différents cas où plusieurs variables explicatives entrent en jeu, on parle de modèle linéaire général² ou modèle de régression linéaire multiple.

1. Exposé du problème

Le modèle linéaire général (ou modèle de régression linéaire multiple) peut s'écrire :

$$Y_t = a_1 X_{1t} + a_2 X_{2t} + \dots + a_k X_{kt} + \varepsilon_t, \quad t = 1, 2, \dots, n$$

où Y est la variable endogène (ou à expliquer), $X_1, X_2, \dots, X_k$ les k variables exogènes (ou explicatives) indépendantes et non aléatoires, $a_1, a_2, \dots, a_k$ des nombres réels inconnus, n le nombre d'observations et ε le terme d'erreur (ou aléa) non observé.

L'erreur ε est une variable aléatoire centrée de variance σ^2 , indépendante des variables X_{it} . Il s'agit d'un terme d'erreur stochastique qui permet de prendre en compte le fait que la variable Y est affectée par d'autres variables que les variables $X_1, X_2, \dots, X_k$.

1- Ce chapitre reprend un certain nombre de développements figurant dans l'ouvrage de F.B. Doucouré (2013) que le lecteur intéressé pourra consulter pour plus de détails.

2- Ce chapitre reprend un certain nombre de développements figurant dans l'ouvrage de F.B. Doucouré (2013) que le lecteur intéressé pourra consulter pour plus de détails.

Généralement, la première variable X_{1t} est égale à 1 pour toutes les observations. Par conséquent a_1 est la constante du modèle qui est de la forme :

$$Y_t = a_1 + a_2 X_{2t} + \dots + a_k X_{kt} + \varepsilon_t, \quad t = 1, 2, \dots, n$$

On désire estimer les paramètres $a_1, a_2, \dots, a_k$ et la variance σ^2 de la variable aléatoire ε à partir d'un ensemble de n observations indépendantes.

2. Notation matricielle du modèle linéaire général

Le modèle linéaire général peut être présenté sous la forme matricielle :

$$\underset{(n,1)}{Y} = \underset{(n,k)}{X} \cdot \underset{(k,1)}{a} + \underset{(n,1)}{\varepsilon}$$

où

$$Y = \begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{pmatrix}; \quad X = \begin{pmatrix} 1 & X_{21} & \dots & X_{k1} \\ 1 & X_{22} & \dots & X_{k2} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & X_{2n} & \dots & X_{kn} \end{pmatrix}; \quad a = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_k \end{pmatrix}; \quad \varepsilon = \begin{pmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_n \end{pmatrix}$$

Y est le vecteur des n observations sur la variable endogène. X est la matrice des variables exogènes, chaque colonne de la matrice X est une variable explicative. a est le vecteur des coefficients de régression et ε est le vecteur des écarts aléatoires.

Illustration 1 : Notation matricielle

Considérons la fonction de consommation keynésienne :

$$C_t = a_1 + a_2 R_t + \varepsilon_t, \quad t = 1, \dots, n$$

où

C_t : consommation au temps t ;

R_t : revenu au temps t ;

a_1 : consommation incompressible ;

a_2 : propension marginale à consommer

Cette fonction est un modèle de régression linéaire simple. Elle s'écrit matriciellement :

$$Y = X \cdot a + \varepsilon$$

où

$$Y = \begin{pmatrix} C_1 \\ C_2 \\ M \\ C_n \end{pmatrix}; \quad X = \begin{pmatrix} 1 & R_1 \\ 1 & R_2 \\ M & M \\ 1 & R_n \end{pmatrix}; \quad a = \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}; \quad \varepsilon = \begin{pmatrix} \varepsilon_1 \\ \varepsilon_2 \\ M \\ \varepsilon_n \end{pmatrix}$$

3. Estimation et propriétés des estimateurs

3.1 Hypothèses d'application de la méthode des moindres carrés ordinaires.

Deux catégories d'hypothèses doivent être faites pour résoudre le problème des moindres carrés. Nous distinguons les hypothèses structurelles ; des hypothèses stochastiques (c'est à dire liées à l'erreur ε).

3.1.1 Hypothèses structurelles

$H_1 : n > k$: le nombre d'observations est supérieur au nombre de séries explicatives, c'est à dire que les degrés de liberté $(n - k)$ doivent être strictement positifs.

H_2 : le rang de la matrice X des variables explicatives est égal à k , le nombre de paramètres à estimer. La matrice X est de plein rang colonne, autrement dit, les variables explicatives sont linéairement indépendantes (absence de colinéarité entre les variables explicatives). Ceci implique que la matrice $X'X$ est inversible c'est à dire que la matrice $(X'X)^{-1}$ existe. X' désigne la transposée de la matrice X .

Ces hypothèses structurelles sont techniques, elles garantissent l'existence d'une solution.

3.1.2 Hypothèses stochastiques

- $H_3 : E(\varepsilon_i) = 0$, l'espérance mathématique de l'erreur est nulle. Cette hypothèse implique que Y est une variable aléatoire d'espérance mathématique x_a . Cette condition permet d'obtenir des estimations sans biais.

- $H_4 : E(\varepsilon_i^2) = \sigma^2$, les erreurs sont homocédastiques (leur variance est constante). Cette hypothèse peut être testée et quand on a repéré l'hétéroscédasticité des erreurs ($E(\varepsilon_i^2) \neq \sigma^2$ pour au moins un i), on peut corriger et obtenir néanmoins de bons estimateurs.

- $H_5 : E(\varepsilon_t \varepsilon_{t'}) = 0, \forall t \neq t'$, les erreurs ne sont pas corrélées (c'est à dire qu'elles sont indépendantes). Cette hypothèse peut elle aussi être testée. Quand H_5 n'est pas vérifiée, on parle alors d'autocorrélation des erreurs.

- $H_6 : \varepsilon_t \rightarrow N(0, \sigma^2)$, les termes d'erreur suivent une loi normale. Cette hypothèse joue un rôle essentiel bien qu'elle soit difficile à vérifier. Il existe des tests de normalité (par exemple le test de Jarque-Bera) mais ils ne sont pas applicables sur les termes d'erreurs, qui par définition demeurent inconnus.

Les hypothèses stochastiques ont pour but de s'assurer que les estimateurs des coefficients a_i jouissent de propriétés statistiques intéressantes.

3.2 Estimation des paramètres par la méthode des moindres carrés ordinaires

La méthode des moindres carrés ordinaires (ci-après MCO) consiste à minimiser la somme des carrés des résidus. On démontre que de l'estimateur des moindres carrés ordinaires est $\hat{A} = (X'X)^{-1} X'Y$. La matrice $(X'X)^{-1}$ existe du fait que X est de rang k (hypothèse H_2).

Remarque 1.1 : L'estimation des paramètres par la méthode des MCO est effectuée par ordinateur grâce à des logiciels comme Eviews, Stata, SPSS (Statistical Package of Social Sciences), RATS, MATLAB, R, GAUSS ou SAS (Statistical Analysis System).

Illustration 2 : Modèle sans constante

Soit le modèle linéaire simple $Y_i = \beta X_i + \varepsilon_i \quad 1 \leq i \leq n$

1) Quelle hypothèse du modèle linéaire simple n'est pas vérifiée a priori dans cette spécification ?

2) Pouvez-vous trouver dans la théorie économique deux exemples où de telles relations se justifient ?

Corrigé :

1) Il n'y a pas de terme constant α . Ou autrement dit, on a supposé une hypothèse supplémentaire $\alpha = 0$.

2) Une telle relation se justifie dans les deux modèles suivants:

a) Dans un modèle d'offre de travail $OT_i = \beta w_i + \varepsilon_i$ où OT est l'offre de travail et w le salaire. Pour un salaire nul, un individu n'offrira pas de travail.

b) Dans la fonction de consommation de Milton Friedman, on a $C = kY_p$ où C est la consommation et Y_p le revenu. Pour un revenu permanent (richesse) nul, l'individu ne pourra pas consommer.

4. Interprétation économique des paramètres

Nous considérons quatre modèles de régression couramment utilisés dans la pratique et donnons l'interprétation économique des paramètres.

4.1 Modèle sans logarithme

Nous considérons le modèle $Y = aX + b + \varepsilon$.

Dans ce cas, le paramètre a est une propension marginale.

$$a = \frac{\partial Y}{\partial X} \Rightarrow \partial Y = a \times \partial X$$

Le paramètre b a une signification économique, c'est la valeur de Y quand $X = 0$.

Illustration 3 :

A l'aide de 35 observations, on souhaite estimer les paramètres du modèle

$\text{Cons} = a \text{ Rev} + b + \varepsilon$, où cons est la consommation de l'individu et Rev son revenu. Les données sont en francs CFA. L'estimation des paramètres par la méthode des MCO donne : $\hat{a} = 0,85$; $\hat{b} = 12\,000$.

On a une fonction de consommation, $\hat{b}$ est l'estimation de la consommation quand le revenu est nul, c'est à dire si $\text{Rev} = 0$. $\hat{b}$ est donc l'estimation de la consommation incompressible (consommation autonome ou minimum de subsistance).

$\hat{a}$ est l'estimation de la propension marginale à consommer : $\partial \text{Cons} = 0,85 \times \partial \text{Rev}$.

$\hat{a} = 0,85$ signifie qu'une augmentation de 100 000 francs CFA du revenu d'un individu implique une augmentation de 85 000 francs CFA de sa consommation.

4.2 Modèle log-linéaire

Le modèle log-linéaire, encore appelé modèle log-log ou modèle double-log, est donné par :

$$\log(Y) = a \log(X) + b + \varepsilon$$

où $\log$ est le logarithme népérien.

Dans ce cas, le paramètre a est une élasticité.

$$a = \frac{\partial \log(Y)}{\partial \log(X)} \Rightarrow \frac{\partial Y}{Y} = a \times \frac{\partial X}{X}$$

a est l'élasticité de Y par rapport à X et le paramètre b n'a pas une signification économique.

Illustration 4 : Les importations du Sénégal (Y) sont mises en relation avec le Produit Intérieur Brut (X) sur la période 1962 à 1995.

$$\log(\text{IMP}) = a \log(\text{PIB}) + b + \varepsilon$$

L'estimation des paramètres par la méthode des MCO donne : $\hat{a} = 0,75$; $\hat{b} = 0,65$.

$\hat{a}$ est l'estimation de l'élasticité des importations par rapport au produit intérieur brut.

$$\frac{\partial \text{IMP}}{\text{IMP}} = 0,75 \times \frac{\partial \text{PIB}}{\text{PIB}}$$

$\hat{a} = 0,75$ signifie qu'une augmentation de 10% du PIB implique une augmentation de 7,5% des importations. Il est à noter que $\hat{b} = 0,65$ n'a pas d'interprétation économique.

Remarque 1.2 : La transformation logarithmique, donnant lieu à un modèle log-linéaire peut être utilisé pour réduire l'hétéroscédasticité des erreurs (voir la sous section 7.7). La réduction de l'hétéroscédasticité provient du fait que la transformation logarithmique « comprime » les échelles dans lesquelles les variables sont mesurées.

4.3 Modèle semi logarithmique

Le modèle semi logarithmique est donné par :

$$\log(Y) = aX + b + \varepsilon$$

Dans ce cas, le paramètre a est une semi élasticité.

$$a = \frac{\partial \log(Y)}{\partial X} \Rightarrow \frac{\partial Y}{Y} = a \times \partial X$$

a est la semi élasticité de Y par rapport à X et le paramètre b n'a pas une signification économique.

Illustration 5 : L'investissement (INV) du Sénégal est mis en relation avec le taux d'intérêt réel (TXINT) sur la période 1972 à 2001.

$$\log(\text{INV}) = a\text{TXINT} + b + \varepsilon$$

L'estimation des paramètres par la méthode des MCO donne : $\hat{a} = -0,06$; $\hat{b} = -2,59$

$\hat{a}$ est l'estimation de la semi élasticité de l'investissement par rapport au taux d'intérêt réel.

$$\frac{\partial \text{INV}}{\text{INV}} = -0,06 \times \partial \text{TXINT}$$

$\hat{a} = -0,06$ signifie qu'une hausse du taux d'intérêt réel d'un point (100% ou cent points) de 1% à 2% par exemple entraîne une diminution de 6% de l'investissement. Il est à noter que $\hat{b} = -2,59$ n'a pas d'interprétation économique.

Si la variable explicative est le temps, le modèle semi logarithmique s'écrit :

$$\log(Y) = at + b + \varepsilon$$

Si t désigne par exemple des mois, des trimestres, des années, on démontre que $\log(1 + g) = a$ où g est le taux de croissance de Y . Le coefficient a s'interprète comme le taux de croissance continu qui donnerait, au bout d'une période, le même résultat qu'une augmentation unique au taux g .

Illustration 6 : Considérons la série macroéconomique masse monétaire (M) du Sénégal sur la période 1971-2007 à fréquence annuelle. Considérons le modèle $\log(M_t) = a + bt + \varepsilon_t$ où t désigne le temps, soit $t = 0, 1, 2, \dots, 36$. L'estimation de ce modèle par la méthode des MCO conduit aux résultats suivants :

$$\log(M_t) = 0,0907t + 4,1888$$

On déduit de cette estimation que $\log(1 + \hat{g}) = 0,0907$ où $\hat{g}$ est le taux de croissance estimé. D'où $\hat{g} = 0,0949$, sur la période 1971-2007, la masse monétaire du Sénégal a augmenté annuellement à un taux de 9,49%.

4.4 Modèle réciproque

Le modèle réciproque s'écrit :

$$Y = a \left(\frac{1}{X} \right) + b + \varepsilon$$

En vertu de ce modèle, lorsque la variable X tend vers l'infini, le terme $a \left(\frac{1}{X} \right)$ tend vers zéro et b est donc la limite asymptotique de Y lorsque X tend vers l'infini. La pente du modèle est donnée par :

$$\frac{\partial Y}{\partial X} = -a \times \left(\frac{1}{X^2} \right)$$

Par conséquent, si $a > 0$, la pente est toujours négative et si $a < 0$, la pente est toujours positive. Ce type de modèle, pour $a > 0$, peut être illustré par la courbe de Phillips. Rappelons que cette dernière relie à l'origine le taux de croissance des salaires nominaux au taux de chômage. Elle a par la suite été transformée en une relation entre taux d'inflation et taux de chômage.

Illustration 7 : A titre d'exemple, supposons que pour un pays donné, la régression du taux d'inflation ($\hat{\pi}_t$) sur l'inverse du taux de chômage (u_t) conduise aux résultats suivants :

$$\hat{\pi}_t = 20,0103 \left(\frac{1}{u_t} \right) - 2,3030$$

Ces résultats montrent que même si le taux de chômage augmente indéfiniment, la plus forte variation des prix sera une baisse du taux d'inflation de l'ordre de 2,3 points.

5. Théorème de Gauss et Markov

L'estimateur $\hat{A} = (X'X)^{-1} X'Y$ de a est «BLU» (Best Linear Unbiased) ou le meilleur estimateur linéaire sans biais de a c'est-à-dire que :

- a) les composantes $\hat{A}_1, \hat{A}_2, K, \hat{A}_k$ de $\hat{A}$ sont des fonctions linéaires des Y_i ,
- b) $\hat{A}$ est un estimateur sans biais de a ,

c) la matrice de variances-covariances de $\hat{\mathbf{A}}$ est $V(\hat{\mathbf{A}}) = \sigma^2 (\mathbf{X}' \mathbf{X})^{-1}$.

d) parmi les estimateurs sans biais des composantes $\hat{\mathbf{A}}_1, \hat{\mathbf{A}}_2, \dots, \hat{\mathbf{A}}_k$ de fonctions linéaires des $\mathbf{Y}_i$, les composantes $\hat{\mathbf{A}}_1, \hat{\mathbf{A}}_2, \dots, \hat{\mathbf{A}}_k$ de $\hat{\mathbf{A}}$ sont celles dont les variances sont minimum.

Le théorème de Gauss et Markov énonce une propriété très importante des MCO : $\hat{\mathbf{A}}$ donne la meilleure connaissance possible de $\mathbf{a}$ puisqu'il est sans biais et à variance minimale. Cette propriété est obtenue sous les hypothèses assez peu exigeantes puisqu'on n'a pas formulé d'hypothèse sur la (ou les lois) suivie(s) par les erreurs ε_t .

Remarque 1.3 : L'estimateur des MCO est appelé estimateur de Gauss-Markov. Il est optimal si les erreurs sont homocédastiques et non corrélées c'est à dire si $V(\varepsilon) = \sigma^2 \mathbf{I}$. Si les erreurs sont hétérocédastiques et/ou corrélées, l'estimation des paramètres se fait par la méthode des moindres carrés généralisés (MCG).

Considérons le modèle $\mathbf{Y} = \mathbf{X} \mathbf{a} + \varepsilon$, pour lequel $V(\varepsilon) = V \neq \sigma^2 \mathbf{I}$, l'estimateur des MCG ou estimateur de Aitken est $\hat{\mathbf{A}} = (\mathbf{X}' \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}' \mathbf{V}^{-1} \mathbf{Y}$. L'estimateur de Aitken généralise donc l'estimateur de Gauss et Markov.

6. Equation d'analyse de la variance et qualité d'un ajustement

6.1 Equation d'analyse de la variance

Ayant évalué par la méthode des moindres carrés ordinaires les paramètres du modèle, on cherche à déterminer si le modèle permet de bien expliquer la variable endogène. Soit $\mathbf{e} = \mathbf{Y} - \hat{\mathbf{Y}}$, le vecteur des résidus et $\hat{\mathbf{Y}} = \mathbf{X} \hat{\mathbf{A}}$, le vecteur des valeurs estimées de $\mathbf{Y}$. On démontre que :

$$\sum_{t=1}^n (Y_t - \bar{Y})^2 = \sum_{t=1}^n (Y_t - \hat{Y}_t)^2 + \sum_{t=1}^n (\hat{Y}_t - \bar{Y})^2$$

$$\text{SCT} = \text{SCR} + \text{SCE}$$

Cette égalité est l'équation d'analyse de la variance. Cette égalité correspond à la décomposition de la variance totale (SCT) en variance résiduelle (SCR) et variance expliquée (SCE).

SCT = Variabilité totale

SCR = Variabilité des résidus

SCE = Variabilité expliquée

6.2 Qualité d'un ajustement

6.2.1 Coefficient de détermination

Un modèle est une représentation simplifiée des déterminants du phénomène économique étudié. Les facteurs qui influencent la variable endogène Y sont certainement beaucoup plus nombreux et complexes que ceux qui sont retenus dans le modèle. Quelle est la part de la variance de Y qui est expliquée par ce modèle réducteur ?

Le coefficient de détermination R^2 est défini à partir de l'équation d'analyse de la variance précédente :

$$R^2 = \frac{\text{variance expliquée}}{\text{variance totale}} = \frac{\sum_{t=1}^n (\hat{Y}_t - \bar{Y})^2}{\sum_{t=1}^n (Y_t - \bar{Y})^2} = \frac{SCE}{SCT}$$

On a :

$$R^2 = \frac{SCT - SCR}{SCT} = 1 - \frac{SCR}{SCT} \quad \text{et} \quad SCR = (1 - R^2) SCT$$

Par construction on a : $0 \leq R^2 \leq 1$. Le coefficient de détermination est donc bien un instrument de mesure de la qualité de l'ajustement. Plus il est proche de 1, mieux cela vaut. Cependant le coefficient de détermination doit être utilisé avec précaution, il est aussi à noter que le coefficient R^2 n'est interprétable que dans un modèle avec terme constant.

6.2.2 Coefficient de détermination corrigé

Le coefficient de détermination est une fonction non décroissante du nombre de variables explicatives incluses dans le modèle. Ainsi, lorsqu'on augmente le nombre de variables explicatives dans le modèle, la valeur du coefficient de détermination augmente. Afin de pallier ce problème, il convient de corriger le coefficient de détermination par le nombre de degrés de liberté. Pour cela, on détermine ce que l'on appelle le coefficient de détermination ajusté de Henry Theil, appelé aussi coefficient de détermination corrigé. Ce dernier noté $\bar{R}^2$ est défini par :

$$\bar{R}^2 = 1 - \frac{n-1}{n-k} (1 - R^2)$$

On a : $\bar{R}^2 < R^2$ et si n est grand $\bar{R}^2 \approx R^2$. $\bar{R}^2$ peut diminuer lorsqu'une variable est ajoutée à l'ensemble des variables exogènes. $\bar{R}^2$ peut même être négatif.

6.2.3 Interprétation du coefficient de détermination

Nous donnons l'interprétation du coefficient de détermination à travers un modèle de croissance économique. On dispose pour un pays donné, des séries macroéconomiques PIB, Investissement (INV), Exportations (EXPT), Taux d'inflation (TINF) et Encours de la dette extérieure (DETTE). L'estimation par la méthode des MCO des paramètres du modèle

$$\text{PIB}_t = a_1 + a_2 \text{EXPT}_t + a_3 \text{TINF}_t + a_4 \text{DETTE}_t + \varepsilon_t$$

donne une valeur de R^2 égale à 0,87. Cette valeur indique que 87% des fluctuations du PIB sont expliquées par les variables EXPT, TINF et DETTE. On peut aussi affirmer que le modèle explique 87% de la variance du PIB. Il faut remarquer que cette valeur obtenue pour le coefficient de détermination R^2 n'indique rien sur la qualité du modèle.

Remarque 1.4 : Les coefficients de détermination et de détermination ajusté permettent de comparer des modèles entre eux. Bien entendu, ces modèles doivent avoir la même variable endogène et comporter le même nombre d'observations. On retiendra le modèle dont le coefficient de détermination – ou le coefficient de détermination ajusté – est le plus élevé.

6.2.4 Utilisation du R^2 pour sélectionner une théorie

Du fait de sa simplicité, le coefficient R^2 invite à plusieurs utilisations discutables. Il ne convient pas d'utiliser le R^2 pour sélectionner une théorie ou une spécification particulière.

Combien de fois ai-je entendu un étudiant, auquel je reprochais la faiblesse économique qui sous-tendait la relation qu'il étudiait me rétorquer : « Mais j'ai un coefficient de détermination de 0,95 ! ».

L'argument lui paraissait de poids. Il n'est pourtant pas recevable. N'hésitons pas à reprendre le propos d'un vieux sage : « L'économétrie est à une théorie ce que le lampadaire est à l'ivrogne : il le soutient plus qu'il ne l'éclaire. » Une théorie parfaitement nébuleuse le reste, en dépit de tous les beaux coefficients que l'on peut exhiber pour la soutenir. La qualité d'un travail économétrique dépend étroitement de la qualité du travail théorique qui l'a précédé. Il convient donc de pas attendre des miracles de l'économétrie : elle n'est que l'art de se servir bien des données dans le cadre d'une théorie, ce n'est déjà pas si mal.

7. Tests économétriques

Une étude économétrique consiste non seulement à estimer des paramètres d'un modèle, mais aussi, à tester des hypothèses afin de valider le modèle économique théorique. Les paramètres estimés $\hat{A}$ sont des variables aléatoires, ce ne sont pas des valeurs certaines, ils ne sont pas exactement identiques à la vraie valeur des paramètres a . On doit effectuer des tests statistiques afin de compléter les résultats des estimations.

7.1 Test d'hypothèses

7.1.1 Définitions

Un test est une procédure qui permet d'accepter ou de rejeter rationnellement une hypothèse relative par exemple à la valeur d'un paramètre de la population ou à sa loi de probabilité. L'hypothèse nulle (ou hypothèse à tester) est notée H_0 tandis que l'hypothèse alternative est notée H_1 .

L'objectif d'un test est de réfuter l'hypothèse nulle, au profit de l'hypothèse alternative. L'hypothèse alternative est acceptée si l'hypothèse nulle est rejetée. Accepter une hypothèse revient donc à dire que l'hypothèse est non rejetée plutôt qu'acceptée.

7.1.2 Types d'erreurs

Tout test d'hypothèses est sujet à des erreurs. Ces dernières surviennent parce que la distribution de ces tests est asymptotique et que l'échantillon sur lequel travaille l'économètre est fini. L'inférence par un test d'hypothèses peut conduire à deux types d'erreur. L'erreur de type I (erreur de première espèce) survient lorsque le test d'hypothèses rejette une hypothèse nulle vraie tandis que l'erreur de type II (erreur de deuxième espèce) ne parvient pas à rejeter une hypothèse nulle fausse.

7.1.3 Niveau de signification d'un test

Le niveau de signification d'un test est la probabilité de l'erreur de première espèce. En théorie, les valeurs critiques d'un test d'hypothèses sont construites selon un niveau désiré. Les seuils statistiques conventionnels sont 1%, 5% ou 10 %. Par exemple, une valeur critique d'un niveau théorique de 5% signifie que la probabilité de rejeter une hypothèse nulle vraie se situe à 5% (on se donne 5 chances sur 100 de se tromper).

Remarque 1.5 : Tests sur les logiciels

Sur les logiciels économétriques, les statistiques des tests sont assorties de leurs probabilités critiques (risque de rejeter à tort l'hypothèse nulle H_0), ce qui évite de se référer aux tables.

La règle de décision suivante doit être appliquée :

- L'hypothèse nulle n'est pas rejetée dès que cette probabilité est supérieure ou égale au seuil α .
- L'hypothèse nulle est rejetée dès que cette probabilité est inférieure au seuil α .

7.2 Test du coefficient de corrélation linéaire

Nous disposons de deux variables quantitatives aléatoires X et Y . Le test du coefficient de corrélation linéaire s'écrit :

H_0 : les variables X et Y ne sont pas corrélées ($\rho = 0$)

H_1 : les variables X et Y sont corrélées ($\rho \neq 0$)

On calcule le coefficient de corrélation linéaire défini par :

$$r(x, y) = \frac{\text{Cov}(X, Y)}{\sigma_X \sigma_Y} = \frac{n \sum_{i=1}^n x_i y_i - \sum_{i=1}^n x_i \sum_{i=1}^n y_i}{\sqrt{n \sum_{i=1}^n x_i^2 - \left(\sum_{i=1}^n x_i \right)^2} \sqrt{n \sum_{i=1}^n y_i^2 - \left(\sum_{i=1}^n y_i \right)^2}}$$

avec :

- $\text{Cov}(X, Y)$ = covariance entre X et Y ;
- σ_X et σ_Y = écart type de X et écart type de Y ;
- n = nombre d'observations.

Pour tester l'hypothèse H_0 , on peut comparer la valeur observée de $r(x, y)$ à la valeur seuil A donnée par la table du coefficient de corrélation pour $v = n - 2$ degrés de liberté.

La règle de décision est alors la suivante :

- Si $|r(X, Y)| < A$, l'hypothèse H_0 n'est pas rejetée et les variables X et Y ne sont pas corrélées.
- Si $|r(x, y)| \geq A$, l'hypothèse H_0 est rejetée et les variables X et Y sont corrélées.

Application 1 : Test du coefficient de corrélation linéaire

Une entreprise commerciale consacre une certaine somme à des opérations publicitaires au début de chaque mois. Dans le tableau 1 sont récapitulés, mensuellement, pour 2012,

- les dépenses de publicité (X),
- les montants des ventes correspondantes (Y)

Les données sont en millions de FCFA.

TABEAU 1 : Dépenses de publicité et ventes

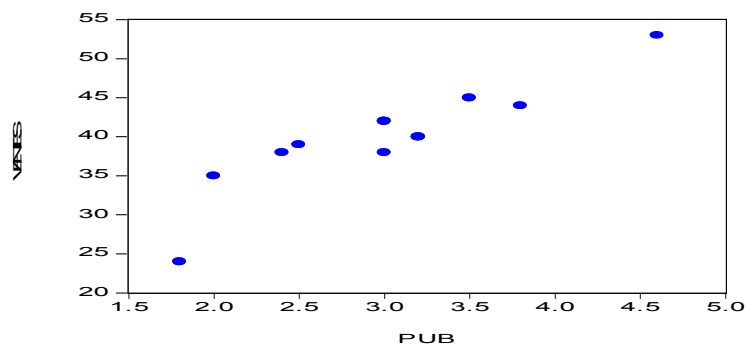
Mois	X	Y
Janvier	2,4	38
Février	3	42
...	...	...
Décembre	4,6	53

- 1) Tracer le nuage de points. Commenter le résultat obtenu.
- 2) Calculer le coefficient de corrélation linéaire.
- 3) Effectuer le test du coefficient de corrélation linéaire. Les ventes sont-elles liées aux dépenses de publicité ?

Corrigé :

- 1) Le nuage de points indique que les couples de valeurs sont approximativement alignés : les deux variables sont corrélées positivement.

GRAPHIQUE 1 : Nuage de points



2) Afin de calculer le coefficient de corrélation linéaire, nous dressons le tableau de calcul suivant :

i	X	Y	X ²	Y ²	XY
1	2,4	38	5,76	1444	91,2
2	3	42	9	1764	126
...	...	...	...	...	...
12	4,6	53	21,16	2809	243,8
Somme	36	480	114,58	19708	1492,4

$$r(x, y) = \frac{\text{Cov}(X, Y)}{\sigma_x \sigma_y} = \frac{n \sum_{i=1}^n x_i y_i - \sum_{i=1}^n x_i \sum_{i=1}^n y_i}{\sqrt{n \sum_{i=1}^n x_i^2 - \left(\sum_{i=1}^n x_i \right)^2} \sqrt{n \sum_{i=1}^n y_i^2 - \left(\sum_{i=1}^n y_i \right)^2}}$$

$$= \frac{12(1492,4) - (36)(480)}{\sqrt{(12)(114,58) - 36^2} \sqrt{(12)(19708) - 480^2}} = 0,906$$

3) Test du coefficient de corrélation linéaire

H_0 : Les variables publicité et ventes ne sont pas corrélées ($\rho = 0$)

H_1 : Les variables publicité et ventes Y sont corrélées ($\rho \neq 0$)

L'hypothèse H_0 signifie que les ventes sont indépendantes des dépenses de Publicité. Le tableau 2 indique les résultats du test du coefficient de corrélation linéaire.

TABLEAU 2 : Test du coefficient de corrélation linéaire, logiciel Eviews6

Included observations: 12		
Correlation		
Probability	PUB	VENTES
PUB	1.000	

VENTES	0.906	1.000
	0.0000	-----

Le coefficient de corrélation linéaire est 0,906 ; sa probabilité critique est nulle. Nous rejetons donc l'hypothèse nulle d'absence de corrélation linéaire au seuil de 1%. Les ventes sont donc liées aux dépenses de publicité.

7.3 Test de normalité de Jarque-Bera

7.3.1 Moment centré d'ordre r

On appelle moment centré d'ordre r ($r \in \mathbb{N}$) associé à la série X_t le nombre :

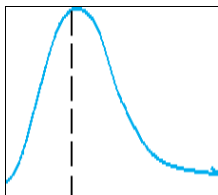
$$\mu_r(X) = \frac{\sum_{t=1}^n (X_t - \bar{x})^r}{n} \quad \text{avec} \quad \bar{x} = \frac{1}{n} \sum_{t=1}^n X_t$$
$$\mu_2 = \sigma^2 = \frac{1}{n} \sum_{t=1}^n (X_t - \bar{x})^2 \quad \mu_3 = \frac{1}{n} \sum_{t=1}^n (X_t - \bar{x})^3$$
$$\mu_4 = \frac{1}{n} \sum_{t=1}^n (X_t - \bar{x})^4$$

Jarque et Bera (1980) ont proposé un test portant sur les deux caractéristiques d'une loi normale : la symétrie et l'aplatissement. La loi normale est caractérisée par un coefficient d'asymétrie (skewness) nul et un coefficient d'aplatissement (kurtosis) égal à 3.

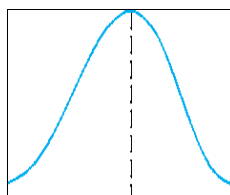
Le coefficient de skewness, noté S , associé à la série X_t est donné par :

$$S = \frac{\mu_3}{\sigma^3} = \frac{\mu_3}{\mu_2^{3/2}}$$

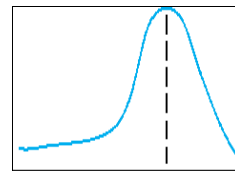
Si le skewness est nul, la distribution est dite symétrique. Lorsque le skewness est non nul, la distribution est dite asymétrique. Plus spécifiquement si $S < 0$, la distribution est étalée vers la gauche (elle a un biais négatif). Si $S > 0$, la distribution est étalée vers la droite (elle a un biais positif).



Distribution étalée
vers la droite



Distribution symétrique

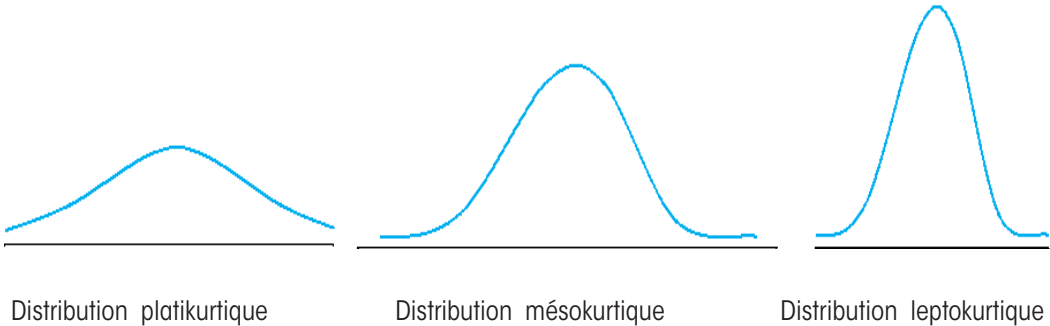


Distribution étalée
vers la gauche

Le coefficient d'aplatissement d'une courbe est caractérisé par la valeur de :

$$K = \frac{\mu_4}{\sigma^4} = \frac{\mu_4}{\mu_2^2}$$

Si $K = 3$, la distribution est normale. Lorsque $K < 3$, la distribution est plus aplatie que la distribution normale (on dit qu'elle est hyponormale ou platykurtique). Si $K > 3$, la distribution présente un excès de kurtosis. Elle est moins aplatie que la distribution normale (on dit qu'elle est hypernormale ou leptokurtique).



7.3.2 Test de Jarque-Bera

Le test d'hypothèses est le suivant :

H_0 : la variable X suit une loi normale $N(m, \sigma)$

H_1 : la variable X ne suit pas une loi normale $N(m, \sigma)$

La statistique de Jarque-Bera est définie par :

$$JB = n \left[\frac{S^2}{6} + \frac{(K - 3)^2}{24} \right]$$

où S est le coefficient de dissymétrie (Skewness) et K le coefficient d'aplatissement (Kurtosis). La statistique JB suit sous l'hypothèse nulle de normalité une loi du Khi-Deux à deux degrés de liberté.

La valeur seuil A pour 5% est égale à 5,991. Par conséquent, si la valeur calculée de la statistique JB est inférieure à A , l'hypothèse nulle de normalité n'est pas rejetée. En revanche, si la valeur de JB est supérieure ou égale à A , l'hypothèse nulle de normalité est rejetée.

Remarque 1.6 : Le test de normalité est utile sur des échantillons de petite dimension mais l'est moins sur de grands échantillons où le « théorème central limite »³ s'applique. Ce test permet notamment de tester l'hypothèse de normalité des résidus.

Application 2 : Test de normalité de Jarque-Bera

Considérons la série macroéconomique masse monétaire du Sénégal sur la période 1971-2007 à fréquence annuelle. Les données sont en milliards de FCFA. Tester la normalité et la lognormalité de la variable masse monétaire.

TABEAU 3 : Masse monétaire du Sénégal

Année	m2
1971	2,4
1972	3
...	...
2007	4,6

Corrigé :

Le tableau 4 présente les résultats du test de Jarque-Bera.

TABEAU 4 : Résultats du test de normalité de Jarque et Bera

	m2	log(m2)
Mean	522.93	5.82
Median	336.52	5.82
Std. Dev.	497.43	1.01
Skewness	1.49	-0.31
Kurtosis	4.34	2.66
Jarque-Bera	16.47	0.77
Probability	0.00	0.68
Observations	37	37

1) Test de normalité

La statistique de Jarque-Bera associée à la variable m2 est 16,47 ; sa probabilité critique est nulle. L'hypothèse nulle de normalité est rejetée. La variable masse monétaire ne suit pas une loi normale.

3 - Le théorème central limite affirme d'une suite de variables aléatoires $(X_1, X_2, \dots, X_n)$ indépendantes et de même loi tend vers une loi normale si le nombre d'observations n est supérieur à 30. Les variables X_i suivent une loi quelconque.

2) Test de lognormalité

Une variable suit une loi lognormale si son logarithme suit une loi normale. La statistique de Jarque-Bera associée à la variable $\log(m_2)$ est 0,77 ; sa probabilité critique est égale à 0,68. L'hypothèse nulle de normalité n'est pas rejetée. La variable masse monétaire suit une loi lognormale de moyenne 5,82 milliards de FCFA et d'écart type 1,01 milliard de FCFA.

Remarque 1.7 : La probabilité critique associée au test de lognormalité vaut 0,68. Autrement dit, si nous rejetons l'hypothèse de lognormalité de la variable masse monétaire, nous aurons 68% de chances de prendre une mauvaise décision.

7.4 Test de Student

7.4.1 Comparaison d'un paramètre a_i à une valeur fixée a

Le test d'hypothèses est le suivant : $H_0 : a_i = a$ contre $H_1 : a_i \neq a$.

Soit $\hat{\sigma}_{\hat{A}_i}$ l'écart type estimé des estimateurs des coefficients de régression. La statistique $\frac{\hat{A}_i - a_i}{\hat{\sigma}_{\hat{A}_i}}$ suit sous l'hypothèse nulle une loi de Student à $(n - k)$ degrés de liberté notée T_{n-k} . On détermine au seuil α , une valeur critique t_α telle que :

$$P\{-t_\alpha \leq T_{n-k} \leq t_\alpha\} = 1 - \alpha$$

La règle de décision est alors la suivante :

- Si $q = \left| \frac{\hat{A}_i - a}{\hat{\sigma}_{\hat{A}_i}} \right| < t_\alpha$, l'hypothèse $H_0 : a_i = a$ n'est pas rejetée,
- Si $q = \left| \frac{\hat{A}_i - a}{\hat{\sigma}_{\hat{A}_i}} \right| \geq t_\alpha$, l'hypothèse $H_0 : a_i = a$ est rejetée.

En pratique, le test le plus utilisé est celui qui consiste à tester l'hypothèse nulle $H_0 : a_i = 0$ contre l'hypothèse alternative $H_1 : a_i \neq 0$. Il s'agit du test de significativité des variables explicatives.

7.4.2 Comparaison d'un paramètre a_i à la valeur $a = 0$

Quelle sont, parmi les variables X_i choisies, celles qui sont réellement explicatives, celles qui ont une influence significative au seuil α . Le test d'hypothèses est $H_0 : a_i = 0$ contre $H_1 : a_i \neq 0$. Les décisions associées étant d'exclure ou de conserver la variable X_i dans le modèle. Dans ce cas, la règle de décision est :

- Si $q = \left| \frac{\hat{A}_i}{\hat{\sigma}_{\hat{A}_i}} \right| < A$, l'hypothèse $H_0 : a_i = 0$ n'est pas rejetée. La variable explicative X_{it} n'est pas significative et n'a aucune influence sur Y_t . On l'élimine du modèle.
- Si $q = \left| \frac{\hat{A}_i}{\hat{\sigma}_{\hat{A}_i}} \right| \geq t_\alpha$, l'hypothèse $H_0 : a_i = 0$ est rejetée. La variable explicative X_{it} est significative et a une influence sur Y_t . On la conserve dans le modèle.

Remarque 1.8 : Le nombre $t = \frac{\hat{A}_i}{\hat{\sigma}_{\hat{A}_i}}$ est appelé ratio de Student. Sur les logiciels, les nombres $\frac{\hat{A}_i}{\hat{\sigma}_{\hat{A}_i}}$ et $\hat{\sigma}_{\hat{A}_i}$ sont notés respectivement t-Statistic et Standard-Error.

7.4.3 Test de l'hypothèse $la = s$

On désire tester l'hypothèse nulle $H_0 : la = s$ contre l'hypothèse alternative $H_1 : la \neq s$ avec

$$(l_{(1,k)}) = (l_1 \ l_2 \ \dots \ l_k), \quad a_{(k,1)} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_k \end{pmatrix}, \quad la = l_1 a_1 + l_2 a_2 + \dots + l_k a_k$$

Dans ce cas, on démontre que la variable aléatoire $\frac{l\hat{A} - la}{\hat{\sigma}_{l\hat{A}}}$ suit une loi Student à $(n - k)$ degrés de liberté. La règle de décision est :

- Si $q = \left| \frac{\hat{A} - s}{\hat{\sigma}_{\hat{A}_i}} \right| < t_{\alpha}$, l'hypothèse $H_0 : a = s$ n'est pas rejetée.
- Si $q = \left| \frac{\hat{A} - s}{\hat{\sigma}_{\hat{A}_i}} \right| \geq t_{\alpha}$, l'hypothèse $H_0 : a = s$ est rejetée.

7.4.4 Intervalle de confiance pour les paramètres a_i

On cherche un intervalle I tel que :

$$P[a_i \in I] = 1 - \alpha$$

I est appelé intervalle de confiance pour le paramètre a_i , au niveau de confiance $1 - \alpha$. Le réel α ($0 < \alpha < 1$) peut être interprété comme le risque que l'intervalle de confiance I ne contienne pas la vraie valeur du paramètre. On utilise le résultat établi précédemment :

$$\frac{\hat{A}_i - a_i}{\hat{\sigma}_{\hat{A}_i}} \rightarrow T_{n-k}$$

On a alors :

$$P[\hat{A}_i - t_{\alpha} \hat{\sigma}_{\hat{A}_i} \leq a_i \leq \hat{A}_i + t_{\alpha} \hat{\sigma}_{\hat{A}_i}] = 1 - \alpha$$

$I = [\hat{A}_i - t_{\alpha} \hat{\sigma}_{\hat{A}_i} ; \hat{A}_i + t_{\alpha} \hat{\sigma}_{\hat{A}_i}]$ est l'intervalle de confiance du paramètre a_i au niveau de confiance $1 - \alpha$. L'intervalle prend aussi la forme $I = [\hat{A}_i - r ; \hat{A}_i + r]$ où $r = t_{\alpha} \hat{\sigma}_{\hat{A}_i}$ est le rayon de l'intervalle de confiance.

Remarque 1.9 : L'intervalle de confiance peut être utilisé pour effectuer le test de Student $H_0 : a_i = a$ contre $H_1 : a_i \neq a$.

- Si $a \in [\hat{A}_i - t_{\alpha} \hat{\sigma}_{\hat{A}_i} ; \hat{A}_i + t_{\alpha} \hat{\sigma}_{\hat{A}_i}]$, l'hypothèse $H_0 : a_i = a$ n'est pas rejetée.
- Si $a \notin [\hat{A}_i - t_{\alpha} \hat{\sigma}_{\hat{A}_i} ; \hat{A}_i + t_{\alpha} \hat{\sigma}_{\hat{A}_i}]$, l'hypothèse $H_0 : a_i = a$ est rejetée.

7.5 Test de significativité globale d'une régression (Fisher)

Dans ce qui suit, nous allons nous interroger sur la signification globale du modèle de régression, c'est-à-dire si l'ensemble des variables explicatives a une influence sur la variable à expliquer Y. Soit le test d'hypothèses :

$$H_0 : a_2 = a_3 = \dots = a_k = 0$$

contre

H_1 : il existe au moins un des coefficients non nul.

Le modèle à tester est $Y_t = a_1 + \varepsilon_t$. Pour le test de signification globale, on démontre que la statistique de Fisher est égale :

$$F^* = \frac{n-k}{k-1} \times \frac{R^2}{1-R^2}$$

Pour effectuer le test de signification globale, nous allons comparer ce F^* calculé au F_{lu} sur la table de Fisher à respectivement $(k-1)$ et $(n-k)$ degrés de liberté. La règle de décision est :

- Si $F^* = \frac{n-k}{k-1} \times \frac{R^2}{1-R^2} < F_{lu}$, nous ne rejetons pas H_0 : le modèle n'est pas globalement significatif au seuil α .
- Si $F^* = \frac{n-k}{k-1} \times \frac{R^2}{1-R^2} \geq F_{lu}$, nous rejetons l'hypothèse H_0 : le modèle est globalement significatif au seuil α .

Application 3 : Estimation des paramètres par les MCO, tests de Student et de Fisher

On dispose pour une entreprise et sur la période 1974 à 2012, des séries production, (Q), facteur capital (K) et facteur travail (L). Nous souhaitons estimer la fonction de production de cette entreprise en utilisant une spécification de type Cobb-Douglas.

TABLEAU 5 : Production, capital et travail

Année	Q	K	L
1974	189.8	87.8	173.3
1975	172.1	87.8	165.4
...	...	...	...
2012	631.1	247.9	305

1) L'expression générale de la fonction de production Cobb-Douglas s'écrit $Q_t = AK_t^{\beta_2} L_t^{\beta_3}$. Pourquoi ce modèle n'est pas estimable économétriquement par les moindres carrés ordinaires ?

2) On considère la forme estimable suivante :

$$\log(Q_t) = \beta_1 + \beta_2 \log(K_t) + \beta_3 \log(L_t) + \varepsilon_t$$

Estimer les paramètres du modèle de Cobb-Douglas par la méthode des moindres carrés ordinaires

a) en utilisant le logiciel Eviews,

b) en utilisant le logiciel Stata,

c) en utilisant le logiciel SPSS

3) Comment s'interprètent économiquement les différents paramètres ?

4) Comment s'interprète la perturbation ε_t du modèle économétrique ?

5) Donner une interprétation du coefficient de détermination R^2 .

6) Testez la significativité des paramètres β_2 et β_3 .

7) Testez l'hypothèse $H_0 : \beta_2 = 0,4$ contre $H_1 : \beta_2 \neq 0,4$.

8) Donnez, au seuil de 5%, des intervalles de confiance pour les paramètres β_2 et β_3 .

9) Déterminer l'importance relative des déterminants de la production en utilisant les coefficients standardisés.

10) Testez la significativité globale du modèle.

11) Testez l'hypothèse selon laquelle les rendements d'échelle constatés pour cette entreprise sont constants.

Corrigé :

1) Le modèle $Q_t = AK_t^{\beta_2} L_t^{\beta_3}$ n'est pas linéaire dans les paramètres, il n'est donc pas estimable économétriquement par les moindres carrés ordinaires. Pour linéariser le modèle, on passe aux logarithmes. D'où l'expression :

$$\log(Q_t) = \beta_1 + \beta_2 \log(K_t) + \beta_3 \log(L_t) + \varepsilon_t$$

2) Estimation des paramètres par les logiciels Eviews, Stata et SPSS

a) L'estimation des paramètres par les logiciels Eviews, Stata et SPSS est donnée respectivement dans les tableaux 6, 7 et 8. Nous remarquons que les résultats d'estimation donnés par les trois logiciels sont identiques.

Les paramètres estimés sont :

$$\hat{\beta}_1 = -3,932 ; \hat{\beta}_2 = 0,385 ; \hat{\beta}_3 = 1,448$$

Outre l'estimation des paramètres des trois coefficients, nous disposons de plusieurs résultats. En particulier, les tableaux d'estimation nous donnent les écarts-types (Standard. Error), les t de Student (t-Statistic) et la statistique de Fisher (F-statistic), le coefficient de détermination, l'intervalle de confiance associé aux paramètres, les coefficients standardisés, ...

TABLEAU 6 :Fonction de Cobb-Douglas.Estimation par les MCO (Logiciel Eviews)

Dependent Variable : LOG(Q) ; Number of observadtions : 39				
Variable	Coef	Standard. Error	t-Statistic	Prob.
C	-3.9321	0.2372	-16.5755	0.0000
LOG(K)	0.3851	0.0480	8.0128	0.0000
LOG(L)	1.4485	0.0833	17.3920	0.0000
R-squared	0.9946	Mean dependent var		5.6874
Adjusted R-squared	0.9943	S.D. dependent var		0.4609
S.E. of regression	0.0347	Akaike info criterion		-3.8055
Sum squared resid	0.0435	Schwarz criterion		-3.6775
Log likelihood	77.2072	Hannan-Quinn criter.		-3.7595
F-statistic	3318.670	Durbin-Watson stat		0.8508
Prob(F-statistic)	0.0000			

b) Logiciel Stata

TABLEAU 7 :Fonction de Cobb-Douglas.Estimation par les MCO (Logiciel Stata)

Source	SS	df	MS	Number of obs = 39		
Model	8.03080693	2	4.01540346	F(2, 36)	=	3318.66
Residual	.043558105	36	.001209947	Prob > F	=	0.0000
Total	8.07436503	38	.21248329	R-squared	=	0.9946
				Adj R-squared	=	0.9943
				Root MSE	=	.03478

lq	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
lk	.3850685	.0480565	8.01	0.000	.2876053	.4825317
ll	1.448587	.0832905	17.39	0.000	1.279666	1.617507
_cons	-3.9321	.2372232	-16.58	0.000	-4.413211	-3.450989

c) Logiciel SPSS

TABLEAU 8 : Fonction de Cobb-Douglas. Estimation par les MCO (Logiciel SPSS)

	Coefficients non standardisés	Coefficients standardisés		
	B	Bêta	t-statistic	prob
(constante)	-3,932		-16,576	0,000
LK	0,385	0,318	8,013	0,000
LL	1,449	0,690	17,392	0,000
Variable dépendante : LQ				
R-deux : 0,995				
R-deux ajusté : 0,994				

3) Interprétation économique des coefficients β_2 et β_3

Comme le modèle est spécifié sous forme log-log, les coefficients β_2 et β_3 sont des élasticités.

-- $\beta_1 = \log(A)$ n'a pas d'interprétation économique intéressante.

-- β_2 est l'élasticité de la production par rapport au capital

$$\beta_2 = \frac{\partial \text{Log} Y}{\partial \text{Log} K} = \frac{\partial Y / Y}{\partial K / K} = e_{Y/K}$$

$\hat{\beta}_2 = 0,385$ signifie que l'on estime qu'une augmentation de 10% du capital implique une augmentation de 3,85% de la production, toutes choses égales par ailleurs. (c'est à dire si le facteur travail reste constant).

-- β_3 est l'élasticité de la production par rapport au travail

$$\beta_3 = \frac{\partial \text{Log} Y}{\partial \text{Log} L} = \frac{\partial Y / Y}{\partial L / L} = e_{Y/L}$$

$\hat{\beta}_3 = 1,448$ signifie que l'on estime qu'une augmentation de 10% du travail implique une augmentation de 14,48% de la production, toutes choses égales par ailleurs. (c'est à dire si le facteur capital reste constant).

4) La perturbation ε_t représente l'effet des pannes de machines, de l'état de santé et de motivation des travailleurs, etc.

5) La valeur du coefficient de détermination est $R^2 = 0,9946$. Ainsi 99,46% des fluctuations de la production sont expliquées par les facteurs capital et travail. On peut aussi dire que les facteurs capital et travail expliquent 99,46% de la variance de la production. Il faut noter que la valeur du coefficient de détermination n'indique rien sur la qualité du modèle.

6) Tests de significativité des variables explicatives

a) On veut tester $H_0 : \beta_2 = 0$ contre $H_1 : \beta_2 \neq 0$.

Le ratio de Student associé au coefficient estimé $\hat{\beta}_2$ est égal à 8,012 ; sa probabilité critique est nulle. On rejette l'hypothèse H_0 au seuil de 1%. Le capital a un impact positif significatif sur la production.

b) On veut tester $H_0 : \beta_3 = 0$ contre $H_1 : \beta_3 \neq 0$. Le ratio de Student associé au coefficient estimé $\hat{\beta}_3$ est égal à 17,392 ; sa probabilité critique est nulle. On rejette l'hypothèse H_0 au seuil de 1%. Le travail a un impact positif significatif sur la production.

7) On veut tester $H_0 : \beta_2 = 0,4$ contre $H_1 : \beta_2 \neq 0,4$.

Le test de Wald effectué sur le logiciel Eviews donne le résultat suivant :

Wald Test: c(2) = 0.4			
Test Statistic	Value	df	Probability
F-statistic	0.096540	(1, 36)	0.7578
Chi-square	0.096540	1	0.7560

Les deux probabilités sont supérieures aux seuils statistiques conventionnels. Nous ne rejetons pas l'hypothèse $H_0 : \beta_2 = 0,4$. L'élasticité de la production par rapport au capital est égale à 0,4.

8) Intervalles de confiance

Le logiciel Stata indique les intervalles de confiance à 95% pour chaque paramètre. Les intervalles sont indiqués par la colonne [95% Conf. Interval].

a) L'intervalle de confiance à 95% de β_2 est $[0,2876 ; 0,4825]$.

On remarque que la valeur zéro n'appartient pas à l'intervalle de confiance. L'hypothèse $H_0 : \beta_2 = 0$ est rejetée. Le capital a un impact positif significatif sur la production.

On remarque aussi que la valeur 0,4 appartient à l'intervalle de confiance. Nous ne rejetons pas l'hypothèse $H_0 : \beta_2 = 0,4$. L'élasticité de la production par rapport au capital est égale à 0,4.

b) L'intervalle de confiance à 95% de β_3 est $[1,2796 ; 1,6175]$.

On remarque que la valeur zéro n'appartient pas à l'intervalle de confiance. L'hypothèse $H_0 : \beta_3 = 0$ est aussi rejetée. Le travail a un impact positif significatif sur la production.

9) Calcul des coefficients standardisés.

L'importance relative des déterminants de la production est évaluée à travers les coefficients standardisés⁴. De deux variables explicatives, la plus influente est celle qui transmet une proportion plus grande de sa variabilité à la variable endogène. Les coefficients standardisés sont calculés par le logiciel SPSS. Les coefficients standardisés associés aux facteurs capital et travail sont respectivement 0,318 et 0,69. La production est donc plus influencée par le facteur travail que par le facteur production.

10) Test de Fisher de significativité globale du modèle

Nous voulons tester

$$H_0 : \beta_2 = \beta_3 = 0$$

contre

H_1 : il existe au moins un coefficient différent de 0

La statistique de Fisher est :

$$F^* = \frac{n-k}{k-1} \times \frac{R^2}{1-R^2} = 3318,67$$

Sa probabilité critique est nulle : nous rejetons l'hypothèse nulle au seuil de 1%. Nous concluons donc que le modèle est globalement significatif. Il est préférable au modèle $\log(Q_t) = \beta_1 + \varepsilon_t$. Le capital et le travail ont globalement un effet significatif sur la production.

4 - Pour obtenir le coefficient standardisé d'une variable explicative donnée, on multiplie, dans une première étape, le coefficient estimé de ladite variable par son écart type. Le résultat ainsi obtenu est rapporté, dans une seconde phase, à l'écart type de la variable endogène, pour obtenir le coefficient standardisé (Pindick et Rubinfeld, page 85).

11) $\beta_2 + \beta_3$ représente les rendements d'échelle : de combien varie la production si on fait varier de la même façon tous les inputs.

Les rendements d'échelle sont décroissants si $\beta_2 + \beta_3 < 1$, la production augmente dans une proportion moindre que les facteurs de production,

Les rendements d'échelle sont constants si $\beta_2 + \beta_3 = 1$, la production augmente dans une proportion identique aux facteurs de production,

Les rendements d'échelle sont croissants si $\beta_2 + \beta_3 > 1$, la production augmente plus vite que les facteurs de production.

Tester l'hypothèse selon laquelle les rendements d'échelle constatés pour cette entreprise sont constants consiste à tester l'hypothèse nulle $H_0 : \beta_2 + \beta_3 = 1$ contre l'hypothèse alternative $H_1 : \beta_2 + \beta_3 \neq 1$.

Ce test effectué sur les logiciels Eviews et Stata donne :

Logiciel Eviews Wald Test : $c(2) + c(3) = 1$			
Test Statistic	Value	df	Probability
F-statistic	425.8555	(1, 36)	0.0000

Logiciel Stata Wald Test : $lk + ltrav = 1$	
Test Statistic	Value
F(1,36)	425.85
Prob > F =	0.0000

La statistique de Fisher est égale à 425,85, sa probabilité critique est nulle. Nous rejetons au seuil de 1% l'hypothèse H_0 de rendements d'échelle constants. Ici, les rendements d'échelle sont croissants : $\beta_2 + \beta_3 > 1$. La production augmente plus vite que les facteurs de production. Le test de cette hypothèse est un test unilatéral.

7.6 Tests d'autocorrélation des erreurs

7.6.1 Définition et causes de l'autocorrélation des erreurs

Nous sommes en présence d'une autocorrélation des erreurs lorsque les erreurs sont liées par un processus à mémoire, donc non indépendantes au cours du temps.

L'autocorrélation des erreurs peut être observée pour plusieurs raisons :

- i) L'absence d'une ou plusieurs variables explicatives dont l'explication résiduelle permettrait de « blanchir » les erreurs.
- ii) Une mauvaise spécification du modèle, les relations entre la variable à expliquer et les variables explicatives ne sont pas linéaires et s'expriment sous une autre forme que celle du modèle estimé (logarithmes, différences premières, etc.).

7.6.2 Estimation en présence d'autocorrélation des erreurs

En présence d'autocorrélation des erreurs, les estimateurs des moindres carrés ordinaires sont sans biais mais ne sont plus de variance minimale. Dans cette situation, on utilise la méthode des moindres carrés généralisés.

7.6.3 Les tests d'autocorrélation des erreurs

7.6.3.1 Test de Durbin et Watson (1951)

Le test de Durbin et Watson (DW) permet de détecter une autocorrélation des erreurs d'ordre un selon la forme :

$$\varepsilon_t = \rho \varepsilon_{t-1} + v_t \quad (1)$$

avec $v_t \rightarrow N(0, \sigma^2)$

Le test d'hypothèses est le suivant :

$H_0 : \rho = 0$ (les erreurs ne sont pas autocorrélées)

contre

$H_1 : \rho \neq 0$ (les erreurs sont autocorrélées)

Considérons le modèle linéaire général $Y = Xa + \varepsilon$, soit $e = Y - X\hat{A}$, le résidu calculé. On suppose que le modèle comporte une constante. La statistique de Durbin et Watson est définie par :

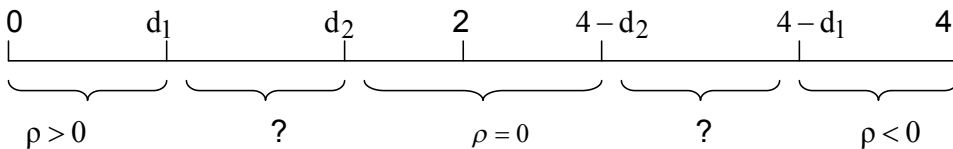
$$DW = \frac{\sum_{t=2}^n (e_t - e_{t-1})^2}{\sum_{t=1}^n e_t^2}$$

où e_t est le résidu calculé au temps t . Pour un grand nombre d'observations c'est-à-dire quand n tend vers $+\infty$, on démontre que $DW \approx 2(1 - \hat{\rho})$ où $\hat{\rho}$ est un coefficient de corrélation linéaire avec $-1 \leq \hat{\rho} \leq 1$.

De par sa construction, la statistique DW varie entre 0 et 4. Ce qui donne $0 \leq DW \leq 4$. Il est intéressant d'avoir en tête l'approximation $DW \approx 2(1 - \hat{\rho})$ qui permet de se faire une idée de la valeur de $\hat{\rho}$ à partir de DW. Ainsi, il est probable que $\rho = 0$ quand $DW \approx 2$. En revanche, quand la valeur de DW est proche de 0, les erreurs sont probablement autocorrélées, avec un $\hat{\rho}$ positif et proche de 1. Quand, à l'inverse, DW est proche de 4, les erreurs sont vraisemblablement affectées d'une forte autocorrélation négative.

Venons-en maintenant à la procédure du test. Les valeurs critiques du test Durbin et Watson tenant compte du nombre d'observations et du nombre de variables explicatives ont été calculées.

La lecture de la table de Durbin et Watson permet de déterminer deux valeurs $d_1 = d_{inf}$ et $d_2 = d_{sup}$ comprises entre 0 et 2 qui délimitent l'espace entre 0 et 4 selon le schéma ci-dessous :



Selon la position du DW empirique dans cet espace, nous pouvons conclure :

- Si $d_2 < DW < 4 - d_2$, on ne rejette pas l'hypothèse H_0 ($\hat{\rho} = 0$), les résidus sont non corrélés.
- Si $0 < DW < d_1$ on rejette l'hypothèse H_0 ($\hat{\rho} > 0$), les résidus sont autocorrélés positivement.
- Si $4 - d_1 < DW < 4$, on rejette l'hypothèse H_0 ($\hat{\rho} < 0$), les résidus sont autocorrélés négativement.

Si $d_1 < DW < d_2$ ou $4 - d_2 < DW < 4 - d_1$, on est dans une zone d'indétermination, ou zone de doute, c'est-à-dire qu'on ne peut pas conclure dans un sens comme dans l'autre.

Remarque 1.10 : Lorsque que la statistique de Durbin et Watson appartient à la région de doute, les économètres conseillent de rejeter l'hypothèse de non autocorrélation des erreurs.

Le test de Durbin et Watson est très fréquemment utilisé. Il est cependant important de préciser ses conditions d'utilisation :

- Le modèle de régression doit comporter impérativement un terme constant ;
- Les variables explicatives doivent être certaines (c'est à dire non aléatoires) ;
- Le nombre d'observations doit être supérieur ou égal à 15 ;
- Le modèle doit être en série temporelle, pour les modèles en coupe instantanée, les observations doivent être ordonnées en fonction de la variable à expliquer ;
- La variable à expliquer ne doit pas figurer parmi les variables explicatives (en tant que variable retardée). Le modèle ne doit pas être autorégressif, si c'est le cas on doit utiliser le test de corrélation des erreurs de Breusch et Godfrey (1978).

Remarque 1.11 : Le test de Durbin et Watson permet uniquement de détecter l'autocorrélation d'ordre 1 des résidus. En d'autres termes, il n'est pas approprié si les résidus présentent de l'autocorrélation à un ordre supérieur à 1.

7.6.3.2 Test de Breusch et Godfrey (1978)

Breusch et Godfrey ont proposé un test permettant de déceler la présence d'autocorrélation d'ordre supérieur à 1 restant valide lorsque le modèle comporte la variable endogène retardée dans les variables explicatives⁵. Ce test est fondé sur un test de Fisher de nullité des coefficients (F-statistic) ou du Multiplicateur de Lagrange (nR^2). L'idée générale de ce test réside dans la recherche d'une relation significative entre le résidu et ce même résidu décalé.

Considérons le modèle linéaire général à erreurs autocorrélées d'ordre p :

$$Y_t = a_1 X_{1t} + a_2 X_{2t} + K + a_k X_{kt} + a_0 + \rho_1 \varepsilon_{t-1} + \rho_2 \varepsilon_{t-2} + K + \rho_p \varepsilon_{t-p} + v_t$$

5- Les modèles linéaires autorégressifs sont étudiés dans le chapitre 2.

Le test de Breusch et Godfrey consiste à tester l'hypothèse H_0 d'absence d'autocorrélation des erreurs, soit $H_0 : \tilde{\alpha}_1 = \tilde{\alpha}_2 = \dots = \tilde{\alpha}_p = 0$. Ce test est mené en trois étapes :

i) On estime par la méthode des moindres carrés ordinaires les paramètres du modèle ci-dessus et on calcule le résidu e_t , puisque les erreurs $\hat{a}_t$ sont inconnues.

ii) On estime par la méthode des moindres carrés ordinaires l'équation intermédiaire

$$e_t = a_1 X_{1t} + a_2 X_{2t} + \dots + a_k X_{kt} + a_0 + \rho_1 e_{t-1} + \rho_2 e_{t-2} + \dots + \rho_p e_{t-p} + v_t$$

et on calcule le coefficient de détermination R^2 associé à cette régression.

iii) On calcule la statistique du multiplicateur de Lagrange de Breusch et Godfrey définie par $BG = nR^2$. Cette statistique suit sous l'hypothèse H_0 de non autocorrélation des erreurs une loi du Khi-Deux à p degrés de liberté. On se fixe un seuil α , et on lit sur la table du Khi-Deux la valeur seuil A .

La règle de décision est la suivante :

Si $BG < A$, l'hypothèse nulle d'absence d'autocorrélation n'est pas rejetée.

Si $BG \geq A$, l'hypothèse nulle d'absence d'autocorrélation est rejetée : au moins un des coefficients ρ_i , $i = 1, \dots, p$, est significativement différent de zéro.

7.6.3.3 Procédures d'estimation en cas d'autocorrélation des erreurs

En cas d'autocorrélation des erreurs, les estimateurs des MCO restent sans biais, mais ne sont plus de variance minimale. Dans ce cas, si la variance du terme d'erreur est connue, on applique la méthode des moindres carrés généralisés. Lorsque la variance du terme d'erreur est inconnue, on utilise les méthodes numériques de Cochrane-Orcutt, de Hildreth-Lu et du maximum de vraisemblance.

Les programmes de régression utilisent fréquemment la méthode de Cochrane-Orcutt. Sargan a montré que la procédure itérative de Cochrane-Orcutt est convergente c'est-à-dire qu'au bout d'un certain nombre d'itérations les valeurs des coefficients vont se stabiliser.

Application 4 : Tests d'autocorrélation des erreurs

On considère la fonction d'importation du Sénégal, à fréquence annuelle sur la période allant de 1962 à 1995.:

$$(1) \log(\text{import}_t) = a + b \log(\text{pib}_t) + \varepsilon_t, t = 1, 2, \dots, 34, n = 34$$

où import désigne les importations, pib le produit intérieur brut et log le logarithme népérien.

TABLEAU 8 : Importations et PIB du Sénégal

Années	Import	PIB
1962	311	761,3
1963	309,1	791,1
...	...	...
1995	466,2	1599,4

Partie 1 : Modèle estimé par les moindres carrés ordinaires

- 1) Estimer les paramètres du modèle (1) par la méthode des moindres carrés
- 2) Tester une éventuelle autocorrélation des erreurs.
 - a) par le test de Durbin-Watson,
 - b) par le test de Breusch-Godfrey.

Partie 2 : Modèle estimé par la méthode de Cochrane-Orcutt

- 3) Estimer les paramètres par la méthode de Cochrane-Orcutt.
- 4) Tester une éventuelle autocorrélation des erreurs du modèle corrigé par la méthode de Cochrane-Orcutt en utilisant le test de Breusch-Godfrey.

Corrigé :

Partie 1

- 1) L'estimation des paramètres par la méthode des MCO conduit aux résultats suivants :

$$\text{Log IMPORT}_t = 0,65 + 0,75 \text{ Log PIB}_t$$

(1,87) (15,19)

$$R^2 = 0,87, \quad \text{SCR} = 0,134, \quad \text{DW} = 1,06, \quad n = 34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés.

2) Tests d'autocorrélation des erreurs

a) Test de Durbin-Watson

Les tests d'hypothèses sont $H_0 : \rho = 0$ contre $H_1 : \rho \neq 0$

Les conditions d'utilisation du test de Durbin et Watson sont bien respectées : le modèle est spécifié en série temporelle, le nombre d'observations ($n = 34$) est supérieur à 15 et, enfin le modèle comporte un terme constant. Le nombre de variables exogènes (constante exclue) est $k = 1$. Sur la table de Durbin et Watson, au seuil de 5%, on lit les valeurs suivantes :

$$\begin{array}{ll} d_1 = 1,39 & d_2 = 1,51 \\ 4 - d_1 = 2,61 & 4 - d_2 = 2,49 \end{array}$$

On a :

$$0 < DW < d_1$$

$$0 < 1,06 < 1,39$$

L'hypothèse nulle est rejetée : on peut donc présumer une autocorrélation positive des erreurs.

Le test de Durbin-Watson effectué sur le logiciel STATA donne les résultats suivants :

Durbin's alternative test for autocorrelation

lags (p)	chi2	df
Prob > chi2		
-----+-----		
1	6.592	1
0.0102		

H0: no serial correlation

Durbin-Watson d-statistic(2,34) =1.064825

La probabilité critique est égale à 0,0102. On rejette l'hypothèse de corrélation des erreurs au seuil de 5%.

b) Le test de Breusch-Godfrey conduit aux résultats suivants :

Breusch-Godfrey Serial Correlation LM Test			
lag : 1			
F-statistic	6.591775	Probability	0.015291
Obs*R-squared	5.961952	Probability	0.014618

Le nombre de décalages (Lag) sur le terme d'erreur est un. On teste donc une autocorrélation d'ordre un des résidus. Le logiciel Eviews effectue 2 tests : le test de Fisher (F-statistic) et le test du multiplicateur de Lagrange (obs*R-squared = nR^2). Les deux probabilités sont inférieures au seuil de 5%. On rejette l'hypothèse de non corrélation des erreurs pour les 2 tests. Les erreurs sont donc autocorrélées.

Partie 2

3) Les erreurs du modèle étant corrélées positivement, l'estimation par la méthode de Cochrane et Orcutt conduit aux résultats suivants :

$$\text{Log IMPORT}_t = 0,47 + 0,77 \text{Log PIB}_t + 0,41 \text{ar} (1) \quad (2)$$

(1,88)
(10,20)
(2,64)

$$R^2 = 0,90 \quad \text{SCR} = 0,09 \quad \text{DW} = 2,18 \quad \hat{\rho} = 0,41$$

Convergence assurée après 6 itérations.

4) Le test de Breusch-Godfrey effectué sur le modèle corrigé par la méthode de Cochrane-Orcutt conduit aux résultats suivants :

Breusch-Godfrey Serial Correlation LM Test:			
lag : 1			
F-statistic	1.182139	Probability	0.285875
Obs*R-squared	1.292506	Probability	0.255587

Les deux probabilités sont maintenant supérieures aux seuils conventionnels. L'hypothèse nulle de non corrélation des erreurs n'est pas rejetée. Le test BG effectué pour les décalages 2 et 3 aboutit au même résultat. **La méthode de Cochrane-Orcutt a donc corrigé l'autocorrélation des erreurs.**

7.7 Tests d'hétéroscédasticité des erreurs

7.7.1 Définition

Les erreurs sont homocédastiques si elles ont la même variance. Dans le cas contraire, on dit qu'elles sont hétéroscédastiques. L'hétéroscédasticité se rencontre plus fréquemment pour les modèles en coupe instantanée, contrairement au problème de l'autocorrélation qui est plutôt spécifique aux modèles en séries temporelles. Sur séries temporelles, les cas les plus fréquents d'hétéroscédasticité concernent les séries financières ; ces dernières étant en général caractérisées par une variance variant au cours du temps.

7.7.2 Tests de détection de l'hétéroscédasticité des erreurs

Les tests statistiques d'homocédasticité des erreurs portent sur l'hypothèse :

$$H_0 : \sigma_1^2 = \sigma_2^2 = \dots = \sigma_n^2$$

c'est à dire sur l'hypothèse d'une variance des erreurs identique pour chaque individu. Il existe toute une batterie de tests permettant de détecter de l'hétéroscédasticité des erreurs dont notamment :

- le test de Goldfeld-Quandt (1965)
- le test de Glejser (1969)
- le test de Breusch-Pagan (1979)
- le test de White (1980)
- le test ARCH (1982)

Dans ce qui suit nous ne revenons que sur les deux derniers tests, qui sont les plus utilisés dans la pratique.

7.7.2.1 Test de White (1980)

Le test de White consiste à estimer dans une première étape le modèle :

$$Y_t = a_1 + a_2 X_{2t} + L + a_k X_{kt} + \varepsilon_t \quad (i)$$

puis à calculer le résidu $e_t = Y_t - \hat{a}_1 - \hat{a}_2 X_{2t} - L - \hat{a}_k X_{kt}$

Dans une deuxième étape, on estime par la méthode des MCO l'équation

$$e_t^2 = \lambda_0 + \beta_1 Z_{1t} + \beta_2 Z_{2t} + L + \beta_p Z_{pt} + u_t \quad (ii)$$

où les variables Z_{kt} , $k = 1, 2, K$, p sont les variables explicatives du modèle, leurs carrés et leurs produits.

De l'estimation par la méthode des MCO de l'équation (ii), on déduit le coefficient de détermination noté R_w^2 . La statistique de White est la statistique du multiplicateur de Lagrange (Lagrange Multiplier : LM) définie par $LM = nR_w^2$.

Cette statistique suit sous l'hypothèse H_0 d'homocédasticité des erreurs une loi du Khi-Deux à p degrés de liberté. On se fixe un seuil α , et on lit sur la table du Khi-Deux la valeur seuil A .

La règle de décision est la suivante :

- Si $W_h = nR_w^2 < A$, l'hypothèse H_0 d'homocédasticité n'est pas rejetée.
- Si $W_h = nR_w^2 \geq A$, l'hypothèse H_0 d'homocédasticité est rejetée.

Remarque 1.12 : Les tests de Glejser, White et ARCH sont disponibles sur le logiciel Eviews. Le test de Breusch-Pagan est exécutable sur le logiciel Stata. L'avantage du test de Glejser par rapport aux autres tests est qu'il permet de spécifier la forme de l'hétéroscédasticité.

7.7.2.2 Test ARCH (1982)

Les modèles ARCH (en anglais AutoRegressive Conditional Heteroscedasticity) ont été introduits par Engle en 1982, pour modéliser le processus d'inflation en Grande Bretagne. Ils permettent de modéliser des chroniques (la plupart du temps financières) qui ont une volatilité instantanée qui dépend du passé. Le test d'hétéroscédasticité de type ARCH est fondé aussi sur le test du Multiplicateur de Lagrange (Lagrange Multiplier Test).

De manière pratique, on procède de la manière suivante :

Première étape : calcul de e_t le résidu du modèle de régression ;

Deuxième étape : calcul des e_t^2 ;

Troisième étape : régression autorégressive des résidus sur p retards où seuls les retards significatifs sont conservés ,

$$e_t^2 = \alpha_0 + \sum_{i=1}^p \alpha_i e_{t-i}^2 .$$

L'hypothèse d'homocédasticité des erreurs est $H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_p = 0$.

Quatrième étape : calcul de la statistique du multiplicateur de Lagrange, $LM = nR^2$ avec :

n = nombre d'observations servant au calcul de la régression de l'étape 3,

R^2 = coefficient de détermination de l'équation intermédiaire de l'étape 3.

La règle de décision est la suivante :

Si $LM < \chi^2(p)$ à p degrés de liberté lu dans la table du Khi-Deux à un seuil α , l'hypothèse H_0 n'est pas rejetée. Les erreurs sont homocédastiques.

Si $LM \geq \chi^2(p)$, l'hypothèse H_0 est rejetée; on considère que les erreurs sont conditionnellement hétéroscédastiques. Le processus suit alors un modèle ARCH(p).

Remarque 1.13 : C'est le test de significativité des coefficients α_i de la régression e_t^2 sur e_{t-p}^2 qui permet de déterminer l'ordre p du processus ARCH sachant qu'un processus ARCH d'ordre 3 est un maximum ($p \leq 3$).

7.7.3 Procédure d'estimation en présence de l'hétéroscédasticité des erreurs

La présence de l'hétéroscédasticité des erreurs a pour conséquence que les estimateurs des moindres carrés restent sans biais, mais ne sont plus de variance minimale. Il peut être pertinent de chercher à corriger l'hétéroscédasticité. Il convient de distinguer les cas où la variance de l'erreur est connue de ceux où elle est inconnue. Lorsque la variance du terme d'erreur est connue, il convient d'appliquer la méthode des moindres carrés généralisés. Dans ces cas où la variance du terme d'erreur est inconnue, les programmes de régression utilisent les corrections suggérées par White (1980) et Newey-West (1987).

Application 5 : Tests d'hétéroscédasticité des erreurs.

Reprenons la fonction d'importation du Sénégal, à fréquence annuelle sur la période allant de 1962 à 1995. L'estimation des paramètres par les MCO conduit aux résultats suivants :

$$\text{Log IMPORT}_t = 0,65 + 0,75 \text{ Log PIB}_t$$

(1,87) (15,19)

$$R^2 = 0,87, \quad \text{SCR} = 0,134, \quad \text{DW} = 1,06, \quad n = 34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés.

Tester une éventuelle hétéroscédasticité des erreurs par les tests suivants :

- a) White
- b) ARCH d'ordre 1
- c) Glejser
- d) Breusch-Pagan

Corrigé :

Pour les 4 tests, le test d'hypothèses est :

$$H_0 : \sigma_t^2 = \sigma^2 \quad \forall t$$

contre

$$H_1 : \sigma_t^2 \neq \sigma^2 \quad \text{pour au moins un } t$$

a) Le test de White conduit aux résultats suivants :

Heteroskedasticity Test: White			
F-statistic	2.848827	Prob. F(2,31)	0.0731
Obs*R-squared	5.278818	Prob. Chi-Square(2)	0.0714

Les deux probabilités sont supérieures au seuil de 5%. Le test de Fisher de significativité du coefficient de régression et la statistique LM sont concordants. On ne rejette pas l'hypothèse d'homocédasticité des erreurs.

b) Le test ARCH d'ordre 1 conduit aux résultats suivants :

Heteroskedasticity Test: ARCH(1)			
F-statistic	0.008791	Prob. F(1,31)	0.9259
Obs*R-squared	0.009355	Prob. Chi-Square(1)	0.9229

Les deux probabilités sont supérieures au seuil de 5%, nous ne rejetons pas l'hypothèse d'homocédasticité des erreurs. Il n'existe pas d'effet ARCH sur les erreurs du modèle.

c) Le test de Glejser conduit aux résultats suivants :

Heteroskedasticity Test: Glejser			
F-statistic	10.12045	Prob. F(1,32)	0.0033
Obs*R-squared	8.169318	Prob. Chi-Square(1)	0.0043

Les deux probabilités sont inférieures au seuil de 5%. L'hypothèse d'homocédasticité des erreurs est rejetée.

L'estimation par les MCO du modèle régressant la série de la valeur absolue de e_t sur une constante et $\log(\text{pib}_t)$ conduit aux résultats suivants :

$$\left| e_t \right| = 0,69 - 0,09 \text{ Log PIB}_t$$

(3,40) (-3,18)

$$R^2 = 0,24 \quad , \quad n = 34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés.

Le coefficient associé à $\log(\text{pi}b_t)$ est significativement différent de zéro car la valeur absolue de son ratio de Student (3,18) est supérieure à la valeur lue dans la table de la loi normale centrée et réduite (1,96 au seuil statistique de 5%). L'hypothèse d'homocédasticité des erreurs est rejetée.

d) Le test de Breusch-Pagan effectué sur le logiciel STATA conduit aux résultats suivants :

Breusch-Pagan / Cook-Weisberg test for heteroskedasticity

Ho: Constant variance

Variables: fitted values of limport

chi2(1) = 4.99

Prob > chi2 = 0.0255

La probabilité critique (0,0255) est inférieure à 5%. L'hypothèse d'homocédasticité des erreurs est rejetée au seuil de 5%.

Les tests ARCH et White concluent au non rejet de l'hypothèse d'homocédasticité des erreurs.

Par contre les tests de Glejser et Breusch et Pagan concluent au rejet de l'hypothèse d'homocédasticité des erreurs.

7.8 Test de spécification de Ramsey

Le test de spécification de Ramsey repose sur la même idée simplificatrice du test de White. Il s'agit d'un test général de variables manquantes. En pratique, le test de Ramsey consiste à procéder en 4 étapes.

1^{re} étape : Estimer les paramètres de la forme linéaire :

$$Y_t = a_1 + a_2 X_{2t} + L + a_k X_{kt} + \varepsilon_t$$

2^e étape : Simuler la variable ajustée $\hat{Y}_t$ définie par :

$$\hat{Y}_t = \hat{a}_1 + \hat{a}_2 X_{2t} + L + \hat{a}_k X_{kt}$$

3^e étape : Estimer les paramètres du modèle linéaire par les MCO

$$Y_t = a_1 + a_2 X_{2t} + L + a_k X_{kt} + \beta_1 \hat{Y}_t^2 + \beta_2 \hat{Y}_t^3 + \beta_3 \hat{Y}_t^4 + u_t$$

4^e étape : Effectuer un test de Fisher de l’hypothèse

$$H_0 : \beta_1 = \beta_2 = \beta_3 = 0 .$$

La règle de décision est la suivante :

- Si l’hypothèse H_0 n’est pas rejetée, le modèle est bien spécifié.
- Si l’hypothèse H_0 est rejetée, le modèle est mal spécifié.

Remarque 1.14 : On peut se demander pourquoi s’arrêter à la puissance 4 dans la troisième étape, d’autant plus la plupart des logiciels permettent de procéder au test en utilisant des puissances supérieures. En fait les expériences de simulation de Monte-Carlo montrent que la puissance du test est maximale dans cette configuration.

Application 6 :Test de spécification de Ramsey

Reprenons la fonction d’importation du Sénégal, à fréquence annuelle sur la période allant de 1962 à 1995. En effectuant la régression sur la période globale ($n = 34$), on obtient :

$$\text{Log IMPORT}_t = 0,65 + 0,75 \text{ Log PIB}_t$$

Tester la spécification du modèle.

Corrigé : Le test de Ramsey conduit aux résultats suivants :

Ramsey RESET Test :	Nombre de termes ajustés :2		
F-statistic	6.366191	Prob. F(2,30)	0.0050
Log likelihood ratio	12.02783	Prob. Chi-Square(2)	0.0024

Le logiciel EvIEWS effectue 2 tests : le test de Fisher et le test du rapport de vraisemblance (Log likelihood ratio).

La statistique de Fisher vaut 6,366 ; sa probabilité critique est nulle. L’hypothèse nulle du test de Ramsey est rejetée : la fonction d’importation est mal spécifiée. Ce résultat n’est pas surprenant

dans la mesure où le pib n'est pas la seule variable explicative des importations. Il manque des variables explicatives comme la production nationale, les prix, etc.

7.9 Tests de stabilité des coefficients

La stabilité des paramètres joue un rôle important lorsqu'on cherche à comprendre les mécanismes économiques et à réaliser des projections. Leur instabilité peut refléter des phénomènes ponctuels dans le temps (choc pétrolier, dévaluation, crise boursière, calamités naturelles, mesures de politiques économiques, nouvelles réglementations, passage d'un régime de change fixe à un régime de change flexible, ...).

Il est souvent intéressant d'évaluer la robustesse du modèle estimé, c'est à dire la stabilité de celui-ci. L'idée sous-jacente est que, sur la période considérée, il peut apparaître un changement structurel ou un changement ponctuel dans la relation entre la variable à expliquer et les variables explicatives. En d'autres termes, il se peut que les valeurs des coefficients du modèle estimé ne soient pas stables sur l'ensemble de la période considérée. Les changements peuvent avoir plusieurs sources. Il existe diverses méthodes pour évaluer la stabilité des coefficients estimés d'un modèle de régression, nous en présentons deux :

- Le test de Chow (1960)
- Les tests Cusum et Cusum Carré de Brown, Durbin et Evans (1975).

7.9.1 Test de Chow

Le test de Chow appelé aussi test de changement structurel, permet d'examiner si les coefficients d'une régression sont stables par rapport à l'observation utilisée.

Sur des séries temporelles, on compare les estimations effectuées sur deux (ou plusieurs) sous ensembles d'observations qui correspondent à un découpage de l'échantillon initial.

On parle dans ce cas de test de stabilité temporelle de la régression.

Sur les données en coupe, on peut comparer les résultats obtenus par exemple sur des pays, des régions, des secteurs industriels différents. Concernant des individus, on peut s'intéresser à des résultats par classe d'âge, par sexe, etc. Dans ce cas, le test de Chow est souvent qualifié de test d'homogénéité des comportements.

On considère le modèle linéaire général suivant :

$$\underset{(n, 1)}{Y} = \underset{(n, k)}{X} \underset{(k, 1)}{a} + \underset{(n, 1)}{\varepsilon} \quad (0)$$

On se pose le problème de la stabilité des coefficients du modèle dans le temps.

Doit-on considérer le modèle comme étant stable sur la totalité de la période ($t = 1, 2, L, n$) ?

Ou doit-on considérer deux sous périodes distinctes d'estimation ?

Supposons qu'on ait deux sous périodes ayant n_1 et n_2 observations :

Sous période 1 :

$$\underset{(n_1, 1)}{Y_1} = \underset{(n_1, k)}{X_1} \underset{(k, 1)}{a_1} + \underset{(n_1, 1)}{\varepsilon_1} \quad (1)$$

Sous période 2 :

$$\underset{(n_2, 1)}{Y_2} = \underset{(n_2, k)}{X_2} \underset{(k, 1)}{a_2} + \underset{(n_2, 1)}{\varepsilon_2} \quad (2)$$

Le principe du test est de voir dans quelle mesure le fait de régresser séparément sur les deux sous périodes améliore le résultat de la régression. Ce test portera sur les sommes des carrés des résidus (variances résiduelles). On commence par régresser sur la population totale soit n observations (modèle (0)), on calcule alors SCR, qui est la somme des carrés des résidus du modèle (0). Puis on refait la même régression sur chacun des deux sous populations séparément et on retient la somme des carrés des résidus de chaque régression.

Soit :

Sur n_1 observations $\rightarrow SCR_1$

Sur n_2 observations $\rightarrow SCR_2$

On pose alors le test d'hypothèses :

$H_0 : SCR = SCR_1 + SCR_2$ (les coefficients du modèle sont stables)
contre

$H_1 : SCR \neq SCR_1 + SCR_2$ (les coefficients du modèle sont instables)

Le Fisher empirique est égal à :

$$F^* = \frac{[SCR - (SCR_1 + SCR_2)]/k}{[SCR_1 + SCR_2]/(n - 2k)}$$

Cette statistique est distribuée sous l'hypothèse nulle de stabilité une loi de Fisher à respectivement k degrés de liberté pour le numérateur et $(n - 2k)$ degrés de liberté pour le dénominateur. On se fixe un seuil de signification α et on lit sur la table de Fisher la valeur seuil F_{lu} .

Si $F^* < F_{lu}$, l'hypothèse H_0 n'est pas rejetée et on retient une estimation par les moindres carrés ordinaires avec l'ensemble des observations. Les coefficients sont dits stables sur l'ensemble de la période.

Si $F^* \geq F_{lu}$, l'hypothèse H_0 est rejetée. Les coefficients ne sont pas stables sur l'ensemble de la période.

Remarque 1.15 : Le test de Chow peut être facilement généralisé à l'existence de plus d'une rupture structurelle. Ainsi, si l'on souhaite tester l'existence de deux ruptures, on procédera en trois sous-périodes, le principe du test restant le même (la somme des carrés SCR étant égale à la somme des carrés des résidus des trois régressions correspondant aux trois sous-périodes). Le test de Chow suppose que l'on connaisse la date à laquelle se produit la (ou les) rupture(s).

7.9.2 Tests Cusum

Le test de Chow suppose que la date de rupture est connue a priori. Quand on travaille sur des séries temporelles, la date à laquelle des changements dans les coefficients interviennent n'est pas toujours facilement repérable. Mais il existe également des tests de stabilité temporelle qui permettent de déterminer les dates de rupture. Brown, Durbin et Evans (1975) ont proposé des tests de stabilité des coefficients basés sur le calcul des résidus récurrents. Ces tests graphiques permettant d'accepter ou non l'hypothèse de stabilité. L'intérêt de ces tests réside dans le fait qu'ils permettent d'étudier la stabilité d'une régression sans définir a priori la date de rupture sur les coefficients.

Il existe deux versions de ce test : le CUSUM fondé sur la somme cumulée des résidus récurrents et le CUSUM SQ (CUSUM of Squares) fondé sur la somme cumulée du carré des résidus récurrents. Les tests Cusum sont ainsi basés sur les résidus récurrents : on procède à une représentation graphique des résidus récurrents cumulés qui permet de tester la stabilité temporelle de la régression et de visualiser les dates d'éventuelles ruptures de comportement. Pour ces deux tests, si on constate une rupture du graphique à la période $t = \pi$ alors on rejette l'hypothèse de stabilité des coefficients de la régression pour cette période π .

Les tests Cusum sont disponibles sur le logiciel EViews.

Remarque 1.16 : Les tests Cusum ne sont valables qu'après une estimation des paramètres par la méthode des MCO.

Application 7 : Tests de stabilité des coefficients d'un modèle

Reprenons la fonction d'importation du Sénégal, à fréquence annuelle sur la période allant de 1962 à 1995. L'estimation des paramètres par la méthode des MCO conduit aux résultats suivants :

$$\text{Log IMPORT}_t = 0,65 + 0,75 \text{ Log PIB}_t$$

(1,87) (15,19)

$$R^2 = 0,87, \quad \text{SCR} = 0,134, \quad \text{DW} = 1,06, \quad n = 34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés.

1) Tester la stabilité de la fonction d'importation du Sénégal en utilisant :

- a) le test de Chow (choisir les dates de rupture 1973 et 1978),
- b) les tests Cusum de Brown-Durbin-Evans.

2) Stabiliser le modèle en utilisant les variables indicatrices.

Corrigé :

1) Tests de stabilité

a1) Pour la date de rupture 1973, le test de Chow aboutit aux résultats suivants :

Chow Breakpoint Test: 1973			
Null Hypothesis: No breaks at specified breakpoints			
Varying regressors: All equation variables			
Equation Sample: 1962 1995			
F-statistic	7.27	Prob. F(2,30)	0.00
Log likelihood ratio	13.44	Prob. Chi-Square(2)	0.00
Wald Statistic	14.54	Prob. Chi-Square(2)	0.00

La statistique de Chow vaut 7,27 ; sa probabilité critique est nulle. L'hypothèse nulle de stabilité est rejetée au seuil de 1%. La fonction d'importation est instable sur la totalité de la période. Le fait de scinder en deux échantillons détériore la qualité du modèle. Le premier choc pétrolier (année 1973) est bien une date d'instabilité pour l'économie sénégalaise.

a2) Pour la date de rupture 1978, le test de Chow aboutit aux résultats suivants :

Chow Breakpoint Test: 1978			
Null Hypothesis: No breaks at specified breakpoints			
Varying regressors: All equation variables			
Equation Sample: 1962 1995			
F-statistic	1.95	Prob. F(2,30)	0.16
Log likelihood ratio	4.15	Prob. Chi-Square(2)	0.12
Wald Statistic	3.89	Prob. Chi-Square(2)	0.14

La statistique de Chow vaut 1,95 ; sa probabilité critique (0,16) est supérieure aux seuils statistiques conventionnels. L'hypothèse de stabilité n'est pas rejetée. La fonction d'importation est donc stable.

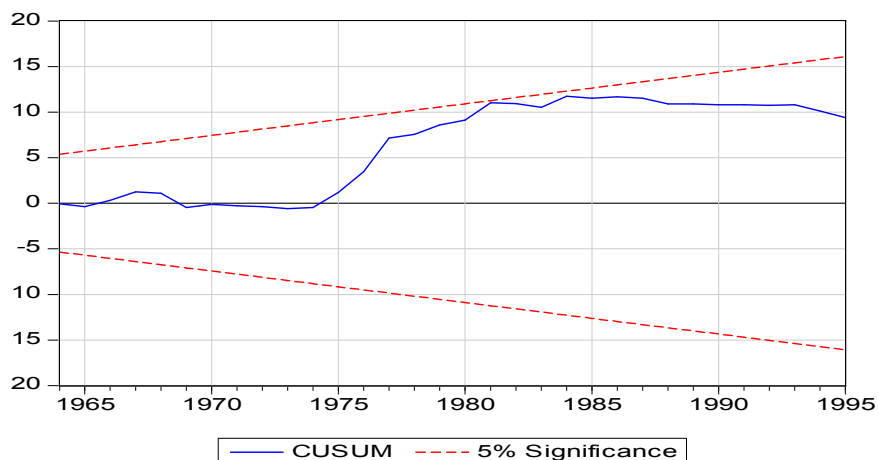
b) Les tests Cusum et Cusum Carré sont effectués avec le logiciel Eviews.

b1) Test Cusum

Ce test permet de détecter les instabilités structurelles.

GRAPHIQUE 5 :Test CUSUM

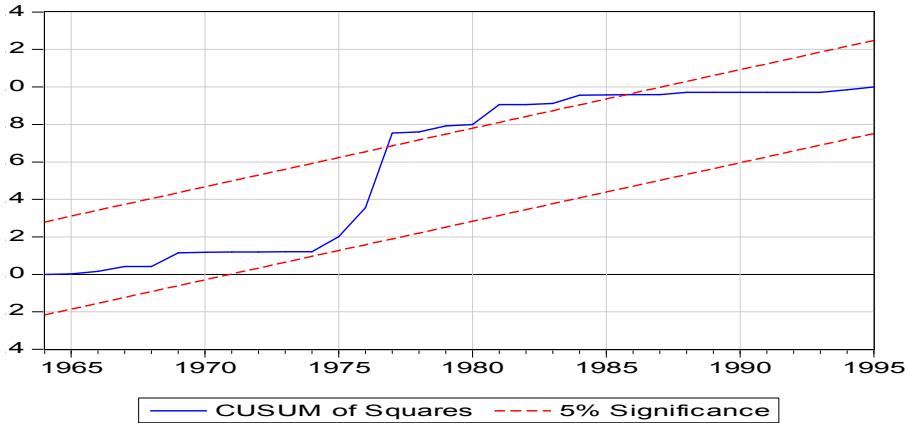
La courbe ne sort pas du corridor représenté en pointillés, la fonction d'importation est structurellement stable.



b2) Test Cusum Carré

Ce test permet de détecter les instabilités ponctuelles.

GRAPHIQUE 6 :Test CUSUM carré



La courbe sort du corridor représenté en pointillés, le modèle est ponctuellement instable. La zone d'instabilité est 1977 à 1986. Cette instabilité peut être expliquée par le deuxième choc pétrolier et les politiques d'ajustement structurel.

Au total, c'est l'hypothèse d'instabilité qui doit être acceptée.

2) Stabilisation du modèle par variable indicatrice

Pour stabiliser le modèle, nous pouvons utiliser une variable indicatrice (appelée aussi variable dummy). La variable indicatrice vaut 1 pendant la ou les zones d'instabilité. Elle vaut 0 ailleurs.

On pose :

$dum_t = 1$ pour les années 1977 à 1986

$dum_t = 0$ ailleurs

L'estimation par les MCO du modèle corrigé conduit aux résultats suivants :

$$\text{Log } IMPORT_t = 0,76 + 0,73 \text{ Log } PIB_t + 0,06 dum_t$$

(16,01) (2,73)

$$R^2 = 0,90, \quad SCR = 0,108 \quad n = 34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés. La variable indicatrice dum est significative au seuil statistique de 5%.

Les tests Cusum effectués sur le modèle corrigé indiquent que la fonction d'importation est maintenant structurellement et ponctuellement stable. **L'utilisation de la variable indicatrice a donc stabilisé le modèle.**

8. La prévision à l'aide du modèle linéaire général

L'un des intérêts pratiques du modèle de régression réside dans la prévision. Ainsi une fois le modèle estimé, il est possible de l'utiliser afin de prévoir l'évolution de la variable endogène Y. Nous considérons le modèle linéaire simple :

$$Y_t = a + bX_t + \varepsilon_t, \quad t=1, 2, \dots, n$$

On suppose connue la valeur X_{n+h} et on cherche à déterminer la prévision de la variable expliquée pour un horizon h. Nous voulons prévoir la valeur Y_{n+h} où n est l'origine de la prévision et h l'horizon de la prévision, $h \in \mathbb{N}^*$.

Le terme prévision a ici un sens différent de celui qu'il reçoit dans le langage courant. Il ne s'agit pas de prévision du futur. On cherche en fait à caractériser les simulations de politique économique que les estimations économétriques rendent possibles.

La valeur de Y_{n+h} est prévue par $\hat{Y}_{n+h} = \hat{a} + \hat{b} X_{n+h}$ où $\hat{a}$ et $\hat{b}$ sont les estimateurs des MCO des paramètres a et b définis par :

$$\hat{b} = \frac{\text{Cov}(X, Y)}{\text{Var}(X)} ; \hat{a} = \bar{Y} - \hat{b} \bar{X}$$

Quand la taille de l'échantillon est grande, l'intervalle de confiance des prévisions - également appelé intervalle de prévision - au niveau $1 - \alpha$ est donnée par :

$$\left[\hat{Y}_{n+h} - 2S ; \hat{Y}_{n+h} + 2S \right]$$

avec

$$S^2 = \frac{1}{n-2} \sum_{t=1}^n e_t^2, \quad e_t = Y_t - \hat{Y}_t = Y_t - \hat{a} - \hat{b} X_t$$

Application 8 : Prédiction d'un modèle linéaire simple

Reprenons l'exemple sur la fonction d'importation du Sénégal, à fréquence annuelle sur la période allant de 1962 à 1995.

L'estimation des paramètres par la méthode des MCO conduit aux résultats suivants :

$$\text{Log } IMPORT_t = 0,65 + 0,75 \text{ Log } PIB_t$$

(1,87) (15, 19)

$$R^2=0,87, \quad SCR=0,134, \quad DW=1,06, \quad n=34$$

où les chiffres entre parenthèses correspondent aux ratios de Student des coefficients estimés.

Supposons que le PIB en 1996 soit fixé à 1606. La valeur des importations prévues pour cette année est :

$$\hat{Y}_{1996} = \exp(0,65 + 0,75 \times \log(1606)) = 485,96$$

Déterminons à présent l'intervalle de prédiction à 95% :

$$\left[\hat{Y}_{1996} - 2S ; \hat{Y}_{1996} + 2S \right]$$

Par ailleurs, nous avons calculé, $S = 33,224$, l'intervalle de prédiction est :

$$[419,512 ; 552,408].$$

La prédiction des importations pour l'année 1996 a 95% de chances de se trouver dans l'intervalle $[419,512 ; 552,408]$.

CONCLUSION

Comme l'indique le sous titre de ce premier tome, cet ouvrage traite l'économétrie du modèle linéaire général. Ce premier tome doit impérativement être lu par l'économètre débutant : il présente un certain nombre de définitions et concepts essentiels. Outre le fait que ce modèle rend compte d'une large palette de phénomènes économiques, l'acquisition des principes de base de l'économétrie du modèle linéaire général est une étape essentielle pour aborder des méthodes plus complexes.

Le second tome traitera les modèles dynamiques. Il débutera avec la présentation des modèles à décalages temporels. Nous verrons comment distinguer le modèle linéaire autorégressif et le modèle à retards échelonnés. Nous verrons ensuite comment déterminer le nombre de décalages du modèle à retards échelonnés par l'utilisation des critères d'information de Akaike et de Schwarz.

Le chapitre 3 du second tome examinera la cointégration et le modèle à correction d'erreur. Nous verrons que lorsque les données ne sont pas stationnaires, la méthode des moindres carrés n'est pas valable. Dans cette situation, il est indiqué d'appliquer les tests de cointégration et d'utiliser le modèle à correction d'erreur dans la situation où l'hypothèse de cointégration n'est pas rejetée.

A l'issue de ces deux tomes, le lecteur aura acquis les techniques de base des méthodes économétriques.

Annexe : Instructions et codes des logiciels Eviews et Stata pour la résolution des applications proposées

« Les machines un jour pourront résoudre tous les problèmes, mais jamais aucune d'entre elles ne pourra en poser un ! »

Albert Einstein

APPLICATION 1 : TEST DU COEFFICIENT DE CORRÉLATION LINÉAIRE

X : dépense de publicité

Y : montant des ventes correspondantes

-- Pour avoir le nuage de points

Instruction Eviews

scat pub ventes

Instruction Stata

scatter ventes pub

-- Pour effectuer le test du coefficient de corrélation linéaire

Instruction EIEWS

Cliquer sur <Quick> puis <Show> Saisir pub ventes dans la fenêtre « Series List » puis OK

Cliquer sur <View> puis <Covariance analysis> puis OK

Instruction Stata

pwcorr pub ventes, sig

APPLICATION 2 : TESTS DE NORMALITÉ ET DE LOGNORMALITÉ

m2 : Masse monétaire

Instruction EViews pour le test de normalité de Jarque-Bera

stats m2 log(m2) ou

Cliquer sur <Quick> puis <Group statistics> puis <Descriptive Statistics> puis <Common Sample>
Saisir ensuite m2 log(m2) dans la fenêtre « Series List » puis OK

Instruction Stata pour le test de normalité

--générer d'abord le logarithme népérien de la variable m2

gen lm2 = log(m2)

-- tester la normalité et la lognormalité

sktest m2 lm2

APPLICATION 3 : ESTIMATION DES PARAMÈTRES D'UNE FONCTION DE PRODUCTION DE TYPE COBB-DOUGLAS.TEST DE SIGNIFICATIVITÉ D'UNE VARIABLE EXPLICATIVE DE STUDENT.TEST DE SIGNIFICATIVITÉ GLOBALE DE FISHER.TEST DE L'HYPOTHÈSE DE RENDEMENT D'ÉCHELLE CONSTANT

Q : production ; K : facteur capital ; L : facteur travail

Instruction Eviews pour l'estimation des paramètres par la méthode des moindres carrés ordinaires

ls log(q) c log(k) log(l) ou

Cliquer sur <Quick> puis <Estimate Equation > Saisir log(q) c log(k) log(l) dans la fenêtre <Equation specification > puis OK .

ls : least squares (en français : moindres carrés ordinaires)

c représente le terme constant du modèle

Instruction Stata pour l'estimation des paramètres par la méthode des moindres carrés ordinaires

-- générer d'abord les logarithmes népériens des variables q , k et l

gen lq = log(q)

gen lk = log(k)

gen ll = log(l)

-- estimer les paramètres du modèle de Cobb-Douglas

regress lq lk ll

le logiciel Stata indique les intervalles de confiance à 95%.

Pour obtenir les intervalles de confiance à 99%,ajouter l'option level(99)

regress lq lk ll ,level(99)

Il est inutile de se préoccuper de la constante,Stata l'ajoute automatiquement à toutes les régressions.Pour forcer Stata à ne pas la rajouter, on doit ajouter l'option nocons.

regress lq lk ll ,nocons

Instruction SPSS pour l'estimation des paramètres par la méthode des moindres carrés ordinaires

Nous devons d'abord générer les logarithmes népériens des variables

Transformer Calculer Variable destination : saisir lq Expression numérique : saisir LN(q) OK

Transformer Calculer Variable destination : saisir Ik Expression numérique : saisir LN(k) OK

Transformer Calculer Variable destination : saisir II Expression numérique : saisir LN(I) OK

Analyse Régression Linéaire variables dépendante : saisir Iq variables explicatives : saisir Ik II
Méthode : choisir Entrée OK

Instruction Eviews pour le test de rendement d'échelle constant

Après l'estimation des paramètres, cliquer sur <View> puis <Coefficients Tests> puis <Wald-Coefficients Restrictions> saisir $c(2) + c(3) = 1$ dans la fenêtre <Coefficients restrictions> puis OK

Instruction Stata pour le test de rendement d'échelle constant

Après l'estimation des paramètres

test Ik + II = 1

APPLICATION 4 : TESTS D'AUTOCORRÉLATION DES ERREURS DE DURBIN-WATSON ET DE BREUSCH-GODFREY SUR LA FONCTION D'IMPORTATION DU SÉNÉGAL. CORRECTION DE L'AUTOCORRÉLATION DES ERREURS PAR LA MÉTHODE DE COCHRANE-ORCUTT

Instruction Stata pour le test de Durbin-Watson

Après l'estimation des paramètres

la statistique de Durbin-Watson est calculée par la commande `dwstat`

le test de Durbin-Watson est effectué avec la commande

`durbina, lags(1)`

Instruction Eviews pour le test de Breusch-Godfrey

Après l'estimation des paramètres

`auto(1)` ou

cliquer sur <View> puis <Residuals Tests> puis <Serial Correlation LM Tests> saisir 1 dans la fenêtre <Lag to include> puis OK.

Cette commande teste une autocorrélation d'ordre 1 des erreurs.

Si nous désirons tester une autocorrélation d'ordre 2 des erreurs, alors on utilise la commande `auto(2)` ou on saisit 2 dans la fenêtre <Lag to include>.

Instruction Stata pour le test de Breusch-Godfrey

Après l'estimation des paramètres

```
bgodfrey(1), lags(1)
```

Instruction Eviews pour la correction de l'autocorrélation des erreurs par la méthode de Cochrane-Orcutt

```
ls log(import) c log(pib) ar(1) ou
```

Cliquer sur <Quick> puis <Estimate Equation> Saisir `log(import) c log(pib) ar(1)` dans la fenêtre <Equation specification> puis OK.

Instruction Stata pour la correction de l'autocorrélation des erreurs par la méthode de Cochrane-Orcutt

-- générer d'abord les logarithmes népériens des variables import et pib

```
gen limport = log(import)
```

```
gen lpib = log(pib)
```

-- estimer par la méthode de Cochrane-Orcutt

prais limport lpib , corc

APPLICATION 5 :TESTS D'HÉTÉROCÉDASTICITÉ DES ERREURS DE WHITE,ARCH(1), GLEJSER ET DE BREUSCH-PAGAN SUR LA FONCTION D'IMPORTATION DU SÉNÉGAL

Instruction Eviews pour le test d'hétérocédasticité des erreurs de White

Après l'estimation des paramètres

cliquer sur <View> puis <Residuals Tests> puis <White Heteroskedasticity Tests > puis OK

Instruction Stata pour le test d'hétérocédasticité des erreurs de Breusch-Pagan

Après l'estimation des paramètres

hettest

Instruction Eviews pour le test ARCH(1) d'hétérocédasticité conditionnelle des erreurs

Après l'estimation des paramètres

archtest(1) ou

cliquer sur <View> puis <Residuals Tests> puis <ARCH LM Tests > saisir 1 dans la fenêtre <Lag to include> puis OK .

Instruction Stata pour le test ARCH(1) d'hétérocédasticité conditionnelle des erreurs

Après l'estimation des paramètres

archlm ,lags(1)

Instruction Eviews pour le test d'hétéroscédasticité des erreurs de Glesjer

Après l'estimation des paramètres, cliquer sur <View> puis <Residuals Tests> puis <Heteroskedasticity Tests> choisir Glejser dans la fenêtre <Test type> puis OK.

APPLICATION 6 :TEST DE SPÉCIFICATION DE RAMSEY SUR LA FONCTION D'IMPORTATION DU SÉNÉGAL

Instruction Eviews pour le test de spécification de Ramsey

Après l'estimation des paramètres

reset(2) ou

cliquer sur <View> puis <Stability Tests> puis <Ramsey Reset Tests> saisir 2 dans la fenêtre <Number of fitted terms> puis OK.

Instruction Stata pour le test de spécification de Ramsey

Après l'estimation des paramètres

ovtest

APPLICATION 7 :TESTS DE STABILITÉ DES COEFFICIENTS CHOW ET DE BROWN-DURBIN-EVANS SUR LA FONCTION D'IMPORTATION DU SÉNÉGAL

Instruction Eviews pour le test de stabilité de Chow

Après l'estimation des paramètres, cliquer sur <View> puis <Stability Tests> puis < Chow Breakpoint Tests> saisir la (les) date(s) de rupture dans la fenêtre <Enter one or more breakpoint dates> puis OK.

Instruction Eviews pour les tests Cusum de stabilité de Brown-Durbin-Evans

Après l'estimation des paramètres, cliquer sur <View> puis <Stability Tests> puis <Recursive Estimates> choisir Cusum Test ou Cusum of Square Test puis OK.

Si on choisit Cusum Test alors le logiciel Eviews effectue le test d'instabilité structurelle.

Si on choisit Cusum of Square Test alors le logiciel Eviews effectue le test d'instabilité ponctuelle (ou conjoncturelle).

Les tests de stabilité ne sont pas disponibles sur le logiciel Stata.

APPLICATION 8 : PRÉVISION DE LA FONCTION D'IMPORTATION DU SÉNÉGAL

Instruction Eviews

Après l'estimation des paramètres, cliquer sur <Forecast>

Choisir Forecast of import

Saisir previmport dans la fenêtre <Forecast name> puis OK.

Instruction Stata pour la prévision

Après l'estimation des paramètres

```
predict prevlimport,xb
```

```
gen previmport = exp(prevlimport)
```



ACHEVÉ D'IMPRIMER SUR LES PRESSES
DE L'IMPRIMERIE DE LA BCEAO
DECEMBRE 2013



BCEAO
BANQUE CENTRALE DES ETATS
DE L'AFRIQUE DE L'OUEST

Avenue Abdoulaye Fadiga
B.P. 3108 - Dakar - Sénégal
www.bceao.int