

La spécification de la fonction de demande de monnaie

Introduction

L'analyse théorique de la demande de monnaie, nous a permis de cerner la nature de la fonction de demande de monnaie. Il nous reste maintenant à envisager les problèmes de la spécification de cette fonction, de sa stabilité dans le temps pour mener à bien des études empiriques.

La fonction de demande de monnaie doit permettre de mettre en vigueur les liens stables dans le temps entre la quantité de monnaie (M1 et M2) et les grandeurs macroéconomiques tels que le produit intérieur brut (PIB), le produit national brut (PNB), le taux d'intérêt (i) le taux d'inflation (π), le taux de change (TC) ...ect

Ainsi, toutes les fonctions de demande de monnaie ont convergé sur les spécifications d'un modèle avec comme variable d'échelle: le revenu, variable de rendement ou d'opportunité: le taux d'intérêt et le taux d'inflation.

La théorie sur la demande de monnaie postule que la demande d'encaisses réelles $(M/P)^d$ est positivement liée à une contrainte budgétaire comme le revenu (y) et négativement liée au rendement des actifs monétaires et financiers. Le taux anticipé de l'inflation est utilisé pour représenter le rendement des actifs non financiers alors que, le taux d'intérêt pour les actifs financiers.

La fonction de demande de monnaie peut être représentée généralement de la façon suivante: $(M/P)^d = f(y, i, \pi, \dots)$

Le présent chapitre sera consacré aux problèmes de spécification qui revêtent trois aspects importants: c'est-à-dire le problème de,

- 1- L'identification de la fonction de demande de monnaie et le problème de simultanéité.
- 2- La définition et le choix des variables qui constituent la fonction de la demande de monnaie.
- 3-processus d'ajustement de la fonction de demande de monnaie.

4.1 Définition de spécification

Un modèle se présente de façon générale, comme un ensemble de relation entre différents variables présentant entre elles un lien de cause à effet. Il ne s'agit pas de n'importe quelles relations, mais, de relations quantitatives entre le niveau de telle variable par exemple un agrégat monétaire M_i ($i=1,2,\dots$) et le niveau de telle autre variable disons le revenu Y ou le taux d'intérêt i .

On peut raisonner sur le niveau de la variable considérée (X) sur sa variation ΔX ou Dx/dt ou sur son logarithme. Pour éviter les problème de mesure on raisonne en valeur relative ($\Delta X/X$). Un des avantage à travailler en logarithme vient de ce que $\ln X$ ou $d \ln X/dt = \Delta X/X$ ou dX/X . Pour effectuer des comparaisons, les économistes préfèrent raisonner en terme d'élasticité qui est un rapport de variation

relative plutôt que de comparer directement des variations absolues. Donc, $(dY/Y)/(dX/X) = d\ln Y/d\ln X = \epsilon_{Y/X}$.

En macroéconomie, on doit agréger des valeurs de biens hétérogènes, ce qui suppose de les pondérer par leurs prix relatifs c'est donc raisonner en valeur. Mais la comparaison des valeurs dans le temps va poser un problème important quant les valeurs unitaires sont mesurées par des prix dont l'évolution peut ne pas refléter un changement de pouvoir d'achat. Donc on va prendre en compte ce phénomène en corrigeant l'influence de l'inflation, c'est-à-dire, travailler en termes réels. La correction à faire subir aux variables consiste à les diviser par un déflateur qui est souvent l'indice de niveau général des prix.

La spécification des relations d'un modèle économétrique est plus ou moins complexe et plus ou moins précise. Avec des spécifications aussi générales, on ne peut guère prétendre exprimer la valeur de la fonction à partir des valeurs des variables. Pour faciliter l'obtention de tels résultats on remplace souvent la fonction par une forme linéaire présentant des propriétés analogues.

Dans notre cas, la fonction de demande de monnaie qui appartient à un modèle théorique de portée générale, on peut se contenter d'une écriture de type $M_d = f(Y, i)$,

Ou encore en opposant les variables endogènes ou explicatives Y, i et les variables exogènes M_d et ou les paramètres $f(Y, i; b, a) = 0$, la fonction est donc mise sous forme implicite. Encore, $M_d = f(Y, i)$ avec $f_Y > 0$ et $f_i < 0$. En logarithme, on écrit $M_d = e^{-ai} Y^b$ en notant le logarithme Y par une minuscule l'équation devient :

$M_d = by + ai$ avec $b > 0$ et $a < 0$.

Les spécifications linéaires se justifient pour plusieurs raisons:

- la simplicité et la facilité à réaliser les calculs même relativement complexes.
- en procédant à des vérifications empiriques on sait guère effectuer des estimations économétriques à partir des modèles non linéaires.
- prendre en compte des phénomènes plus complexes et des spécifications plus précises en liens avec les enseignements de la théorie qui à leur tour, peuvent être testés au niveau empirique.

En macroéconomie, il y'a plusieurs théories concurrentes comme nous les avons précisé dans les chapitres précédents. Ces théories reposent sur la nature des

procédures utilisées en terme de modélisation du système économique. L'adéquation de la théorie aux faits est essentielle. Donc, il est nécessaire de savoir s'il existe des critères permettant de préférer telle théorie à telle autre. Toute fois, les méthodes économétriques ne permettent guère de tester autre chose que la comptabilité d'une théorie avec les faits observés qui, par ailleurs, peuvent être également compatibles avec autres théories. Les tests empiriques permettant de discriminer véritablement entre différentes théories sont en pratique extrêmement rares.

On peut toutefois fonder des spécifications macroéconomiques différentes en modifiant certaines hypothèses de base utilisées en microéconomie et notamment l'hypothèse de maximisation de l'utilité ou encore celle de la concurrence pure et parfaite.

La nouvelle économie Keynésienne, développée à partir des années 70 fait largement appel à cette méthode en utilisant en particulier les résultats de la théorie microéconomique de la concurrence imparfaite(asymétrie d'information) ou encore en proposant une analyse microéconomique des marchés fonctionnant hors de leurs situation d'équilibre. En conséquence, les économistes vont chercher dans les approches microéconomiques la justification des spécifications utilisées en macroéconomie.

La spécification de certaines relations de comportement peut conduire à l'introduction des variables endogènes retardées. Ce ci afin de tenir compte de délais d'ajustement. Dans notre cas, on fait répondre la fonction de demande de monnaie en plus du revenu et du taux d'intérêt et autres variables explicatives, la quantité de monnaie demandée pendant la période précédente (M_{t-1}). La variable endogène retardée est donc donnée et invariante, elle est donc considérée comme une variable exogène en raisonnant purement en statique. Toute fois, le modèle est beaucoup plus riche.

4.1.1 Identification et simultanéité

4.1.1.1 Problème d'identification:

Théoriquement, et conformément aux analyses des chapitres précédents, la fonction de demande de monnaie suggère que la quantité de monnaie désirée par les individus dépend essentiellement d'un groupe de variables. Parmi ces variables, une

variable d'échelle(le revenu réel) , un coût d'opportunité de détention de monnaie (taux d'intérêt) et d'autres variables spécifiques à l'approche envisagée, par exemple le taux d'inflation ou le niveau général des prix , le taux de change, l'innovation financière...

La spécification de la fonction de demande de monnaie est le choix des variables rentrant dans l'analyse dépendent du genre de relation à la quelle le chercheur s'intéresse.

La première difficulté qui se pose à la formalisation de demande de monnaie concerne l'identification de cette fonction. Il n'est pas toujours facile dans la pratique de distinguer dans le stock nominal de monnaie la partie désirée de celle qui est transitoire (non désirée)¹⁷⁷. De plus, il est difficile d'appréhender les fonctions d'offre et de demande de monnaie, lorsque la demande de monnaie est mesurée en termes nominaux. Il devient difficile de distinguer les deux fonctions puisqu'on ne sait pas exactement que se sont les variables tels, le revenu, le taux d'intérêt, le taux de change et le taux d'inflation qui provoquent les fluctuations de demande de monnaie nominale ou inversement. Si se sont les variations de l'offre de monnaie qui influent les déterminants de la fonction de demande de monnaie, comme il le précise D. LAIDLER 1974¹⁷⁸, la quantité de monnaie demandée n'est pas une variable que l'on peut observer. Seule la quantité de monnaie offerte peut être mesuré, et ce n'est qu'en supposant l'équilibre sur le marché monétaire que se second concept peut servir à évaluer le premier. En général, la fonction de demande de monnaie est une fonction décroissante du taux d'intérêt et l'offre de monnaie est considérée comme une fonction positive de ce taux. Graphiquement elle représentée de la manière suivante.

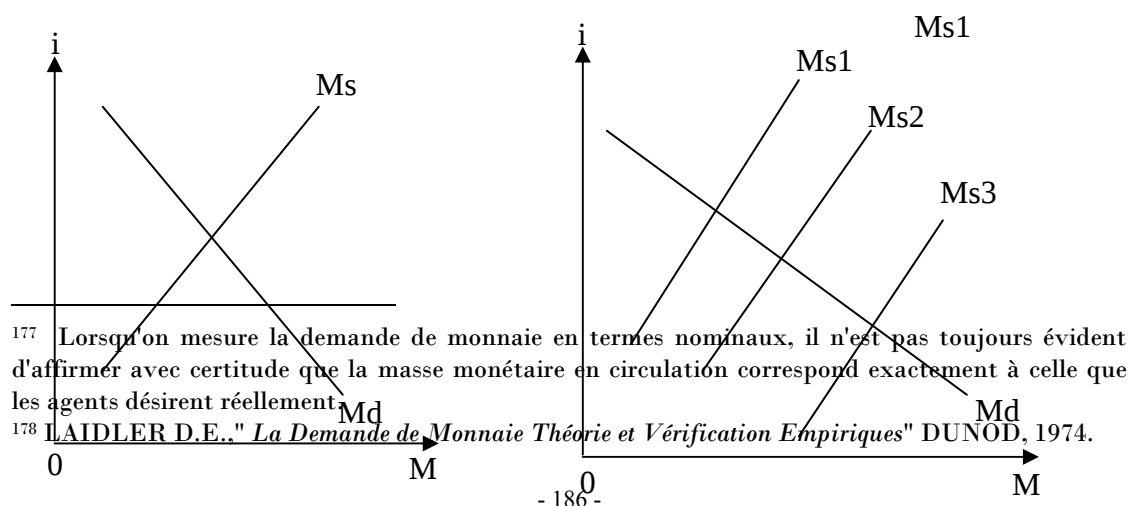


Figure 4.1 et 4.2 : représentent le problème de simultanéité cas 1

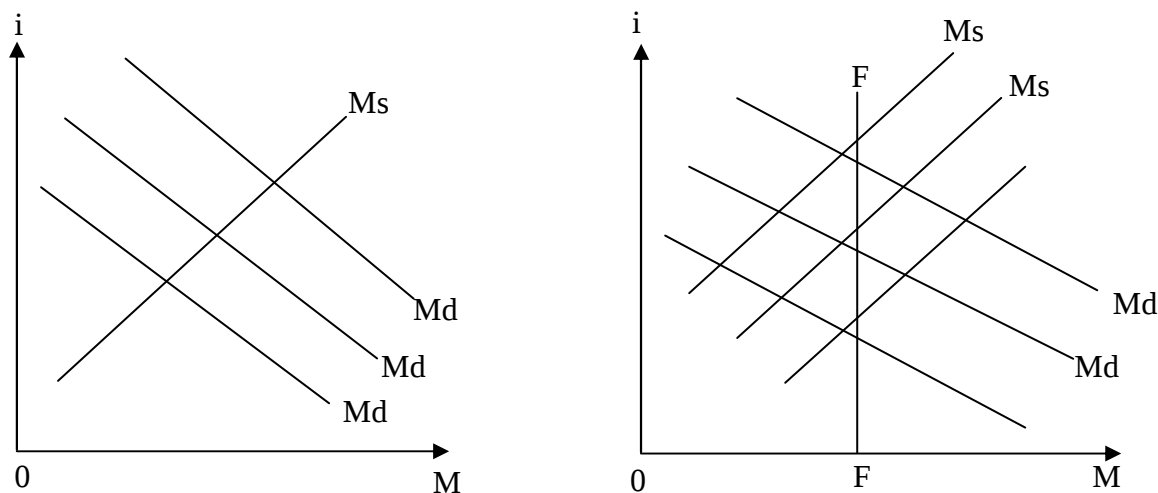


Figure 4.3 et 4.4 : représentent le problème de simultanéité cas 2
Source: D.E. LAIDLER, " *La Demande de Monnaie Théorie et Vérification Empiriques*", DUNOD, 1974,

Le problème reste le même dès qu'il y a plus d'une variable dans l'analyse de la demande de monnaie. Pour surmonter le problème d'identification deux directions de recherches ont été entreprises.

4.1.1.2 L'approche à dominante économétrique

L'existence de deux fonctions (offre et demande) pose un problème au niveau de leurs variables explicatives. La première démarche consiste à s'assurer que la fonction d'offre varie en toute indépendance par rapport à la fonction de demande de monnaie. Si au moins l'une des variables exogène de la fonction d'offre de monnaie n'apparaît pas dans la fonction de demande de monnaie, il est alors possible d'identifier les relations entre la demande de monnaie et ses déterminants.

En suivant cette démarche, il faut connaître les caractéristiques du système économique qui produit ces observations. Il faut déterminer les facteurs exogènes

qui causent les déplacements de demande et l'offre de monnaie. Ces facteurs doivent être indépendants les uns des autres. En plus, il faut s'assurer que la variable exogène qui n'apparaît pas que dans la fonction d'offre de monnaie est strictement indépendante des facteurs déterminants et résiduels de la fonction de demande de monnaie.

Si Par hypothèse, on suppose que la demande de monnaie nominale soit une fonction de revenu national réel, du taux d'intérêt, de niveau général des prix et du taux de change, que l'offre de monnaie nominale dépend de ces mêmes variables et de la base monétaire en circulation et si les fluctuations de la base monétaire ne dépendent ni des fluctuations du revenu national réel et des prix, ni des fluctuations des facteurs aléatoires, alors les paramètres de la fonction de demande de monnaie sont identifiables.

A l'inverse si les autorités monétaires ont l'intention d'accommoder la variation de la base monétaire aux variations des facteurs exogènes de la fonction de la demande de monnaie, tels que le revenu, le taux d'intérêt et les prix, l'identification de la fonction de demande de monnaie sera rendu impossible. En plus, comme Leadler l'affirme, ce n'est pas difficile d'établir ce point puisque le niveau de réserve monétaire (base monétaire) mis à la disposition du système bancaire par la banque centrale figure de façon prééminente dans toute la théorie de l'offre alors, qu'elle n'apparaît dans aucune théorie de la demande¹⁷⁹.

La seconde démarche est de s'assurer que toutes ces observations se situent bien sur la même fonction de demande. C'est-à-dire, il faut supposer de plus que cette dernière demeure fixe entre les observations. Certains auteurs ont d'autre part critiqué l'insuffisance d'identification dans la mesure ou la politique monétaire ne peut apparaître totalement indépendante du niveau de revenu et du prix¹⁸⁰

4.1.2 L'approche empirique

Cette approche est pratique et simple. Il consiste à contourner le problème de l'identification en traitant la relation en termes réels plutôt que nominaux. La quantité de monnaie réelle est prise comme une variable dépendante.

¹⁷⁹ LAIDLER D.E., " *La Demande de Monnaie Théorie et Vérification Empirique*", op. cit. p,

¹⁸⁰ Pour une discussion détaillée, voir T.F.Cooley et S.F. Leroy (1981), dans leur description des voies de recherches et des résultats concernant le problème d'identification et de simultanéité.

Le raisonnement essentiel est que l'offre de monnaie nominale est considérée comme une donnée exogène, donc susceptible d'être contrôlée par les autorités monétaires dont la banque centrale. Par contre, la demande réelle de monnaie est déterminée par les agents économiques qui ajustent leurs encaisses monétaires réelles via les mouvements des prix.

En effet, dans un système économique où les prix varient d'une façon endogène, il n'existe pas de fonction d'offre de monnaie réelle. Mais, seule existe la fonction de demande de monnaie en termes réels.

Si on utilise l'hypothèse selon laquelle la demande réelle de monnaie est indépendante du niveau général des prix et où les variables exogènes de la fonction de demande de monnaie en termes réels sont indépendantes de cette demande et elles sont déterminées ailleurs, on peut alors, identifier les paramètres de la fonction de la demande de monnaie et le problème d'identification de la fonction de demande de monnaie peut être résolu.

4.1.3 Le problème de simultanéité

La méthode de régression considère que les fluctuations inexplicables d'une variable dépendante doivent être attribuées à un facteur d'erreur résiduelle. La méthode de simultanéité se pose quant les variables estimées indépendantes apparaissent et font subir une dépendance réciproque du fait en particulier de l'existence de cross corrélation entre les erreurs résiduelles du système d'équation.

En général, lorsqu'on formalise la fonction de demande de monnaie, on peut appliquer de manière simple la méthode des moindres carrés (MCO). Cependant, il peut arriver que les estimateurs de MCO ne soient pas convergents du fait de l'existence du problème de simultanéité. C'est-à-dire l'existence d'une corrélation croisée entre les perturbations et les variables explicatives. Car, les variables explicatives (exogènes) de la fonction de demande de monnaie tels que le revenu, le taux d'intérêt, les prix et le taux de change sont aussi des variables endogènes du système économique. Ce caractère endogène des variables explicatives pose une série de conséquences pour l'analyse de la méthode des moindres carrés ordinaire. Par exemple si on va prendre la structure du modèle de la consommation dans une économie (Algérie) on a :

$$C = f(Y). \quad (4.1)$$

$$Y = C + I. \quad (4.2)$$

$$I = I^*. \quad (4.3)$$

En général, un modèle économique va intégrer plusieurs relations mais pas une seule. Les variables explicatives pour chacune d'entre elles ne sont pas nécessairement des variables exogènes de ce modèle.

Dans la première relation, la consommation est expliquée par le revenu. Il s'agit là d'une variable explicative ayant à la fois des fondements empiriques et théoriques. Cette relation reflète la façon dont les agents économiques se comportent pour dépenser leurs revenus. On parle de relation de comportement dont la théorie économique et les travaux empiriques vont permettre de préciser certaines caractéristiques minimales de cette fonction. Par exemple, la consommation augmente quand le revenu augmente, donc, $dC/dY > 0$ ou encore que la consommation augmente proportionnellement moins que le revenu soit $d^2 C/d^2 Y < 0$.

La seconde relation explique que la variable expliquée, le revenu Y est égal à la somme de la consommation et de l'investissement. La qualité de l'explication est faible et on parle souvent de relation comptable ou de définition. La troisième relation, n'est là que pour indiquer la limite du modèle qui ne cherche pas à expliquer l'investissement celui-ci est exogène.

Dans ce modèle on trouve une forme d'interdépendance, car, il y'a trois équations. Deux variables expliquées ou endogènes, et on sait que la variable expliquée dans la seconde relation est aussi la variable explicative dans la première équation. Alors on doit procéder à une résolution qui s'explique de la manière suivante.

On doit expliquer uniquement le modèle on fonction de la ou des variables exogènes, par le remplacement dans l'équation de Y la variable C et la variable I fournies par leurs relations, on obtient donc;

$$Y = f(Y) + I^* \text{ et } Y - f(Y) = I^*. \quad (4.4)$$

On trouve une forme linéaire en remplaçant la relation $C = f(Y)$ par $C = cY + a$ ou a représente un constant. Alors cette résolution nous donne l'équation suivante:

$$Y = 1/(1-c) (I^* + a). \quad (4.5)$$

On sait que le multiplicateur Keynésien $k = 1/1-c$, alors $Y = k \cdot (I^* + a)$. Le résultat c'est que la variable endogène ne dépend maintenant que de la variable exogène et des coefficients (c) qui est que la propension à consommer supposée constante.

De la même façon on obtient la valeur de la variable $C = c/1-c (I^* + a)$. Finalement; la résolution du modèle se présente sous forme d'une ou plusieurs relations entre les variables endogènes du modèle C et Y et exclusivement les variables exogènes et les paramètres I^* et a et c . on dit que ce modèle est mis sous la forme réduite par opposition à la forme structurelle.

Dans ce cas il faut révéler Trois remarques:

- L'interprétation des relations de la forme réduite d'un modèle n'est jamais immédiate puisqu'elle synthétise des relations de nature différente et souvent en économie des relations d'équilibre et des relations de comportement. Il s'agit donc d'être prudent.
- Souvent, on limite la forme réduite à une équation concernant la variable considérée comme la plus importante.
- Il est commode de travailler uniquement sur les formes réduites. Mais un problème évident se pose, Compte tenu des paramètres de la forme réduite, peut-on déterminer un et un seul ensemble de paramètres de la forme structurelle ? Dans le cas contraire, le modèle est dit non identifiable et il y a danger à raisonner sur sa forme réduite.

Quand on cherche à résoudre un modèle, la solution peut inclure une variable endogène comme variable explicative dans la forme réduite. Dans le modèle précédent par exemple, si on veut en élargir le champ en considérant par exemple que l'investissement n'est pas purement exogène, mais dépend aussi d'une variable qu'on doit pouvoir expliquer, ici le taux d'intérêt, en écrivant $I = h R + I^*$, la forme réduite sera :

$$Y = k (a + I^*) + h k R \quad (4.6)$$

Le modèle est alors dit sous déterminé et, par exemple, la condition d'équilibre ne nous fournira ici que des couples de variables endogènes compatibles entre elles, mais en aucun cas, le niveau d'une de ces variables associé aux seules variables exogènes. Le niveau des variables exogènes étant donné, on peut toutefois obtenir le niveau d'une des variables, ici Y , pour chaque niveau de l'autre, ici R . On dit

qu'on paramétrise Y par rapport à R . On aura reconnu dans notre exemple la relation dite IS qui donne les couples (Y, R) compatibles avec l'équilibre du marché des biens et services.

A l'opposer, une relation du modèle structurel, en général une relation comptable ou une relation d'équilibre, ne peut être satisfaite que s'il existe une relation bien spécifiée entre les variables exogènes et/ou les paramètres. On dit alors que le modèle est sur déterminé. L'exemple le plus connu est ici le modèle de Harrod dans lequel la condition de croissance équilibre impose une relation parfaitement définie entre les paramètres qui caractérisent la forme structurelle du modèle. On doit en effet avoir $s/Cr = n$.

Il est raisonnable de penser qu'un modèle doit être juste déterminé. Dans le cas contraire, en effet, il n'est pas complet et ce qui devrait pouvoir être expliqué ne l'est pas. Une façon de s'en sortir est d'immobiliser les variables en excès, ici soit R , soit Y . Ceci n'est cependant que formel et n'est satisfaisant que si on peut fournir de bonnes raisons pour opérer de cette manière. On dira par exemple :

« A l'équilibre le revenu se situe à son niveau de plein emploi, Y^* ». Une autre façon de procéder consistera à rajouter une ou des relations au modèle de telle sorte que l'ensemble soit juste déterminé. Mais en procédant comme cela, on est souvent condamné à élargir le champ concerné par la modélisation. Dans l'exemple précédent, on peut rajouter des relations caractérisant le marché monétaire et écrire:

$$M_s = M_d \text{ l'offre de monnaie est égale à la demande de monnaie} \quad (4.7)$$

$$M_s = M^* \text{ l'offre de monnaie est exogène} \quad (4.8)$$

$M_d = u Y + v R$ la demande de monnaie dépend du revenu et du taux d'intérêt (4.9). Si on résout ce modèle propre au marché monétaire, on s'aperçoit qu'il est également sous déterminé et on obtient :

$$R = -u/v Y + M^*/v \quad (4.9)$$

qui correspond à la relation dite LM. En rapprochant cette relation de la forme réduite IS on obtient un modèle complet mis sous une forme semi réduite :

$$Y = k(a + I^*) + hk R \quad (4.10)$$

$$R = -u/v Y + M^*/v, \quad (4.11)$$

et qui correspond au modèle dit IS-LM. qui ici est juste déterminé.

Encore, les problématiques en économie, Quand un modèle est sous ou sur déterminé, il est nécessaire de chercher à lever cette sur ou sous détermination. Dans le premier cas, nous avons noté qu'on pouvait le faire en complétant le modèle afin d'arriver à un modèle juste déterminé, ou encore immobiliser une ou des variables endogènes, ce qui revient à intégrer autant d'équations supplémentaires. Toutefois, il n'y a pas nécessairement accord entre tous les économistes sur la façon de procéder. Chaque méthode, pour laquelle on donnera des justifications, conduira à une problématique spécifique ou encore à une théorie spécifique.

Un bon exemple peut être trouvé dans le modèle IS-LM où d'évidence, la présentation qui en est faite plus haut ne prend pas en compte une variable économique importante, à savoir le niveau général des prix. Ceci est sans influence au niveau du marché des biens et services où toutes les relations sont spécifiées en valeur réelle. Par contre, s'agissant du marché monétaire, la demande de monnaie est une demande d'encaisses réelles, alors que seule l'offre nominale de monnaie peut être considérée comme exogène. La forme réduite du modèle du marché monétaire s'écrit donc:

$$M^*/P = u_y + v r \quad (4.12)$$

où y et r représente respectivement le revenu réel et le taux d'intérêt. La sous détermination est donc plus importante et il

existe une infinité de relations de type LM entre r et y . On dit aussi que LM est paramétrée par rapport au niveau général des prix. Le modèle complet IS-LM s'écrit alors (les minuscules représentants des variables réelles ou des paramètres et les majuscules des variables nominales):

$$y = k (a + l^*) + h k r \quad (4.13)$$

$$r = -u/V y + 1/v (M^*/P) \quad (4.14)$$

Ce modèle est sous déterminé puisqu'il se ramène à deux équations et trois variables endogènes r , y , et P .

Deux visions du fonctionnement de l'économie vont alors s'opposer pour lever la sous détermination :

- Les néoclassiques vont considérer que y est intégralement déterminé par l'offre qui se situe à un niveau assurant le plein emploi, y^* . Sur cette base on peut

alors compléter le modèle en ajoutant une fonction de production de type $y = f(n)$ et les équations permettant de modéliser le marché du travail;

$$n_d = d(W/P) ; n_s = s(W/P) ; n_d = n_s. \quad (4.15)$$

Le modèle complet est alors juste déterminé et comprend 6 équations et 6 variables endogènes $r ; y ; P ; n_d ; n_s$ et w le taux de salaire réel. L'hypothèse fondamentale sur laquelle repose cette façon de voir est la parfaite flexibilité des prix et des salaires nominaux qui vont garantir un retour instantané à l'équilibre sur le marché du travail, donc l'absence de chômage involontaire.

- Keynes, au contraire, considère que les prix ne sont pas flexibles à la baisse et raisonne pour un niveau général des prix donné. Le modèle IS-LM est donc juste déterminé et la prise en compte du marché de travail est inutile pour déterminer le niveau de y qui dépend des conditions de la demande globale. La prise en compte du marché du travail ne sert alors qu'à montrer la compatibilité entre l'équilibre et un chômage involontaire. Lorsqu'un modèle est sur déterminé, c'est à dire qu'il existe une équation de trop qui peut toujours se ramener à une contrainte supplémentaire entre les paramètres et/ou les variables exogènes, il s'agira d'endogénéiser l'une de ces variables ou paramètres. Là aussi apparaissent des différences de conception entre les économistes reposant sur des visions différentes du fonctionnement de l'économie.

Nous rencontrons ce problème à propos de la théorie de la croissance. Harrod obtient une condition de croissance équilibrée assurant l'équilibre sur le marché des biens et le marché du travail qui s'écrit : $s/Cr = n$

où s est la propension à épargner ; Cr le coefficient de capital, qui sont deux paramètres de son modèle et n le taux de croissance naturel qui est exogène. Le modèle est donc sur déterminé et, pour lever cette sur détermination, on peut soit endogénéiser Cr en intégrant une possibilité de substitution des facteurs de production, ce que feront les néoclassiques, soit endogénéiser s , Cr restant donné, ce que feront les néo-keynésiens, soit enfin endogénéiser n dans une optique néo-malthusienne ou néo-marxiste.

Si le système économique dans son ensemble supporte des chocs aléatoires, ces derniers pourront causer des fluctuations aléatoires et simultanées de toutes les variables endogènes du système. Ces fluctuations seront corrélées entre elles. Cette

corrélation ne résulte pas des relations causales entre les variables (à savoir la fluctuation d'une variable endogène engendre la fluctuation de l'autre), mais elle réside tout simplement dans le fait d'une cause commune dans le système. Quant les estimateurs régressent la demande de monnaie sur par exemple le revenu et le taux d'intérêt, ils considèrent en effet que toutes les variations de la demande de monnaie résultent de celle de ses déterminants. Dans le cas où les éléments du système subissent des chocs communs que les estimateurs ignorent, les estimations des paramètres de la fonction de demande de monnaie seront biaisées par l'influence de l'erreur résiduelle commune à toutes les variables (endogènes et exogènes)¹⁸¹.

En effet, lorsqu'on régresse la demande de monnaie sur le revenu et le taux d'intérêt, on suppose que toutes les variations de la demande de monnaie sont provoquées par ses déterminants.

Par contre si les séries sont affectées les unes les autres par des chocs aléatoires, alors les estimations des paramètres de la fonction de demande de monnaie seront biaisées par la contamination d'une ou toutes les variables explicatives provoquées par l'erreur résiduelle. L'ampleur du biais dépend de celui de l'erreur résiduelle. Alors, si leur contribution est faible le biais sera mineur.

Les études empiriques réalisées sur des échantillons suffisamment grands, fondées sur des données à l'intervalle long souffrent relativement moins du biais d'équation simultanée, car les données ont un degré d'agrégation temporelle relativement élevé. L'utilisation de ses données permet de réduire le problème de simultanéité du fait que les fluctuations aléatoires ont tendance à se compenser au sien même de chaque série.

Par exemple, les variations du revenu réel durant une longue période sont données par l'effet de tendance, quant aux fluctuations résiduelles provoquées par les chocs exogènes frappant le système économique, elles produisent un effet relativement faible.

En revanche, lorsqu'on effectue des travaux empiriques sur des petits échantillons le problème de simultanéité tend à s'accroître. Dans ce cas de figure, on utilise des techniques d'estimations alternatives à celle de la MCO. L'une des méthodes souvent utilisée par les économètres est celle des doubles moindres carrés (DMC) ou

¹⁸¹ Voir, Laidler D., " *The Demand For Money*", HAPER et ROW, New York, 1985.

encore celle des variables instrumentales (VI) qui permet de supprimer la contamination des variables explicatives. La méthode de DMC a été utilisée pour l'estimation de la fonction de demande de monnaie à court terme. Les écarts de résultats observés par rapport à l'utilisation de la méthode de MCO apparaissent forcément moins, ceci pourra l'expliquer par le fait que le biais des équations simultanées n'est pas un problème majeur dans le contexte de la fonction de demande de monnaie. Cooley et Leroy (1981)¹⁸² ont mentionné que les instruments exogènes ne sont pas correctement choisis. Beaucoup d'études fondées sur des données à l'intervalle long, utilisent les encaisses réelles ou la vitesse de circulation de la monnaie comme variables dépendantes. Les arguments développés plus hauts montrent que ces études ont traité d'une façon adéquate mais implicite les problèmes d'identification et de simultanéité. A long terme, les variables touchées par les chocs tendent à atteindre leur chemin d'équilibre et l'effet des fluctuations aléatoires est moins important. Ceux ci contribuent à l'atténuation des problèmes d'identification et de simultanéité. Plusieurs économètres se sont explicitement attaqués à ce problème¹⁸³. Ils ont testé la fonction de demande de monnaie dans des modèles macroéconomiques concernant principalement les Etat Unis d'Amérique Ainsi que d'autres pays tels que le Royaume Unis, l'Allemagne et la France. Les résultats de ces études approchent des ceux obtenus par la méthode de MCO. L'idée à retenir c'est qu'il n'existe pas une grande différence au niveau des résultats entre les différentes méthodes.

Pour notre étude expérimentale on opte pour des méthodes économétriques récentes pour surmonter ces problèmes, en utilisant pour un processus VAR la théorie de la cointegration. Plusieurs travaux en cette direction ont été menés à partir des années quatre-vingt-dix en utilisant la procédure de Johansen¹⁸⁴, comme

¹⁸² Cooley T. et Leroy S., "Identification And Estimation of Money Demand," *American Economic Review* 1981, vol.

¹⁸³ On peut consulter les travaux de Brunner K. et Meltzer H. (1963), Walter A. (1965), Jonson P. (1976), Boghton J. (1979), Lieberman C. (1980), Laidler D. (1974,1980,1982,1985), Carr J. et Darby (1981), Judd J. et Scadding J.L. (1982).

¹⁸⁴ Johansen S., " Statistical analysis of cointegration vectors", *Journal of Economic Dynamics and Control*, 1988,12, p. 231-254.

par exemple, Johansen et Juselius (1990)¹⁸⁵, Hendry et Ericsson (1991)¹⁸⁶, Hansen(1992)¹⁸⁷, Baba, Hendry et Starr (1992)¹⁸⁸, Hoffman et Rasche (1996)¹⁸⁹, Ericsson et Sharma(1998)¹⁹⁰, Fagan et Henry(1999)¹⁹¹, et récemment par, Avouyi-Dovi et al 2002b)¹⁹², Stracca (2001)¹⁹³, Diop, Jacquinot et Sahuc (2002a)¹⁹⁴, Bruggeman, Calza et Sousa (2002)¹⁹⁵, Dubois E., Lardic S. Et Mignon (2003)¹⁹⁶.

Les techniques économétriques standard supposent que les séries étudiées soient stables au cours du temps (hypothèse de stationnarité). La plupart des séries économiques ne vérifient pas cette hypothèse, ce qui suppose, lorsqu'on souhaite étudier les relations qui les lient, de mettre en œuvre des techniques spécifiques. Après une présentation des conséquences de la non stationnarité et des tests qui permettent de la détecter, la formation présente la notion de cointegration qui caractérise des séries non stationnaires (intégrées) dont une combinaison linéaire est stationnaire.

Depuis les travaux de Johansen, cette approche est généralement présentée dans le cadre d'une analyse multi variée et permet de spécifier des relations stables à long

¹⁸⁵ Johansen S., Juselius K., "Maximum likelihood estimation and inference on cointegration with application to the demand for money", *Oxford Bulletin of Economics and Statistics*, 1990, 52, p. 169-210.

¹⁸⁶ Hendry D.F. et Ericsson N.R., "An Econometric Analysis of U.K. Money Demand in Monetary Trends...by M. Friedman et Anna J. Schwartz", *American Economic Review*, 1991, 81, 8-38.

¹⁸⁷ Hansen B.E., "Tests for Parameter Instability in Regressions with I(1) Processes", *Journal of Business & Economic Statistics*, 1992, 10, 321-335.

¹⁸⁸ Baba Y., Hendry D. et Starr R., "The Demand for M1 in the USA, 1960-1988", *Review of Economic Studies*, 1992, 59, 25-61.

¹⁸⁹ Hoffman D.L. et Rasche R.H., "Aggregate Money Demand Functions", Kluwer Academic Publisher, 1996.

¹⁹⁰ Ericsson N.R. et Sharma S., "Broad Money Demand and Financial Liberalisation in Greece 1998", *Empirical Economics*, 1998, 23, 417-436.

¹⁹¹ Fagan G. et Henry J., "Long-Run Money Demand in the EU: Evidence for Area-Wide Aggregates", in *Money Demand in Europe*, Lutkepohl H. and Wolters J. ed., Heidelberg, Physica-Verlag, 1999, 217-40.

¹⁹² Avouyi-Dovi S., Diop A., Jacquinot P. et Sahuc J.-G., "Estimation d'une fonction de demande de monnaie pour la zone euro", DGEI-DEER, *Note interne*, réseau demande de monnaie, r02068, Banque de France, 2002b.

¹⁹³ Stracca L., "The Functional Form of the Demand for Euro Area M1", *ECB Working Paper*, 2201, 51.

¹⁹⁴ Diop A., Avouyi-Dovi S., Fonteny E., Gervais E., Jacquinot P. et Sahuc J.-G., "Bases de données pour la demande de monnaie en zone euro", DGEI-DEER, *Note interne*, réseau demande de monnaie, r02054, Banque de France, 2002.

¹⁹⁵ Bruggeman A., Calza A. et Sousa J., "An Evaluation of Different Estimates of the Real Money Gap for the Euro Area", *ECB mimeo*, 2002.

¹⁹⁶ Dubois E., Lardic S. et Mignon V., "The Exact Maximum Likelihood Based-Test for Fractional Cointegration: Critical Values, Power and Size", *Document de travail THEMA*, Université Paris X, 2003.

terme tout en analysant dans le même temps la dynamique de court terme des variables considérées. Ce sera donc également l'occasion de présenter les modèles VAR qui constituent une généralisation des processus AR au cas multi varié.

4.2 Définition des variables de la fonction de demande de monnaie.

Pour la réalisation des tests empiriques, il est nécessaire de procéder à une définition précise de toutes les variables qui rentrent dans la fonction de demande de monnaie. Sur la base des hypothèses adoptées nous allons examiner successivement les problèmes concernant les définitions de ces variables c'est-à-dire le concept de monnaie, la variable d'échelle, les variables de taux de rendements, les prix et autres variables explicatives.

4.2.1 La variable expliquée: La monnaie

Les multiples fonctions de la monnaie donnent une certaine hétérogénéité aux services rendus par les actifs monétaires. La monnaie est un actif particulier. Elle constitue une forme sous laquelle un agent peut concrétiser la richesse¹⁹⁷. La monnaie génère pour ses détenteurs un flux de services. Mais le problème de définition d'actif dont les services rendus sont assez similaires entre eux et suffisamment différents de ceux des autres.

En fonction de la nature des services rendus par la monnaie, du fait que les fonctions de la monnaie sont multiples et qu'il existe des nuances institutionnelles entre les pays et les innovations au cours du temps, la définition de la monnaie est parfois rendue difficile. Nous pouvons regrouper ces actifs monétaires et déterminer ainsi une définition au sens étroit ou bien au sens large comme nous l'avons vu dans le premier chapitre.

La définition étroite de la monnaie ou bien M1 intègre la monnaie fiduciaire plus la monnaie scripturale (c'est la disponibilité monétaire). En ajoutant les dépôts à terme (la quasi monnaie) à M1 on obtient l'agrégat monétaire M2.

En règle générale, de la meilleure définition de la monnaie doit prendre en compte le degré de substitution des actifs, la stabilité de la fonction de la demande de monnaie et l'homogénéité du comportement de l'offre.

¹⁹⁷ Voir, Friedman M. et Schwartz D.J. (1982), F. Aftalion et Poncet P. (1987), Poncet P. et Portrait R. 1980).

La détention des encaisses monétaires expliquées pour le motif de transaction, s'appuie sur le fait que la monnaie est un moyen d'échange (monnaie active) dans la mesure où elle doit être définie par le critère de la liquidité, sa définition en tant que disponibilité monétaire M1 apparaît comme le meilleur choix.

Les motifs de précaution et de spéculation (monnaie oisive) mettent l'accent sur l'incertitude des transactions et les anticipations sur le taux d'intérêt. Dans ce cas là ils peuvent être détenus pour des dépenses probables non seulement les disponibilités monétaires mais aussi autres actifs financiers dont le coût de conversion en disponibilité monétaire M1 est très faible même négligeable tels que les dépôts à terme ou épargne, donc ici l'agrégat monétaire M2 apparaît comme le meilleur choix.

Il faut noter qu'il n'est pas toujours aisé de justifier l'introduction des dépôts à terme dans la définition de la monnaie, si on adopte l'hypothèse selon laquelle la monnaie représente un instrument d'échange. Mais pour estimer la demande de monnaie spéculative une définition plus large que M1 en tant que variable dépendante est certainement préférable pour des raisons suivantes:

- la caractéristique d'actif financier dont la valeur ne varie pas avec les fluctuations du taux d'intérêt s'applique non seulement à l'agrégat M1 à l'ensemble de liquidité.
- Les liquidités économiques, autre que M1 tout en demeurant non risquées M2 produisent à leurs détenteurs des rendements supérieurs et elles sont plus attrayantes que les disponibilités monétaires.

Des travaux réalisés par D. Laidler (1985)¹⁹⁸ ont montré que l'intégration des dépôts à terme ne change en rien la stabilité de la fonction de demande de monnaie. Dans certains cas, elle peut même améliorer sensiblement leurs comportements sur des tests. D'autres auteurs ont testé pour le même pays à la fois l'agrégat monétaire au sens étroit M1 et l'agrégat aux large M2 (KOMERAK et MELECKY : 2001)¹⁹⁹. Les problèmes de définition de l'agrégat monétaire s'aggravent à partir des années quatre-vingt, du fait essentiellement de la libéralisation financière et des diverses

¹⁹⁸Laidler D., (1985),op. cit.

¹⁹⁹ Komerak L. et Melecky M., " Demand For Money in The Transition Economy: The Case of CZECH Republic (1993-2001)," *Warwick Economic Research Papers*, 2001.

innovations financières²⁰⁰, c'est donc, le développement des institutions financières et du système financier (l'apparition de la dénomination des trois D; la déréglementation, désintermédiation et décloisonnement).

D'après Gowland (1991)²⁰¹ l'innovation peut se définir comme « l'introduction d'un nouveau produit sur le marché ou la production d'un produit existant, mais d'une nouvelle manière ». Dans les établissements financiers, le rôle de la technologie n'est pas vraiment clair, car les changements que l'on doit apporter à « la méthode de production » pour introduire un nouveau produit ne sont pas énormes : les prêts sont, globalement, toujours les mêmes avant et après. L'innovation Les banques sont constamment à la recherche d'innovations financières pour accroître leurs profits. En ce sens, l'innovation financière ne peut être considérée comme un événement contemporain. Cela a toujours existé depuis que les banques sont banques. A la fin des années 1980 et au début des années 1990 suite à la vague intensive de déréglementation des marchés financiers, de même que sur certains développements des secteurs des communications et des technologies informatiques. Ces deux facteurs ont profondément transformé l'activité de l'intermédiation financière entre emprunteurs et prêteurs, activité traditionnellement aux mains des banques. En conséquence, d'autres intermédiaires financiers et même certaines institutions non financières (par exemple, les compagnies d'assurances) ont commencé à s'approprier ce rôle. Devant le rôle grandissant des marchés financiers, et confrontées au déclin de leur activité traditionnelle, les banques ont commencé à se comporter de plus en plus comme des courtiers. Afin d'améliorer leur position de courtiers et de mieux faire face à la concurrence croissante sur leur marché traditionnel, les banques se sont lancées activement dans la recherche d'innovations financières. L'environnement économique leur a également été favorable, à ce nouveau stade de leur évolution, car la fluctuation grandissante des taux d'intérêt et des taux de change était devenue une source intéressante de profits en raison des possibilités d'arbitrage qui

²⁰⁰ Gozland D., " Financial Innovation in Theory and Practice", in Green C. J. and D.T. Llewellyn (eds.), *Surveys in Monetary Economics*, Oxford, Basil Blackwell, 1991 vol. 2

leurs étaient offertes. Gowland (idem, p.92), par exemple, considère six caractéristiques comme déterminants de cette nouvelle vague d'innovation financière : a - le remplissage de la coquille du marché financier , b- les marchés secondaires, c- la négociabilité, d- le rechargement, e- l'internationalisation, et f- la concurrence et le marketing. Toutes ces caractéristiques aident à comprendre l'ampleur de cette vague. Plus généralement, l'innovation financière désigne la création de nouveaux produits financiers destinés à combler les vides du marché afin de remplir l'ensemble des marchés et, par conséquent, permettre les transferts de fonds de prêteurs à emprunteurs de manière plus efficace. Silber (1983, p.89)²⁰² argumente également en faveur de l'innovation comme moyen de desserrer les contraintes financières imposées aux établissements, c'est-à-dire comme tentative pour desserrer les contraintes auxquelles doit faire face un établissement lorsqu'il essaie d'atteindre ses objectifs. L'un des points importants de l'innovation en matière de finance est, comme le souligne Gowland (ibid., p.98), que de telles innovations font rarement l'objet de brevets, contrairement à ce qui se passe dans l'industrie. D'après cet auteur, il existe trois raisons pouvant expliquer ce nombre relativement faible de dépôts de brevets dans le milieu de la finance : a) c'est juridiquement impossible, b) ce qui importe, c'est d'être le premier ou, du moins, d'être perçu comme étant le premier, c) ce n'est pas souhaitable pour des raisons de « réseaux ». La stratégie optimale consiste à mettre sur le marché des produits qui doivent être bientôt offerts par d'autres établissements. Goodhart (1989, p.377)²⁰³ voit également que le processus d'innovation implique une série d'étapes imposées par la concurrence, l'évolution technologique et la déréglementation qui ont fortement réduit le coût de l'intermédiation (c'est-à-dire l'écart effectif entre le taux de rémunération des dépôts et le taux de crédit) . Autrement dit, l'innovation représente une incitation pour les banques et autres intermédiaires financiers car elle leur permet de compenser la réduction de l'écart entre le taux de rémunération des dépôts et le taux de crédit.

²⁰² Silber.I.W., " The Process of Financial Innovation", *American Economic Review*, 1983 73(2), Mai, pp. 89-95.

²⁰³ Goodhart C., "*Money, Information And Uncertainty*", Londres, Macmillan 1989.

De nouveaux produits sont apparus tout en gardant leur ancienne caractéristique d'actif monétaire. Aux Etats Unis Amérique on trouve le "*Now Account*" (compte cheque rémunéré), au Royaume Unis c'est le découvert autorisé pour la carte de crédit. Le découvert c'est généralisé maintenant dans la plus part des pays développés.

Le Canada offre l'exemple d'une des économies qui a connu la plus grande déréglementation depuis de nombreuses années. Les taux créditeurs ont pu évoluer en fonction des taux du marché et il n'y a pas eu, aussi bien dans les années 70 que dans les années 80 de mesures d'encadrement du crédit ou de contrôle des mouvements de capitaux. Des instruments négociables du marché monétaire ont été créés dès les années 60.

Il n'est guère surprenant, dans un contexte de constante innovation financière, d'observer des signes d'instabilité de la demande de monnaie. L'innovation financière est devenue très importante dans la détermination et la spécification de la fonction de demande de monnaie et ses fluctuations et surtout pendant les périodes de l'inflation (Arrau, De Gregorio, Reinhart 1991)²⁰⁴. L'apparition de nouveaux produits et de nouveaux intermédiaires tend à estomper la distinction entre actifs monétaires et non monétaires et à modifier le comportement financier des agents économiques. Il est juste de dire que le rôle précis de l'innovation financière à cet égard n'est pas encore connu, mais nous avons toutes les raisons de penser que sa prise en compte devrait permettre d'expliquer cette évolution de la demande de monnaie et de mieux évaluer ses conséquences pour les tensions inflationnistes futures.

Cette constatation implicite de surveiller l'évolution potentielle de la définition de la monnaie; surtout pour tester la stabilité de la fonction de demande de monnaie. Ainsi, un effort de recherche sur les nuances institutionnelles est nécessaire pour comparer et généraliser les résultats obtenus dans différents communautés économiques et surtout dans les pays en voie de développement ou se domaine enregistre un retard énorme car peu d'études étaient consacrées à ce domaine et ou le système économique monétaire est toujours dans un état embryonnaire.

²⁰⁴ Arrau P., J. Gregorio De, Reinhart C., "Financial Innovation and Money Demand: Theory and Empirical Implementation," *Working paper N0.585, World Bank*, 1991.

Il convient toutefois, de souligner que relatives à la monnaie peuvent changer d'un pays à un autre en fonction des réalités institutionnelles particulières et de la réglementation monétaire en vigueur. Ce pendant, le rôle institutionnel de conseil que joue le fond monétaire international et la banque mondiale en matière de politique monétaire tend à harmoniser les définitions de la monnaie et des agrégats monétaires dans les différents pays.

4.3 Les variables explicatives

4.3.1 La variable d'échelle

La définition de la variable d'échelle concerne essentiellement le choix entre deux groupes de variables: le revenu et la richesse.

Dans la typologie des différents types de revenus, il existe le revenu courant et le revenu permanent. On définit le revenu permanent comme le revenu moyen qu'un individu touchera jusqu'à sa retraite. Dans la littérature économique on considère le plus souvent le revenu courant comme une approximation de volume de transaction dans l'économie.

Le revenu joue un rôle primordial dans les tests empiriques. La mesure du revenu ne présente pas beaucoup de difficulté, bien qu'il existe des mesures différentes, tels que le produit national brut (PNB)²⁰⁵ le produit national net (PNN)²⁰⁶ ou encore le produit intérieur brut (PIB)²⁰⁷. On peut dire que les résultats empiriques obtenus en utilisant l'une ou l'autre de ses variables d'échelles dans la fonction de demande de monnaie ne sont pas très différents car elles n'apportent pas des divergences sensibles.

La notion de la richesse a été récemment précisée. Les premières mesures de richesse utilisées dans les travaux empiriques sont de celles de la seule richesse monétaire. Cependant il n'est pas facile de constituer cette série à cause de l'indisponibilité des données détaillées permettant de construire une série temporelle longue comme avec le revenu.

²⁰⁵ Le PNB est la valeur totale de la production finale de biens et services des acteurs économiques d'un pays donné au cours d'une année donnée. A la différence du PIB, le PNB inclut les produits nets provenant de l'étranger, c'est-à-dire le revenu sur les investissements nets réalisés à l'étranger.

²⁰⁶ PNN = PNB - amortissements

²⁰⁷ le PIB est égal à la somme des emplois finals intérieurs de biens et de services (consommation finale effective, formation brute de capital fixe, variations de stocks), plus les exportations, moins les importations

En outre il est difficile de dire quel degré de consolidation on doit appliquer pour disposer d'une série agrégée en partant des séries désagrégées. Le problème consiste à calculer un agrégat qui contient la richesse non humaine de tous les agents économiques, ou de consolider les bilans des propriétaires finaux.

A cause de ces inconvénients, la richesse non humaine est donnée par une conception plus large. La richesse au sens de Friedman qui est plus couramment utilisée dans les travaux empiriques dont la mesure directe de cette conception pose énormément de problèmes. Mais indirectement, elle peut être remplacée par le revenu anticipé ou le revenu permanent. Selon la logique suivante, la richesse est la valeur présente du revenu futur anticipé et si le taux d'actualisation reste approximativement constant, la richesse variera de la même façon que le revenu futur.

Ainsi, les individus peuvent souvent former des anticipations au niveau de leurs revenus. Alors la question qui se pose c'est quelle formation des anticipations de revenu? Anticipations adaptives ou anticipations rationnelles?

Comment anticiper puisque le temps et l'incertitude sont deux inconnus. Les agents décident sur la base de variables (revenu ou prix) dont ils ne connaissent pas la valeur future. L'incertitude chez Keynes est un facteur d'instabilité. La formation des anticipations reste inexpliquée donc, anticipations exogènes. Keynes se laisse guider par un phénomène extérieur pour savoir si le climat est propice "épaisseurs du Times".

Les modèles d'anticipation formulés par les économistes peuvent être regrouper en deux grandes catégories:

- le premier groupe considère chaque variable par exemple le taux d'inflation, comme un processus stochastique et ne prend en compte que les valeurs passées et présentes de cette variable. C'est le concept commun sous le nom d'anticipation extrapolative et adaptive.
- le deuxième groupe se compose des anticipations dites rationnelles (au sens Muth) ou toutes les informations accessibles sont prises en compte pour la formation des anticipations.

4.3.1.1 Les anticipations adaptives

L'idée des anticipations adaptives c'est que l'évolution passée d'une variable permet de présenter de façon linéaire l'évolution anticipée de cette série, sur la base des travaux économétriques sur les retards échelonnée. La théorie a mis en relief des exemples comme la dynamique de l'hyperinflation. CAGAN (1956)²⁰⁸ montre que les agents prévoient l'évolution des prix, à partir de l'évolution passée, corrigée des erreurs du passé. Les agents n'apprennent que les erreurs du passé et ne concertent pas d'autres sources d'information. Il semble, de plus, que les agents répètent certaines erreurs. Donc la rationalité est limitée puisque l'erreur existe ²⁰⁹. Les anticipations extrapolatives constituent l'exemple suivant:

$$X_t^* = X_{t-1} + b(X_{t-1} - X_{t-2}) \quad \text{avec } 0 \leq b < 1 \quad (4.16)$$

Où X_t^* présente la valeur anticipée à la période courante. Cette dernière dépend de la deuxième période X_{t-1} et de l'avant dernière période - X_{t-2} . On extrapole ainsi l'évolution future à partir de l'évolution passée. Si $b = 0$ alors $X_t^* = X_{t-1}$, cette hypothèse correspond aux cas des anticipations naïves. Les agents s'attendent à voir dans l'avenir ce qu'ils ont connu dans le passé jusqu'à maintenant.

La prévision optimale de X_{t-1} si est seulement si X_t suit une marche au hasard, autrement dit un processus autoregressive integrated moving average (AR IMA).

Cela signifie un processus autorégressif intégré de moyenne mobile (0, 1,0). Si la série X_t représente le taux d'inflation, alors la représentation extrapolative indique que le taux d'inflation anticipé est égal au taux d'inflation observé à la période précédente corrigée de l'écart d'inflation entre la dernière et l'avant dernière période. La correction est une fonction du paramètre b . on peut tenir compte d'un passé plus ancien en le pondérant de manière arithmétique suivante:

$$X_t^* = \sum_{j=1}^p b_j X_{t-j} \quad \text{avec } 0 < b_j < 1 \text{ et } \sum_{j=1}^p b_j = 1 \quad (4.17)$$

Où alors géométrique:

$$X_t^* = \sum_{j=1}^p b_j^2 X_{t-j} \quad \text{avec } 0 < b_j < 1 \text{ et } \sum_{j=1}^p b_j^2 = 1 \quad (4.18)$$

²⁰⁸ CAGAN P. (1956), op. cit .

²⁰⁹CAGAN P.,idem.

Dans l'analyse de la fonction de demande de monnaie on utilise souvent le processus d'apprentissage pour mesurer le revenu permanent. Ce processus implique que les agents économiques tiennent compte de leur expérience passée pour saisir le futur en donnant un poids plus important à leur expérience récente. Donc, les anticipations adaptives tiennent compte de ces erreurs du passé en révisant l'anticipation d'une fraction α de la dernière erreur d'anticipation ($\alpha \neq 1$). On écrit donc;

$$X_t^* - X_{t-1}^* = \alpha (X_t - X_{t-1}^*) \quad (4.19)$$

$$\text{Ou } X_{t-1}^* = \alpha X_{t-1} + (1 - \alpha) X_{t-1}^* \quad (4.20)$$

Si cette équation est exacte, donc elle est pour la période passée et la période d'avant aussi exacte. Algébriquement on peut dériver de cette équation les équation suivantes:

$$X_{t-1}^* = \alpha X_{t-2} + (1 - \alpha) X_{t-2}^* \quad (4.21)$$

$$X_{t-2}^* = \alpha X_{t-3} + (1 - \alpha) X_{t-3}^* \quad (4.22)$$

$$X_{t-3}^* = \alpha X_{t-4} + (1 - \alpha) X_{t-4}^* \quad (4.23)$$

Et ainsi de suite, en substituant, on arrive à écrire :

$$X_t^* = \alpha X_{t-1} + \alpha (1 - \alpha) X_{t-2} + \alpha (1 - \alpha)^2 X_{t-3} + \alpha (1 - \alpha)^3 X_{t-4} + \dots \quad (4.24)$$

Plus α est grand, plus en accord de l'importance à la valeur effectivement observé dans le passé et moins la valeur anticipée dans le passé et significative.

A cause de leur simplicité, les modèles adaptifs ont été largement utilisés dans les recherches empiriques concernant la fonction de la demande de monnaie. Ces modèles manquent cependant de fondement économique précis d'une part et produisent des erreurs systématiques d'anticipations d'autre part. D'où il manque de la rationalité ou ils ne sont pas rationnels.

4.3.1.2 Les anticipations rationnelles

Les années soixante constituent une période charnière dans le développement des écoles de pensée. Chacune d'elle annonce une orientation différente au concept de rationalité. Simon (1956), s'emploie à préciser ses idées sur la rationalité. Pour lui les individus sont souvent amenés à prendre des décisions dans un univers où ils ne disposent que des possibilités de calculs limitées. En revanche, Muth

(1961)²¹⁰aborde le problème dans une perspective complètement différente. Il indique qu'il faut que les anticipations des individus soient compatibles avec les modèles qui sont supposés expliquer à priori les comportements. Si les agents sont rationnels pour maximiser le profit et l'utilité, ils sont rationnels lorsque' ils anticipent l'avenir. Puisque l'information est rare (faits et enseignements de la théorie économique), alors, ils l'utilisent au mieux sans la gaspiller. Les anticipations sont rationnelles lorsqu'elles suivent les prévisions de la théorie économique dominante. Lorsque l'information est parfaite, les incertitudes sont éliminées (valeurs anticipées= valeurs réelles).L'hypothèse des anticipations rationnelles est un élément essentiel de la nouvelle macroéconomie comme nous l'avons déjà vue dans le chapitre précédent (crédibilité et efficacité de la politique économique). L'idée de base est que le processus de formation des anticipations repose sur le "vrai modèle" structurel de l'économie²¹¹.

L'anticipation de la valeur future d'une variable est dite rationnelle si , elle est égale à l'espérance mathématique de celle-ci calculée sur la base de toutes les informations pertinentes et disponible. Formellement cette condition est souvent écrite sous la forme suivante:

$$X_t^* = E(X_t / I_t) \quad (4.25)$$

Ou X_t^* représente la valeur anticipée de la variable X en début de la période (t). X_t égale la valeur effective et aléatoire de X pendant la période (t). I_t représente l'ensemble des informations pertinentes et disponibles en début de la période (t). or cette définition implique que X_t et X_t^* ne diffèrent que d'un terme d'erreur de prévision aléatoire u_t . C'est-à-dire que:

$$X_t - X_t^* = X_t - E(X_t / I_t) = u_t \quad (4.26)$$

Cette erreur de prévision possède plusieurs propriétés importantes. La propriété centrale est que $E(u_t / I_t) = 0$, non corrélée avec l'anticipation $cov(u_t / I_t) = 0$, de plus, elle est non corrélée avec toute variable appartenant à l'ensemble informationnel.

²¹⁰ Muth J.F., op.cit

²¹¹Muth est l'auteur qui a développé en premier la notion des anticipations rationnelles telle qu'on l'utilise aujourd'hui. Ainsi on peut considérer qu'il représente la source de la nouvelle macroéconomie classique. Toutefois, l'on attribue à Lucas 1976 et Sargent et Wallas 1973 la paternité des premières applications du concept des anticipations rationnelles aux modèles macroéconomiques.

L'hypothèse d'anticipations rationnelles permet de prendre en considération la moindre information contenue dans l'erreur dans l'erreur d'anticipations u_t , si u_t suit un processus autorégressif d'ordre 1 : $u_t = \lambda u_{t-1} + e_t$. Ainsi les agents rationnels seront en mesure de calculer la valeur à l'instant (t) de u_t suivant ce processus la différence entre la valeur prévue u_t^* et la valeur courante u_t sera égale à e_t soit une erreur d'anticipation purement aléatoire. Du point de vue pratique, il n'existe pas de description générale du processus de pondération rationnelle. Les modèles varient en fonction des comportements temporels de la variable en question. Compte tenu des difficultés et des coûts importants de fonctionnement de processus (coûts de collecte de données dans l'anticipation de la valeur future de la variable), les anticipations rationnelles sont plutôt utilisées pour les recherches concernant l'inflation que pour celles concernant la fonction de demande de monnaie.

4.3.2 Le choix de la variable d'échelle

Le choix de la variable d'échelle entre le revenu courant et la richesse est de nature à la fois théorique et empirique. Dans le cadre théorique, deux courants de pensée s'affrontent; la théorie des transactions et la théorie de portefeuille.

La théorie de transactions s'articule sur le fait que la monnaie est avant tout un moyen d'échange. C'est la vraie nature de la monnaie qui permet de distinguer la monnaie de tous les autres biens et actifs financiers. Sous cet angle, le revenu courant est une variable plus appropriée que la notion de la richesse pour expliquer la fonction de demande de monnaie.

La théorie de portefeuille suppose que la monnaie est un actif parmi tant d'autres qui rentrent dans la composition du portefeuille des agents. En tant que substitut des autres actifs, la monnaie doit être expliquée par une contrainte plus large que celle des transactions. Donc, la richesse mesurée par le revenu anticipé (permanent) doit rentrer dans la fonction de demande de monnaie.

Les résultats des tests économétriques sur cet aspect sont mitigés. Certains favorisent l'approche des transactions telles que Mayer et Neri (1975). D'autres privilégient l'approche du portefeuille comme Brunner et Meltzer (1963), Chow (1966) et Laidler (1971-1985).

Une étude comparée réalisée par Gupta (1980) en utilisant des données trimestrielles évalue les performances relatives du revenu courant et du revenu permanent. Les tests réalisés avec le revenu courant ont donné de meilleurs résultats. Mais ces résultats ne sont pas pertinents dans la mesure où les statistiques de Durbin-Watson (D W) montre qu'il existe une forte corrélation non corrigée.

Par contre, Spenilli (1980) a testé la fonction de demande de monnaie en Italie et parmi les quatre modèles de base, ceux qui incluent le revenu permanent ont les résultats les plus satisfaisants. Pour l'Afrique, des modèles de la fonction de demande de monnaie ont utilisé des données annuelles pour estimer le modèle à correction d'erreur dynamique (MCE) Domowitz et El Badaoui (1987) et Simmons 1992. Ces modèles ont imposé a priori une élasticité d'unité à long terme sur le revenu et une élasticité égale à zéro sur les autres variables. La variable de richesse apparaît donc mieux refléter le comportement de la fonction de demande de monnaie. Toutefois, cette conclusion est fondée sur des données annuelles, alors que l'analyse de la fonction de demande de monnaie avec des données à l'intervalle courte n'apportera peut être pas de résultats aussi satisfaisants. Face à ces tests non convainquant et très mitigés, il n'est pas évident de tirer une conclusion définitive pour tel ou tel choix de variable d'échelle.

4.3.3 Les variables des coûts d'opportunité

Il convient d'introduire dans la fonction de demande de monnaie une variable représentant le coût d'opportunité des encaisses monétaires. Cela pour mieux tester les hypothèses alternatives. La variable qui représente le coût d'opportunité mesure l'écart de rendement entre la monnaie et les autres actifs qui la substituent. Le problème est donc, de choisir les mesures de rendement d'actifs alternatifs à la monnaie.

4.3.3.1 Le rendement des actifs alternatifs

Il existe un nombre considérable d'actifs sur les marchés financiers. En général on les classe en deux catégories en fonction du critère de maturité; les actifs de long terme et les actifs de court terme. Du fait que les rendements ou les taux d'intérêts de ces deux catégories d'actifs sont étroitement liés, il est probable que la fonction de demande de monnaie s'explique aussi bien par un seul taux d'intérêt qu'avec deux ou plus. Cependant, le choix entre le taux d'intérêt court et le taux d'intérêt

long est fondé sur des divergences de nature analytique et aussi sur des critères théoriques.

En effet, l'approche du portefeuille démontre que le taux d'intérêt long est une mesure plus convenable, puisqu'il représente mieux le rendement moyen des actifs financiers à tout moment.

En revanche, pour les partisans de l'approche par motif de transaction les rendements des produits financiers à court terme sont les meilleures substitutions de la monnaie grâce à leurs maturités courtes. Dans la pratique, les tests et les résultats empiriques n'ont pu fournir une réponse décisive²¹².

Certains économistes ont essayé de concilier cette différence à l'aide de la théorie de la structure à terme des taux qui décrit l'interrelation de rendement des actifs ayant des maturités différentes²¹³. Sur la base d'hypothèse des marchés efficients²¹⁴, la théorie de la structure des taux démontre que, après l'ajustement du risque, le rendement anticipé des actifs financiers ayant des maturités différentes a tendance à s'égaliser pour une même période de détention.

En conséquence, le choix entre le taux court et le taux long relève essentiellement du choix d'une période de détention des encaisses monétaires. Si la période de détention est courte le bon de trésor à court terme reflète mieux le rendement d'actifs alternatifs de la monnaie que les obligations (ou d'autres actifs à long terme) à 15, 20 ans ou plus de maturités. Du fait que le choix de la période de détention est effectué par les individus, les auteurs tels que Laidler 1971, Karvi 1972 et Friedman 1977 ont préféré de mesurer le rendement alternatif par l'ensemble de la structure des taux.

L'idée de l'ensemble de la structure des taux de à été développé par autres économistes tels que Heler et Khan 1979 Friedman et Schwartz 1982 dont les

²¹² Plusieurs tests et travaux empiriques, tels que ceux réalisés par Bougton 1979 et Laidler 1980 ont montré que les performances de la fonction de demande de monnaie sont peu sensible à l'utilisation du taux court ou du taux long.

²¹³ Il existe une relation entre les taux d'intérêt d'échéances différentes : en avenir certain, le taux à long terme à n années est la moyenne géométrique des taux à court terme pour les périodes futures. On observe en général une relation positive entre le taux d'intérêt de tout actif financier et sa durée, d'où une courbe des taux ascendante. Mais la courbe des taux peut aussi s'inverser, notamment en période de récession

²¹⁴ Un marché est efficient lorsque le prix des titres financiers y reflète à tout moment toute l'information pertinente disponible. Dans un tel marché, il est impossible de prévoir les rentabilités futures, et un titre financier est à tout moment à son prix ". Un tel marché est également appelé marché à l'équilibre ou marché parfait.

résultats étaient très convenables. Bien qu'ils n'expliquent pas mieux la fonction de demande de monnaie qu'un taux d'intérêt simple. Il est à noter que les paramètres estimés pouvaient être sensible à la maturité utilisée.

Pour l'approche de la théorie du portefeuille, les actifs financiers sont des substitues de la monnaie, mais aussi à la fois des actifs physiques. Le rendement réel des actifs physiques peut être mesuré par exemple, par la relation du revenu généré par l'actif sur le prix de l'actif. Dans le cas inflationniste il faut majorer le rendement des actifs réels par le taux d'inflation anticipé.

Comme dans le cas du revenu anticipé, les recherches sur la formation des anticipations ont suivi deux directions. Dans l'hypothèse des anticipations adaptatives, l'inflation anticipée a été mesurée comme un moyenne pondérée des taux d'inflation présents et passés. L'erreur systématique dont souffre cette approche est éliminée par l'approche des anticipations rationnelles qui tient compte de plus d'information que la première direction. Dans une économie de plein emploi, le taux d'inflation anticipé peut être mesuré par le taux de variation de l'offre de monnaie à la condition que la demande de monnaie soit stable. Toutefois, dans une économie où la stabilité de la fonction de demande de monnaie n'est pas assurée, cette formule des anticipations n'est pas adaptée.

Dans une économie où le système financier est bien développé et où les marchés monétaire et financier fonctionnent bien, alors on peut mesurer le taux d'inflation anticipé par la variation du taux d'intérêt nominal, puisque le taux d'intérêt réel est à peu près indépendant de l'inflation. Quant les agents anticipent de l'inflation, le taux d'intérêt nominal est alors ajusté vers le haut pour compenser l'érosion de la valeur réelle anticipée et pour stabiliser la demande intérieure.

Dans l'hypothèse où le taux d'inflation anticipé est parfaitement incorporé dans le taux d'intérêt nominal, l'inclusion de ce taux dans la fonction de demande de monnaie est équivalente à l'inclusion d'une mesure d'inflation anticipée. Ceci permet de détourner les difficultés de mesure des anticipations.

Dans une économie ouverte, il existe naturellement d'autres outils potentiels de mesurer l'inflation anticipée, comme par exemple le taux de change ou l'écart entre le taux d'intérêt domestique et le taux d'intérêt étranger. Il se trouve que dans un environnement international, la monnaie et les actifs financiers échangés

deviennent des substituts des encaisses de monnaie domestique. Si leur degré de substitution est assez élevé pour formaliser le taux d'inflation anticipé avec ces variables et si les variables ont un pouvoir explicatif important, donc la nécessité de les intégrer dans la fonction de demande de monnaie s'impose. Sur la base des travaux réalisés par Frenkel²¹⁵ portant sur l'hyperinflation, conclut que le taux de change est un rendement alternatif de la monnaie et qu'il est un bon indicateur à terme de taux change au comptant et futur. Le taux d'inflation anticipé peut être directement mesuré par la prime de change à terme. Cette approche n'a pas été très suivie à cause de la disponibilité des données d'une part et des résultats non satisfaisants obtenus par des études ultérieures d'autre part. par exemple Salin²¹⁶ a effectué les mêmes hypothèse que Frenkel, mais qui les a cependant amené à des résultats différents. Ces résultats concluent que le marché des changes est inefficace et qu'on ne peut pas seulement utiliser la prime de change à terme comme mesure de l'inflation.

Briten 1981 a étudié le cas de l'Allemagne et du Royaume Uni, en suivant une approche de portefeuille. Il a montré que l'écart entre le taux d'intérêt domestique et étranger a un effet significatif et négatif sur demande de monnaie domestique.

En revanche, les résultats de certaines recherches semblent dire que le taux de change n'est pas significatif dans la fonction de demande de monnaie (Bourrgton1979) et par contre, le taux d'inflation à un effet prononcé après 1970 dans cette étude.

4.3.3.2 Le rendement de la monnaie

Généralement, dans la plus part des études on considère que le rendement de la monnaie comme nul, ou au moins constant. Certains auteurs ont réalisé des études empiriques en vue de mesurer cette variable, ou bien la négliger dans les recherches sur les déterminants de la fonction de demande de monnaie. Selon ces études, le rendement de la monnaie a une nature à la fois explicite et implicite.

Explicitement, certains catégories de dépôts touchent un taux d'intérêt, par exemple au Canada, le taux d'intérêt attaché au compte courant existe, aux Etats-Unis d'Amérique, l'intérêt sur les dépôts à vue fut prohibé en 1933. Cette

²¹⁵ Frenkel J.,1977

²¹⁶Salin M.K., 1980

prohibition a perdu son sens par l'introduction des Now Account. Les dépôts à terme quant à eux sont toujours rémunérés bien que leurs taux ont été plafonné par des réglementations dans certains pays jusqu'au la déréglementation du système bancaire fut apparu dans les années 80 et 90. Les variations explicites de ces taux d'intérêts ont une influence considérable sur la fonction de demande de monnaie.

Implicitement, les établissements financiers, sous la pression de la concurrence viennent par fois à accorder des intérêts non réglementés à leur clientèle, tels que la baisse des charges des services des prêts à taux préférentiels, cadeaux et certains prévalent. Donc, la mesure de ces intérêts implicites est moins évidente. Selon des études, cette mesure est réalisée selon deux méthodes:

La première, suppose que les banques influencent sur le taux d'intérêt payé à leurs clients par le changement des frais de services imputés aux dépôts à vue. Dans cette optique, on peut prendre la variation du ratio du total des charges sur le total des dépôts à vue comme le rendement de la monnaie. Ce ratio doit normalement, varier positivement avec la demande des encaisses monétaires.

La deuxième méthode, suppose d'introduire la concurrence parfaite du marché monétaire. La pression de la concurrence entre les établissements bancaires force les banques à transmettre le rendement marginal du dépôt à vue à leurs clients. Sur ces idées on calcule le rendement d'agrégat de la monnaie en pondérant les différentes parties de la monnaie (les billets bancaires, les dépôts à vue et les dépôts à terme) par leurs poids dans la masse monétaire totale.

4.3.4 Le processus d'ajustement des encaisses monétaires

Au-delà des divergences entre les différentes théories de la demande de monnaie, la fonction de demande de monnaie est généralement spécifiée comme une fonction à élasticité de la variable d'échelle et du taux de rendement. Soit:

$$(M/P)_t = \alpha_1 X_t^{B1} \cdot R_t^{B2} \cdot u_t \quad (4.27)$$

Ou α_1 et est une constante,

X_t = variable d'échelle exprimée en terme réel,

$B1$ = l'élasticité de la demande des encaisses réelles par rapport au revenu,

$B2$ = l'élasticité de la demande des encaisses réelles par rapport au taux d'intérêt,

u_t = ou signifie l'erreur résiduelle.

L'équation (4.12) peut s'écrire de la manière suivante:

$$M_t^d = \alpha_0 + \alpha_1 Y_t + \alpha_2 r_t + \alpha_3 \Delta P^e + u_t \quad (4.28)$$

Où Y est le revenu en terme réel $y = Y/P$, r est la variable du coût d'opportunité et ΔP^e représente le taux d'inflation anticipé.

Sous la forme logarithmique (ln) nous avons;

$$\text{Log } (M/P)_t = \log \alpha_1 + B_1 \log X_t + B_2 \log R_t + u_t \quad (4.29)$$

Pour les études qui utilisent des séries temporelles à intervalle à un an, cette formule apparaît satisfaisante puisque les données sont lissées dans le temps²¹⁷.

Il n'est pas par contre évident que les encaisses monétaires à tout moments soient toujours à leurs niveaux désiré ou bien souhaité. Plusieurs études ont été consacrées à la recherche d'une forme conventionnelle de définition de la demande de monnaie à court terme. Généralement, les différentes interprétations sont classées en deux groupes;

- 1- la méthode d'ajustement partiel.
- 2- Les différentes versions de la méthode d'ajustement partiel.

4.3.4.1 L'ajustement partiel

La méthode de l'ajustement partiel met l'accent sur le fait que les agents économiques peuvent s'éloigner des encaisses désirées à un moment donné. Il est donc, de procéder à un ajustement partiel entre la demande de monnaie détenue et la de la demande de monnaie désirée en fonction d'une fraction Φ qui définit la différence entre elles.

$$M_t - M_{t-1} = \Phi (M_t^* - M_{t-1}) \quad \text{avec} \quad 0 \leq \Phi \leq 1 \quad (4.30)$$

Sous cette hypothèse, l'ajustement partiel des encaisses monétaires repose sur des suppositions principales telles que:

1-Qu'en s'éloignant du niveau désiré à long terme, les agents économiques rencontrent deux types de coûts;

Le premier, est une perte de l'utilité résultant de la différence entre l'encaisse désiré e et l'encaisse réelle (effective).

Le deuxième, est relatifs au frais de transaction engendrés par l'ajustement de l'encaisse réelle vers l'encaisse d'équilibre.

²¹⁷ Sur ce point voir Friedman M. et Schwartz A.J. 1982 ou les études empiriques sont fondées sur des données cycliques.

Le coût total (C) peut être exprimé comme une fonction quadratique des ces deux écarts. Formellement nous avons:

$$C = \alpha_1 [\log(M/P)_t^* - (M/P)_t]^2 + \alpha_2 [\log(M/P)_t - \log(M/P)_{t-1}]^2 + u_t \quad (4.31)$$

Mais si le coût d'ajustement donné où Φ tend vers 0 cette équation devient $M_t = M_{t-1}$ et l'ajustement partiel ne prend pas place.

2-Supposant que les agents économiques sont rationnels, donc il contrôlent leurs encaisses monétaires de telle façon que le coût total soit à son niveau minimum.

Sous forme mathématique, ce comportement implique que la dérivé partielle du coût C par rapport à $\log(M/P)_t$ nulle:

$$\log(M/P)_t - \log(M/P)_{t-1} = (1 - \Phi) [\log(M/P)_t^* - \log(M/P)_{t-1}]^2 + u_t \quad (4.32)$$

Où $(1 - \Phi) = \alpha_1 / (\alpha_1 + \alpha_2)$.

$\Phi = \alpha_1 / (\alpha_1 + \alpha_2)$ si le coût de déséquilibre est grand que le coût d'ajustement, donc

α_1 donne α_2 et Φ tend vers unité (1), donc si $\Phi=1$ l'équation

$M_t - M_{t-1} = M_t^* - M_{t-1}$ et $M_t = M_t^*$, dans ce cas l'ajustement est instantané (4.33).

On suppose que l'équation (4.13) soit une forme correcte de demande de monnaie individuelle à long terme. En remplaçant $\log(M/P)_t^*$ par sa valeur telle que définie dans l'équation 4.14 nous avons :

$$(M/P)_t = \Phi (M/P)_{t-1} + (1 - \Phi) \alpha_1 + (1 - \Phi) B_1 (Y/P)_t + (1 - \Phi) B_2 R_t + u_t \quad (4.34)$$

En introduisant la variable dépendante retardée à une étape comme variable explicative cette équation nous permet de distinguer l'encaisse monétaire à un moment donné de son niveau d'équilibre.

$$M_t^d = \alpha_0 + \alpha_1 Y_t + \alpha_2 r_t + \alpha_3 \Delta P^e + \alpha_4 M_{t-1} + u_t \quad (4.35)$$

De plus, elle permet de retrouver la fonction de long terme de la demande de monnaie, en divisant l'équation (4.15) par le terme $(1 - \Phi)$. Grâce à ces caractéristiques cette équation a fréquemment été utilisée dans les travaux empiriques. Bien que les résultats obtenus semblent être satisfaisants, l'emploi de cette équation doit être prudent tant pour des raisons économiques que pour des raisons économétriques.

De point de vue théorique, l'hypothèse d'ajustement des coûts ne correspond pas à certains modèles théoriques comme les modèles de transactions. Dans ces derniers la demande de monnaie de long terme intègre à priori le concept de coûts. Alors, il n'est pas cohérent d'introduire d'autres coûts pour des raisons de court terme. Milbourne et al (1983) ont indiqué que, si les coûts sont tous importants, il faut les intégrer dans l'analyse dès le début. Mais ceci déduira la simplicité de l'hypothèse d'ajustement partiel.

Sur le plan économétrique, le problème a son origine dans l'autocorrélation liée aux résidus. Dans une régression telle que (4.15) la méthode de l'estimateur des moindres carrés ordinaires (MCO) produira des estimations sans biais pour les paramètres d'ajustement et de demande de monnaie à long terme mais si et seulement si l'espérance mathématique du terme d'erreur résiduel u_t est égale à 0 ($E(u_t) = 0$).

En revanche, si le résidu contient des erreurs systématiques le coefficient estimé de la variable dépendante récurrente sera biaisé vers le haut ou vers le bas en fonction de l'autocorrélation positive ou négative.

4.3.4.2 La justification théorique

L'ajustement partiel suppose que la variable de contrôle converge avec un certain retard vers sa valeur d'équilibre. Si on prend $(1 - \Phi)$ comme proportion d'ajustement la formule générale partielle peut être écrite comme suit:

$$M_t - M_{t-1} = (1 - \Phi) (M_t^* - M_{t-1}). \quad (4.36)$$

L'application de cette formule produit sous certaines hypothèses de la variable M_t , des modèles différents. Nous examinons deux versions du modèle d'ajustement des coûts : l'ajustement réel ou nominal et l'ajustement des prix.

a- l'ajustement réel ou l'ajustement nominal:

L'ajustement est fondé sur la contribution de G. CHOW 1966 pour une comparaison nous rappelons les deux équations de cette approche:

$$\log(M/P)_t - \log(M/P)_{t-1} = (1 - \Phi) [\log(M/P)_t^* - \log(M/P)_{t-1}] + u_t. \quad (4.37)$$

$$\log(M/P)_t = (1 - \Phi) f(Z) + \Phi \log(M/P)_{t-1} + u_t. \quad (4.38)$$

$$\text{Avec } f(z) = \alpha_1 + B_1 \log X_t + B_2 \log R_t. \quad (4.49)$$

Cette représentation implique que les individus ajustent progressivement leurs encaisses réelles aux changements de la variable d'échelle et du taux d'intérêt.

Toute fois ils ajustent instantanément leurs encaisses nominales au variable de prix afin de conserver une encaisse constante. Ce cas de figure présente un cas parmi tant d'autres dans le raisonnement possible des agents (voir aussi W.L.Coal1982).quant il y a illusion monétaire, à savoir qu'il existe un décalage dans le temps entre la variation des prix et la perception de cette dernière, les ajustements des encaisses monétaires par rapport au variation des prix ne sont plus instantanés. Dans ce cas l'ajustement réel n'est plus soutenable. En tenant compte de ce point, des économistes, notamment Goldfeld(1976)²¹⁸ ont choisi l'approche de l'ajustement nominal. En exprimant la fonction de demande de monnaie à long terme en termes réels, ils maintiennent la fonction du coût d'ajustement en termes nominaux.

L'approche d'ajustement nominal met l'accent sur le fait qu'un agent ne peut contrôler que ses encaisses nominales et par conséquent, il ajuste partiellement sa détention de monnaie nominale vers le point d'équilibre. Formellement on a

$$M_t^* = f(Z) + P_t + u_t . \quad (4.40)$$

$$\text{Log} (M_t - M_{t-1}) = (1 - \Phi) \log (M_t^* - M_{t-1}) + u_t . \quad (4.41)$$

$$\log (M/P)_t = (1 - \Phi) f(Z) + \Phi \log (M/P)_{t-1} + u_t \quad (4.42)$$

Cependant, l'ajustement nominal soulève des questions d'ordre à la fois empirique et analytique. Empiriquement, l'expression de l'ajustement nominal s'est révélée très proche de celle de l'ajustement réel. La seule différence entre l'équation 9 et 10 et la substitution des prix retardés par le prix courant dans le terme de variable dépendante retardée d'une étape. La formulation 11 donne des résultats aussi de ceux de l'équation 9, bien que l'hypothèse d'ajustement réel soit vraie et inversement. Sous l'angle analytique, trois points sont à retenir.

1-l'ajustement nominal, raisonnable au niveaux de l'analyse individuelle devient inadapté au niveaux agrégé. Dans le dernier cas, la quantité de monnaie nominale est déterminée du coté de l'offre. Ainsi, dans l'économie à l'échelle globale, la quantité de monnaie nominale ne peut être simultanément une variable exogène qui influence déterminants de la fonction de demande de monnaie et une variable endogène par rapport à ces mêmes déterminants.

²¹⁸ Goldfeld S., "The Case of Missing money", *Brooking Papers on Economic activity*, 1976.
- Goldfeld S. et Sichel D.E., "The Demand For Money", in B. Friedman et F. Hanhn, 1990.

2- il est vrai que l'offre de monnaie n'est pas absolument exogène par rapport à la demande de monnaie. En fonction de la politique monétaire adaptée. Il est tout à fait possible que l'offre de monnaie s'ajuste lentement aux variations de la demande de monnaie. En effet, une offre de monnaie endogène n'est pas une raison suffisante pour introduire un terme récurrent de l'offre dans la fonction de demande de monnaie. De ce fait, comme l'a indiqué D. Laidler (1985), il ressort que le problème d'identification n'est pas correctement résolu par l'adoption d'une mesure de l'offre de monnaie en termes réels.

3- la fonction d'ajustement nominal n'est pas nécessairement applicable en termes économiques. Si on prend par exemple, le niveau général des prix comme variable endogène, cela implique que les prix s'ajustent instantanément aux variations du revenu et du taux d'intérêt. D'autre part, il existe des surajustements des prix par rapport à l'offre de monnaie. Pour atténuer l'effet de sur ajustements, il faut que les variables réelles réagissent en même temps. Cela viole la règle d'indépendance des déterminants dans une équation de régression. Cependant, cela nous révèle que l'approche de l'ajustement nominal doit être dans le meilleur cas limitée au niveau individuel comme nous l'avons déjà montré. Ce cas a été supposée par Patinkin (1956) dans l'hypothèse d'absence d'illusion monétaire, quant il y'avait absence de proportionnalité de demande d'encaisses nominales par rapport au revenu.

4.3.4.3 L'ajustement du coût et l'ajustement de prix

Si par hypothèse, on considère que la demande de monnaie en terme réel est bien identifiée, le modèle de l'ajustement de prix paraît alors au niveau de l'économie globale cohérent et adéquat. Dans une structure macroéconomique à prix flexible, l'équilibre entre l'offre de monnaie et sa demande est assuré par la variation des prix. Dans cette optique, Walker(1965) et autres économistes ont privilégié l'approche d'ajustement par les prix. Lorsque les s'ajustent partiellement vers le niveau d'équilibre, la demande de monnaie à court terme peut être testé sous la formule suivante:

$$\text{Log} (P_t - P_{t-1}) = (1 - \Phi) \log (P_t^* - P_{t-1}) + u_t. \quad (4.43)$$

$$\log (M/P)_t = (1 - \Phi) \log f (Z) + \Phi \log (M/P)_{t-1} + u_t . \quad (4.44)$$

bien que cette spécification ressemble à celle de l'ajustement réel, ce sont deux modèles différents. Le modèle d'ajustement réel est fondé sur la notion d'ajustement

de coût des encaisses monétaires réelles, tandis que, le modèle d'ajustement par les prix repose sur l'existence d'une rigidité des prix à la suite d'une augmentation de la masse monétaire. A partir des tests empiriques, les deux formulations montrent des divergences, excepté le cas où l'on peut s'assurer que la valeur constante de la quantité de monnaie nominale dans l'équation d'ajustement réel et de la capacité du coefficient $(1 - \Phi)$ à capturer ou à absorber tout le mécanisme de transmission c'est-à-dire son intégrité.

La question de savoir si le processus d'ajustement correspond plutôt aux variations de l'offre de monnaie qu'au souci de coût d'ajustement, révèle au phénomène de la stabilité ou l'instabilité de la fonction de demande de monnaie qu'il faut évaluer. Ce phénomène prend aussi de l'importance pour la conduite de la politique monétaire.

Dans la critique de différentes approches analytiques utilisés pour application de la politiques économique et surtout la politique monétaire, la question était toujours si la fonction de demande de monnaie est stable est prédictible principalement dans les pays en voie de développement. Des résultats des études empiriques ont montré que après les années 70, la spécification traditionnelle de la fonction de demande de monnaie dans des pays industrialisés et en voie de développement était temporairement instable (Goldfeld 1976)²¹⁹. Les paramètres estimés étaient souvent non plausibles ; et montre des autocorrections bien élevés des erreurs résiduelles. Les paramètres estimés de la fonction de demande de monnaie à long terme ne correspondront plus à leurs vraies valeurs. Dans ce contexte, d'aucuns affirment que ces résultats mitigées seraient dus à des problèmes de spécification (Baba, Hendry et Starr (1992))²²⁰, en particulier à l'utilisation d'un modèle d'ajustement partiel et l'omission de certaines variables. Ainsi, depuis plus d'une décennie, la recherche s'est orientée dans deux principales directions : en élargissant, d'une part, la gamme des variables de coût d'opportunité utilisée dans une fonction de demande de monnaie; d'autre part, en substituant au modèle d'ajustement partiel, une représentation à correction d'erreur. Cette dernière dont la construction découle de la vérification de l'hypothèse de cointégration, présente l'avantage de regrouper dans un même modèle les effets de court terme et de long terme ainsi qu'un

²¹⁹ Goldfeld S. (1976) op.cit.

²²⁰ Baba, Hendry et Starr (1992), op.cit.

mécanisme d'ajustement à la relation d'équilibre de long terme. Ainsi, l'on est passé d'un équilibre statique à un équilibre dynamique. Parallèlement à l'essor de cette nouvelle spécification issue de la théorie de la Cointégration.

4.3.5 Economie des séries temporelles

Les séries temporelles, appelées aussi séries chronologiques (ou même chroniques), occupent une place importante dans tous les domaines de l'observation ou de la collection de données macroéconomiques et surtout financières. Une série temporelle est une suite d'observations indexées par les entiers relatifs tels que le temps. Pour chaque instant du temps, la valeur de la quantité étudiée Y_t est appelé variable aléatoire. L'ensemble des valeurs Y_t quand t varie est appelé processus aléatoire:

$\{ Y_t, t \in \mathbb{Z} \}$. Une série temporelle est ainsi la réalisation d'un processus aléatoire. La date à laquelle l'observation est faite est une information importante sur le phénomène observé²²¹.

Au cours de la seconde moitié des années soixante-dix, les grands modèles Macro économétriques, et surtout leur utilisation pour évaluer les effets de la politique économique, ont subi de nombreuses critiques qui ont conduit à leur abandon progressif, du moins au niveau académique.

La première critique majeure est connue sous le nom de *critique de Lucas*. Selon Lucas (1976)²²², les paramètres des modèles économétriques peuvent évoluer sous l'influence de la politique économique lorsque les agents intègrent les changements de règles de politique économique dans leur comportement. Dans la mesure où les modèles économétriques existants à cette époque ne tenaient pas compte de manière adéquate des anticipations des agents, les modèles dont les paramètres avaient été évalués sur la base des données passées – c'est-à-dire pour des politiques économiques passées – ne permettaient pas une évaluation correcte des effets des nouvelles politiques économiques à venir. Pour Lucas et Sargent (1979)²²³, l'analyse des politiques économiques ne peut s'effectuer dès lors qu'en abandonnant le cadre

²²¹ Lardic S. et Mignon V., " *Econométrie Des Séries Temporelles Macroéconomiques et Financières*", ECONOMICA, Paris, 2002.

²²² Lucas R., 1976, op.cit.

²²³ Lucas R., Sargent T., " After Keynesians Macroeconomics ", Federal Reserve Bank of Minneapolis, *Quarterly Review*, , 1979, 3(2), p. 1-16.

théorique de référence au profit d'une modélisation structurelle cohérente et de la dérivation explicite des règles de décisions compatibles avec un ensemble de restrictions associées aux conditions d'équilibre et aux schémas d'anticipation. Si la *critique de Lucas* s'avère théoriquement fondée, il n'en reste pas moins que sa portée empirique peut être discutée et évaluée. La question est alors de savoir – sans remettre en cause la logique de la critique – si les erreurs de prévisions issues d'un modèle économétrique soumis à la critique sont ou non importantes (Fair 1984). De nombreuses études empiriques suggèrent que la *critique de Lucas* a une importance faible en pratique (Favero et Hendry (1992)²²⁴ et Ericsson et Irons (1995)²²⁵. Par exemple, Favero et Hendry (1992) ont montré sur données britanniques, que suite à un changement significatif dans l'offre de monnaie, les équations décrivant la demande de monnaie ne présentaient pas de ruptures significatives. Dans les modèles agrégés avec anticipations rationnelles, les formes réduites mettent en évidence des restrictions inter équations sur les paramètres. Puisque ces restrictions inter équations n'apparaissaient pas satisfaites, Hendry et Favero concluaient que l'hypothèse d'anticipations rationnelles n'était pas vérifiée.

Comme par ailleurs, les modèles dynamiques avec anticipations rationnelles, immuns à la critique, imposent des restrictions en si grand nombre qu'ils sont difficiles à spécifier correctement, Fair (1993)²²⁶ en vient à se demander s'il n'est pas préférable de conserver les modèles macro économétriques traditionnels, quitte à vérifier si les paramètres des formes réduites restent invariants à la politique économique.

Une seconde critique majeure des modèles macro économétriques s'appuie sur l'approche développée par Sims (1980)²²⁷. Sims remet en cause ces modèles dans la mesure où ils imposent une structure d'identification sur la base d'*a priori* théoriques. D'une part, cette structure peut ne pas être en cohérence avec les propriétés statistiques des observations disponibles. D'autre part, il n'est pas

²²⁴Favero C., Hendry D.F., " Testing the Lucas Critique ", *Econometric Review*, 1992,11(3), p. 265-306.

²²⁵Ericsson T., Irons j., " The Lucas Critique in Practice: Theory without Measurement " in *Macroeconometric: Developments, Tensions and Prospects*, K. Hoover ed., 1995, chap. 8, Kluwer Academic Publisher.

²²⁶Fair R. "Testing Macroeconometric Models ", *American Economic Review Papers and Proceedings*, 1993,83(2), p. 287-293.

²²⁷Sims C., " Macroeconomics and Reality ", *Econometrica*, 1980, 48(1), p. 1-48.

toujours possible de justifier simplement ces choix d'identification, en particulier dans le cas de modèles dynamiques. Sims propose de donner à l'économetre le rôle d'étudier sans *a priori* la structure des données économiques pour analyser les conditions dans lesquelles elles ne soient pas en contradiction avec la théorie économique retenue par le modélisateur. A l'aide d'outils statistiques et de procédures d'inférence, l'économetre essaie de déterminer la structure dynamique des variables économiques qu'il considère.

Puis, en imposant différentes restrictions identifiantes sur la base d'un *a priori* théorique acceptable, il spécifie une forme structurelle à partir de laquelle différents exercices de prévisions et de réponses à des chocs peuvent être menés. Il n'en reste pas moins que cette approche ne prémunit pas les modèles obtenus de la *critique de Lucas*, dans le sens où ils ne permettent pas d'étudier les effets liés à des changements de politique économique (si ces derniers n'ont pas été explicitement modélisés). Face à l'approche développée par Sims, les partisans de la calibration (Kydland et Prescott 1996)²²⁸ opposent la prédominance de la théorie sur l'économétrie. Dans cette approche, le modèle théorique – non soumis à la *critique de Lucas* – permet à lui seul de déterminer la structure de l'économie et sa parcimonie rend plus explicite les principaux mécanismes économiques.

Quelle que soit la portée empirique de la *critique de Lucas*, il faut noter que celle-ci a donné lieu à un développement important de nouvelles méthodes d'estimation et d'inférence. Par exemple, la méthode des moments généralisés développée par Hansen (1982)²²⁹ permet dans un cadre structurel d'estimer les paramètres résumant les préférences et la technologie, sous l'hypothèse de leur stabilité, de modèles dynamiques avec anticipations rationnelles. De même, les différentes stratégies de tests de stabilité temporelle ou encore les tests d'exogénéité (Engle, Hendry et Richard 1983)²³⁰ sont venus enrichir la panoplie d'outils à la disposition de l'économétrie de la politique économique.

²²⁸Kydland F., Prescott E., " The Computational Experiment: An Econometric Tool ", Journal of Economic Perspectives, 1996, 10(1), p. 69-85.

²²⁹Hansen L., " Large Sample Properties of Generalized Method of Moments Estimators", *Econometrica*, 1982, 50(5), p. 1269-1286.

²³⁰Engle R., Hendry D.F., Richard J.F., " Exogeneity ", *Econometrica*, 1983, 52(2), p. 277-304.

Enfin, le développement considérable de la micro économétrie structurelle ou descriptive, notamment grâce à la constitution de banques de données individuelles de taille très importante, a permis d'aborder la question de l'évaluation des changements de politique en analysant quantitativement les modifications qu'ils induisent dans les comportements individuels hétérogènes, qu'il s'agisse des ménages ou des entreprises. L'optique privilégiée dans ce cadre peut être néanmoins sensiblement différente de ce qui a été brièvement évoqué jusqu'ici. En particulier la question de la sensibilité des schémas d'anticipation des agents au choix des objectifs et des priorités de la politique économique n'est pas abordée. Deux approches peuvent être exemplifiées.

Dans un premier type d'approche, dans la mesure où il n'est en général pas possible de réaliser une expérience contrôlée dans le domaine de l'économie, l'économetre a développé des techniques qui visent à comparer au mieux les comportements des agents qui ont été touchés par la modification de la politique aux comportements de ceux qui ne l'ont pas été en contrôlant autant que faire se peut les effets de l'environnement. Suivant la nature de la réforme économique, les résultats empiriques obtenus sont plus ou moins probants et peuvent apporter une évaluation *ex post* instructive.

Dans un second type d'approche, l'analyse des effets d'une modification d'une politique est faite sur un état stationnaire d'un modèle dynamique qui traite explicitement la présence d'une hétérogénéité à l'équilibre. La mesure des effets des changements de politique s'appuie alors sur l'estimation des effets de la modification des comportements à l'équilibre déduits de l'estimation des paramètres structurels du modèle. La transition entre états stationnaires n'est cependant pas étudiée.

Étant donné les controverses qui agitent l'économétrie de la politique économique, mais aussi le foisonnement des méthodes économétriques, il n'est pas forcément facile de repérer les différents enjeux méthodologiques, théoriques et appliqués. Hendry étudie la modélisation des anticipations lorsque les erreurs de prévision résultent de modifications importantes et non anticipées des facteurs déterministes, tels que le terme constant ou encore la tendance linéaire. Les travaux étudient différentes implications de ce type d'erreur de prévision, notamment au regard des

anticipations rationnelles. En premier lieu, la prévision basée sur un ensemble de variables causales n'est pas forcément la meilleure. En effet, dans le cas où les facteurs déterministes changent de façon importante, tout exercice de prévision même s'il est mené à partir d'un jeu de variables causales nécessairement incomplet n'est pas réellement utile tant que les changements futurs ne peuvent pas être clairement anticipés pour tout l'horizon de la prévision. En second lieu, les meilleures méthodes de prévision sont celles qui sont relativement immunisées vis-à-vis d'une mauvaise modélisation des changements déterministes. Ces méthodes de prévision – qualifiées de robustes – diffèrent des anticipations rationnelles. Si les agents adoptent alors de telles règles robustes de prévision, elles auront la propriété de ne pas se modifier face à une famille de nouvelles politiques économiques implicitement envisagées, rendant la *critique de Lucas* ineffective et non pertinente empiriquement. Krozlig et Toro reprennent cette idée de modifications des termes déterministes. Leur article définit dans un premier temps la notion de ruptures déterministes communes puis la relie au concept de super exogénéité. Les ruptures déterministes communes apparaissent lorsque les amplitudes des ruptures frappant les termes déterministes de plusieurs variables sont linéairement reliées si bien que certaines combinaisons linéaires des variables étudiées n'en sont pas affectées. La super exogénéité établit un ensemble de conditions telles que les paramètres d'un modèle conditionnel soient invariants à des changements dans les paramètres d'un modèle marginal. Si les changements dans les paramètres du modèle marginal se réduisent à ceux déterministes – de termes déterministes communs, les concepts de super exogénéité et l'occurrence de ruptures déterministes communes sont étroitement reliés.

Différentes simulations mettent en évidence de très bonnes propriétés en terme de puissance pour des échantillons de petite taille. L'étude reprend ensuite le concept de super exogénéité en plaçant ces modifications déterministes communes dans un modèle conditionnel. Différentes simulations en petits échantillons viennent confirmer la puissance du test. Carrasco et Gregoir introduisent un modèle avec coefficients aléatoires afin de capturer la relation possible entre les paramètres résumant la politique économique et ceux issus du comportement des agents privés. Le cadre statistique retenu est celui des modèles doublement stochastiques. Ces

modèles étendent les modèles linéaires avec coefficients aléatoires car ils incorporent des variables endogènes retardées. Ce cadre statistique s'avère plus adéquat car il permet de traiter la dépendance temporelle. De plus, les paramètres – aléatoires – définissant les comportements des agents privés peuvent dépendre de ceux déterminant la politique économique. Le cadre statistique permet alors de tester la *critique de Lucas*. Lorsque l'invariance des paramètres des comportements privés est rejetée, le modèle permet alors d'illustrer les conséquences possibles d'une modification de la politique économique sur le comportement des variables agrégées. Une application sur données françaises met en évidence une dépendance significative des paramètres de la fonction de consommation aux changements observés de la politique économique au début des années quatre-vingt-dix, supportant la portée empirique de la *critique de Lucas*.

L'analyse quantitative de la politique économique est un champ d'investigation théorique et empirique particulièrement actif de l'analyse économique. Ainsi, les modèles structurels envisagés sont aujourd'hui suffisamment sophistiqués, tant au niveau macro que micro, pour qu'il soit possible de ne plus les considérer comme des curiosités théoriques mais plutôt comme des cadres utiles et pertinents d'analyse de la politique économique. De même, les méthodes statistiques de l'économétrie des séries temporelles attirent notre attention sur les possibilités de développer des modèles statistiques parcimonieux permettant d'évaluer convenablement les effets de la politique économique. Il convient cependant de souligner à ce stade que la macro et la micro économétrie structurelle n'abordent pas ces problématiques sous des angles toujours comparables. L'approche macro économétrique insiste surtout sur la sensibilité des règles de prévisions des agents représentatifs aux modifications des priorités des politiques et sur le problème de l'instabilité des modèles. L'approche micro économétrique tente plutôt soit d'apprendre à mieux décrire le comportement des agents en exploitant les modifications des comportements induits par les changements de politique, soit de représenter parcimonieusement les sources structurelles de l'hétérogénéité des agents pour analyser les conséquences des modifications de politique, en comparant les états stationnaires de l'économie. Un enrichissement mutuel de ces approches doit être possible.

4.3.6 Propriétés de base des séries temporelles

4.3.6.1 Stationnarité

La stationnarité est la clef de l'analyse des séries temporelles²³¹. La série $\{Y_t\}$ est dite strictement stationnaire si la distribution conjointe de $(Y_{t_1}, \dots, Y_{t_k})$ est identique à celle de $(Y_{t_1+t}, \dots, Y_{t_k+t})$, quel que soit t , ou k est un entier positif arbitraire et $(t_1, \dots, t_k)$ une liste de k entiers positifs arbitraires. Autrement dit, la stationnarité stricte dit que la distribution conjointe de $(Y_{t_1}, \dots, Y_{t_k})$ est invariante quand on fait glisser le temps. Cette condition est difficile à vérifier et on utilise en général, une version plus faible de stationnarité. On dit qu'une série temporelle $\{Y_t\}$ est faiblement stationnaire si la moyenne de Y_t et la covariance entre Y_t et Y_{t-l} sont invariantes par translation du temps. Précisément, $\{Y_t\}$ est faiblement stationnaire si : (a) $E(Y_t) = \mu$ ou μ est une constante indépendante de t , (b) $\text{cov}(Y_t, Y_{t-l})$ ne dépend que de l , entier. La stationnarité faible (ou du second ordre) implique que le Figure de la série en fonction du temps montre des fluctuations autour d'un niveau moyen, fluctuations qui se ressemblent, quelque soit la date autour de laquelle on examine la série.

4.3.6.2 Autocovariance

La covariance $\gamma_l = \text{cov}(Y_t, Y_{t-l})$ est appelée autocovariance d'ordre l (ou de décalage) (lag- l autocovariance). Pour chaque décalage l , il y a une autocovariance. La fonction : $l \rightarrow \gamma_l$ est la fonction d'autocovariance de (Y_t) .

Cette fonction a trois propriétés importantes :

$$(a) \gamma_0 = \text{var}(Y_t), \quad (4.45)$$

$$(b) \gamma_l = \gamma_{-l}, \text{ car :} \quad (4.46)$$

$$\gamma_{-l} = \text{cov}(Y_t, Y_{t-(-l)}) = \text{cov}(Y_{t-(-l)}, Y_t) = \text{cov}(Y_{t+l}, Y_t) = \text{cov}(Y_{t+l}, Y_{(t+l)-l}) = \gamma_l \quad (4.47)$$

²³¹ Consulter: R. Bourbonnais, "Econométrie", 6^e édition, DUNOD, Paris, 2005.

- Araujo C., Brun J.F. et Combes J.L., "Econométrie", BREAL, 2004.

- Pirotte A., "L'économétrie, Des Origines Aux Développements Récents", CNRS Economie, Paris, 2004.

- Lardic S. et Mignon V., "Econométrie Des Séries Temporelles Macroéconomiques Et Financières", ECONOMICA, Paris, 2002.

Autre notation. On écrit aussi $\gamma_Y(l)$, en particulier pour distinguer la fonction d'autocovariance d'une série (Y_t) , de celle d'une autre série.

4.3.6.3 Corrélation et fonction d'autocorrélation

Le coefficient de corrélation entre deux variables X et Y de moyennes μ_X et μ_Y est défini par :

$$\rho_{X,Y} = \frac{\text{cov}(X,Y)}{\sqrt{\text{var}(X)\text{var}(Y)}} = \frac{E[(X - \mu_X)(Y - \mu_Y)]}{\sqrt{E(X - \mu_X)^2 E(Y - \mu_Y)^2}} \quad (4.47)$$

Ce coefficient est compris entre -1 et 1. Il mesure la force de la dépendance linéaire entre X et Y . Si on dispose d'un échantillon (x_t, y_t) , $t = 1, \dots, T$ d'observations indépendantes de (X, Y) , on peut estimer de façon convergente le coefficient de corrélation par le coefficient de corrélation empirique :

$$\hat{\rho}_{X,Y} = \frac{\sum_{t=1}^T (x_t - \bar{x})(y_t - \bar{y})}{\sqrt{\sum_{t=1}^T (x_t - \bar{x})^2 \sum_{t=1}^T (y_t - \bar{y})^2}} \quad (4.48)$$

Ou, $\bar{x} = \sum_{t=1}^T x_t / T$ et, $\bar{y} = \sum_{t=1}^T y_t / T$ sont les moyens empiriques de X et Y .

Considérons maintenant une série temporelle x_t , $t = 1, \dots, T$ de valeurs numériques, sans nous interroger sur son modèle mathématique, et formons la série retardée : $y_t = x_{t-1}$, $t = 2, \dots, T$.

On peut calculer le coefficient de corrélation entre les deux séries :

$$r = \frac{\sum_{t=2}^T (x_t - \bar{x})(y_t - \bar{y})}{\sqrt{\sum_{t=1}^{T-1} (x_t - \bar{x})^2 \sum_{t=2}^T (y_t - \bar{y})^2}} \quad (4.49)$$

Si la série observée x_t , $t = 1, \dots, T$ est la réalisation d'une série (ou processus) stationnaire, ce coefficient mesure la liaison entre la valeur de la série en une date et en la date voisine.

Observons que $\bar{x} = \sum_{t=1}^{T-1} x_t / T - 1$ et $\bar{y} = \sum_{t=2}^T y_t / T - 1$ ne diffèrent que par les valeurs

x_1 et x_T . Le paraFigure suivant formalise cette observation : on y définit d'abord

la covariance décalée de 1, de 2, . . . d'une série stationnaire, puis le coefficient de corrélation décalé de 1, de 2, . . . On introduit ensuite la version empirique de ce coefficient de corrélation. On pourra observer qu'elle est plus simple que le r donné par la formule ci-dessus.

4.3.6.4 Fonction d'autocorrélation (ACF)

Considérons une série (faiblement) stationnaire (Y_t). On est souvent intéressé par décrire la dépendance de (Y_t) par rapport à son passé, notamment pour expliquer le niveau actuel de la série par le niveau à une date précédente. On sait que si une dépendance est linéaire, elle est bien décrite par le coefficient d'autocorrélation. Par définition, le coefficient d'autocorrélation d'ordre l est:

$$\rho_l = \frac{\text{cov}(Y_t, Y_{t-l})}{\sqrt{\text{var}(Y_t) \text{var}(Y_{t-l})}} \quad (4.50)$$

Mais, $\text{var}(Y_{t-l}) = \text{var}(Y_t) = \gamma_0$ donc,

$$\rho_l = \frac{\text{cov}(Y_t, Y_{t-l})}{\text{var} Y_t} = \frac{\gamma_l}{\gamma_0} \quad (4.51)$$

Enfin en terme d'espérance mathématique et notant que par la stationnarité : $E(Y_t) = \mu$ indépendant de t , on a :

$$\rho_l = \frac{E[(Y_t - \mu)(Y_{t-l} - \mu)]}{E[(Y_t - \mu)^2]} \quad (4.52)$$

ρ_l est une mesure de la dépendance de la valeur Y en une date par rapport à sa valeur à une date décalée de l'intervalles de temps.

La fonction : $l \rightarrow \rho_l, l = 0, 1, 2, \dots$ est appelée fonction d'autocorrélation (théorique), FAC (ou ACF en anglais) de la série $\{Y_t\}$. De la définition on voit que : $\rho_0 = 1, -1 \leq \rho_l \leq 1$

Etant donné un échantillon $y_t, t = 1, \dots, T$, de $\{Y_t\}$ stationnaire, notons la moyenne empirique, $\bar{y} = \sum_t^T y_t / T$. Le coefficient d'autocorrélation empirique d'ordre l est:

$$\hat{\rho}_1 = \frac{\sum_{t=2}^T (y_t - \bar{y})(y_{t-1} - \bar{y})}{\sum_{t=1}^T (y_t - \bar{y})^2} \quad (4.53)$$

Le coefficient d'autocorrélation empirique d'ordre $l \geq 1$ est:

$$\hat{\rho}_l = \frac{\sum_{t=l+1}^T (y_t - \bar{y})(y_{t-l} - \bar{y})}{\sum_{t=1}^T (y_t - \bar{y})^2}, 0 \leq l \leq T-1 \quad (4.54)$$

Sous des conditions générales, voir par exemple Brockwell et Davis (1997)²³², $\hat{\rho}_l$ est un estimateur convergent de ρ_l .

La fonction : $l \rightarrow \hat{\rho}_l, l = 0, 1, 2, \dots$ est appelée fonction d'autocorrélation empirique de la série (Y_t) .

Si (Y_t) est une suite de v.a. i.i.d., de moment d'ordre 2 fini, $E(Y_t^2) < \infty$, alors les coefficients d'autocorrélation $\hat{\rho}_l$ sont approximativement indépendants et normalement distribués, de moyenne 0, de variance $1/T$.

4.3.6.5 Autocorrélation partielle

L'autocorrélation (ou le coefficient d'autocorrélation) partielle d'ordre K de la série Y_t est égale au coefficient de corrélation entre :

$$Y_t - E(Y_t / Y_{t-1}, Y_{t-2}, \dots, Y_{t-k+1}) \text{ et } Y_{t-k} - E(Y_{t-k} / Y_{t-1}, Y_{t-2}, \dots, Y_{t-k+1})$$

Ce coefficient, noté $r(k)$ mesure la corrélation entre Y_t et Y_{t-k} , lorsqu'on a éliminé les parties de Y_t et Y_{t-k} , expliquées par les variables intermédiaires.

a) Modèles de séries stationnaires

Maintenant nous examinons les exemples classiques de séries stationnaires et calculons leurs fonctions d'autocorrélation. Un bruit blanc est une série stationnaire, si $\{Z_t\} \sim \text{BB}(0, \sigma_z^2)$, sa fonction d'autocovariance est :

$$\gamma_z(k) = \{\sigma_z^2, k = 0, 0, k = 0\} \quad (4.55)$$

Etant donné une série empirique, sa modélisation revient souvent à trouver, c'est-à-dire identifier et estimer le mécanisme qui fait passer d'un BB à la série.

1-Série linéaire

Une série Y_t est dite linéaire si elle peut s'écrire :

²³²Brockwell P.J., Davis R.A., " *Introduction to Time Series and Forecasting* ", Springer, 1997.

$$Y_t = \mu + \sum_{i=-\alpha}^{\alpha} \psi_i Z_{t-i} \quad (4.56)$$

ou Z_t est un $BB(0, \sigma_z^2)$, $\psi_0 = 1$ et la suite $\{\psi_i\}$ est absolument sommable, c'est-à-dire $\sum_i |\psi_i| < \infty$. On admettra qu'une série linéaire est stationnaire. Une série est dite linéaire et causale si elle est linéaire et $i \geq 0$, autrement dit elle ne dépend que du BB passé.

Si Y_t est linéaire et causal on obtient :

$$EY_t = \mu, \text{var}(Y_t) = \sigma_z^2 \sum_{i=0}^{\alpha} \psi_i^2. \quad (4.57)$$

L'autocovariance d'ordre k est :

$$\gamma_k = \text{cov}(Y_t, Y_{t-k}) = E \left[\sum_{i=0}^{\alpha} \psi_i Z_{t-i}, \sum_{j=0}^{\alpha} \psi_j Z_{t-k-j} \right] \quad (4.58)$$

$$= E \left(\sum_{i,j=0}^{\alpha} \psi_i \psi_j Z_{t-i} Z_{t-k-j} \right) \quad (4.59)$$

$$= \sum_{j=0}^{\alpha} \psi_{j+k} \psi_j E(Z_{t-k-j}^2) = \sigma_z^2 \sum_{j=0}^{\alpha} \psi_j \psi_{j+k} \quad (4.60)$$

Si la série est linéaire et causale et si de plus $\psi_i = 0$ pour $i > q$ on dit que Y_t est une moyenne mobile d'ordre q (MA(q)). Une série linéaire causale est un MA(α).

4.3.7 Processus autorégressif d'ordre p

4.3.7.1 Processus autorégressif d'ordre 1

On dit que $\{Y_t\}$ est un processus autorégressif d'ordre 1 s'il obéit à une équation:

$$Y_t = \phi_0 + \phi_1 Y_{t-1} + Z_t, \quad t \in \mathbb{Z} \quad (4.61)$$

Moments d'ordres 1 et 2 d'un AR(1)

Supposons $\{Y_t\}$ est stationnaire alors, sa moyenne μ , est constante et prenant l'espérance mathématique des deux cotés de l'équation Y_t on obtient:

$$\mu = \phi_0 + \phi_1 \mu \quad (4.62)$$

et si $\phi_1 \neq 1$

$$E(Y_t) = \mu = \frac{\phi_0}{1 - \phi_1} \quad (4.63)$$

On pose $\dot{Y}_t = Y_t - \mu$. C'est le processus centré. Avec l'opérateur de retard, on a:

$$(1 - \phi_1 B) \dot{Y}_t = Z_t \quad (4.64)$$

Par substitutions successives on obtient que $\dot{Y}_t$ peut être exprimé comme une moyenne mobile infinie :

$$\dot{Y}_t = Z_t + \phi_1 Z_{t-1} + \phi_1^2 Z_{t-2} + \dots \quad (4.65)$$

pourvu que $-1 < \phi_1 < 1$. Cette condition est suffisante pour que le processus soit stationnaire.

On appelle $\dot{Y}_t$ la représentation MA(1) de Y_t . L'écriture de Y_t comme une somme de variables non corrélées permet de calculer facilement les variances et autocovariances comme suite:

Elevons au carré les deux cotés de $\dot{Y}_t$, il vient :

$$\text{Var}(Y_t) = \sigma_Z^2 (1 + \phi^2 + \phi^4 + \dots) \quad (4.66)$$

$$= \frac{\sigma_Z^2}{1 - \phi^2} \quad (4.67)$$

Enfin écrivons $\dot{Y}_t$ en $t - k$ et calculons les espérances des deux cotés de :

$$Y_t Y_{t-k} = (Z_t + \phi Z_{t-1} + \phi^2 Z_{t-2} + \dots)(Z_{t-k} + \phi Z_{t-k-1} + \phi^2 Z_{t-k-2} + \dots) \quad (4.68)$$

ou, Z_t étant un BB, $E(Z_t Z_{t-m}) = 0$, $m \neq 0$. On obtient pour $k > 0$,

$$\gamma_k (\phi^k + \phi^{k+2} + \phi^{k+4} + \dots) \sigma_Z^2 = \phi^k \gamma_0 \quad (4.69)$$

La fonction d'autocorrélation de l'AR(1) est:

$$\rho_k = \phi^k, k = 0, 1, 2, \dots \quad (4.70)$$

Observons enfin que $\dot{Y}_t$ est l'écriture d'un AR(1) comme une moyenne mobile infinie.

Modèle AR(2)

Soit Y_t stationnaire, obéissant à l'équation:

$$Y_t = \phi_0 + \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + Z_t, \quad t \in \mathbb{Z} \quad (4.71)$$

Prenant l'espérance des deux cotés on obtient :

$$E(Y_t) = \mu = \frac{\phi_0}{1 - \phi_1 - \phi_2} \quad (4.72)$$

pourvu que $1 \neq \phi_1 + \phi_2$. Comme $\phi_0 = \mu(1 - \phi_1 - \phi_2)$, on a:

$$Y_t - \mu = \phi_1 (Y_{t-1} - \mu) + \phi_2 (Y_{t-2} - \mu) + Z_t \quad (4.73)$$

et on va étudier la fonction d'autocovariance sur le processus centré, $\dot{Y}_t = Y_t - \mu$ qui vérifie :

$$\dot{Y}_t = \phi_1 \dot{Y}_{t-1} + \phi_2 \dot{Y}_{t-2} + Z_t \quad (4.74)$$

Multiplions les deux cotés de cette équation par $\dot{Y}_{t-l}, l \geq 0$

$$\dot{Y}_t \dot{Y}_{t-l} = \phi_1 \dot{Y}_{t-1} \dot{Y}_{t-l} + \phi_2 \dot{Y}_{t-2} \dot{Y}_{t-l} + Z_t \dot{Y}_{t-l} \quad (4.75)$$

et prenons les espérances mathématiques. Nous obtenons:

$$\gamma_l = \phi_1 \gamma_{l-1} + \phi_2 \gamma_{l-2}, l \geq 0 \quad (4.76)$$

En effet, par substitution successive de Y_{t-1} en fonction de $Y_{t-2}, Z_{t-2}, \dots$ on voit que $\text{cov}(\dot{Y}_{t-l}, Z_t) = 0, l \geq 0$. On appelle cette dernière équation, l'équation de moments d'un AR(2). La fonction d'autocorrélation d'un AR(2) est :

$$\rho_1 = \frac{\phi_1}{1 - \phi_2} \quad (4.77)$$

$$\rho_l = \phi_1 \rho_{l-1} + \phi_2 \rho_{l-2}, l \geq 1 \quad (4.78)$$

Nous avons supposé Y_t , stationnaire. Nous examinons maintenant les conditions sur les ϕ_i qui assurent cette stationnarité. L'équation aux différences correspondant à $\dot{Y}_t$

est:

$$1 - \phi_1 B - \phi_2 B^2 = 0 \quad (4.79)$$

C'est le polynôme caractéristique de l'équation de récurrence qui décrit l'AR(2).

Cette équation du second degré a deux racines réelles ou complexes : $\frac{1}{w_1}$ et $\frac{1}{w_2}$:

$$1 - \phi_1 B - \phi_2 B^2 = (1 - w_1 B)(1 - w_2 B) \quad (4.80)$$

Pour aller plus loin, examinons ce qu'on a fait pour le processus AR(1). Le processus AR(1) centré obéit à:

$$(1 - \phi_1 B) \dot{Y}_t = Z_t \quad (4.81)$$

La substitution a donné $\dot{Y}_t = Z_t + \phi_1 Z_{t-1} + \phi_1^2 Z_{t-2} + \dots$ ou $\dot{Y}_t = (1 - \phi_1 B)^{-1} Z_t$. Elle revient à développer en série la fraction rationnelle $(1 - \phi_1 B)^{-1}$, opération possible

car $|\phi_1| < 1$. Pour l'AR(2), on veut développer en série : $(1 - \phi_1 B - \phi_2 B^2)^{-1}$. On peut décomposer cette opération:

$$(1 - w_1 B)(1 - w_2 B)\dot{Y}_t = Z_t \quad (4.82)$$

Donne

$$(1 - w_2 B)\dot{Y}_t = (1 - w_1 B)^{-1} Z_t \quad (4.83)$$

Puis

$$\dot{Y}_t = (1 - w_2 B)^{-1}(1 - w_1 B)^{-1} Z_t. \quad (4.84)$$

Ces opérations sont possibles si $|w_1| < 1$ et $|w_2| < 1$ c'est-à-dire si les racines du polynôme caractéristique sont en module > 1 .

Donc un processus qui vérifie Y_t est stationnaire si et seulement si les racines du polynôme caractéristique $1 - \phi_1 B - \phi_2 B^2$ sont > 1 en module.

Une extension du modèle AR(1) est le modèle AR(p). Soit un (Z_t) un BB. Un processus (Y_t) est dit autorégressif d'ordre p s'il s'écrit:

$$Y_t = \phi_0 + \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \dots + \phi_p Y_{t-p} + Z_t \quad (4.85)$$

Avec l'opérateur retard on peut écrire cette autorégression à l'ordre p comme :

$$(1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p)Y_t = \phi_0 + Z_t \quad (4.86)$$

$$\phi(B)Y_t = \phi_0 Z_t \quad (4.87)$$

Nous inspirant de ce qu'on a obtenu pour un AR(2), nous admettrons qu'un processus autorégressif d'ordre p est stationnaire si les racines de l'équation : $1 - \phi_1 z - \phi_2 z^2 - \dots - \phi_p z^p = 0$ sont en module > 1 . C'est la condition de stationnarité d'un processus autorégressif d'ordre p.

4.3.8 Processus Moyenne mobile

4.3.8.1 Processus MA(1)

On dit que (Y_t) est un processus moyenne mobile d'ordre 1 (MA(1)), s'il obéit à une équation :

$$Y_t = \mu + Z_t - \theta Z_{t-1} \quad (4.88)$$

ou $Z_t \sim \text{BB}(0, \sigma_Z^2)$. Cette équation s'écrit encore :

$$Y_t - \mu = (1 - \theta B)Z_t \quad (4.89)$$

Moments d'ordres 1 et 2 d'un MA(1) En prenant l'espérance mathématique des deux cotest de (4.89) , on voit que

$$E(Y_t) = \mu \quad (4.90)$$

La variance de Y_t est la variance d'une combinaison affine de variables non corrélées donc $\text{var}(Y_t) = (1 + \theta^2)\sigma_z^2$. De même, $\text{cov}(Y_t, Y_{t-1}) = \text{cov}(\mu + Z_t - \theta Z_{t-1}, \mu + Z_{t-1} - \theta Z_{t-2}) = -\theta\sigma_z^2$. On voit que $\text{cov}(Y_t, Y_{t-k}) = 0, k > 1$.

En résumé, $\forall \theta$ le processus MA(1) défini par(4.89) est stationnaire, de moyenne μ , de fonction d'autocorrélation :

$$\begin{aligned} \rho_k &= 1 \quad \text{si } k=0 \\ \rho_k &= \frac{-\theta}{1+\theta^2} \quad \text{si } k=1 \\ \rho_k &= 0 \quad \text{si } k > 1 \end{aligned} \quad (4.91)$$

On aimerait pouvoir exprimer le processus MA(1) en fonction de son passé (observé) et pas seulement en fonction d'un bruit non observé. Introduisons le processus centré, $\dot{Y}_t = Y_t - \mu$ correspondant à (4.87). On voit que si $|\theta| < 1$, on peut développer $(1 - \theta B)^{-1}$ en série entière. Ceci nous amène à une définition.

On dit qu'un processus est inversible si on peut l'écrire comme une autorégression infinie. Ainsi, un MA(1) est inversible si la racine de l'équation $1 - \theta z = 0$ est > 1 en module. On observe que la condition d'inversibilité d'un MA (1) est techniquement parallèle à la condition de stationnarité d'un autorégressif d'ordre 1.

4.3.8.2 Processus MA(q)

Un processus (Y_t) est dit processus moyenne mobile d'ordre q (MA(q)) si :

$$Y_t = \mu + Z_t - \theta_1 Z_{t-1} - \theta_2 Z_{t-2} - \dots - \theta_q Z_{t-q} \quad (4.92)$$

ou $Z_t \sim \text{BB}(0, \sigma_z^2)$. On peut noter de façon équivalente :

$$Y_t = \mu + (1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q) Z_t \quad (4.93)$$

$$= \mu + \Theta(B) Z_t \quad (4.94)$$

Un MA(q) est un cas de série linéaire dont les propriétés sont:

1. Un MA(q) est un processus stationnaire.

2. La fonction d'autocorrélation d'un processus MA(q) est nulle à partir de l'ordre q + 1.

Cette dernière propriété est utile pour deviner (identifier) l'ordre de moyenne mobile convenable pour modéliser une série. En effet, en présence d'un corrélogramme empirique non significativement différent de 0 à partir d'un certain ordre (k), on pensera à modéliser la série correspondante par un MA(k - 1). Et un MA(q) est inversible si les racines de $1 - \theta_1 z - \theta_2 z^2 - \dots - \theta_q z^q = 0$ sont, en module, > 1 .

4.3.8.3 Processus ARMA (p,q)

(Y_t) est un processus ARMA(p,q) s'il est stationnaire avec une partie MA et une partie AR :

$$Y_t - \phi Y_{t-1} - \phi Y_{t-2} - \dots - \phi Y_{t-p} = \theta_0 - Z_t + \theta_1 Z_{t-1} - \theta_2 Z_{t-2} - \dots - \theta_q Z_{t-q} \quad (4.95)$$

Ou $Z_t \sim BB(0, \sigma_Z^2)$. On voit que:

$$\mu = E(Y_t) = \frac{\theta_0}{1 - \phi_1 - \dots - \phi_p} \quad (4.96)$$

Un ARMA(p,q) peut se noter :

$$Y_t = \mu + \frac{1 - \theta_1 B - \dots - \theta_q B^q}{1 - \phi B - \dots - \phi_p B^p} Z_t \quad (4.97)$$

Dans ces expressions, il faut bien voir que μ est la moyenne et que l'autre terme est une erreur de moyenne nulle, autocorrélée. On pourrait envisager une moyenne fonction du temps avec toujours un modèle ARMA de moyenne nulle pour l'erreur. Plus, on appelle équations de Yule-Walker, les équations que vérifient les autocovariances ou les autocorrélations d'un processus AR(p), formule (4.74) ou un ARMA(p,q) (formule (4.93))

4.3.8.4 Saisonnalité: Saisonnalité multiplicative

Décrivons brièvement la modélisation de la saisonnalité dans l'approche de Box-Jenkins. Soit une série mensuelle observée (pour simplifier) sur un nombre entier d'années, à partir d'un mois de janvier. On note y_{ij} l'observation du mois j de l'année i; $j = 1, \dots, 12$, $i = 1, \dots, n$.

Supposons qu'on modélise la dépendance d'un mois sur un ou deux mois précédents (sans s'occuper de l'effet saisonnier) et qu'on adopte un ARMA (p,q):

$$\Phi(B)Y_t = \Theta(B)b_t \quad (4.98)$$

Il est fort probable, si la série présente une saisonnalité, que le résidu $\hat{b}_t$ ne sera pas blanc mais aura une structure de corrélation saisonnière. On peut envisager deux traitements de cette "non blancheur". Ou bien on ajoute des termes de retard dans les polynômes Φ et Θ , ou bien on modélise b_t par un ARMA dont l'unité de temps est l'année :

$$b_t = \frac{\Theta_s(B^s)}{\Phi_s(B^s)} Z_t \quad (4.99)$$

où s désigne la période. Ce qui donne :

$$\Phi_s(B^s)\Phi(B)Y_t = \Theta(B)\Theta_s(B^s)Z_t \quad (4.100)$$

avec $Z_t \sim \text{BB}$, où $\Phi(B)$, $\Phi_s(B^s)$, $\Theta(B)$, $\Theta_s(B^s)$ sont respectivement des polynômes de degrés p , q en B et P, Q en B^s . On dit que Y_t est un SARMA(p, q)(P, Q)s'il vérifie (4.95) et est stationnaire. Les conditions de

– stationnarité de Y_t sont : les racines des polynômes $\Phi(B)$ et $\Phi_s(B^s)$ sont en module > 1 .

– inversibilité de Y_t sont : les racines des polynômes $\Theta(B)$ et $\Theta_s(B^s)$ sont en module > 1

4.3.8.5 Processus VAR (vecteur auto régressif)

Le VAR (Vector Auto-Regression) est une extension à plusieurs variables de l'approche développée par Box et Jenkins²³³. Chaque variable est expliquée non seulement par son propre passé comme dans le cas d'un AR traditionnel, mais également par le passé des autres variables du système.

L'utilisation d'un modèle VAR est méthodiquement justifiée par le fait que les modèles VAR autorisent des simulations permettant de saisir les modifications des variables objectifs suite à un choc sur les variables instruments.

Les modèles VAR comportent trois avantages : en premier ils permettent d'expliquer une variable par rapport à ses retards et en fonction de l'information

²³³ Box G.E.P. et Jenkins G.M., " *Times Series Analysis, Forecasting And Control*", San Francisco, Holder Day, 1970.

contenue dans d'autres variables pertinentes ce qui soulève des problèmes de cointégration, en second on dispose d'un espace d'information très large et enfin, cette méthode est assez simple à mettre en œuvre et comprend des procédures d'estimation et des tests.

Dans un modèle VAR l'on ne se donne pas de modèle théorique à priori. Tout modèle peut en effet être arbitraire si les variables apparaissent à la fois à la droite et à gauche des équations et si nous n'avons pas de causalité. Dans ce cas l'on peut chercher à régresser chaque variable endogène sur l'ensemble des autres variables, endogènes et exogènes. Le modèle sera déterminé en ne retenant que les coefficients significatifs.

Les modèles VAR doivent être estimés à partir de séries stationnaires. C'est une propriété d'invariance des caractéristiques statistiques des processus pour toutes les translations dans le temps. Il est impossible d'identifier un processus si toutes ses caractéristiques changent au cours du temps. Ces modèles permettent, d'une part, d'analyser les effets d'une variable sur l'autre au travers de simulations de chocs aléatoires.

Un modèle VAR à k variables et à p décalages, noté VAR (p) s'écrit de la manière suivante:

$$\Delta Y_t = \Gamma_1 \Delta Y_{t-1} + \dots + \Gamma_k \Delta Y_{t-k} + \Pi Y_{t-1} + \mu + \Phi D_t + \varepsilon_t \quad (4.101)$$

avec $t = 1, \dots, T$; k est le nombre de retards , Y_t est un vecteur de p variables coïntégrées, et $\Delta Y_{t-1} \dots \Delta Y_{t-k}$ sont des vecteurs de leurs variations, μ est une constante, D_t est un vecteur de variable non stochastiques (coefficients saisonniers, trend temporel, variables auxiliaires) ou de variables stochastiques exclues de l'espace de co-intégration (variables incluses dans la dynamique de court terme, mais pas dans l'espace de co-intégration) « dummy » ou *faiblement exogènes*, les termes d'erreur ε_t sont indépendants et identiquement distribués *niid* (0, Σ), $\Gamma_1 \dots \Gamma_k$, sont des vecteurs de constantes, et Π est un produit matriciel .donc un modèle VAR passe les étapes suivantes: 1.Représentation, 2.Estimation 3.Causalité et 4. Chocs impulse réponse.

Conclusion

L'évaluation quantitative des effets de la politique économique a toujours représenté un enjeu majeur de l'analyse économique en mobilisant tous les champs de la discipline : théorie, mesure, statistique et économétrie. Cet aspect de l'analyse économique apparaît naturellement dès que l'on aborde des travaux concernant des études macroéconomiques. Très rapidement apparaissent dans l'analyse des paramètres simples dont les valeurs dépendent du comportement des agents privés. Estimer la valeur de ces paramètres à partir d'observations semble en général possible, mais il importe de connaître le degré de confiance que l'on peut attribuer à ces évaluations si on souhaite en faire usage dans une comparaison de simulations de différentes politiques économiques. La quantification de ces effets pose de nombreux problèmes liés aux propriétés du modèle retenu : est-il bien spécifié ?

Est-il une approximation acceptable de l'économie ? ses paramètres sont-ils invariants à la politique économique envisagée ?...

L'objectif de l'économétrie de la politique économique est précisément de fournir un cadre statistique et des méthodes adaptées pour traiter ces questions et ainsi clarifier certains enjeux de la modélisation

les rôles de la théorie et de l'économétrie ont peu donné lieu à controverses jusqu'au début des années soixante-dix. À cette époque, l'économétrie

a principalement été au service de la théorie (essentiellement de la théorie macro-économique), en quantifiant différents effets multiplicateurs de court et de long terme. Le partage du travail entre théorie et économétrie s'est maintenu pour deux raisons. D'une part, la théorie gardait une forme d'autonomie vis-à-vis de la mesure et son but principal était de dériver des modèles formels. D'autre part, le modèle de référence – le modèle offre globale/demande globale – utilisé pour l'analyse de la politique économique faisait l'objet d'un large consensus.

Les rôles respectifs de la théorie et de l'économétrie ont considérablement évolué depuis le début des années soixante-dix. Identifier la majorité des variables clés dans les relations économétriques et le travail de l'économetre consistait à fournir des estimations de ces relations postulées *a priori*. Ainsi, la théorie déterminait *a priori* les variables endogènes et exogènes et l'économétrie ne cherchait pas à évaluer la pertinence empirique d'une telle distinction. Ce partage du travail s'est révélé fructueux. L'économétrie s'est surtout développée autour des méthodes d'estimation et d'identification. Bon nombre des estimateurs couramment employés aujourd'hui ont été introduits à cette époque.