

L'économie retrouvée

Renaud Fillieule

L'école autrichienne d'économie

Une autre hétérodoxie

Sciences Sociales

Septentrion
PRESSES UNIVERSITAIRES

L'école autrichienne d'économie

Une autre hétérodoxie

La collection
Économie retrouvée
est dirigée par
Laurent Cordonnier et Franck Van de Velde

Cet ouvrage est publié après
l'expertise éditoriale du comité
Sciences sociales composé de :

Jean-Pierre Bourgois (Lille 2), Vincent Caradec, (Lille 3),
Francis Danvers (Lille 3), Xavier Labbée (Lille 2),
Rémi Lefebvre (Reims/Lille 2), Nicolas Postel (Lille 1),
Loïc Sallé (Lille 2), Helga-Jane Scarwell (Lille 1), Franck
Van de Velde (Lille 1), Pierre-Yves Verkindt (Lille 2),
Bruno Villalba (IEP/Lille 2)

Renaud Fillieule

L'école autrichienne d'économie

Une autre hétérodoxie

Publié avec le soutien
du Centre Lillois d'Études et de Recherches
Sociologiques et Économiques
(CLERSE, Lille 1)

Presses Universitaires du Septentrion
www.septentrion.com

Les Presses Universitaires du Septentrion

sont une association de six universités :

- Université des Sciences et Technologies de Lille, Lille 1,
- Université du Droit et de la Santé, Lille 2,
- Université Charles-de-Gaulle – Lille 3,
- Université du Littoral – Côte d'Opale,
- Université de Valenciennes et du Hainaut-Cambrésis,
- Fédération Universitaire Polytechnique de Lille.

La politique éditoriale est conçue dans les comités éditoriaux.

Six comités et la collection « Les savoirs mieux de Septentrion » couvrent les grands champs disciplinaires suivants :

- Acquisition et Transmission des Savoirs
- Lettres et Arts
- Lettres et Civilisations Étrangères
- Savoirs et Systèmes de Pensée
- Temps, Espace et Société
- Sciences Sociales

Publié avec le soutien

de l'Agence Nationale de la Recherche,

du Conseil Régional Nord-Pas de Calais

© Presses Universitaires du Septentrion, 2010

www.septentrion.com

Villeneuve d'Ascq

France

Toute reproduction ou représentation, intégrale ou partielle, par quelque procédé que ce soit, de la présente publication, faite sans l'autorisation de l'éditeur est illicite (article L 122 4 du Code de la propriété intellectuelle) et constitue une contrefaçon.

L'autorisation d'effectuer des reproductions par reprographie doit être obtenue auprès du Centre Français d'Exploitation du Droit de Copie (CFC) 20 rue des Grands-Augustins à Paris.

ISBN : 978-2-7574-0163-7

ISSN : 1773-8814

Livre imprimé en France

Remerciements

L'auteur tient tout d'abord à remercier Franck Van de Velde qui l'a incité à entreprendre ce projet et lui a offert, avec Laurent Cordonnier, cette occasion de publication aux Presses Universitaires du Septentrion. Il est également redevable aux membres du *Séminaire de recherche en économie autrichienne* de l'Université Paris 2, qui lui ont adressé des questions et des remarques lors d'une présentation du contenu d'une partie de l'ouvrage. Des remerciements particuliers s'adressent à Laurent Carnis, Alain Fillieule, Guido Hülsmann, à nouveau Franck Van de Velde et Laurent Cordonnier, et Pierre-Édouard Visse pour leurs commentaires et suggestions détaillés sur ce texte. Le contenu de ce livre reste bien sûr de la seule responsabilité de son auteur.

Sommaire

Préface des éditeurs	11
Introduction	15
Chapitre 1 : biens et valeur.....	21
Chapitre 2 : échange et prix.....	45
Chapitre 3 : monopole et concurrence.....	69
Chapitre 4 : la production et sa structure.....	89
Chapitre 5 : capital et intérêt	111
Chapitre 6 : la monnaie et son pouvoir d'achat.....	133
Chapitre 7 : inflation et crise	155
Chapitre 8 : État et marché.....	181
Conclusion.....	217
Bibliographie.....	223
Table des matières détaillée	235

Le grand public, les journalistes, et très probablement la plupart des étudiants qui terminent leurs études d'économie et de gestion aujourd'hui, n'ont sans doute jamais entendu parler des économistes Autrichiens. À certains d'entre eux, peut-être, le nom de Friedrich Hayek dira quelque chose, sans qu'ils puissent mieux situer son œuvre qu'en le classant dans le camp des néolibéraux et des fervents opposants à la doctrine keynésienne – ce qui est profondément juste, d'ailleurs. Mais Hayek n'est sans doute pas le plus autrichien des auteurs autrichiens, et ce n'est sans doute pas le plus économiste d'entre eux. On pourrait même dire que c'est l'arbre qui cache la forêt, et qui pourrait dispenser à trop bon compte les curieux d'aller voir ce que ce grand courant de pensée a apporté à l'histoire des idées économiques, depuis presque un siècle et demi, à travers les grands auteurs que sont Carl Menger, Eugen von Böhm-Bawerk, Ludwig von Mises et Murray Rothbard.

Se dispenser d'aller voir... Avec la publication de cet ouvrage de Renaud Fillieule, c'est ce que ne pourront plus se permettre les amateurs d'économie qui se piquent d'avoir l'esprit large, en arguant du fait que la littérature autrichienne est faite de « pavés » rebutants, ne comportant presque jamais d'équations, et qui n'aurait au surplus pas fait avancer beaucoup l'ingénierie des politiques économiques – mais c'est justement ce qu'elle ne veut pas faire. La prouesse de Renaud Fillieule est en effet de mettre à la disposition des curieux et des esprits larges une véritable « somme » sur l'école autrichienne... une somme qui a le bonheur de tenir en 200 pages et qui prend le lecteur par la main du début à la fin, en lui donnant à voir, dans un récit très simple et parfaitement ordonné, la plupart des grandes notions qui constituent le cadre intellectuel de cette école, ainsi que les éléments de doctrine qui fondent ses positions dans la plupart des grandes controverses qui animent la discipline. Et ce n'est pas encore tout : l'auteur prend soin de montrer que l'école autrichienne n'est pas une relique à ranger dans les manuels d'histoire des idées, mais que les grands auteurs sur lesquels s'appuie cette tradition lui valent d'être toujours bien vivante, parce qu'ils fournissent des clés d'interprétation pour traiter des « grandes questions économiques

contemporaines », comme le montrent les recherches récentes qui s'appliquent à fournir une grille d'analyse « autrichienne » de la crise économique et financière que nous traversons.

Les lecteurs familiers de la collection « L'économie retrouvée » s'étonneront peut-être de trouver, à côté des ouvrages d'inspiration keynésienne, marxiste, institutionnaliste, régulationniste qui commencent à forger l'identité de cette collection, un ouvrage sur les économistes autrichiens, qui plus est bien fait, c'est-à-dire qui présente le risque d'emporter certaines convictions. Les Autrichiens ne sont-ils pas en effet les pires pourfendeurs du keynésianisme, et ses pires accusateurs... eux qui ont vu dès le départ dans les idées de Keynes le terrassement sur lequel serait pavée « la route de la servitude », pour reprendre le titre d'un ouvrage de Hayek (1944) ? Ne sont-ils pas de ceux qui n'ont jamais cessé de dénoncer l'interventionnisme des keynésiens, les manipulations budgétaires et monétaires pratiquées par leurs émules, les distorsions de la structure productive provoquées par ces politiques, leurs conséquences inflationnistes et leur rôle éminent dans le retour des bulles spéculatives, des booms artificiels et des récessions qui s'ensuivent nécessairement ? On veut rassurer le lecteur : ce sont bien eux, et l'on retrouvera tout cela clairement explicité ici. Alors quoi ?

Il ne s'agit pas de répondre à cet étonnement en disant simplement, pour tenir une bonne raison, que les économistes autrichiens sont également les ennemis déclarés des économistes néoclassiques – à qui ils reprochent violemment leur constructivisme qui mène tout droit à la tyrannie – pour tenir une bonne raison. Les ennemis de nos ennemis ne sont pas forcément nos alliés, et il n'y a pas de front commun anti-néoclassique entre les hétérodoxies keynésienne, institutionnaliste, marxiste, régulationniste, d'un côté, et l'école autrichienne de l'autre. Mais il y a tout de même des points communs importants dans l'approche de l'économie qui valent la peine d'être notés. À l'opposé du modèle canonique qui sert d'idéal type unificateur de la pensée néoclassique, il n'existe pas, ni chez les Autrichiens ni chez les keynésiens et leurs compagnons de route, de système complet de marchés, sur lesquels interagissent des agents rationnels maximisant leur intérêt sur l'ensemble de leur cycle de vie, aidés dans leurs choix intertemporels par des marchés financiers efficients, et mis à l'abri des

fluctuations de la demande sous le voile monétaire. Les hétérodoxes autrichiens, comme les keynésiens et les institutionnalistes, admettent que les agents sont plongés dans un univers radicalement incertain, fait d'un ensemble de marchés carrément incomplet ; que dans cet univers leur rationalité est limitée et se déploie dans un temps qui a de l'épaisseur ; un univers où la monnaie n'est pas neutre (même à long terme) et dont les effets dépendent de la manière dont elle est mise en circulation. Par-dessus tout, les uns et les autres voient l'économie comme un processus qui comporte des dynamiques d'ajustement, et non comme un état d'équilibre. Ce sont ces points communs qui méritent à nos yeux d'être soulignés, pour faire comprendre du même coup combien les économistes néoclassiques sont les seuls à voir le monde comme ils le voient – même s'ils sont encore aujourd'hui hégémoniques. Et ce sont encore ces points communs qui font ressortir de manière saillante, par un contraste saisissant, le gouffre abyssal qui sépare en retour la pensée autrichienne des autres hétérodoxies. Ce gouffre, c'est bien entendu le rejet de ce qui a constitué le cœur de la révolution keynésienne, à savoir : le principe de la demande effective. Et pourtant la pensée autrichienne débouche sur une authentique construction macroéconomique, mais une macroéconomie sans demande effective ! En plaçant au cœur de leur représentation de l'économie un marché des fonds prêtables, les Autrichiens prolongent la vision classique d'un monde dans lequel l'épargne gouverne l'investissement, et passent à côté – diraient les keynésiens – de ce qui fait la spécificité d'une économie capitaliste, dans laquelle c'est l'inverse qui est vrai. Outre l'impossibilité de concevoir les crises de débouchés, qui rendent compte de la permanence de la pauvreté (du chômage involontaire) dans nos économies d'abondance, le maintien de la fiction d'un marché sur lequel s'ajusteraient les souhaits d'épargne et d'investissement inverse à peu près tout ce qui fait l'essentiel des processus économiques, selon les keynésiens. La macroéconomie autrichienne aboutit donc à des conclusions, la plupart du temps, complètement opposées aux conclusions keynésiennes, des conclusions qui fondent une position très hostile à l'intervention publique ou, plus clairement encore : des positions néolibérales.

En donnant à voir aussi clairement cela, en montrant l'économie autrichienne dans son plus simple appareil, dirions-

nous, Renaud Fillieule rend peut-être autant service aux futurs économistes autrichiens – qui disposeront enfin d’un bréviaire sans équivalent, en français comme en anglais – qu’il rend service aux autres économistes hétérodoxes, lesquels pourront y trouver, pour une fois clairement balisé à leurs yeux, à côté de l’école néoclassique, un autre sentier de la perdition.

FRANCK VAN DE VELDE
LAURENT CORDONNIER

Quels sont les principes fondamentaux et les développements majeurs de l'analyse économique ? Ce livre a pour but de présenter la réponse apportée à cette question par *l'école autrichienne*. Si l'on considère que les piliers de l'économie « autrichienne » ont été posés par son fondateur Carl Menger à la fin du XIX^e siècle, alors ils se composent d'une théorie des biens, d'une théorie de la valeur et d'une théorie de l'échange. Ses successeurs ont à leur tour développé des théories des prix, de la concurrence, de la production et des revenus, de la monnaie, des cycles d'affaires et des interventions de l'État. Ce texte s'adresse aux lecteurs qui souhaitent découvrir, ou se familiariser avec, ces théories. Il peut aussi servir d'introduction aux traités de référence de cette école, beaucoup plus complets et volumineux, comme celui de von Mises (1985 [1949]) et celui de Rothbard (1962). Plus généralement, il intéressera toutes celles et tous ceux qui recherchent une présentation synthétique, rigoureuse et néanmoins non mathématisée, de la science économique.

L'école autrichienne d'économie constitue un paradigme à part entière. Elle diffère en profondeur du paradigme néo-classique standard ou « orthodoxe » qui forme aujourd'hui la base de l'enseignement universitaire de l'économie. Cette différence apparaît dès la théorie de la valeur. Structure de production, imputation, spécificité ou convertibilité des facteurs de production : il n'existe aucune trace de ces notions essentielles, ni des analyses qu'elles permettent de développer, dans les manuels standards. La théorie autrichienne des prix s'appuie sur le principe de l'adaptation aux chocs dynamiques par réallocation des capitaux et disparition des profits et pertes entrepreneuriaux engendrés par ces chocs. Là aussi, l'analyse théorique est conduite de façon très différente de celle des modèles standards d'équilibre général, et plus encore de celle des modèles d'équilibre partiel. Quant aux conceptions orthodoxes de la concurrence – concurrence parfaite et concurrence monopolistique –, elles sont sévèrement critiquées par les « Autrichiens » pour leur incapacité à rendre compte de la nature du processus compétitif de l'économie de marché. Ces derniers leur substituent une conception de la concurrence comme rivalité, qui prend en

compte l'incertitude radicale de l'avenir et la fonction entrepreneuriale de recherche et d'exploitation des occasions de profit.

Dans les domaines qui relèvent de la macroéconomie – structure du capital, taux d'intérêt, monnaie, cycle –, la divergence est tout aussi marquée. Théorie du détour de production, triangles hayékiens, théorème de la régression monétaire, rejet de l'équation des échanges ($MV = PT$), explication de l'intérêt par la théorie de l'actualisation, explication du cycle par la théorie du crédit de circulation : autant d'éléments fondamentaux du paradigme autrichien qui sont totalement absents de la présentation standard. Mais c'est avec le paradigme keynésien que l'opposition est ici la plus frontale, avec une réfutation du « paradoxe de l'épargne », et un rejet des politiques dites d'investissement public et des politiques monétaires inflationnistes.

L'objet de ce livre n'est pas d'analyser en détail les points de discordance entre le paradigme autrichien et les autres – même si certains seront évoqués dans les pages qui suivent. Il s'agit, plus modestement, d'offrir *une présentation brève et synthétique des concepts et théories économiques fondamentaux de l'école autrichienne*. Celle-ci s'est construite à partir d'une réflexion collective qui a rassemblé plusieurs générations successives d'économistes. Cet ouvrage entend restituer de façon concise mais fidèle cette riche histoire conceptuelle et théorique.

Les principaux économistes de l'école autrichienne (jusqu'aux années 1970)

Le paradigme autrichien s'est peu à peu constitué grâce aux apports d'économistes successifs qui ont élaboré les différents éléments d'un système complet d'analyse économique. L'axe principal se compose de quatre auteurs, tous d'origine autrichienne :



S'y ajoutent Friedrich von Wieser, Frank Fetter, Murray Rothbard et Israel Kirzner, les trois derniers étant d'origine américaine, dont les apports sont plus limités mais néanmoins importants. Enfin, deux économistes américains doivent aussi être cités parce que certains de leurs travaux jouent un grand rôle dans le paradigme autrichien : John Bates Clark et Frank Knight. À partir des années 1970, de nombreux autres économistes ont renforcé les rangs de l'école autrichienne, mais seuls les principaux auteurs « classiques », si l'on peut dire, sont évoqués ici. Voici leurs apports théoriques les plus importants.

- Carl Menger (1840-1921) : théorie des biens, théorie de la valeur, théorie de l'échange, théorie du prix d'enchères (*Grundsätze der Volkswirtschaftslehre – Principes d'économie*, 1871)
- Friedrich von Wieser (1851-1926) : théorie de la valeur des facteurs de production convertibles (*Der natürliche Werth – La valeur naturelle*, 1889)
- Eugen von Böhm-Bawerk (1851-1914) : théorie du capital (dé-tour de production), réinterprétation de la « loi des coûts », théorie de la préférence pour le présent, théories de l'intérêt (*Kapital und Kapitalzin – Capital et intérêt*, 1884-1912)
- [● John Bates Clark (1847-1938)] : théorie des changements dynamiques, notion de distribution fonctionnelle (*The Distribution of Wealth – La distribution des richesses*, 1899)] (auteur influent mais qui ne fait pas partie de l'école autrichienne)
- Frank Fetter (1863-1949) : théorie de l'actualisation (*Economic Principles – Principes d'économie*, 1915)
- Ludwig von Mises (1881-1973) : théorie de la monnaie, théorie du crédit, théorie du cycle, théorie du calcul économique, critique du collectivisme (argument du calcul économique), théorie de l'interventionnisme étatique (*Theorie des Geldes und der Umlaufsmittel – Théorie de la monnaie et du crédit*, 1912 ; *Die Gemeinwirtschaft – Le socialisme*, 1922 ; *Human Action – L'action humaine*, 1949)

[● Frank Knight (1885-1972) : théories de l'incertitude et du profit entrepreneurial (*Risk, Uncertainty, and Profit – Risque, incertitude, et profit*, 1921)] (auteur influent mais qui ne fait pas partie de l'école autrichienne)

● Friedrich von Hayek (1899-1992) : macroéconomie de la structure de production (triangles hayékiens), critique du collectivisme (argument de la complexité), théorie de la concurrence dynamique (*Prices and Production – Prix et production*, 1931 ; *Individualism and Economic Order – Individualisme et ordre économique*, 1948)

● Murray Rothbard (1926-1995) : critique de la théorie du prix de monopole, théorie des impôts et dépenses de l'État (*Man, Economy, and State – L'Homme, l'économie, et l'État*, 1962 ; *Power and Market – Pouvoir et marché*, 1977)

● Israel Kirzner (1930-) : théorie de la concurrence entrepreneuriale (*Competition and Entrepreneurship – Concurrence et entrepreneurial*, 1973)

Il peut sembler étonnant que le nom de Joseph Schumpeter (1883-1950) ne figure pas dans cette liste. Bien que fortement marqué par l'école autrichienne dont il est issu, ce dernier développe des théories qui s'éloignent de l'édifice progressivement bâti par Menger, Böhm-Bawerk, von Mises et Hayek (Schumpeter 1999 [1911], 1939, 1951 [1942]). Cet édifice est d'ailleurs loin d'être parfaitement cohérent. Des fissures existent, notamment en ce qui concerne les théories du capital, de l'intérêt et du collectivisme, et elles seront indiquées en leur temps.

Les relations entre les différents auteurs ayant contribué à constituer le cœur de l'école autrichienne, telle qu'elle est définie dans ce livre, sont précisées dans le schéma de la figure I.1 (les auteurs influents mais ne faisant pas partie de cette école sont représentés entre crochets). Le personnage central de ce schéma est von Mises : son traité sur *L'action humaine* réalise en effet la synthèse théorique et épistémologique de l'école autrichienne, un point de confluence avec certains des apports les plus significatifs de l'école néo-classique américaine – mais aussi avec un rejet explicite et argumenté de l'économie mathématique de l'école de

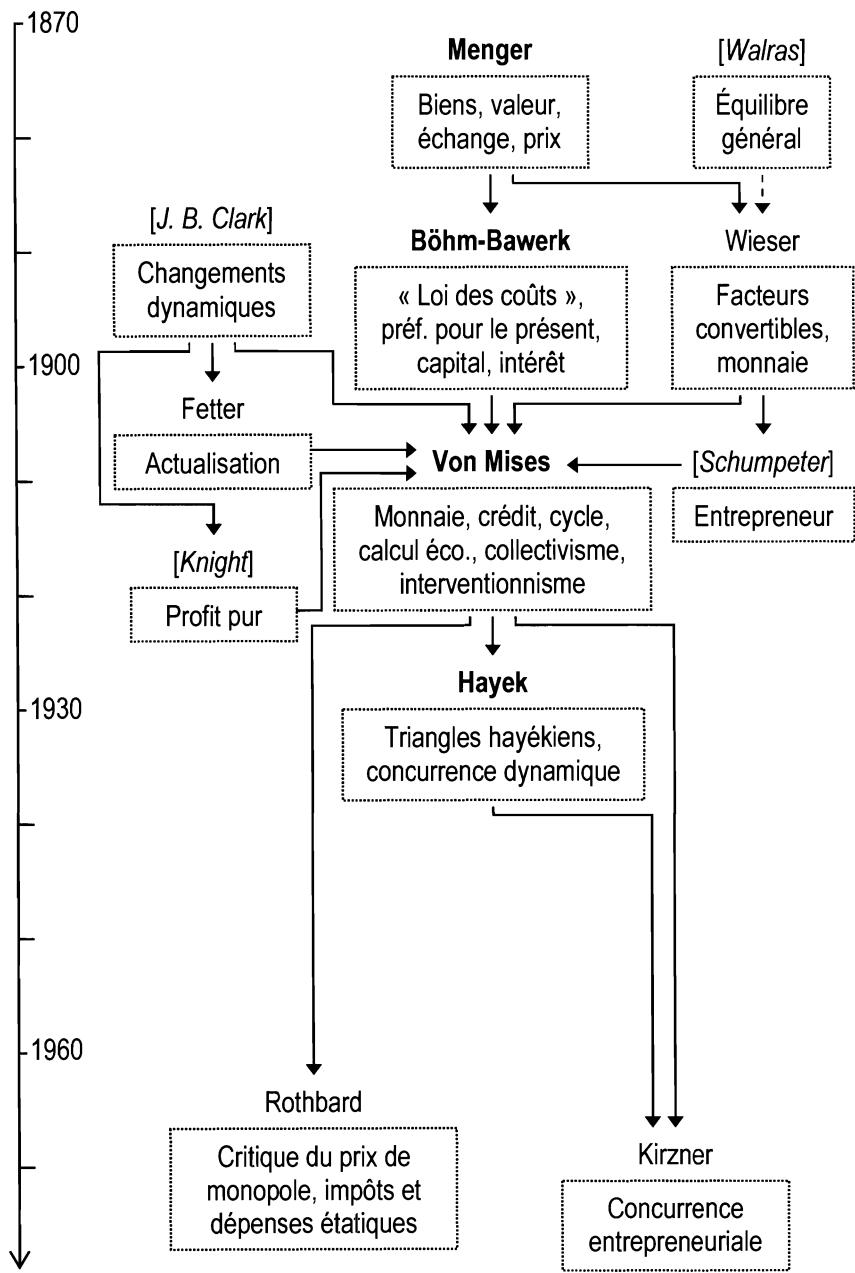


Figure I.1. Présentation succincte de l'école autrichienne

Lausanne. De nombreux auteurs poursuivent aujourd'hui l'étude et l'approfondissement de la tradition autrichienne. Ils ne sont pas placés dans ce schéma qui se limite, comme il a été dit, aux auteurs classiques.

Il reste à préciser que seules les questions de *théorie* économique seront abordées dans les chapitres qui suivent. Des travaux majeurs ont été consacrés à la méthodologie et l'épistémologie de la science sociale par Menger (1985 [1883]), von Mises (1981 [1933], 1985 [1949], 2007 [1957], 1962), Hayek (1952, 1964, 1974), et d'autres. Il n'y sera fait presque aucune allusion. Il en va de même des nombreux textes et ouvrages de philosophie politique consacrés à la défense et l'illustration du libéralisme : ils se situent hors du champ de ce livre (voir par exemple von Mises 2005 [1927], Hayek 1960, Rothbard 1991 [1982]).

Un très vaste corpus de littérature autrichienne se trouve en libre accès sur le site internet du *Mises Institute* (<http://mises.org>). Les lecteurs souhaitant approfondir leurs connaissances peuvent aussi consulter des revues scientifiques (en anglais) comme la *Review of Austrian Economics*, le *Quarterly Journal of Austrian Economics* et *New Perspectives on Political Economy*, ces deux dernières étant accessibles gratuitement en ligne. Les analyses des économistes de l'école autrichienne contemporaine sur les questions actuelles de politique et de conjoncture économiques peuvent être suivies au jour le jour sur trois blogs (en anglais), celui du *Mises Institute* (<http://blog.mises.org/>), le blog *Coordination Problem* (<http://austrianeconomists.typepad.com/>), et enfin *ThinkMarkets* (<http://thinkmarkets.wordpress.com/>).

Renaud Fillieule
La Madeleine, juin 2010

Les bases de l'analyse économique autrichienne ont été posées par son fondateur Carl Menger (1976 [1871]) en réponse à deux premières séries de questions. (1) Sous quelles conditions une chose est-elle utile à un individu, c'est-à-dire sous quelles conditions est-elle un *bien* pour cet individu ? (2) Comment expliquer qu'un bien ait de la *valeur* pour un individu, et comment expliquer à quel niveau s'établit cette valeur ? Ses réponses prennent respectivement la forme d'une théorie des biens et d'une théorie de la valeur qui constituent les deux piliers de l'édifice autrichien.

1.1 La théorie des biens

1.1.1 *Les quatre pré-requis d'un bien.* La théorie des biens de Menger part du concept de *besoin*. Si un individu passe d'un état où il ressent un besoin à un état où ce besoin est satisfait, cela implique, nous dit-il, qu'une « chose » a interféré avec cet individu. Cette chose est alors définie comme *utile*. Et si, en outre, l'individu l'a appliquée en toute connaissance de cause à la satisfaction de son besoin, alors cette chose est un *bien*. Quatre conditions doivent donc selon lui être simultanément présentes pour qu'une chose soit un bien :

- [1] Un besoin humain ;
- [2] La capacité causale de cette chose à satisfaire ce besoin ;
- [3] La connaissance de cette relation causale ;
- [4] Le contrôle de la chose permettant de l'appliquer à la satisfaction du besoin.

1.1.2 *Biens et subjectivité.* De nouveaux besoins, de nouvelles connaissances ou de nouvelles possibilités de contrôle peuvent transformer en biens des choses qui n'en étaient pas auparavant. Une nappe de pétrole, par exemple, devient un bien à partir du moment où l'on connaît son utilité et où l'on dispose des techniques et des moyens de l'extraire du sous-sol. Le fait pour une chose d'être un bien n'est donc pas une propriété « inhérente » à

cette chose, pour reprendre l'expression de Menger, mais une *relation* entre un individu et cette chose. Supposons que le goût pour le tabac disparaisse soudainement et totalement. Dans ces conditions, nous dit Menger, toutes les choses qui ne peuvent servir qu'à satisfaire le besoin de tabac perdraient instantanément leur caractère de bien : les paquets de cigarettes des consommateurs, les stocks entreposés par les producteurs, les feuilles de tabac déjà cueillies, les semences servant à produire le tabac, etc. En revanche, les choses qui pourraient être réutilisées en vue de servir d'autres besoins ne perdraient pas leur caractère de bien, par exemple les terres agricoles qui pourraient être reconverties pour produire d'autres biens.

Böhm-Bawerk (1962 [1881]) caractérise les biens à l'aide du concept de *subjectivité*. Ce sont les connaissances de l'acteur et ses besoins tels qu'il les ressent, qui font qu'une chose est ou non un bien. Même les caractéristiques objectives d'un bien ne peuvent être perçues, appréhendées, qu'à travers la subjectivité des individus. Une chose est un bien, non pas dans l'absolu, mais toujours pour un certain individu, placé à un certain moment dans une certaine situation. Deux biens matériels identiques du point de vue physique mais situés à des endroits différents, ou disponibles à deux dates différentes, sont deux biens *différents* puisque leurs capacités causales à satisfaire des besoins ne sont pas les mêmes.

1.1.3 *Les services*. Menger décrit les biens comme des « choses », mais il précise que ce ne sont pas nécessairement des objets matériels. Il existe un type de biens immatériels : les *services* du travail. Certains biens de consommation sont aussi des services, comme par exemple le diagnostic d'un médecin, le concert d'un musicien, la plaidoirie d'un avocat. Les économistes autrichiens ne distinguent donc pas les biens et les services, comme on le fait habituellement aujourd'hui, mais plutôt les « biens matériels » et les « services » (biens immatériels), à l'intérieur de la catégorie la plus générale qui est celle des « biens ».

1.1.4 *Biens d'ordre supérieur et étapes de production*. Menger classe les biens en deux grandes catégories, d'une part ceux qui satisfont *directement* les besoins, à savoir les *biens de consommation*, et d'autre part ceux qui jouent un rôle *indirect* dans la satisfaction des besoins, à savoir les *facteurs de production*. Mais il

introduit surtout une distinction plus fine entre les biens, selon l'étape plus ou moins éloignée de la consommation finale à laquelle ils se situent. La baguette de pain disponible pour une consommation directe est un bien d'ordre 1 (bien de consommation), la farine qui a servi à la produire est un bien d'ordre 2, le blé qui a servi à produire la farine est un bien d'ordre 3, les semences qui ont servi à produire le blé sont un bien d'ordre 4, et ainsi de suite.

Les différents ordres de biens sont représentés verticalement (figure 1.1) : les ordres plus proches de la consommation se situent vers le bas (aval) et les plus éloignés vers le haut (amont). Les biens d'ordre supérieur ou égal à 2 sont les facteurs de production, et ils peuvent appartenir à des étapes « basses » (proches de la consommation finale) ou à des étapes « hautes » (éloignées de la consommation finale). La *structure de production* d'un système économique est constituée de l'ensemble plus ou moins complexe et entrelacé des étapes de production qui conduisent finalement à la production de l'ensemble des biens de consommation.

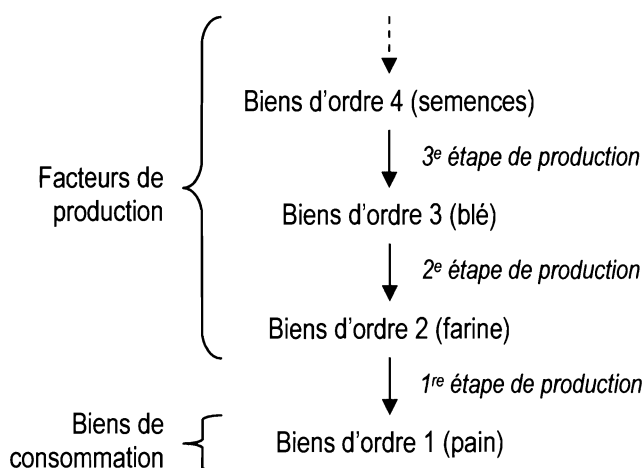


Figure 1.1. Ordres de biens et étapes de production

1.1.5 *Biens complémentaires ou substituables, convertibles ou spécifiques.* Pour qu'un bien d'ordre supérieur puisse servir à satisfaire des besoins, il faut que l'acteur économique dispose d'un

certain nombre *d'autres* biens que Menger appelle les biens *complémentaires*, et qui sont indispensables pour mener la production jusqu'à son terme. Si l'individu dispose de blé, mais pas des biens complémentaires qui vont lui permettre de transformer ce blé en pain et de satisfaire un besoin, alors ce blé n'est pas un bien pour lui (en supposant ici qu'il ne puisse pas l'utiliser autrement qu'en produisant du pain).

Menger s'intéresse surtout à la relation de causalité entre les biens complémentaires d'ordre différent (1976 [1871], p. 63). Il énonce le principe selon lequel le caractère de bien est rétrocedé des biens complémentaires d'ordre inférieur vers les biens d'ordre supérieur qui ont servi à les produire. En d'autres termes, c'est parce que les choses qu'il contribue à produire sont des biens qu'un bien d'ordre supérieur est lui-même un bien : c'est parce que le pain est un bien que la farine qui sert à le produire en est un aussi, parce que la farine est un bien que le blé aussi en est un, et ainsi de suite en remontant la chaîne du processus de production. Ce principe, d'après lequel la transmission du caractère de bien s'opère des étapes inférieures vers les étapes supérieures, préfigure celui de l'imputation de la valeur (voir § 1.3.1 ci-dessous).

Même s'il ne pousse pas son analyse au-delà du concept de complémentarité, il est facile de la prolonger pour retrouver les distinctions qui seront effectuées plus tard entre biens complémentaires et substituables d'une part, et entre biens convertibles et spécifiques d'autre part. Considérons un bien d'ordre supérieur *A*. Supposons que l'un de ses biens complémentaires *B* devienne indisponible. *A* perd-il son caractère de bien ? Pas nécessairement. Si le bien complémentaire manquant *B* est remplacé par un bien *substituable* *C*, alors *A* conserve son caractère de bien. Si aucun bien substituable n'est disponible, *A* peut demeurer un bien à condition qu'il soit *convertible*, c'est-à-dire utilisable dans un autre processus de production (dont tous les biens complémentaires sont disponibles). Si en revanche *A* est un bien *spécifique* – utilisable dans un seul processus de production – et si l'un de ses biens complémentaires *B* dépourvu de substitut vient à manquer, alors il perd son caractère de bien et devient une simple chose.

1.1.6 *Les biens « non économiques »*. Menger opère une distinction entre les biens « économiques » d'une part, qui sont ceux dont

la disponibilité est inférieure aux besoins et qui doivent être économisés, et les biens « non économiques » d'autre part, qui sont surabondants par rapport aux besoins. L'exemple le plus simple d'un bien « non économique » ou gratuit est l'air que nous respirons. L'air permet de satisfaire le besoin de respirer, qui est à l'évidence un besoin vital, et il est disponible en quantités qui dépassent très largement les besoins (sauf dans des situations exceptionnelles). Pour von Mises cependant, l'air n'est *pas* un bien, parce que le besoin qu'il satisfait ne nécessite pas la moindre action de la part des individus (1985 [1949], p. 99). Lorsque l'acteur bénéficie d'une chose sans avoir à poser un acte conscient et intentionnel, cette chose ne doit selon lui pas être définie comme un bien mais plutôt comme une « condition générale du bien-être humain » (*general condition of human welfare*).

Cette objection de von Mises ne se réduit pas à une simple querelle terminologique où « condition générale du bien-être » aurait la même signification que « bien non économique ». En effet, alors que le raisonnement de Menger est fondé sur la notion de *besoin*, celui de von Mises s'appuie sur la notion d'*action*. Toute action vise la satisfaction d'un besoin, mais toute satisfaction d'un besoin n'est pas la conséquence d'une action, comme le montre l'exemple de la respiration. Une action présuppose une fin visée par l'acteur et des moyens mobilisés pour l'atteindre. Dans ce cadre, un bien n'est autre qu'un *moyen* employé dans une action. Et un moyen est nécessairement rare car s'il était surabondant il n'aurait pas besoin d'être l'objet d'une action (Rothbard 1962, p. 4).

1.1.7 *Les biens « imaginaires »*. Menger établit une autre distinction, entre les « vrais » biens et les biens « imaginaires ». Ces derniers sont les biens pour lesquels la relation de causalité de la condition [2] (énoncée au § 1.1.1) est objectivement fausse. Les gens *croient* que la chose a la capacité causale de satisfaire leur besoin, mais cette capacité est en fait purement imaginaire, comme dans le cas des remèdes de charlatans qui sont censés guérir certaines maladies alors qu'ils sont en réalité inefficaces. Le développement de la civilisation s'accompagne d'une meilleure connaissance du monde qui nous entoure, ce qui tend selon Menger à faire augmenter le nombre de « vrais » biens et diminuer le nombre de biens « imaginaires ». Mais cette distinction est-elle pertinente pour

l'analyse économique ? Elle est évidemment cruciale pour les sciences naturelles pures et appliquées, puisque ces dernières recherchent les vraies causes des phénomènes. Pour l'économiste qui étudie la notion de bien, en revanche, il importe peu que la relation causale soit objectivement valide ou non. Dès lors que les gens croient que cette relation existe, la chose *est un bien* pour eux, même si, pour d'autres qui n'y croient pas, elle n'est pas un bien (von Mises 1985 [1949], p. 99). Les conditions [2] et [3] peuvent donc être reformulées en une seule condition [2'] qui tient mieux compte du rôle de la subjectivité des acteurs dans la définition du bien :

[2'] la *croyance* en une capacité causale de cette chose à satisfaire ce besoin.

1.1.8 *Les droits de propriété et les « relations » sont-ils des biens ?*

Les droits de propriété sont des relations immatérielles entre les acteurs économiques et les biens. Böhm-Bawerk (1962 [1881]) a rédigé une étude approfondie sur la question de savoir si ces droits sont eux-mêmes des biens, et il aboutit à une réponse négative. Un droit de propriété sur un bien n'est rien d'autre que le renforcement de la probabilité de contrôler ce bien, grâce à l'appui de la police et de la justice étatiques. Ces forces étatiques ont pour fonction de faire en sorte que les biens soient possédés par leurs propriétaires légitimes et leurs soient rendus s'ils ont été volés. Le droit de propriété lui-même n'est pourtant pas un bien puisqu'il ne permet pas au propriétaire de parvenir à la satisfaction de ses besoins. Seul le bien possédé – et donc contrôlé – le permet. Si un propriétaire se fait voler son bien, le perd ou le détruit par inadvertance, même si son droit de propriété subsiste le besoin ne peut plus être satisfait (la condition [4] n'est plus respectée). Böhm-Bawerk remarque en outre que si les droits de propriété étaient des biens, alors le propriétaire et possesseur d'un bien aurait toujours *deux* biens, le bien lui-même et le droit sur ce bien. On assisterait donc à une multiplication par deux, tout à fait injustifiée, du nombre de biens.

Menger fait entrer dans la catégorie des biens ce qu'il appelle les « relations », dont il cite les exemples suivants : les entreprises, les clientèles, les monopoles, les droits d'auteur, les brevets et les privilèges commerciaux exclusifs, et même les relations familiales,

amicales et sentimentales. Pourtant, si l'on accepte l'argument de Böhm-Bawerk, il faut conclure que ces relations ne sont pas des biens. Le soutien que nous apporte un ami lorsque nous sommes dans la détresse est-il un bien ? Oui, mais c'est ce soutien – ce service – qui est un bien, et non pas la relation amicale en elle-même. De même, le droit d'auteur n'est pas un bien : ce sont les sommes qu'il rapporte (s'il en rapporte) qui sont des biens, et ainsi de suite. L'argument avancé par Menger pour affirmer que les entreprises, clientèles, etc., sont des biens est qu'il est possible de les acheter. Comme ces relations font l'objet d'un commerce, nous dit-il, elles sont des biens. Pourtant, ces relations ne permettent pas de satisfaire les besoins et ne sont donc pas des biens selon la définition de Menger lui-même.

1.1.9 *Le temps*. La réflexion sur la prise en compte du temps par l'action humaine tient une grande place dans l'école autrichienne, on le verra avec les théories de la production et de l'intérêt. Von Mises évoque la rareté du temps et la nécessité de l'économiser (1985 [1949], p. 107), et Rothbard considère le temps lui-même comme un bien puisqu'il constitue un moyen indispensable à toute action (1962, p. 11).

1.1.10 *Critique de la notion de « bien public »*. L'économie néo-classique standard utilise la notion de « bien public », notamment à partir des travaux de Samuelson (1954). Un bien est dit « public » si le fait qu'un individu le consomme ne réduit pas la consommation de ce même bien par les autres individus (principe de non-rivalité) et si aucun individu ne peut être exclu de sa consommation (principe de non-exclusion). Ce type de bien est aux antipodes du bien privé classique qui est rival et exclusif.

Rothbard (1962, p. 884-885) propose une critique de la notion de bien public dans la perspective autrichienne. L'exemple le plus évident de bien public est censé être celui de la *défense nationale*, rendu par les forces militaires du pays. Mais ces forces militaires sont-elles un bien pour les citoyens ? La condition [1] de Menger requiert un besoin humain à satisfaire. Lequel ? Disons celui de vivre dans un pays en paix, protégé des agressions militaires extérieures. Or, certains individus peuvent souhaiter la guerre, pour des raisons diverses ; pour eux, la condition [1] n'est pas satisfaite et

les forces militaires armées, dans la mesure où elles dissuadent efficacement les agressions extérieures, ne sont pas un bien. Admettons maintenant que la population d'un pays souhaite unanimement la paix. La condition [2'], c'est-à-dire la croyance en la capacité causale de la chose à satisfaire le besoin, est-elle vérifiée ? Les pacifistes peuvent penser que les forces militaires ont un effet, non pas dissuasif, mais au contraire provocateur vis-à-vis des autres pays, qui conduit ces derniers à s'armer à leur tour et risque de déclencher la guerre au lieu de l'empêcher. Pour ces individus, la condition [2'] n'est pas satisfaite et les forces de défense nationale ne sont pas non plus un bien. Rothbard souligne une autre difficulté : les ressources militaires sont inégalement réparties sur le territoire et ne permettent pas de défendre aussi bien certains endroits que d'autres, ce qui contredit tout autant le principe de non-rivalité que celui de non-exclusion. La possibilité même de l'existence de biens publics se trouve ainsi remise en cause.

1.2 La valeur subjective

1.2.1 *Définition et origine de la valeur.* Menger définit la valeur d'un bien comme *l'importance* que l'individu attribue à ce bien du point de vue de la satisfaction directe ou indirecte de ses besoins. Les biens « non économiques », c'est-à-dire ceux qui sont disponibles en surabondance par rapport aux besoins, n'ont aucune valeur : posséder ou non une unité du bien ne change rien à la satisfaction des besoins de l'individu, puisque d'autres unités similaires sont disponibles en surnombre. Une unité d'un bien « non économique » n'a aucune importance pour l'individu, et elle est donc sans valeur pour lui. Une bouffée d'air est dénuée de valeur, dans les conditions habituelles de notre existence, puisque si l'individu en était privé il satisferait tout aussi bien son besoin de respirer à l'aide de l'une quelconque des autres unités d'air surabondantes qui l'avoisinent. Seuls les biens « économiques », dont les dotations sont inférieures aux besoins, ont de la valeur. En d'autres termes, la *rareté* est un pré-requis de la valeur. Si, en suivant von Mises, on rejette la notion de bien « non économique » et que l'on considère qu'un bien est nécessairement rare (et donc économisé), alors on peut simplement dire que tout bien a une valeur (puisque

s'il n'en avait pas il serait surabondant ou inutile et ne serait donc pas un bien).

1.2.2 *Valeur et subjectivité.* Menger insiste d'emblée sur le fait que la valeur n'est pas inhérente aux biens, mais dépend de la relation que l'individu entretient avec eux. Un bien qui possède une très grande valeur pour un individu dans une certaine situation (une outre d'eau dans le désert), peut en avoir très peu dans une autre situation (la même quantité d'eau près d'une source abondante). Comme le précise Böhm-Bawerk, la valeur ainsi définie est un phénomène « subjectif », une *valeur subjective* qui dépend de la hiérarchie des besoins et des dotations d'un certain individu à un certain moment. Si ses goûts changent ou si sa situation se modifie, alors les valeurs qu'il attribue aux biens changent elles aussi. Menger prend l'exemple de la transformation des goûts qui s'opère lors du passage à l'âge adulte, qui fait disparaître la valeur que l'individu attribuait à ses jouets d'enfant et qui confère une valeur à des biens qui n'en avaient pas auparavant.

Le terme « valeur » a bien sûr plusieurs significations très différentes. Böhm-Bawerk distingue la valeur subjective et la *valeur objective*. Lorsque l'on dit que tel bien a pour « valeur » telle somme de monnaie, il s'agit d'une valeur objective qui est tout simplement le prix du bien. Dans ce chapitre, pour éviter les confusions, le terme « valeur » employé sans autre qualificatif signifiera « valeur subjective ».

1.2.3 *L'échelle des valeurs.* Les valeurs respectives qu'un individu attribue aux biens sont hiérarchisées, certaines sont plus élevées et d'autres plus basses. Elles forment donc une échelle, qui peut bien sûr changer d'un moment à l'autre et d'une situation à l'autre. La place d'un bien sur cette échelle de valeurs s'explique selon Menger par un principe simple : un bien a d'autant plus de valeur pour un individu que le besoin qu'il permet de satisfaire est jugé important par cet individu. En d'autres termes, l'individu hiérarchise ses besoins en fonction de leur importance relative, puis il attribue – il *impute* – à chaque bien une valeur plus ou moins élevée qui correspond à l'importance plus ou moins grande des besoins qu'il permet de satisfaire. Ainsi, la hiérarchisation des besoins conduit aussi à hiérarchiser les biens susceptibles de les satisfaire.

Les besoins classés par ordre de préférence par un individu sont des besoins *concrets* et non pas des *catégories abstraites* de besoins. Même si la catégorie de besoin « nourriture » est placée plus haut que la catégorie « divertissement » pour la survie biologique de l'organisme, un individu dont les besoins concrets de nourriture sont déjà correctement satisfaits peut conférer une plus grande importance à la satisfaction d'un besoin concret de divertissement qu'à la satisfaction d'un besoin concret supplémentaire de nourriture : dans cette situation, l'individu attribuera davantage de valeur à un spectacle de divertissement (par exemple) qu'à une portion supplémentaire de telle ou telle nourriture.

1.2.4 *La résolution du « paradoxe de la valeur »*. Un bien a d'autant plus de valeur subjective pour un individu que les besoins qu'il lui permet de satisfaire sont jugés importants par cet individu. Ce principe fondamental de l'évaluation subjective des biens semble se heurter au fameux « paradoxe de la valeur » qui avait été noté par les économistes classiques (Smith 1976 [1776], p. 32-33). Si la valeur d'un bien est d'autant plus grande que les besoins qu'il satisfait sont importants, comment se fait-il que l'eau n'ait que très peu de valeur alors que le diamant, qui ne satisfait aucun besoin vital, a une très grande valeur ? Il est beaucoup plus important de pouvoir étancher sa soif que de posséder un bijou en diamant. L'eau ne devrait-elle donc pas avoir considérablement plus de valeur que le diamant ? La réponse de Menger est que, dans les conditions habituelles de notre existence, l'eau est disponible en très grande quantité. La valeur d'un verre d'eau est très faible parce que si l'individu en était privé, une grande quantité d'eau resterait disponible pour satisfaire son besoin d'eau (dans les situations de la vie courante). La quantité disponible de diamant est en revanche très limitée, et la perte d'un diamant obligerait l'individu à renoncer à satisfaire des besoins beaucoup plus importants que la perte d'un verre d'eau. Le principe fondamental de l'évaluation subjective n'est pas réfuté, mais au contraire confirmé par cet exemple.

1.2.5 *Unité et stock*. Le paradoxe de la valeur provient d'une confusion entre les concepts d'unité et de stock d'un bien. Lorsque l'on dit que « l'eau est beaucoup plus utile que le diamant et devrait donc avoir beaucoup plus de valeur », on raisonne implicite-

ment sur la totalité du stock d'eau disponible pour l'humanité. Et il est vrai que si un individu devait choisir entre la totalité de l'eau potable disponible (stock d'eau) et la totalité des diamants disponibles (stock de diamants), il choisirait l'eau et attribuerait donc plus de valeur à celle-ci qu'au diamant. Or, en l'occurrence, les individus n'ont pas à choisir entre la totalité des stocks de ces deux biens, mais entre des *unités concrètes* de ces biens, comme le dit Menger, qui ne représentent que de très faibles quantités par rapport à la totalité des stocks. Une explication correcte de l'attribution des valeurs doit partir des unités de bien concernées par les choix concrets auxquels sont confrontés les individus : quelle est la valeur pour tel individu à tel moment de telle unité de bien ? Pourquoi est-il disposé à échanger une unité de tel bien contre une unité de tel autre bien ? Pourquoi choisit-il d'employer une unité de tel bien d'ordre supérieur dans telle ligne de production plutôt que dans telle autre ? L'unité du stock de bien est donc l'élément de base de l'analyse de la satisfaction des besoins.

La notion d'unité, tout comme celles de bien et de valeur, présente une dimension subjective (Rothbard 1991 [1954], p. 118). Lorsque quatre clous sont nécessaires pour accrocher une affiche au mur, l'unité de bien se compose d'un ensemble de quatre clous. L'unité prise en compte par l'individu dépend de l'objectif poursuivi et donc de l'action entreprise. Cette considération conduit aussi Rothbard à critiquer les modèles mathématiques standards dans lesquels l'unité de bien est considérée comme infiniment petite pour pouvoir faire l'objet d'un calcul différentiel. L'action humaine ne porte jamais sur des infiniment petits, mais toujours sur des quantités « discrètes » au sens mathématique du terme.

1.2.6 *La loi de l'utilité marginale*. Les notions d'unité et de stock d'un bien conduisent à la loi de l'utilité marginale, que Böhm-Bawerk considère comme la principale loi de la théorie économique. L'expression *utilité marginale* a été inventée par Wieser (1884) pour désigner la valeur subjective d'une unité d'un stock de bien (le terme original en allemand *Grenznutzen* a été traduit en 1889 par *marginal utility*). Soit un individu qui possède un stock d'unités d'un bien :

– les unités de ce stock sont par définition *identiques* aux yeux de l'individu, sans quoi elles représenteraient pour lui des biens

différents ;

- comme elles sont interchangeables, chacune de ces unités a pour l'individu la *même valeur* subjective que chacune des autres ;

- l'utilité marginale d'un bien est définie comme la valeur subjective d'une unité du stock de ce bien (si le stock se réduit à une seule unité, alors l'utilité marginale est égale à l'utilité totale) ;

- *loi de l'utilité marginale* : l'utilité marginale d'un bien est l'importance du besoin *le moins urgent* que l'individu compte satisfaire avec une unité de son stock de ce bien ;

- cette loi est une loi *praxéologique* au sens de von Mises (1985 [1949], p. 130) : ce n'est pas une loi observable, empirique, expérimentale ou testable, mais une loi formelle et a priori de l'action, nécessairement vraie par définition même des termes employés.

Böhm-Bawerk (1959 [1889], p. 143) illustre la loi de l'utilité marginale en reprenant un exemple de Menger, celui d'un fermier isolé qui dispose après sa récolte de cinq sacs de blé. Par ordre d'importance : le 1^{er} sac va servir à assurer son minimum de subsistance jusqu'à la prochaine récolte, le 2^e à accroître la quantité de ses repas quotidiens pour le maintenir en bonne santé, le 3^e sera utilisé pour nourrir des volailles qui lui donneront de la viande qui diversifiera ses repas, le 4^e servira à produire par distillation des boissons alcoolisées, et le 5^e à nourrir des animaux de compagnie. Les cinq sacs sont supposés identiques : ils contiennent la même quantité de blé, sans aucune différence de qualité. Comme ils sont interchangeables, chacun d'eux a pour le fermier exactement la même valeur que chacun des autres. Quelle est cette valeur ? Pour la déterminer, il faut tout simplement se demander quelle est la satisfaction à laquelle le fermier choisirait de renoncer s'il était privé de l'un d'eux. Il est clair que si l'un des sacs était détruit, il ne renoncerait pas à se nourrir. Il choisirait de renoncer à satisfaire le besoin *le moins important*, en l'occurrence maintenir en vie ses animaux de compagnie. L'utilité marginale de son stock de blé correspond donc à l'importance du « dernier » besoin qui serait satisfait avec une unité de son stock.

1.2.7 *Variation de stock et utilité marginale*. Menger et Böhm-Bawerk illustrent l'effet des variations de stocks en approfondissant l'exemple du fermier isolé. Supposons que ce fermier ne pos-

sède que quatre sacs de blé au lieu de cinq. La valeur de chaque sac augmenterait puisque le dernier besoin satisfait avec quatre sacs (boissons) est plus important que celui satisfait avec cinq sacs (animaux de compagnie). Inversement, s'il possédait six sacs au lieu de cinq, la valeur de chaque sac diminuerait puisqu'il attribuerait le 6^e sac à un besoin moins important que celui de nourrir ses animaux de compagnie. Il est clair que plus les sacs dont dispose le fermier sont nombreux, plus leur valeur subjective est faible ; moins ils sont nombreux, plus leur valeur est forte. Plus généralement et sous la condition « toutes choses égales par ailleurs » : si le stock d'un bien dont dispose un individu augmente, alors son utilité marginale (utilité d'une unité) diminue ; et si ce stock diminue, alors son utilité marginale augmente.

1.2.8 *La mesure des valeurs*. Dans les illustrations où il expose sa théorie, Menger présuppose que les valeurs sont mesurables. Il attribue à chaque bien un nombre qui représente sa valeur pour l'individu concerné, puis il *additionne* ces nombres pour définir la valeur totale de cet ensemble de biens pour l'individu. Wieser (1893 [1889]) adopte une conception plus simpliste encore, qui consiste à calculer la valeur totale du stock d'un bien par le produit du nombre d'unités du stock et de l'utilité marginale. Böhm-Bawerk (1959 [1889], p. 147) calcule lui aussi la valeur d'un stock en ajoutant les valeurs des biens qui le composent. Mais il critique à juste titre la procédure de Wieser : les unités du stock d'un même bien sont attribuées successivement à la satisfaction de besoins de moins en moins importants, et la somme de ces valeurs est donc nécessairement plus élevée que la somme des utilités marginales (puisque l'utilité marginale est la plus faible de ces valeurs).

Von Mises (1980 [1912], p. 51-55) va opérer une rupture définitive dans l'école autrichienne, en affirmant que les valeurs subjectives ne sont *pas* mesurables. Elles sont une simple échelle ou gradation qui permet de dire que ce bien est placé plus haut ou qu'il est placé plus bas que cet autre bien, mais pas « de combien » il est plus haut ou plus bas. En d'autres termes, il est impossible d'ajouter des valeurs comme on ajoute des longueurs. Les valeurs subjectives peuvent seulement être ordonnées, elles ne peuvent pas faire l'objet d'un calcul (von Mises 1985 [1949], p. 128) : il est impossible de connaître la valeur d'un stock pour un acteur à partir

de celles de ses unités, ou de connaître la valeur d'une unité à partir de celle d'un stock (sauf bien sûr si le stock se limite à une seule unité). Von Mises a pris ici le parti de Čuhel (1907), qui n'était pas membre de l'école autrichienne, contre celui de Böhm-Bawerk dans le débat qui a opposé ces derniers sur cette question (Hülsmann 2007, p. 218). Dans la terminologie standard, reprise par von Mises et par Rothbard, on dit que les valeurs subjectives constituent une représentation *ordinale*, et non pas une mesure *cardinale*, des préférences entre les biens.

1.2.9 *Valeur et indifférence*. La microéconomie standard formalise les choix économiques en s'appuyant sur la notion de courbe d'indifférence (Pareto 1981 [1909], p. 168-169). Une courbe d'indifférence est l'ensemble des paniers de biens entre lesquels l'individu est indifférent, c'est-à-dire des paniers qui lui apportent exactement la même satisfaction. Or, Rothbard (1991 [1954], p. 123) fait remarquer que l'indifférence est incompatible avec l'action : agir, c'est choisir. Un individu est toujours en mesure de dire, entre deux biens différents, lequel a le plus de valeur pour lui : c'est celui qu'il préfère, c'est-à-dire celui qu'il choisirait de garder s'il devait renoncer à l'un des deux. Les valeurs subjectives définissent donc selon lui un ordre strict entre les biens.

1.2.10 *La comparaison interpersonnelle des valeurs*. Böhm-Bawerk termine son analyse de la loi de l'utilité marginale en affirmant qu'une quantité donnée d'un certain bien a davantage de valeur pour un pauvre que pour un riche. En effet, nous dit-il, le riche dispose de beaucoup plus de biens que le pauvre, et le gain ou la perte d'un bien n'aura donc que peu d'importance pour lui (faible valeur) alors qu'ils auront un fort impact sur la satisfaction des besoins du pauvre (forte valeur).

Dans son ouvrage de réflexion sur l'analyse économique, fortement influencé par des économistes de l'école autrichienne comme Wicksteed et von Mises, Robbins (1947 [1932], p. 134-138) montre que cette application de la loi de l'utilité marginale est *illégitime* d'un point de vue scientifique (sur la source de ces considérations chez Čuhel 1907, voir Hülsmann 2007, p. 219). Car cette loi ne s'applique, en toute rigueur, qu'à un acteur unique : si son stock d'un bien augmente (par exemple sa quantité de mon-

naie), alors toutes choses égales par ailleurs la valeur d'une unité diminue. Mais cette loi ne permet pas de comparer les valeurs attribuées aux unités de monnaie (par exemple) par un acteur à celles attribuées par un *autre* acteur. Il est donc impossible de déduire de cette loi qu'en taxant le riche et en subventionnant le pauvre, l'État augmentera la satisfaction globale de la société. Le choix d'opérer une redistribution plus égalitaire des richesses doit s'appuyer sur certains postulats éthiques et ne peut pas se justifier scientifiquement à partir de la loi de l'utilité marginale.

1.2.11 *La préférence pour le présent.* La relation entre les valeurs subjectives imputées aux biens et l'écoulement du temps est un prélude aux théories autrichiennes de la production, du capital et de l'intérêt. Menger remarque qu'une satisfaction présente a généralement plus d'importance pour un individu qu'une satisfaction similaire dans le futur (1976 [1871], p. 153-154). Il donne à ce phénomène de préférence pour le présent un fondement émotionnel que l'on peut considérer comme tautologique : la peur de manquer de biens est plus forte pour les périodes proches que pour les périodes plus éloignées, mais ceci n'est autre que le phénomène à expliquer. Il fournit aussi une explication rationnelle un peu limitée, qui consiste à dire que la satisfaction des besoins futurs requiert nécessairement la satisfaction préalable des besoins présents, car notre santé future peut être gravement compromise par un trop grand sacrifice des besoins présents. En d'autres termes, il faut bien survivre dans le présent et le futur proche si l'on veut pouvoir bénéficier de la satisfaction des besoins plus lointains. Par imputation, on peut en déduire qu'en général un bien présent a plus de valeur que le même bien disponible dans le futur, et c'est exactement en ces termes que Böhm-Bawerk (1959 [1889], p. 259) va formuler le principe qui se trouve au fondement de sa théorie de l'intérêt. Il explique cette préférence pour le présent par trois causes :

(1) les individus qui estiment être moins bien pourvus dans le présent qu'ils ne le seront dans le futur attribuent plus de valeur aux biens présents, cela est évident ; mais ceux qui se trouvent dans la situation inverse (mieux pourvus dans le présent que dans le futur) peuvent en général transférer leurs ressources présentes vers le futur, et donc préféreront eux aussi les biens présents aux

biens futurs.

(2) Les individus sous-estiment systématiquement leurs satisfactions futures par rapport à leurs satisfactions présentes, à cause de l'incertitude de l'avenir (on n'est même pas sûr d'être encore vivant pour profiter des satisfactions futures) et à cause d'un défaut de volonté (lorsque la tentation de profiter du moment présent est trop forte) ; cette explication est aussi celle de Wieser (1893 [1889], p. 17), qui évoque le conflit entre la pulsion et la raison.

(3) Disposer d'un facteur de production dans le présent est toujours préféré à en disposer dans le futur, car ce facteur contribue à produire des quantités de biens de consommation d'autant plus grandes qu'il est utilisé dans un processus de production plus long (voir § 4.1.8).

Fetter (1915, p. 240), l'un des principaux successeurs de Böhm-Bawerk sur ce thème, retient lui aussi l'explication par l'incertitude de l'avenir, ainsi qu'une explication biologique – mais qui semble tautologique – selon laquelle « l'impulsion à rechercher la gratification immédiate est profondément enracinée dans la nature biologique de l'homme ».

1.2.12 *La loi de la préférence temporelle.* Von Mises est le premier à avoir énoncé la loi de la préférence temporelle comme un principe de pure logique de l'action (un principe praxéologique), débarrassé de toute considération biologique ou psychologique : « La satisfaction d'un besoin dans un avenir rapproché est, toutes choses égales d'ailleurs, préférée à la même dans un avenir distant » (1985 [1949], p. 508). Sa démonstration (par l'absurde) est la suivante. Imaginons une suite d'intervalles de temps entre lesquels la condition « toutes choses égales par ailleurs » est respectée. Si l'on suppose que l'acteur choisit de ne *pas* consommer ce bien lors du 1^{er} moment, alors il n'y a aucune raison, puisque toutes choses sont égales par ailleurs, qu'il le consomme lors du 2^e moment, ni lors du 3^e, et ainsi de suite : l'acteur ne consommera jamais le bien, ce qui est absurde. Il le consommera donc nécessairement lors du 1^{er} moment.

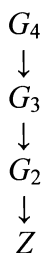
La formulation misésienne de cette loi permet de rendre compte de ses exceptions apparentes. Lorsque par exemple un acteur situé en hiver préfère un stock de glaçons l'été suivant à un stock immédiatement disponible, il compare des biens différents :

la glace étant abondante en hiver et rare en été, la condition « toutes choses égales par ailleurs » n'est évidemment pas respectée et la loi de la préférence temporelle ne s'applique pas.

1.3 La valeur des facteurs de production

1.3.1 *Le principe d'imputation.* Pour Menger, les facteurs de production, c'est-à-dire les biens d'ordre supérieur, acquièrent leur valeur selon le même principe que les biens de consommation, à savoir en fonction de l'importance que l'individu leur attribue du point de vue de la satisfaction de ses besoins. Mais, par définition, un bien d'ordre supérieur ne satisfait pas directement les besoins. Il participe d'abord à un processus de production qui donne naissance à des biens d'ordre inférieur, et finalement des biens d'ordre 1 (biens de consommation). Le bien d'ordre supérieur ne peut donc avoir de la valeur que si ces biens d'ordre inférieur ont eux aussi de la valeur. En effet, s'ils n'en avaient pas – s'ils n'avaient pas d'importance pour l'individu –, alors le bien d'ordre supérieur qui est utilisé pour les produire ne pourrait pas non plus en avoir. Comme, en outre, le processus de production prend du temps, au moment où l'individu impute une valeur au bien d'ordre supérieur, ces biens d'ordre inférieur n'existent pas encore et leur valeur au moment de leur apparition doit être *anticipée*. Menger en conclut logiquement que la valeur des biens d'ordre supérieur est « déterminée par la valeur anticipée des biens d'ordre inférieur à la production desquels ils contribuent » (1976 [1871], p. 150).

Böhm-Bawerk (1959 [1889], p. 170) prend l'exemple de groupes de facteurs de production G_4 , G_3 et G_2 utilisés au cours d'étapes successives pour produire un bien final Z :



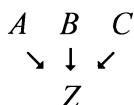
Si l'on suppose que ces groupes de facteurs sont irremplaçables (non substituables) et spécifiques (ils ne peuvent pas servir à produire un bien autre que Z), et si l'on néglige la durée du processus de production (pour éviter la complication due à la préférence pour le présent), alors tous ces groupes ont pour l'individu la *même* valeur que le bien produit. Par imputation, le bien Z confère sa valeur au groupe de facteurs G_2 (privé de G_2 , l'individu serait privé de la satisfaction correspondant à Z), le groupe G_2 confère sa valeur au groupe de facteurs G_3 (privé de G_3 , l'individu serait privé de la satisfaction correspondant à G_2), etc.

D'après le principe menguérien de l'imputation, un bien reçoit donc sa valeur, non pas de ses biens d'ordre supérieur (biens qui servent à le produire) mais au contraire de ses biens d'ordre inférieur (qu'il contribue à produire). Ce principe constitue l'un des traits les plus importants et les plus profonds de la révolution autrichienne. Il s'oppose radicalement à toutes les théories qui voudraient expliquer la valeur des biens par celle des facteurs qui ont servi à les produire. En particulier, les efforts ou les sacrifices – parfois appelés les « coûts réels » – qui ont été consentis pour produire un bien ne lui confèrent d'après Menger pas la moindre valeur, puisque seule sa capacité à satisfaire directement ou indirectement des besoins (futurs) peut expliquer sa valorisation subjective.

1.3.2 La valeur des facteurs complémentaires. Compte tenu de la nécessaire complémentarité entre les facteurs de production, on peut se demander s'il est possible pour l'acteur d'évaluer ces biens séparément les uns des autres. La réponse apportée par Menger est positive : la valeur d'un facteur de production correspond, comme celle d'un bien de consommation, à l'importance des besoins auxquels l'individu devrait renoncer s'il était privé de ce facteur. Menger remarque que si ce facteur est remplaçable, l'individu pourra réduire la perte de satisfaction – mais pas la supprimer – en réorganisant la production, c'est-à-dire en répartissant différemment les facteurs entre les lignes de production. Pour approfondir le raisonnement, plusieurs cas doivent donc être distingués selon que le facteur est remplaçable (substituable) ou non, et selon que ses compléments sont convertibles ou non.

Inspirons nous de l'analyse de Böhm-Bawerk (1959 [1889],

p. 162-163), plus détaillée et plus systématique que celle de Menger, et considérons trois facteurs A , B et C complémentaires que l'individu choisit d'appliquer à la production d'un bien de consommation Z (nous pouvons par exemple imaginer que A est une lampe de poche, B la pile de cette lampe, et C le travail pour mettre la pile dans son emplacement). Ces facteurs sont utilisés ensemble pour effectuer le processus de production :



Pour déterminer la valeur du bien A , supposons que l'individu en soit privé.

1^{er} cas : A est *irremplaçable*, c'est-à-dire qu'aucun autre bien ne peut lui être substitué. Si B et C sont *spécifiques* (ne peuvent être utilisés que dans ce processus), alors la perte du bien A prive B et C de toute valeur, et l'acteur subit la perte de satisfaction maximale qui correspond à la perte du bien de consommation Z produit grâce à ce trio. Si, en revanche, B et C sont *convertibles* (peuvent être utilisés dans d'autres processus), alors la perte du bien A ne prive pas B et C de valeur puisque ces derniers peuvent être réutilisés dans d'autres processus où ils contribuent à produire d'autres biens de consommation, et donc à réduire la perte de satisfaction par rapport au cas précédent.

2^e cas : A est *remplaçable*, c'est-à-dire qu'un bien D possédé par l'acteur peut lui être substitué (à la limite, l'acteur pourrait réaffecter certains de ses facteurs à la production d'un nouveau bien A). Si B et C sont *spécifiques*, alors l'acteur a le choix entre (a) retirer le bien D du processus où il aurait été affecté et le substituer à A (c'est-à-dire remplacer la combinaison ABC par la combinaison DBC), et (b) ne pas remplacer le bien A et utiliser le bien D comme initialement prévu, auquel cas B et C perdent toute valeur ; il choisit celle de ces deux possibilités qui minimise la diminution de satisfaction due à la perte de A . Si B et C sont *convertibles*, l'acteur a deux possibilités : (a) substituer le bien D au bien A , (b) ne pas opérer de substitution et réaffecter les biens B et C à d'autres processus de production ; comme précédemment, il choisit l'option qui lui fait perdre le moins de satisfaction.

Les cas où le bien A est remplaçable et où les biens B ou C sont convertibles sont plus favorables à l'acteur puisqu'ils lui offrent davantage de possibilités pour atténuer sa perte de satisfaction s'il est privé de A .

1.3.3 Proportions fixes ou variables des facteurs. Le raisonnement précédent a été conduit implicitement jusqu'ici dans un cas de *proportions fixes* : ABC produisent, dans ces proportions, un bien Z , $2A2B2C$ produisent deux biens Z , $3A3B3C$ produisent trois biens Z , etc., les biens A , B et C devant toujours être combinés dans la même proportion pour que le processus de production puisse avoir lieu. Menger montre que le raisonnement peut aussi être conduit dans le cas où les proportions sont *variables*, c'est-à-dire quand le bien de consommation Z peut être produit avec des proportions différentes des biens A , B et C . Si l'individu a par exemple choisi d'affecter $3A$, $5B$ et $4C$ à la production d'unités du bien Z , alors en perdant une unité de A , et même si B et C sont spécifiques, le processus de production peut être effectué en combinant $2A$, $5B$ et $4C$. La production du bien de consommation Z s'en trouve néanmoins réduite et la valeur de l'unité de A – l'utilité marginale de A – correspond à l'importance pour l'individu des besoins qu'il ne peut plus satisfaire à cause de cette baisse de la production (on peut bien sûr imaginer des cas plus complexes si les biens B et C sont convertibles).

1.3.4 Imputation et facteurs convertibles. Menger a énoncé et démontré le principe d'imputation de la valeur des produits vers leurs facteurs, mais il revient à Wieser (1884, 1893 [1889]) d'avoir montré comment il s'applique dans le cas très important où les facteurs sont *convertibles*. En effet, si un facteur est spécifique, c'est-à-dire ne peut être utilisé que dans un seul type de processus et dans aucun autre, il est facile de conclure qu'il reçoit sa valeur de celle des besoins qu'il permet de satisfaire. Böhm-Bawerk nous le fait comprendre avec l'exemple du vignoble de Tokay : ce n'est pas parce que ce vignoble (bien spécifique) a une grande valeur que le vin qu'il produit a une grande valeur, mais au contraire parce que ce vin a une grande valeur que son vignoble en a une aussi. Si le principe d'imputation s'applique sans difficulté au cas des biens spécifiques, celui des biens convertibles est plus délicat.

Considérons un groupe de facteurs G qui permet de produire le bien X , mais aussi le bien Y et le bien Z . Supposons que l'acteur possède trois unités de G et qu'il décide de les affecter respectivement à la production d'un bien X , d'un bien Y et d'un bien Z :

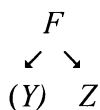
$$\begin{array}{ccc}
 G & G & G \\
 \downarrow & \downarrow & \downarrow \\
 X & Y & Z \\
 \text{(Valeur : 5)} & \text{(Valeur : 7)} & \text{(Valeur : 12)}
 \end{array}$$

Supposons maintenant que la valeur du bien Z soit la plus élevée, et attribuons-lui arbitrairement le nombre 12 ; la valeur du bien X est la plus faible, attribuons-lui le nombre 5 ; la valeur de Y est intermédiaire, égale à 7. La valeur d'une unité de G est 5, puisque si l'individu en est privé il choisira de produire Y et Z avec ses deux unités restantes de facteurs, et renoncera à la satisfaction la plus faible, c'est-à-dire celle obtenu grâce au bien X et de valeur égale à 5.

Si la valeur du produit Z est 12 et que celle de son facteur G est 5, le principe d'imputation n'est-il pas contredit ? Non, car la valeur de Z n'atteint 12 que lorsque les trois processus de production ($G \rightarrow X$, $G \rightarrow Y$, $G \rightarrow Z$) ont été effectués, c'est-à-dire lorsque les unités de G n'existent plus et n'ont donc plus aucune valeur pour l'individu. Mais lorsque ce dernier vient de produire le bien Z et qu'il lui reste deux unités de G , la valeur de Z est égale à 5. En effet, s'il est privé de Z il utilisera ses deux unités restantes de G pour produire Y et un nouveau Z , et ne devra donc renoncer qu'à la valeur de X : la valeur du produit Z (= 5) est bien égale à la valeur de son facteur G (= 5).

On pourrait avoir l'impression que le facteur G impose sa valeur 5 au produit Z . Mais l'analyse montre que le bien G reçoit, en fin de compte, sa valeur des biens de consommation qu'il sert à produire, en l'occurrence du bien X . Le principe d'imputation est donc respecté dans le cas des facteurs convertibles, et l'attribution de la valeur s'opère là aussi des produits vers leurs facteurs, même si elle emprunte un chemin détourné. Comme l'écrit Böhm-Bawerk, les facteurs convertibles « reflètent » la valeur qu'ils reçoivent de leur produit marginal sur leurs autres produits (1959 [1889], p. 176).

1.3.5 *Le coût d'opportunité*. Böhm-Bawerk (1959 [1884], p. 185) prend l'exemple d'un individu qui utilise un groupe de facteurs F pour produire un bien Z . Que lui a « coûté » ce bien Z ? Ou en d'autres termes : quel est le *sacrifice* qu'il a dû consentir pour obtenir Z ? Le coût « direct » est évident, nous dit Böhm-Bawerk, puisque pour obtenir le bien Z l'individu a dû sacrifier les facteurs F . Mais que signifie ce sacrifice du point de vue subjectif? Pour le savoir, il faut déterminer à quelle satisfaction ou valeur l'individu a dû renoncer pour avoir Z . S'il n'avait pas produit Z , il aurait utilisé les facteurs F pour produire un autre bien Y (ce qui suppose implicitement que ces facteurs sont convertibles) :



L'obtention du bien Z a donc conduit l'individu à sacrifier la satisfaction que lui aurait procurée le bien Y : ce coût « indirect », comme le nomme Böhm-Bawerk, est de nature subjective et sera appelé ultérieurement le « coût d'opportunité ». Si les facteurs F sont purement spécifiques et ne peuvent donc pas servir à produire un bien autre que Z , alors leur coût d'opportunité est nul puisqu'ils n'ont aucun usage alternatif. Seuls les biens convertibles sont des biens « coûteux » puisque leur utilisation dans tel processus ou tel autre nécessite un choix et donc un sacrifice. Pour von Mises (1985 [1949], p. 306), toute action vise à atteindre un *profit psychique*, c'est-à-dire un écart positif entre la valeur de l'objectif poursuivi et son coût d'opportunité, entre la valeur du bien produit (Z) et celle de l'usage alternatif des moyens qui ont servi à le produire (Y).

Thirlby (1981 [1946], p. 138-141) résume cette théorie du coût d'opportunité subjectif en disant que lorsqu'un acteur est face à une alternative entre deux actions, le choix de l'une au détriment de l'autre lui coûte ce qu'il aurait obtenu en effectuant cette autre action. Ce coût n'est pas constitué par des biens, et en particulier il ne consiste pas en des sommes de monnaie dépensées de telle ou telle façon. Il est subjectif puisqu'il n'est connu que par l'acteur lui-même, et dépend des options envisagées par celui-ci. Comme il n'est pas omniscient, de nombreuses options lui sont inconnues et

n'interviennent donc pas dans ses choix : un changement de situation ou la découverte de nouvelles options peut l'amener à devoir choisir, et donc à subir de nouveaux coûts. Le coût est « éphémère » en ce sens que dès que l'action est effectuée il n'a plus la moindre importance pour l'acteur qui va faire face à de nouvelles alternatives et donc à de nouveaux coûts. Aussi longtemps qu'une décision de production se déroule selon le plan prévu, elle ne nécessite aucun choix et l'acteur ne subit plus de coût subjectif, même s'il est amené à dépenser – c'est-à-dire à détruire au cours du processus – toutes sortes de ressources.

La troisième grande contribution de Menger, qui constitue une application de sa théorie de la valeur, est une théorie de l'échange. Elle permet d'expliquer pourquoi les individus échangent des biens, selon le principe de la double inégalité des valeurs, et comment se déterminent les taux d'échange, c'est-à-dire les prix. Il reviendra à ses successeurs, et d'abord à Böhm-Bawerk (1959 [1889]), d'élaborer une théorie du système des prix qui réinterprète la vieille « loi des coûts » des économistes classiques dans le cadre entièrement renouvelé du principe d'imputation, en expliquant les coûts par les prix et non plus les prix par les coûts. Les apports de l'école américaine, en particulier ceux de Clark (1899) sur les changements dynamiques et ceux de Knight (1921) sur le profit, seront intégrés par von Mises (1985 [1949]) à la théorie autrichienne du système des prix. Celle-ci diffère en profondeur, aussi bien de la théorie de l'équilibre partiel que de celle de l'équilibre général, et elle transcende la distinction standard entre microéconomie et macroéconomie.

2.1 La théorie de l'échange

2.1.1 Les formes d'interaction sociale. Dans son traité, Menger ne considère qu'une seule forme de circulation des biens entre les individus, à savoir l'échange. Von Mises (1985 [1949]), en revanche, élabore un cadre général qui distingue *l'action autistique* d'une part, et les relations de coopération d'autre part, elles-mêmes divisées en deux catégories, les *relations contractuelles* et les *relations hégémoniques* :

(1) une action autistique ne vise aucune réciprocité et s'accomplit, soit dans l'isolement, soit sans attendre en retour un bien de la part d'autrui (c'est l'exemple du *don*) ;

(2) la relation contractuelle est celle de *l'échange*, dans laquelle un acteur cède volontairement un bien à un autre acteur, en vue d'obtenir un bien volontairement cédé par ce dernier ;

(3) la relation hégémonique est une relation asymétrique dans

laquelle un acteur décide, sous une contrainte qui peut aller jusqu'à la violence physique, d'obéir aux ordres d'un autre acteur ou groupe, sans contrepartie définie et acceptée par avance. Von Mises cite comme exemples de liens hégémoniques les relations entre un État et ses citoyens, entre les parents et leurs enfants, entre le seigneur et ses serfs, entre le maître et ses esclaves.

2.1.2 *L'explication de l'échange.* Menger commence par analyser le type d'échange le plus simple, dans lequel les avantages mutuels sont les plus frappants, qui est celui où un acteur dispose d'une quantité surabondante (excédant ses besoins) d'un bien *A* mais pas de bien *B*, et où un autre acteur dispose d'une quantité surabondante d'un bien *B* mais pas de bien *A*. Si le premier cède au second ses unités excédentaires du bien *A*, et si le second cède au premier ses unités excédentaires du bien *B*, alors chacun d'eux pourra grâce à cet échange améliorer la satisfaction de ses besoins sans subir de sacrifice (sauf ceux directement liés à l'échange lui-même, qui seront évoqués plus bas).

En général, cependant, l'échange impose un sacrifice aux acteurs car ils auraient pu satisfaire certains de leurs besoins grâce aux unités cédées. S'ils choisissent dans ce cas de procéder à l'échange, c'est parce que les unités reçues leur permettent de satisfaire des besoins qu'ils estiment plus importants que ceux auxquels ils doivent renoncer en cédant des unités de leur bien initial. En d'autres termes, pour chacun des deux participants à l'échange, le bien reçu a une valeur subjective supérieure à celle du bien cédé. Menger énonce les trois conditions qui doivent être réunies pour que deux acteurs procèdent à un échange :

(1) ils sont informés de la situation (chacun sait ce que l'autre veut échanger),

(2) ils peuvent effectuer le transfert de biens (il n'y a pas d'obstacle physique ou légal à l'échange),

(3) chacun d'eux attribue davantage de valeur au bien reçu qu'au bien cédé (principe de « double inégalité des valeurs », selon l'expression de Rothbard).

Si une seule de ces trois conditions n'est pas remplie, alors l'échange – le transfert volontaire des biens – n'a pas lieu. La condition (3) montre que l'échange s'explique tout simplement par le fait que les acteurs recherchent une satisfaction plus complète de

leurs besoins. Elle montre aussi qu'il n'y a pas d'égalité des valeurs subjectives dans l'échange : les gens n'échangent pas des biens de même valeur (subjective), mais au contraire des biens dont les valeurs sont différentes. Si les valeurs étaient égales, ils n'auraient aucune raison de procéder à l'échange.

2.1.3 *L'échange à prix fixé.* Menger étudie les limites de l'échange dans le cas très simple où le prix entre les deux biens échangés est fixé. Soit un individu qui possède un stock d'unités d'un bien A et un autre individu un stock d'unités d'un bien B . Supposons que ces unités s'échangent au prix de une contre une, et que ces individus décident de commencer à procéder à des échanges. Au fur et à mesure que cette série d'échanges se déroule, le stock de A du premier individu diminue et son stock de B augmente, et la situation est inverse pour le second individu. D'après la loi de l'utilité marginale : pour le premier individu l'utilité marginale de A augmente et celle de B diminue, et pour le second l'utilité marginale de A diminue et celle de B augmente. L'échange se poursuit aussi longtemps que, pour chacun des deux individus, l'utilité marginale du bien cédé reste inférieure à celle du bien reçu. Si les stocks sont suffisamment importants de part et d'autre, il arrive un moment où, pour l'un des deux acteurs, l'utilité marginale du bien cédé finit par dépasser celle du bien reçu (si son stock est faible, l'acteur peut avoir cédé tout son stock avant d'avoir atteint sa limite de l'échange). Or, un tel échange n'apporterait aucun surcroît de valeur à cet acteur, qui décide donc de ne pas l'effectuer. L'échange trouve alors sa limite et il s'arrête là, même si l'autre acteur souhaite continuer les transferts d'unités.

2.1.4 *Échange et maximisation de la valeur.* Le cas précédent peut être généralisé à l'échange avec plusieurs autres biens, de façon à comprendre comment un individu retire la plus grande satisfaction possible de ses ressources. Soit un individu qui possède un stock d'un bien A (ressources), et qui peut substituer à ses unités du bien A , par l'échange ou la production (peu importe ici), des unités d'autres biens B , C et D . Les taux d'échange entre ces biens sont supposés fixés au taux de 1 contre 1 : $1A \leftrightarrow 1B$, $1A \leftrightarrow 1C$ et $1A \leftrightarrow 1D$. L'individu va pouvoir augmenter la satisfaction qu'il retire de ses ressources en procédant à des échanges successifs entre ses

unités de A et les unités de B , C et D . Il commencera par l'échange qui lui rapporte le plus de satisfaction contre une unité de A , puis, compte tenu des nouvelles utilités marginales, à nouveau l'échange qui lui apporte le plus de satisfaction, et ainsi de suite. Au fur et à mesure que se déroulent les échanges successifs, l'utilité marginale de A augmente et celles respectivement de B , C et D diminuent, par application de la loi de l'utilité marginale. Lorsqu'arrive le point où plus aucun échange n'est avantageux, les utilités marginales de A , B , C et D se trouvent aussi proches les unes des autres qu'elles peuvent l'être pour cet individu. Si les biens étaient infiniment divisibles, on aboutirait à l'égalité des utilités marginales, mais compte tenu de l'indivisibilité des biens l'égalité n'est jamais atteinte et il subsiste des différences entre les utilités marginales (Böhm-Bawerk 1959 [1889], p. 173).

Les économistes néo-classiques standards ont développé des formalisations mathématiques poussées de l'optimisation de la valeur subjective dans l'échange (c'est la théorie du consommateur présentée dans tous les manuels). Mais il faut remarquer que ces exemples d'échanges à prix fixés sont assez artificiels car ils supposent que les préférences subjectives font face à des prix préexistants, alors que, comme on le verra, les prix sont en fait déterminés par les préférences subjectives et les choix de l'ensemble des individus (voir § 2.2.15).

2.1.5 Productivité et coûts de l'échange. L'échange est productif, il est créateur de richesses, au même titre qu'un accroissement du nombre d'unités d'un bien matériel, puisqu'il permet aux individus de mieux satisfaire leurs besoins. Les professions qui servent d'intermédiaires dans les échanges sont donc elles aussi productives, tout comme les professions agricoles ou industrielles. Cette conception de Menger concernant la productivité de l'échange repose bien sûr sur une analyse en termes de valeurs subjectives. Une analyse purement objective conduirait à la conclusion erronée que l'échange ne crée aucune richesse puisque les biens qui sont échangés existaient déjà, toute la richesse était « déjà là ». Or, comme leurs préférences ne sont connues que par les acteurs eux-mêmes, et peuvent changer d'un instant à l'autre, seule une analyse subjectiviste permet d'appréhender correctement la productivité de l'échange.

Menger montre ensuite que l'échange impose toujours des sacrifices, et parmi ces coûts – que l'on nomme aujourd'hui les coûts de transaction – il compte le temps pris pour effectuer la transaction, les assurances, les commissions. Il compte aussi des coûts qui devraient plutôt être considérés comme des coûts de production : les coûts de transport, les coûts d'emballage et de stockage, et les droits de douane et autres impôts sur les échanges. Il est évident que l'échange ne sera effectué que si, pour les deux participants, les avantages surpassent les coûts.

2.1.6 *Valeur d'usage et valeur d'échange*. Pour un individu isolé, les biens ne peuvent contribuer à la satisfaction des besoins que d'une seule façon : par leur utilisation directe. Mais lorsque l'échange devient possible, l'individu peut, soit utiliser le bien lui-même, soit l'utiliser pour se procurer un autre bien en échange. L'utilisation directe permet de satisfaire des besoins dont la valeur (subjective) est définie par Menger comme la *valeur d'usage* du bien. En cédant son bien, l'individu peut satisfaire, grâce au bien reçu en échange, des besoins dont la valeur (subjective) la plus élevée est définie par Menger comme la *valeur d'échange* du bien initialement possédé.

Il existe des biens qui ont une valeur d'usage mais pas de valeur d'échange, par exemple une photo de famille. Inversement, et ce cas est beaucoup plus répandu dans un système économique développé, il existe des biens qui ont une valeur d'échange mais pas de valeur d'usage, comme par exemple le stock de lunettes de vue possédé par un opticien. Enfin, des biens ont à la fois une valeur d'usage et une valeur d'échange, et dans ce cas la valeur du bien pour l'individu qui le possède est tout simplement la plus élevée des deux. L'individu décide d'échanger si la valeur d'échange de son bien est supérieure à sa valeur d'usage, c'est-à-dire si les besoins qu'il peut satisfaire avec le bien obtenu par l'échange sont plus importants pour lui que ceux qu'il satisferait avec le bien qu'il cède. Si la valeur d'usage est la plus haute, alors il conserve son bien pour l'utiliser lui-même. D'un moment à l'autre, si les préférences ou les stocks de l'acteur se modifient, alors ce que Menger appelle le « centre de gravité » d'un bien peut changer, c'est-à-dire que sa valeur d'échange peut passer au-dessus, ou au contraire au-dessous, de sa valeur d'usage.

2.1.7 *Échange et division du travail*. Soit deux acteurs qui produisent les biens de consommation *A* et *B*. S'ils ont la possibilité d'échanger une partie de leur production, ils ont intérêt, pour accroître leur consommation, à se spécialiser chacun dans la production de l'un des deux biens, c'est-à-dire à se consacrer à produire uniquement ce bien. Cette efficacité productive de la division du travail est évidente si l'un des individus est plus efficace pour produire *A* et l'autre pour produire *B* (loi des avantages absolus). Comme Ricardo (1951 [1817]) l'a montré dans sa théorie du commerce international, la division du travail est avantageuse même dans le cas où l'un des deux acteurs est plus efficace que l'autre dans les deux activités : l'acteur le plus efficace a intérêt à se spécialiser dans l'activité où sa supériorité productive est la plus forte. Cette loi des avantages comparatifs sera rebaptisée *loi d'association* par von Mises. Ainsi, la possibilité d'échanger permet la division du travail, qui permet à son tour l'accroissement de la production et le développement des échanges (Rothbard 1962, p. 80-81). La spécialisation sera d'autant plus complète que l'acteur peut s'attendre à échanger une plus grande part de sa production : selon l'expression célèbre d'Adam Smith, « la division du travail est limitée par l'étendue du marché ».

2.1.8 *Marchandises et « échangeabilité »*. Dans le système économique le plus élémentaire, le ou les individus produisent des biens pour eux-mêmes, c'est-à-dire pour leur consommation personnelle. Lorsque le système se complexifie grâce à l'échange et à la division du travail, les individus commencent à produire pour autrui en vue d'obtenir en échange des biens produits par les autres. Mais ils ne produisent que sur commande de leurs clients, lorsqu'ils sont (presque) sûrs de pouvoir effectuer cet échange au terme du processus de production. Enfin, dans la forme développée du système d'échanges, les acteurs produisent des *marchandises*, c'est-à-dire des biens qu'ils tiennent prêts, à la disposition d'acheteurs éventuels. Les producteurs n'ont pas attendu d'être sûrs d'avoir un client, ils ont mené à son terme le processus de production et stockent ces biens produits – ces marchandises – dans l'attente que des clients viennent les acheter (Menger 1976 [1871], p. 238).

La notion de marchandise, comme celles de valeur et de bien,

doit être appréhendée dans le cadre subjectiviste : elle n'est pas une propriété objective des biens, mais une relation entre les individus et les biens. Menger prend l'exemple d'un chapelier qui décide de conserver l'un des chapeaux de son stock pour son usage personnel : ce chapeau perd dès cet instant son caractère de marchandise pour devenir un bien de consommation. Ainsi, dès qu'une marchandise arrive entre les mains de son utilisateur final, elle n'en est plus une et devient un bien de consommation.

L'échangeabilité (en anglais : *marketability*) d'une marchandise est la capacité de cette marchandise à pouvoir être vendue, c'est-à-dire échangée contre d'autres biens. Cette échangeabilité sera plus ou moins forte en fonction du plus ou moins grand nombre d'acheteurs potentiels, de la plus ou moins grande mobilité de la marchandise, de sa plus ou moins forte valeur, de l'extension plus ou moins grande de la zone géographique et de la période de temps où elle peut être vendue, de l'organisation plus ou moins efficace des marchés, etc. Ainsi, un stock de blé a beaucoup plus d'« échangeabilité » qu'un stock de fourrure, puisqu'il intéresse beaucoup plus de gens, dans des zones et périodes plus étendues, et qu'il s'échange sur des marchés mieux organisés.

2.1.9 *Du troc à l'échange monétaire.* L'échange prend tout d'abord la forme du troc, dans lequel les individus cherchent à obtenir, en échange du bien qu'ils cèdent, un bien qu'ils utiliseront directement eux-mêmes, soit pour le consommer, soit pour produire d'autres biens. Comme le note Menger, ce type d'échange se heurte à des limites étroites, car chaque individu doit trouver un partenaire qui possède exactement ce que veut l'individu et qui souhaite exactement ce qu'il propose. Cette double coïncidence des besoins sera difficile à réaliser, surtout pour les biens ayant une faible échangeabilité comme par exemple les biens indivisibles de grande taille. Le propriétaire d'un cheval, qui souhaite l'échanger contre des outils et de la nourriture aura de grandes difficultés à trouver un partenaire qui, non seulement possède tout ce qu'il désire, mais en outre souhaite l'échanger contre ce cheval. La difficulté du troc, c'est-à-dire de l'échange *direct*, devient alors presque insurmontable.

Il existe néanmoins un moyen de contourner cette difficulté en recourant à *l'échange indirect*. Le possesseur du cheval, s'il ne

trouve pas de partenaire direct pour échanger, peut en effet essayer d'atteindre par un chemin détourné les individus qui possèdent les biens qu'il recherche. Il lui faut pour cela essayer de céder son cheval contre des biens qui ont davantage d'échangeabilité. En effet, même si ces biens ne lui sont pas directement utiles, ils ont plus de chances d'être demandés par les individus qui possèdent les outils et la nourriture qui l'intéressent. S'il parvient à se procurer des biens à forte échangeabilité, il se sera rapproché de son objectif initial, puisqu'il lui sera plus facile de les échanger à leur tour contre les biens qu'il recherchait au début. Il aura ainsi débloqué une situation qui ne trouvait pas d'issue par l'échange direct.

Comme chaque acteur a intérêt à céder ses biens contre des biens plus « échangeables », cette procédure d'échange indirect va se généraliser. Elle va permettre d'intensifier la division du travail, mais elle va aussi conférer aux biens les plus « échangeables » du système économique un surcroît de demande, car ils ne vont plus seulement être demandés pour un usage direct mais aussi pour un usage indirect dans l'échange. Cette hausse de demande, en accroissant la valeur de ces biens les plus échangeables, va encore augmenter leur degré d'échangeabilité. Au terme de ce processus qui se renforce sous l'effet de ses propres conséquences (processus de causalité circulaire), un petit nombre de biens vont être systématiquement utilisés dans l'échange indirect : ils finissent par être acceptés dans l'échange par tous les acteurs économiques, qui savent qu'ils pourront ensuite les échanger contre ce qu'ils désirent. Ces biens sont alors devenus des *monnaies*. Menger montre ainsi comment l'échange direct laisse place à l'échange indirect puis monétaire, ce qui porte la division – et donc la productivité – du travail à des niveaux qui seraient impossibles à atteindre en l'absence de monnaie.

2.2 Le système des prix de marché

2.2.1 *La notion de prix*. Un prix est un taux d'échange entre les biens. Si deux acteurs échangent 3 pommes contre 2 poires, alors le prix d'une pomme (exprimé en poires) est $2/3$ de poire et le prix d'une poire (exprimé en pommes) est $3/2$ pommes. En économie monétaire, les prix des biens sont exprimés en monnaie, mais on

pourrait aussi bien exprimer le prix de la monnaie en biens : si le prix d'un téléviseur en euros est 449, alors le prix d'un euro en ce téléviseur est $1/449$.

Pour Menger, le phénomène du prix est dérivé de celui de l'échange, qui provient à son tour du phénomène économique le plus fondamental, à savoir la recherche par les individus d'une meilleure satisfaction de leurs besoins. Il ajoute que le prix n'est en aucun cas le signe d'une quelconque égalité de valeur entre les biens. Les économistes classiques du XIX^e siècle tentaient d'expliquer les prix par l'égalité de la valeur des biens : une table a un prix 5 fois supérieur à celui d'une chaise, donc elle « vaut » 5 chaises. Ils cherchaient ensuite à expliquer cette égalité de valeur par l'égalité des quantités de travail nécessaires pour produire les biens (théorie de la valeur-travail) ou par l'égalité des coûts de production des biens. Ils commettaient là, nous dit-il, une grave erreur qui a causé des « dommages incalculables » à la science économique. Le prix est la conséquence de l'échange, qui suppose, non pas une égalité, mais au contraire une *inégalité* des valeurs subjectives des biens échangés.

2.2.2 Le marché d'enchères. Les bases de la théorie des prix, aussi bien chez Menger (1976 [1871]) que chez Böhm-Bawerk (1959 [1889], p. 220) et plus tard chez Rothbard (1962, p. 206), sont présentées à partir du modèle du marché d'enchère. Ce modèle bien connu est exposé dans tous les manuels actuels d'introduction à l'économie, il est inutile de s'y attarder. La *demande* est la courbe qui relie chaque prix monétaire du bien à la quantité que les acheteurs seraient disposés à acheter pour ce prix, toutes choses égales par ailleurs : elle est déterminée par les préférences subjectives des acheteurs entre la monnaie et ce bien. Compte tenu de la loi de l'utilité marginale, la demande est nécessairement décroissante (*loi de la demande*) : quand le prix augmente, alors la quantité demandée pour ce prix diminue, toutes choses égales par ailleurs. L'*offre* est la courbe qui relie chaque prix du bien à la quantité que les vendeurs sont disposés à céder pour ce prix, toutes choses égales par ailleurs. Contrairement à la demande, son sens de variation est quelconque : l'offre peut être croissante, mais aussi décroissante (l'exemple classique est celui du travail : si son prix horaire augmente suffisamment, le travailleur peut décider de réduire sa quan-

tité offerte de travail pour bénéficier de davantage de loisir). Dans un souci de simplification, les courbes d'offre et de demande seront représentées ci-dessous comme des droites.

Considérons le cas typique où l'offre est croissante (voir figure 2.1). Le prix a tendance à se fixer à son niveau *d'équilibre*, qui égalise la quantité offerte et la quantité demandée : si le prix de marché est supérieur au prix d'équilibre, alors la quantité offerte est supérieure à la quantité demandée et la concurrence entre les vendeurs fait baisser le prix ; si le prix est inférieur au prix d'équilibre, alors ce sont les acheteurs qui sont en concurrence et font monter le prix. Dans les deux cas, le prix tend donc à s'approcher du prix d'équilibre. Ce modèle d'enchère permet de retrouver, dans un cadre général et rigoureux, les résultats de l'ancienne *loi de l'offre et de la demande* : si l'offre augmente (la courbe d'offre se déplace vers la droite), alors le prix d'équilibre tend à baisser ; si la demande augmente (la courbe de demande se déplace vers la droite), alors le prix d'équilibre tend à augmenter, etc.

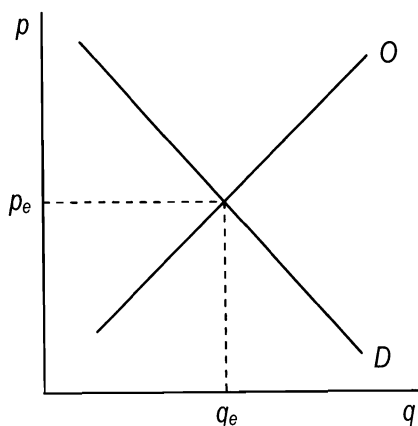


Figure 2.1. Détermination du prix d'enchères

2.2.3 La version classique de la loi des coûts. Lorsqu'un bien est produit à l'aide de facteurs de production, son prix monétaire est déterminé par le modèle d'enchères qui vient d'être présenté, mais dans la mesure où le processus de production se répète d'une pé-

riode à l'autre, le prix de vente du bien entre aussi en relation avec son coût de production, c'est-à-dire avec les prix des facteurs qui sont utilisés pour le produire. Böhm-Bawerk (1959 [1889], p. 248-256) propose des considérations détaillées sur cette relation entre prix et coût monétaire (coût unitaire ou moyen). Il accepte la loi des coûts, selon laquelle le prix d'un bien reproductible à volonté ne peut rester durablement au-dessus ou au-dessous de son coût de production. L'argument d'Adam Smith (1976 [1776], chap. 7) et de Ricardo (1951 [1817], chap. 4) lui paraît tout à fait valide. Il le résume ainsi : si le prix du bien passe au-dessus des coûts de production, alors les producteurs réalisent des profits exceptionnels, ce qui les incite à produire en plus grandes quantités, et ce qui attire de nouveaux producteurs sur ce marché ; la quantité produite augmente, ce qui tend à faire baisser le prix et donc aussi à faire disparaître les profits exceptionnels. Si le prix passe au contraire au-dessous des coûts de production, les pertes subies par les producteurs les conduisent à restreindre leur production ou même à quitter ce marché, l'offre baisse, le prix de vente tend à monter, et les pertes tendent à disparaître. Le modèle d'enchères s'applique donc bien à chaque période, mais la courbe d'offre se déplace, ce qui conduit à ramener l'écart relatif entre prix et coût unitaire vers la rentabilité moyenne (voir figure 2.2).

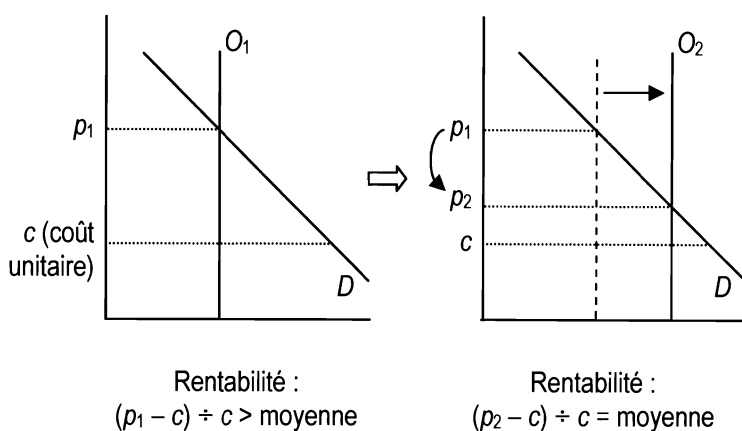


Figure 2.2. La loi des coûts (version classique : le coût unitaire c reste constant pendant l'ajustement)

Böhm-Bawerk (1962 [1894]) note cependant qu'il y a une confusion chez les Classiques et chez Alfred Marshall (1920 [1890]) à propos de la notion de coût, tantôt interprétée comme une dépense monétaire, et tantôt comme un sacrifice en travail (« coût réel »). Or, l'interprétation par le coût réel est incorrecte. Un travailleur qualifié ne subit pas de plus grands sacrifices en travaillant qu'un travailleur non qualifié ; son travail coûte pourtant plus cher (en monnaie), et c'est bien ce salaire monétaire plus élevé qui doit être pris en compte dans la loi des coûts.

2.2.4 L'inversion de la causalité classique : la détermination des coûts par les prix. Böhm-Bawerk a surtout montré que la loi des coûts monétaires devait être complètement réinterprétée, en parallèle avec le principe d'imputation de Menger. Contrairement à ce que pensaient les Classiques, ce ne sont pas les coûts qui déterminent les prix, mais au contraire les prix qui déterminent les coûts.

Il prend l'exemple des produits dont le fer est le facteur de production, en négligeant pour simplifier tous les facteurs complémentaires (travail, etc.) ainsi que le taux d'intérêt. Les consommateurs forment des demandes pour les produits en fer, qui font face aux offres des vendeurs, ce qui détermine les prix (d'équilibre) de ces produits. Les producteurs se basent ensuite sur ces prix pour former à leur tour leur demande de fer. Si le prix d'un produit est suffisamment élevé, par exemple 10 €, ses producteurs sont prêts à payer jusqu'à 10 € par unité de fer pour se procurer la quantité de fer qui leur servira à fabriquer la quantité d'équilibre de leur produit (sous l'hypothèse que chaque unité de produit nécessite une seule unité de fer). Si le prix du produit n'est que de 9 €, ses producteurs n'iront pas au-delà de 9 €, et ainsi de suite pour des produits de prix de plus en plus faibles jusqu'à un produit de 1 €. La confrontation entre ces demandes de fer et l'offre des producteurs de fer détermine le prix du fer, par exemple à 3 € l'unité. Tous les producteurs dont le produit se vend à 3 € ou plus se procurent ainsi le fer dont ils ont besoin et peuvent produire les quantités d'équilibre de leur produit, qu'ils écoulent auprès des consommateurs.

Comme le précise Böhm-Bawerk, ce prix du fer (3 €) est la conséquence d'une chaîne causale qui part des évaluations subjectives par les consommateurs des produits en fer, puis se dirige vers

les prix de ces produits en fer, et enfin aboutit au prix du fer lui-même. Cette séquence conduit bien de la valeur subjective des produits en fer vers leur prix, puis vers leur coût. Elle constitue aussi une illustration de la loi de l'utilité marginale puisque le prix du fer est déterminé par son utilisation « marginale », c'est-à-dire ici par la plus faible valeur monétaire que les consommateurs qui parviennent à se procurer des produits en fer leur attribuent (3 €). Ce raisonnement illustre aussi la nature *causale* de la théorie autrichienne (Rothbard 1962, p. 277) : les valeurs et les prix ne se déterminent pas mutuellement dans un équilibre général, mais plutôt selon un ordre de causalité qui part des besoins subjectifs des individus, puis remonte vers les biens qui servent à les satisfaire directement (biens d'ordre 1), remonte d'une étape vers les biens d'ordre 2 (qui servent à produire les biens d'ordre 1), puis vers les biens d'ordre 3, et ainsi de suite jusqu'aux facteurs originaires.

2.2.5 La convergence vers l'équilibre par résorption des profits et pertes entrepreneuriaux. Le raisonnement de Böhm-Bawerk ne s'arrête pas là. Car si les produits en fer se vendent entre 10 € et 1 €, alors que chacun ne contient qu'une seule unité de fer à 3 €, certains producteurs font des profits entrepreneuriaux (par exemple ceux qui vendent à 8 € un produit qui ne leur coûte que 3 €) et d'autres font des pertes entrepreneuriales (ceux qui vendent à 1 ou à 2 €). Les investisseurs capitalistes vont donc :

(1) réduire les investissements dans les branches en déficit, ce qui va faire diminuer les offres et faire monter les prix de ces produits,

(2) augmenter les investissements dans les branches profitables, ce qui va accroître les offres et faire diminuer les prix.

Ce processus de réallocation des capitaux se poursuit aussi longtemps que des différences de rentabilité subsistent entre ces branches de production, jusqu'à ce que, finalement, les prix des produits soient tous fixés à 3 € et que la loi des coûts soit strictement respectée (prix = coûts). Ainsi, tous les consommateurs disposés à payer au moins 3 € pour un produit en fer peuvent se le procurer. Böhm-Bawerk note que ce processus d'ajustement peut modifier le prix du fer. Si par exemple les consommateurs prêts à payer 4 € pour les produits en fer épuisent exactement la totalité des stocks, alors par imputation le prix du fer va monter de 3 € à

4 €, ce qui constitue une illustration supplémentaire du fait que ce ne sont pas les coûts qui déterminent les prix, mais bien le contraire.

2.2.6 Un cas apparent de détermination des prix par les coûts. Supposons que l'offre de fer augmente, soit parce que des mines plus abondantes sont découvertes, soit parce qu'un nouveau procédé technique accroît l'efficacité de la production de ce métal. Le prix du fer va tendre à baisser, ce qui va faire baisser à leur tour les prix des produits en fer. N'y a-t-il pas là un cas où c'est la baisse du coût d'un facteur de production (le fer) qui se répercute sur les prix de ses produits ? Böhm-Bawerk répond par la négative, et il montre – en remontant jusqu'au facteur originaire travail – que même dans ce cas la chaîne causale va bien des prix des produits vers leurs coûts de production.

Il suppose que la journée de travail est payée 1 €, ce prix étant lui-même formé par imputation à partir des prix des produits finaux du travail (le travail étant un facteur originaire, non produit, il n'a pas de coût de production et il est donc impossible de lui appliquer la loi des coûts). Si, avant l'augmentation de l'offre de fer, il fallait 3 jours de travail pour produire une unité de fer, et que désormais 2 jours suffisent, alors le coût (en travail) des produits en fer n'est plus que de 2 €, mais leur prix est resté à 3 € (comme le prix des produits n'a pas changé, ce raisonnement suppose que la production n'a pas varié, et donc que la quantité de travail utilisée pour produire le fer a diminué de façon à maintenir la production constante). La profitabilité de cette branche (fer) attire des capitaux qui servent à embaucher des travailleurs supplémentaires, ce qui conduit à accroître la production de fer jusqu'à ce que cette profitabilité exceptionnelle disparaisse, c'est-à-dire jusqu'à ce que son prix unitaire soit ramené à son coût de 2 € (2 jours de travail à 1 € la journée de travail = 2 € par unité de fer). Sous l'effet de l'augmentation de leur production, les produits en fer voient aussi leur prix baisser de 3 € à 2 €. Ainsi, les prix des biens de consommation déterminent le prix du travail, qui détermine à son tour le prix du fer. Le prix du fer est donc déterminé par les prix des biens de consommation, tels que ceux-ci sont reflétés dans le prix du travail. La détermination fondamentale s'effectue aussi dans ce cas des prix des biens finaux vers les coûts des facteurs, et l'on re-

trouve ici en théorie des prix un principe déjà présenté dans le domaine de la valeur subjective (voir § 1.3.4).

2.2.7 Le produit monétaire marginal. Le raisonnement qui vient d'être effectué repose sur deux hypothèses très simplificatrices. La première est que le fer est le seul facteur de production. Böhm-Bawerk a parfaitement conscience de cette limite, et il précise que, comme un bien est toujours produit par la combinaison de plusieurs facteurs, chaque facteur se voit imputer une fraction seulement du prix de vente. Rothbard (1962) a proposé une analyse plus générale de la détermination du prix des facteurs de production, qui prend en compte cet aspect grâce au concept de « produit monétaire marginal » (*marginal value product*).

Un producteur évalue une unité de facteur (les stocks des autres facteurs étant fixés) grâce à son produit monétaire marginal, qui est le revenu monétaire auquel il devrait renoncer s'il était privé de cette unité. Connaissant sa fonction de production, il sait que s'il dispose de cette unité il produira une quantité q_2 de bien, et que s'il n'en dispose pas il produira une quantité moindre q_1 ($q_1 < q_2$). Compte tenu de la demande pour son produit, il estime que s'il produit q_2 il pourra vendre les unités de son bien au prix p_2 , et que s'il produit q_1 il pourra vendre au prix plus élevé p_1 ($p_1 > p_2$) puisque l'offre sera réduite. Le produit monétaire marginal de cette unité est donc $p_2q_2 - p_1q_1$, et il constitue la somme maximale que le producteur est disposé à payer pour acheter cette unité. Rothbard démontre que la production ne peut avoir lieu que dans une zone où le produit monétaire marginal est décroissant en fonction de la quantité produite : comme le produit monétaire marginal est aussi la demande du facteur, cette dernière est décroissante pour chaque producteur individuel. La demande de marché, agrégation de toutes les demandes individuelles, est donc elle aussi décroissante. La rencontre de cette demande avec l'offre du facteur détermine son prix (prix d'équilibre du marché de ce facteur). Ce prix est un coût monétaire d'opportunité puisqu'il indique à chaque producteur ce que les autres sont prêts à payer pour se procurer une unité supplémentaire du bien.

2.2.8 L'escompte. Le raisonnement de Böhm-Bawerk repose sur une seconde simplification qui consiste à négliger le déroulement

du temps et donc le taux d'intérêt. Or il sait fort bien que les producteurs tiennent compte de la durée plus ou moins longue du processus dans lequel un facteur est utilisé, en escomptant la valeur de son produit au taux d'intérêt. Supposons par exemple que le taux d'intérêt soit de 5 %. Si la durée qui s'écoule entre l'achat de l'unité de facteur et la vente du bien qu'elle contribue à produire est de 1 an, et si le producteur estime que cette « dernière » unité du facteur rapportera 105 € dans 1 an (au terme du processus), alors il est prêt à se procurer cette unité aujourd'hui à condition que son prix courant ne dépasse pas 100 € ($100(1 + 5 \%) = 105$). En effet, si le prix dépasse 100 €, par exemple 101 €, alors l'emploi de cette unité rapportera un intérêt de 3,96 % seulement ($(105 - 101)/101$), moins avantageux que le taux d'intérêt courant de 5 %. Le producteur n'aurait donc pas intérêt à placer 101 € dans ce processus puisque cette somme pourrait lui rapporter davantage dans un autre processus. Rothbard lève les deux hypothèses simplificatrices de Böhm-Bawerk grâce au concept de « produit monétaire marginal escompté » (*discounted marginal value product*), et il parvient au principe général selon lequel les prix des facteurs sont déterminés par imputation du produit monétaire marginal escompté.

2.2.9 Des « obstacles frictionnels » aux « changements dynamiques ». Il est bien évident que dans le monde réel, la loi des coûts n'est jamais respectée de façon stricte. Dans certaines branches ou certaines entreprises, les prix excèdent les coûts, et dans d'autres les prix se trouvent au contraire en deçà des coûts pendant des périodes plus ou moins longues. Ces écarts proviennent de ce que Böhm-Bawerk appelle des « obstacles frictionnels » : les progrès techniques et les changements imprévus d'offre et de demande font constamment apparaître et réapparaître des écarts positifs ou négatifs entre prix et coûts. Ces écarts constituent les sources des profits et des pertes que l'activité entrepreneuriale contribue à résorber par la réallocation des capitaux entre les branches et entre les étapes de production. L'expression « obstacle frictionnel » n'est pourtant pas tout à fait appropriée, car elle ne distingue pas les forces qui tendent à faire diverger les prix des coûts et celles qui se contentent de ralentir l'ajustement entre prix et coûts.

La deuxième génération des économistes autrichiens, celle de von Mises et de Schumpeter, va s'inspirer de l'analyse plus satisfaisante de ces phénomènes proposée par John Bates Clark. Ce dernier a fondé la théorie néo-classique de la dynamique du capitalisme en distinguant les différents types de *changements dynamiques* qui provoquent des recompositions du système productif, et plus précisément, selon ses propres termes, qui entraînent « une altération des tailles relatives des différents groupes industriels » (Clark 1899, p. 56). Il en dénombre cinq : (1) l'accroissement de la population, (2) l'augmentation du capital, (3) l'amélioration des méthodes de production, (4) la sélection des établissements industriels les plus efficaces (la disparition des moins efficaces), et (5) la multiplication (diversification) des besoins des consommateurs. Von Mises (1981 [1922]) propose une typologie en six catégories explicitement inspirée de celle de Clark (changements naturels, changements de la quantité et de la qualité de la population, changements de la quantité et de la qualité des biens du capital, changements des techniques de production, changements de l'organisation du travail, changements de demande). Schumpeter (1999 [1911]) a popularisé une classification un peu différente, qui a pour but de mettre en lumière ce qu'il appelle les « nouvelles combinaisons », c'est-à-dire les innovations qualitatives qui transforment la structure de production : nouveau bien, nouvelle méthode de production, nouveau débouché, nouvelle source de matières premières, nouvelle organisation (par exemple l'apparition d'un trust ou d'un monopole).

2.2.10 Profits et pertes monétaires. Le profit entrepreneurial monétaire est une différence positive entre le prix de vente et le coût de production monétaire au sens large, c'est-à-dire incluant le taux d'intérêt moyen (en toute rigueur, l'intérêt n'est pas un coût mais un revenu de la production : voir § 5.2.1). Si le revenu de la vente ne couvre pas ces coûts au sens large, alors l'entrepreneur subit une perte.

Concernant la question de l'origine des profits/pertes, Knight (1921) a eu une influence majeure sur l'école autrichienne. Il montre que, contrairement à ce que pensait Clark, les changements dynamiques sont à eux seuls insuffisants pour faire naître des profits et des pertes entrepreneuriaux dans le système économique. En

effet, si un changement a été correctement anticipé par les entrepreneurs, alors ces derniers effectuent à temps les ajustements qui font disparaître les sources de profits/pertes : si les entrepreneurs anticipent correctement une forte hausse du prix du blé dans un mois, cette source de profit futur va disparaître puisqu'ils vont dès maintenant réserver des quantités de blé pour les revendre dans un mois, ce qui va faire baisser le prix futur (le prix courant va aussi augmenter, ce qui fait apparaître des profits pour les propriétaires actuels de blé, mais ce profit vient précisément de ce que la hausse du prix futur n'avait pas été correctement anticipée auparavant, car si cela avait été le cas l'occasion de profit futur n'aurait même pas existé dans le présent : elle aurait disparu à une date antérieure). Knight en conclut que ce qui cause les profits/pertes est une divergence entre la situation courante et la situation qui était attendue par les agents économiques. Cette erreur d'anticipation provient elle-même de *l'incertitude du futur*. L'incertitude de l'avenir qui fait naître les profits/pertes doit être distinguée du « risque », qui peut être mesuré en décomptant les occurrences d'un type d'événement. L'incertitude, en ce sens, n'est pas mesurable par un calcul de fréquence car elle correspond à une configuration unique de facteurs de causalité dont certains sont inconnus. Von Mises (1985 [1949], p. 310) reprend la thèse de Knight en affirmant que l'incertitude sur les offres et les demandes futures est la « source ultime » d'où proviennent les profits et les pertes d'entrepreneur.

2.2.11 *La tendance à l'égalité des taux de rentabilité*. Lorsque des profits et des pertes monétaires apparaissent dans le système économique suite à une erreur d'anticipation, les actions des entrepreneurs vont tendre à les faire disparaître et à égaliser ainsi les taux de rentabilité (définis comme l'écart relatif du prix de vente et du coût unitaire : $(p - c) \div c$). Rothbard (1962) et Reisman (1996) exposent ce raisonnement et l'illustrent par de nombreux exemples. Un profit entrepreneurial pour une firme, une branche ou une étape de production attire les investissements, ce qui entraîne deux conséquences :

- (1) tout d'abord, la production du bien va augmenter, ce qui va accroître l'offre et faire *baisser le prix* de vente,
- (2) et ensuite, ces investissements constituent une demande accrue des facteurs de production nécessaires, ce qui tend à élever

le prix de ces facteurs, ou en d'autres termes à *augmenter les coûts de production*.

Sous ce double effet de baisse du prix et de hausse des coûts, le profit tend à disparaître et le taux de rentabilité de l'entreprise, de la branche ou de l'étape de production tend à revenir vers la moyenne. Si ce sont des pertes entrepreneuriales qui apparaissent, un processus symétrique se déroule – hausse du prix et baisse des coûts –, qui fait disparaître ces pertes et ramène lui aussi le taux de rentabilité vers la moyenne des taux du système économique. Il existe donc une tendance à l'égalisation des taux de rentabilité au sein d'une économie de marché, une tendance constamment à l'œuvre face aux chocs dynamiques non anticipés.

En toute rigueur, comme les acteurs visent à maximiser leur satisfaction subjective et non pas leur revenu monétaire, les taux de rentabilité peuvent être différents selon les branches, même à l'équilibre final. Si les gens répugnent pour des raisons morales à investir dans certaines industries (armement, tabac, etc.), le taux de rentabilité d'équilibre de ces dernières sera plus élevé, et il sera au contraire plus faible pour les branches dont le financement apporte en soi un surcroît de satisfaction aux investisseurs (Rothbard 1962, p. 378).

2.2.12 *L'économie en rotation uniforme*. Si, par hypothèse, plus aucun changement dynamique ne se produisait dans l'économie – ni dans les préférences des individus, ni dans les techniques, ni dans la disponibilité des ressources –, alors le système convergerait vers un état d'équilibre statique que von Mises (1985 [1949]) nomme « l'économie en rotation uniforme » (*evenly rotating economy*). Dans cette situation de rotation uniforme, une construction imaginaire qui ne peut pas survenir dans le monde réel, il n'y a ni profit ni perte entrepreneuriaux puisque les mêmes activités sont effectuées de la même façon période après période. Les prix des biens produits sont alors les prix « naturels » des Classiques, les prix « statiques » de Clark (1899) et les prix « finaux » de von Mises : ils sont tous égaux aux coûts de production en facteurs plus l'intérêt sur le capital investi. L'économie en rotation uniforme est un outil théorique qui permet d'une part d'analyser les effets de long terme d'un changement dynamique, et d'autre part de concevoir par contraste la situation réelle d'une économie de

marché comme un système caractérisé à chaque instant par de multiples déséquilibres qui se manifestent par des écarts entre prix et coûts au sens large (c'est-à-dire incluant l'intérêt sur le capital investi).

2.2.13 *Les processus d'adaptation aux grands types de changements dynamiques.* Les trois principaux types de chocs dynamiques du marché sont les chocs de demande, les chocs techniques, et les chocs de ressources. Les processus d'adaptation à ces chocs illustrent le principe de la tendance à l'égalité des taux de rentabilité, qui est aussi celui de la tendance à la disparition des profits et pertes entrepreneuriaux.

(1) *Choc de ressource.* Hayek (1948) analyse quelles seraient les conséquences d'une raréfaction de l'étain. Comme, par hypothèse, l'offre de cette matière première s'est restreinte, ses propriétaires peuvent en obtenir des prix plus élevés ; les entreprises qui utilisent de l'étain tendent à se reporter en partie sur des substituts proches, ce qui accroît la demande de ces substituts et fait augmenter leur prix. L'augmentation des coûts des entreprises qui utilisent l'étain ou ses substituts est donc inévitable (toutes choses égales par ailleurs), ce qui réduit leur taux de rentabilité : les capitaux – et donc aussi les facteurs de production – ont tendance à être réalloués vers des branches qui utilisent moins d'étain ou de ses substituts. Les firmes qui font usage de ces ressources raréfiées les économisent en se contractant, ce qui leur permet de restaurer leur taux de rentabilité.

(2) *Progrès technique.* Hazlitt (2006 [1946]) suppose qu'une entreprise, dans une certaine branche de production *A*, utilise un nouveau type de machine qui permet de réduire les coûts de production. Dans un premier temps, elle réalise un profit, ce qui va d'une part attirer les capitaux et d'autre part inciter les autres entreprises de la branche *A* à l'imiter : sous l'effet de cette concurrence et de l'afflux de capitaux, la production augmente et la baisse des coûts va finir par être répercutée sur le prix de vente du bien *A*. (a) Si la demande qui s'adresse à la branche *A* est *élastique*, alors par définition l'augmentation de la production s'accompagne d'une hausse du revenu global par période (chiffre d'affaires) : la branche *A* bénéficie d'un taux de rentabilité supérieur à la moyenne, les capitaux affluent, son taux de rentabilité

revient vers la moyenne pendant qu'elle s'étend au détriment des branches qui produisent des substituts ; ces dernières voient leur demande baisser, leur prix de vente diminuer, elles subissent des pertes et se contractent. (b) Si la demande est *inélastique*, alors la branche *A* va au contraire se contracter puisque l'augmentation de la production totale s'accompagne d'une baisse de son chiffre d'affaires : son taux de rentabilité passe au-dessous de la moyenne, les capitaux et les facteurs refluent pour s'investir dans d'autres branches où ils contribueront à accroître la production d'autres biens.

(3) *Changement de demande*. Reisman (1996, p. 174, p. 184, p. 209) examine les conséquences d'un changement des préférences des consommateurs. S'ils décident d'acheter davantage d'un certain bien *A*, et moins d'un autre bien *B*, alors le prix de *A* tend à augmenter puisque sa demande augmente, des profits apparaissent et les capitaux affluent ; le prix de *B* tend au contraire à baisser, les producteurs de *B* subissent des pertes et les capitaux refluent. La branche *A* se développe et celle de *B* se contracte, pendant que leurs taux de rentabilité respectifs reviennent vers la moyenne. Si les facteurs de production sont facilement convertibles d'une branche à l'autre, alors les effets restent circonscrits aux deux branches initialement concernées. Mais si la convertibilité est difficile, voire impossible, alors les effets se propagent en amont de la structure : le prix des facteurs qui servent à produire *A* tend à augmenter et les branches en amont se développent, alors que les prix des facteurs qui servent à produire *B* tendent à baisser et les branches en amont se contractent.

Ces modèles de théorie des prix transcendent la distinction standard entre micro- et macroéconomie, puisqu'ils analysent les conséquences des changements dynamiques sur l'organisation globale du système économique en tenant compte des interactions horizontales et verticales entre les marchés.

2.2.14 *La rationalité de l'économie de marché*. La théorie autrichienne des prix montre comment le système de l'échange marchand coordonne l'ensemble des activités économiques en rationalisant l'utilisation des facteurs de production. Quand la demande d'un bien augmente (ou diminue), le fonctionnement de l'économie de marché conduit à un développement (respective-

ment, une contraction) de la branche qui le produit, ainsi que des branches situées en amont. Lorsque survient dans une branche un progrès technique, et si la demande est élastique, alors le fonctionnement du marché entraîne une croissance de cette branche et de sa filière amont, au détriment des branches produisant des substituts qui n'ont pas bénéficié d'un tel surcroît d'efficacité productive ; si la demande est inélastique, alors la branche et ses filières se contractent en valeur au profit des branches de substituts et de leurs filières, ce qui permet d'accroître la production dans l'ensemble du système. Quand un facteur originaire se raréfie, les producteurs sont conduits à l'économiser, à le remplacer par ses substituts, et à réduire leur activité au profit de ceux qui n'en utilisent pas. L'économie de marché s'adapte ainsi aux différents types de chocs en allouant les ressources productives comme le ferait un acteur rationnel unique (Hayek 1948, p. 85).

2.2.15 Les entrepreneurs face à la souveraineté des consommateurs. Les données du marché subissent des changements incessants auxquels s'adapte la structure de production. Mais cette adaptation n'a rien d'automatique : *elle est voulue et imposée par les consommateurs, pilotée par les entrepreneurs, et appliquée par les propriétaires de facteurs* (travailleurs, propriétaires de « terre » et capitalistes).

Cette réalité est évidente dans le cas d'un changement de la demande des consommateurs. Mieux et plus vite les entrepreneurs anticipent ce changement, plus ils réaliseront de profits en appliquant des facteurs de production dans les filières nouvellement valorisées par les consommateurs, plus ils éviteront les pertes en retirant des facteurs des filières récemment dépréciées par ces derniers. Les entrepreneurs qui perçoivent tardivement ce changement ou refusent de se soumettre au diktat des consommateurs subissent les plus fortes pertes et peuvent être « démis » de leur fonction entrepreneuriale (en perdant leurs capitaux et en n'inspirant plus assez confiance pour pouvoir en emprunter de nouveaux). La même réalité est à l'œuvre dans les autres cas. Supposons que dans une branche de production un progrès technique rentable soit susceptible de réduire les coûts et donc le prix de vente. Les entrepreneurs qui intégreront au plus vite cette nouvelle technique feront des profits au détriment de ceux qui tarderont ou refuseront de les

imiter. Pourquoi ? Parce que les consommateurs veulent acheter au meilleur rapport qualité/prix. Ils vont donc réorienter leurs dépenses vers les entreprises qui leur proposent des termes de l'échange plus avantageux, c'est-à-dire vers les firmes innovantes et au détriment des firmes attentistes qui vont subir des pertes. Là encore, ce sont les entrepreneurs qui anticipent correctement les desiderata des consommateurs, ou s'y conforment rapidement, qui obtiendront des profits et éviteront les pertes subies par les entrepreneurs moins attentifs ou moins réactifs.

Les consommateurs sont bien les souverains du processus de marché en ce sens que ce sont eux qui, en dernier recours, valident ou font échouer les plans des entrepreneurs. Ils contrôlent indirectement l'allocation des capitaux entre les branches de production par leurs choix de dépense, et obligent les entrepreneurs recherchant le profit à réduire les coûts de production (ils déterminent aussi, par leurs choix de consommation intertemporelle, l'épargne disponible et donc la longueur de la structure de production : voir § 4.1.11). Lorsque les entrepreneurs se trompent sur la volonté des consommateurs, ou veulent lui résister, ils le payent par des pertes d'autant plus élevées que leur erreur a été plus grave ou leur résistance plus acharnée.

Le processus de marché présente des similitudes avec le processus politique de l'élection démocratique par lequel la volonté majoritaire de la population porte un responsable politique au pouvoir (von Mises 1985 [1949], p. 287). Il y a néanmoins des différences notables entre le processus de l'élection politique et celui de l'élection économique des entrepreneurs par achat ou abstention d'achat des consommateurs. L'élection économique a lieu en continu, non à des intervalles de plusieurs années : le verdict du jour peut être remis en cause dès le lendemain. Dans un scrutin politique, la majorité qui l'emporte récupère tous les rênes du pouvoir. Dans un scrutin économique, en revanche, les minorités influencent le résultat. Pour reprendre l'exemple utilisé par von Mises, les éditeurs ne publient pas que des romans policiers pour le grand public, mais aussi des œuvres philosophiques ou artistiques difficiles pour un public très restreint. Enfin, le nombre de voix de chaque électeur économique dépend du montant de sa dépense de consommation. Un individu plus riche, qui dépense davantage, pèse plus lourd sur le résultat du

scrutin qu'un individu moins riche (en revanche, la dépense de consommation de la classe moyenne dans son ensemble excède de loin celle du groupe des plus riches, et a donc un impact global beaucoup plus important sur l'organisation de la structure de production). Mais l'inégalité des revenus est elle-même un résultat du processus de marché dirigé par les choix des consommateurs, qui déterminent les prix des biens de consommation à partir desquels sont déterminés à leur tour par imputation les revenus des facteurs de production et en particulier les salaires des différentes qualités de travail (von Mises 1985 [1949], p. 286-287). Les empiètements sur la souveraineté des consommateurs proviennent essentiellement des interventions de l'État. Lorsque ce dernier attribue des privilèges monopolistiques à certains producteurs ou finance par l'impôt des établissements de services publics (voir chap. 8), il se substitue aux consommateurs pour décider quoi produire et en quelle quantité.

Chapitre 3

MONOPOLE ET CONCURRENCE

La théorie du monopole fait l'objet de développements substantiels dès la première formulation du paradigme autrichien par Menger (1976 [1871]). Sa théorie du prix de monopole sera reprise par Fetter (1915) et par von Mises (1998 [1944]), mais plus tard critiquée et rejetée par Rothbard (1962). À partir des années 1940, les économistes autrichiens se montrent très critiques vis-à-vis des modèles standards de la concurrence comme celui de la concurrence parfaite et celui de la concurrence monopolistique. Kirzner (1973) élaborera la conception autrichienne de la concurrence entrepreneuriale, en réalisant une synthèse entre la théorie de la concurrence comme processus de Hayek (1946 [1948]) et la théorie de l'entrepreneur de von Mises (1985 [1949]).

3.1 La théorie du prix de monopole

3.1.1 *Monopole économique et monopole politique.* Les monopoles économiques sont ceux qui émergent du fonctionnement du marché libre, dans le respect des droits de propriété privée et de la liberté contractuelle. Les monopoles politiques, en revanche, sont ceux qui bénéficient d'un privilège d'État limitant ou interdisant, par la menace de sanction judiciaire, l'entrée de concurrents sur ce marché. Seuls les monopoles économiques seront évoqués dans ce chapitre. Les monopoles politiques relèvent de la théorie des interventions de l'État traitée au chapitre 8.

3.1.2 *La relativité du monopole.* Menger définit très classiquement le monopoleur comme le vendeur exclusif d'un certain bien (comme le note Reisman 1996, si le bien vendu est défini de façon suffisamment précise en termes de qualité, de localisation, de date de disponibilité, etc., alors tous les producteurs ou presque peuvent être considérés comme des monopoles économiques). À l'aide d'une illustration simple où un monopoleur cherche à vendre son bien aux enchères à un ensemble d'acheteurs, il montre que : (1) si le monopoleur décide de vendre une certaine quantité, alors il ne

pourra pas fixer le prix unitaire de vente du bien au-dessus d'une certaine limite, et (2) s'il décide au contraire de fixer le prix de son bien, alors il ne pourra pas en vendre au-delà d'une certaine quantité (1976 [1871], p. 199-210). Comme le diront Fetter (1915, p. 79) et plus tard von Mises (1998 [1944]), le monopole est *relatif* et non pas absolu. Le monopoleur n'échappe pas aux forces concurrentielles, puisque les consommateurs établissent leur demande pour le bien monopolisé en tenant compte de leur demande des autres marchandises. Les producteurs de tous ces biens, y compris le monopoleur, sont donc en concurrence pour se procurer leurs revenus dans le flux de la dépense totale des consommateurs (voir § 3.2.2).

3.1.3 *La restriction monopolistique.* Bien que les concepts dont dispose Menger soient encore assez rudimentaires, il se rend compte que la situation de monopole ne pose un problème spécifique à l'analyse économique que dans un cas particulier : si la restriction de la quantité vendue permet d'accroître suffisamment le prix de vente pour que le monopoleur puisse augmenter son revenu, c'est-à-dire en d'autres termes si la demande est *inélastique* dans cette zone de prix. Menger n'emploie évidemment pas cette terminologie – « élasticité » et « inélasticité » – puisqu'elle ne sera forgée que plus tard (par Marshall 1920 [1890]). Il se contente d'un exemple numérique très simple, avec un monopoleur en possession de 1 000 unités d'un bien :

- s'il vend la totalité du stock, le prix unitaire s'établira à 6 onces d'argent, d'où un revenu total de 6 000 onces,
- mais s'il ne vend que 800 unités, le prix s'établira à 9 onces par pièce, pour un revenu total de 7 200 onces.

Le monopoleur a donc bien intérêt dans ce cas – demande inélastique – à ne mettre en vente que 800 unités, même si les 200 unités restantes sont des biens périssables qu'il ne pourra pas conserver pour les vendre plus tard. Dans l'histoire des marchés des épices et du tabac, il est arrivé que les producteurs détruisent ainsi une partie de leur récolte pour augmenter leurs revenus par restriction de l'offre. Menger poursuit son illustration en expliquant que si deux producteurs étaient en concurrence, chacun disposant de la moitié du stock de 1 000 unités, alors aucun d'eux n'aurait intérêt à réduire son stock de 200 unités. En situation de concurrence

(c'est-à-dire en l'absence d'entente) la totalité du stock serait mise en vente, alors qu'en situation de monopole une partie du stock seulement sera proposée à la vente pour un prix plus élevé. Dans le cas d'une demande inélastique, la restriction monopolistique de la quantité offerte s'effectue au détriment des consommateurs qui doivent payer plus cher un bien raréfié, mais à l'avantage du monopoleur puisque son revenu s'en trouve augmenté.

3.1.4 Le prix de monopole. Tous les éléments de la théorie du prix de monopole sont déjà présents chez Menger. Il ne reste plus qu'à la formuler de façon plus complète, comme le fait par exemple Fetter (1915) en distinguant « prix de monopole » et « prix concurrentiel », demande inélastique et demande élastique.

Dans le tableau 3.1, la demande se compose (a) d'une partie élastique, lorsque le nombre d'unités vendues est inférieur à 4, puisque l'augmentation du nombre d'unités vendues s'accompagne alors d'un accroissement du revenu brut, et (b) d'une partie inélastique, au-delà de 4 unités vendues, où le revenu brut diminue avec le nombre des ventes.

Nombre d'unités	Prix unitaire	Revenu brut	
1	7	7	Demande élastique
2	6	12	
3	5	15	
4	4	16	Revenu brut maximal
5	3	15	Demande inélastique
6	2	12	
7	1	7	

Tableau 3.1. Prix de monopole et maximisation du revenu brut
(d'après Fetter 1915, p. 82)

Dans la zone de prix/quantités où la demande est *élastique*, le monopoleur n'a aucun intérêt à réduire la quantité offerte puisque cela ne ferait que diminuer son revenu brut : le prix concurrentiel est dans ce cas égal au prix de monopole et la monopolisation du bien n'entraîne aucune restriction de l'offre. Dans la zone de

prix/quantités où la demande est au contraire *inélastique*, un prix de monopole apparaît, qui se situe au-dessus du prix concurrentiel puisque le monopoleur a dans ce cas intérêt à réduire la quantité offerte sur le marché. S'il dispose de 6 unités, par exemple, il a intérêt à n'en mettre que 4 en vente pour un prix unitaire de 4 € (prix de monopole). Si les 6 unités étaient possédées par des producteurs en concurrence, elles seraient vendues au prix unitaire de 2 € (prix concurrentiel).

Il est par ailleurs évident que si un producteur en monopole a intérêt à renoncer à vendre une partie de son stock, c'est qu'il a commis une erreur de prévision. S'il a produit 6 unités et n'a intérêt à en vendre que 4, alors il aurait pu économiser sur ses coûts – et donc augmenter son revenu net – en ne produisant que 4 unités. Un producteur fixe sa production en sorte de maximiser son revenu *net*. Si cette quantité se situe sur la partie inélastique de la courbe de demande, alors il restreint la production jusqu'au point où la demande devient élastique. Si la quantité qui maximise le revenu net se trouve sur une partie élastique de la courbe de demande, alors le producteur a intérêt à fixer un « prix concurrentiel ».

3.1.5 *Les types de monopoles économiques chez von Mises*. Lorsqu'il énumère les types de monopoles, von Mises (1998 [1944]) abandonne la définition étymologique (producteur unique) pour adopter une définition plus large. Le monopole peut être un *producteur unique*, plusieurs producteurs agissant de concert dans le cadre d'un *cartel*, mais aussi un groupe de producteurs indépendants les uns des autres dont le nombre est limité par l'État (Rothbard parle dans ce dernier cas d'un *quasi-monopole*). Les monopoles non étatiques, qui seuls nous intéressent dans ce chapitre, entrent selon lui dans trois catégories :

(1) les « services publics » (*public utilities*), qui sont ici définis comme des producteurs *privés* chargés de la fourniture locale de l'eau, de l'électricité, du gaz naturel, etc. ; bien souvent, aujourd'hui, ces producteurs sont des monopoles politiques, bénéficiant de licences d'exclusivité octroyées par les autorités politiques nationales ou municipales ; avant ces interventions étatiques, ces « services publics » étaient en général fournis par des compagnies en concurrence, mais parfois par des compagnies en monopole

économique (1998 [1944], p. 26),

(2) la production de certains biens volumineux, dont les coûts de transport sont suffisamment élevés pour procurer un monopole s'il n'y a qu'un producteur local,

(3) les ressources naturelles concentrées en un endroit ou un petit nombre d'endroits ; il cite l'exemple du diamant, qui lui semble constituer à l'époque le seul exemple d'un monopole mondial instauré sans aucun soutien des gouvernements.

3.1.6 La nature du revenu de monopole. D'après la théorie du prix de monopole, un monopoleur peut obtenir un gain spécifique – s'il fait face à une demande inélastique – en restreignant la production par rapport au niveau qu'elle atteindrait si plusieurs producteurs se concurrençaient sur ce marché. Ce revenu de monopole est égal à la différence entre le revenu brut du monopole et la somme des revenus bruts qui seraient obtenus par plusieurs producteurs en concurrence. Il est le prix de l'élément qui se trouve à l'origine du monopole et qui selon von Mises peut être : un bien de consommation, un facteur matériel de production, un facteur travail ou une technique tenue secrète (ou bien sûr un privilège légal dans le cas d'un monopole ou quasi-monopole étatique). Le revenu de monopole est imputé à la possession, par le monopoleur, de cette marchandise ou de ce privilège (von Mises 1998 [1944], p. 6).

Le revenu de monopole n'est pas un profit entrepreneurial : il est nécessairement imputé à un facteur et constitue donc le revenu de ce facteur. Les profits et les pertes entrepreneuriales résultent de l'incertitude du futur (voir § 2.2.10). Le gain de monopole ne dépend pas de l'incertitude de l'avenir et il subsisterait, en tant que revenu de facteur, dans l'économie en rotation uniforme (Rothbard 1962, p. 596-599).

Kirzner (1973, p. 23) adopte un point de vue plus nuancé. Selon lui, si un entrepreneur a perçu avant les autres le gain qui pouvait être retiré de la monopolisation d'un certain bien, et s'il a pu se procurer la totalité du stock de ce bien, alors son revenu doit être considéré dans une perspective de long terme comme un profit entrepreneurial, même si à court terme il apparaît comme un revenu de monopole imputable à la possession exclusive d'un facteur. La nature même du revenu dépend ici de la perspective temporelle dans laquelle il est analysé.

3.1.7 *Les conséquences systémiques d'un prix de monopole.* Quelles sont les répercussions sur le système économique de la fixation d'un prix de monopole sur un marché ? Supposons qu'une branche de production dont la demande est inélastique passe d'une organisation concurrentielle à une organisation monopolistique, suite à une fusion ou à une cartellisation. La production va alors être restreinte pour instaurer un prix de monopole qui accroîtra le revenu du producteur. Une partie des facteurs spécifiques, qui par définition ne peuvent être utilisés pour produire d'autres biens, va devenir oisive. Une partie des facteurs convertibles, qui font les frais de la restriction de production, devient en revanche disponible pour les autres branches de production : leurs propriétaires les offrent sur d'autres marchés où leur prix va donc tendre à baisser. Ainsi, la contraction de la branche monopolisée s'accompagne d'un agrandissement de certaines autres branches. En situation de concurrence les entreprises poussent la production jusqu'à la limite de la rentabilité, au-delà de laquelle l'abondance du produit ferait passer son prix au-dessous de son coût au sens large (qui est le coût moyen + l'intérêt sur le capital investi) : tant que le taux de rentabilité est supérieur au taux moyen, les investissements affluent, ce qui fait baisser le prix (puisque l'offre augmente) et monter les coûts (puisque la demande de facteurs s'accroît), jusqu'au point où ce taux est ramené à la moyenne. En situation de prix de monopole, en revanche, la restriction de la production monopolisée s'opère à l'avantage de la production de certains autres biens, ce qui constitue une *distorsion de la production* par rapport aux souhaits des consommateurs (von Mises 1998 [1944]).

3.1.8 *Les prix de monopole ont-ils tendance à remplacer les prix concurrentiels ?* Une critique récurrente a été adressée au système capitaliste au cours du XX^e siècle, à savoir que son évolution se caractériserait par un remplacement graduel des prix concurrentiels par des prix de monopole. Ainsi, un « capitalisme monopolistique » préjudiciable aux consommateurs mais bénéfique aux capitalistes se substituerait peu à peu, et par une tendance inéluctable, au capitalisme concurrentiel d'origine. Von Mises rejette cette thèse. Il estime que la multiplication des monopoles est due, non pas à une tendance naturelle, mais aux interventions des États qui ont octroyé des privilèges monopolistiques de plus en plus nom-

breux aux producteurs en place et aux producteurs nationaux (sous forme de licences et de cartels nationaux ou internationaux). Si les prix de monopole se multipliaient naturellement, alors les gouvernements n'auraient nul besoin de recourir à des législations de ce type.

Rothbard, de son côté, défend la proposition selon laquelle ce sont les prix concurrentiels qui ont tendance à remplacer les prix de monopole dans l'évolution du capitalisme (1962, p. 596). En effet, le développement du marché s'accompagne d'une augmentation du nombre de biens produits. Or, plus le choix entre les biens de consommation s'élargit, plus les possibilités de reporter sa dépense d'un bien vers un autre sont nombreuses, plus les demandes sont élastiques, et moins il y a de possibilités d'établir des prix de monopole.

3.1.9 Prix de monopole ou alignement sur les coûts de production ? Reisman (1996, p. 411) reconnaît que les biens de consommation dont la demande est inélastique sont peu nombreux, mais il note que le cas des biens intermédiaires est différent. Un producteur serait par exemple disposé à payer très cher les pièces détachées sans lesquelles il ne peut pas terminer ses processus de production en cours, ce qui signifie que les vendeurs d'un type de pièce détachée font globalement face à une demande inélastique. Or, ces biens intermédiaires (pièces détachées, etc.) sont le plus souvent vendus à leur coût de production, c'est-à-dire très au-dessous de ce que les acheteurs seraient prêts à payer pour les obtenir si ces biens venaient à se raréfier. Pourquoi les vendeurs ne profitent-ils pas de l'inélasticité de la demande pour augmenter fortement leurs prix ? Parce que sur le moyen terme, et plus encore sur le long terme, ils n'y ont pas du tout intérêt. En effet, s'ils cherchaient à exploiter cette inélasticité de la demande en restreignant leur production, ils ouvriraient toute grande la porte à la concurrence : de nouveaux producteurs feraient leur apparition en proposant des prix proches des coûts de production, et même en les garantissant par contrat à leurs acheteurs. Ces derniers n'hésiteraient évidemment pas à se tourner vers ces nouveaux fournisseurs et à laisser tomber les anciens qui leur faisaient subir la menace de hausses des prix monopolistiques. En pratique, conclut Reisman, même si la demande est inélastique la limite supé-

rière du prix de vente est presque toujours déterminée par les coûts de production des concurrents potentiels.

3.1.10 *Le prix de monopole : une illusion ?* Bien que Rothbard offre une présentation détaillée de la théorie du prix de monopole, il n'adhère pas lui-même à cette théorie puisqu'il considère que le concept de prix de monopole n'est qu'une « illusion » (1962, p. 586). Selon lui, ni un producteur ni un observateur extérieur ne peut distinguer un prix de monopole d'un prix concurrentiel. Ni une restriction de la production, ni le gain particulièrement élevé d'un facteur, ni l'existence de ressources oisives ne permettent d'établir qu'un producteur a fixé un prix de monopole. Et comme il n'y a aucun moyen de savoir si un prix est monopolistique ou concurrentiel, cette distinction n'a pas d'existence sur un marché libre où il ne peut y avoir que des « prix de marché libre ». Seul un privilège étatique de monopole ou de quasi-monopole fait apparaître, sans la moindre ambiguïté, un prix restrictionniste (voir § 8.1.7).

3.2 La concurrence entrepreneuriale

3.2.1 *Concurrence biologique et concurrence catallactique.* Le concept de concurrence n'a pas fait l'objet d'une élaboration très poussée chez les fondateurs de l'école autrichienne. Menger et Böhm-Bawerk qualifient un marché de « concurrentiel » si plusieurs participants cherchent, indépendamment les uns des autres, à se procurer les mêmes biens par l'échange. Fetter (1915, p. 73) définit lui aussi la concurrence comme « la tentative de deux personnes ou plus d'obtenir la même chose ».

Von Mises (1985 [1949], p. 289) propose une discussion plus approfondie dans laquelle il oppose la concurrence « catallactique », c'est-à-dire celle qui caractérise une économie de marché, à la concurrence biologique. Cette dernière désigne l'antagonisme « implacable » qui oppose les animaux confrontés aux limitations de leurs moyens de subsistance. La concurrence catallactique, en revanche, est un aspect du fonctionnement du système de coopération sociale fondé sur la propriété privée, l'échange, et la division du travail. Il la définit comme une « émulation entre des gens qui

désirent se surpasser l'un l'autre », les vendeurs cherchant à offrir des biens moins chers et de meilleure qualité, et les acheteurs surenchérissant pour se procurer les biens. Elle n'est pas une lutte pour la survie où le plus fort triomphe en faisant disparaître le plus faible, mais plutôt un processus de reclassement dans lequel ceux qui échouent sont reportés à une « place plus proportionnée à leur efficacité que celle à laquelle ils avaient voulu parvenir ». Ainsi, les métaphores biologiques ou militaires ne s'appliquent pas à la concurrence catallactique.

3.2.2. *Une conception large de la concurrence.* La concurrence (catallactique) ne se limite pas au marché d'un seul type de biens, ni même aux marchés de biens similaires les uns des autres. Elle constitue un processus global dans lequel :

(1) les consommateurs sont en concurrence pour se procurer les marchandises et les services des entreprises qui produisent les biens de consommation, et ces entreprises sont elles aussi en concurrence (en aval) pour obtenir les dépenses des consommateurs ;

(2) toutes les entreprises sont en concurrence (en amont) pour obtenir les capitaux des épargnants ;

(3) les propriétaires de facteurs originaires (travail, « terre ») et de biens du capital (*capital goods*) sont en concurrence pour vendre leurs ressources aux entreprises, et ces dernières sont en concurrence pour les acheter ; la concurrence pour vendre ou acheter des facteurs est d'autant plus large que ces facteurs sont plus convertibles, et d'autant plus étroite qu'ils sont plus spécifiques – mais elle ne se limite que très rarement à un seul marché.

Les conceptions orthodoxes de la concurrence – modèle de la concurrence pure et parfaite et modèle de la concurrence monopolistique – sont, au contraire, des conceptions très étroites de la concurrence (elles seront critiquées ci-dessous, entre autres pour cette raison).

3.2.3 *La concurrence comme processus dynamique de découverte.* Hayek (1946) introduit une distinction importante entre la concurrence comme processus qui conduit vers l'équilibre d'une part, et la concurrence comme résultat de ce processus d'autre part. Dans un modèle standard comme par exemple celui de la concurrence parfaite, un marché est considéré comme étant en situation de

« concurrence » lorsqu'il a *atteint* son équilibre final. Or, à l'équilibre, les décisions de tous les acteurs sont compatibles, personne ne peut proposer de réelle innovation ni surprendre les autres : la concurrence comme émulation, c'est-à-dire comme véritable compétition dont l'issue n'est pas connue à l'avance, est impossible. Il faut donc distinguer la *concurrence dynamique*, dans laquelle les acteurs découvrent et utilisent peu à peu les informations qui vont leur permettre de proposer à autrui de meilleurs termes de l'échange, et la *concurrence statique* issue de ce processus de découverte. En se concentrant uniquement sur les conditions mathématiques caractérisant la concurrence statique (c'est-à-dire l'équilibre concurrentiel), les économistes standards passent à côté du phénomène de la concurrence dynamique qui est le plus important pour comprendre le fonctionnement d'une économie de marché, et qui correspond aussi à la signification du terme « concurrence » dans le langage courant et le langage des affaires.

3.2.4 *L'élément entrepreneurial de l'action humaine*. Dans un processus de concurrence dynamique, l'action humaine ne peut pas se réduire à une pure logique du choix (Hayek 1948). Au départ du processus, le système est en déséquilibre et les plans des différents acteurs sont incohérents : ces plans ne pourront pas être tous menés à bien comme prévu, et les acteurs devront les rectifier en intégrant les informations non anticipées obtenues dans le cours même du déroulement du processus. Kirzner (1973) approfondit cette analyse hayékienne en distinguant deux aspects très différents de l'action humaine. Le premier est l'aspect *maximisateur*, qui correspond à la pure logique du choix, c'est-à-dire à la rationalité instrumentale : l'acteur adapte les moyens dont il dispose aux fins qu'il vise, compte tenu des techniques qu'il connaît. Le second est l'aspect *entrepreneurial*, qui consiste pour l'acteur à découvrir de nouvelles fins, de nouveaux moyens, ou de nouvelles techniques, et même, en amont, à être attentif à ces possibilités de découverte. Kirzner définit donc cet élément entrepreneurial comme la *vigilance* (*alertness*) que manifeste l'acteur vis-à-vis de la possibilité que surgissent de nouveaux buts, moyens ou techniques, et qui permet de concevoir l'action comme active et créative plutôt que comme passive et mécanique (1973, p. 35).

3.2.5 *La concurrence entrepreneuriale*. Pour les acteurs interagissant dans le cadre d'un processus de marché, cet élément entrepreneurial consiste à être attentif à l'apparition de meilleures possibilités d'achat et de vente. Or, cette vigilance aux occasions d'acheter moins cher et de vendre plus cher est précisément ce qui donne naissance à la concurrence au sens d'une véritable compétition : l'entrepreneuriat constitue donc la source du processus de la dynamique concurrentielle. Kirzner en conclut que la concurrence (dynamique) et l'entrepreneuriat sont deux concepts *indissociables* l'un de l'autre, tout comme les deux faces d'une même médaille. À l'équilibre individuel ou systémique, les acteurs prennent leurs décisions dans un cadre donné de fins, de moyens et de techniques. Il n'y a aucune place pour l'entrepreneuriat, ni non plus pour la concurrence au sens d'une compétition. Ce n'est qu'en dehors de l'équilibre, lorsque les actions des uns et des autres s'entrechoquent, que les acteurs qui font preuve de vigilance entrepreneuriale rectifient leurs plans en identifiant de nouvelles fins, de nouveaux moyens ou de nouvelles techniques. Ces rectifications leur permettent de faire des offres plus adaptées à l'état du marché (vendeurs) ou d'exprimer des demandes apportant davantage de satisfaction (acheteurs). Dans cette perspective, le processus de marché se caractérise à la fois par son aspect concurrentiel et par son aspect entrepreneurial.

Tous les acteurs exercent une vigilance de type entrepreneurial, les consommateurs en étant attentifs aux nouveaux produits qui seraient susceptibles de les satisfaire, ou attentifs aux occasions d'acheter moins cher, les propriétaires de facteurs en restant vigilants aux possibilités de céder leurs facteurs à des conditions plus avantageuses pour eux, etc. Mais l'activité entrepreneuriale par excellence, dans le cadre du processus de marché, est l'activité spéculative *d'arbitrage* qui consiste à acheter pour revendre : acheter et revendre les mêmes types de biens, ou acheter des facteurs pour les combiner et revendre leurs produits. Le « pur » entrepreneur au sens de Kirzner exerce donc sa vigilance pour acheter moins cher et revendre plus cher, en vue d'obtenir un profit, et il ne doit être confondu ni avec les capitalistes (propriétaires de ressources) ni avec les managers. Si un entrepreneur investit ses propres fonds, ou dirige sa propre entreprise, alors il faut considérer qu'il se loue à lui-même ces ressources en capital ou en travail.

L'une des originalités de la conception de Kirzner est de considérer que l'activité entrepreneuriale, en tant que telle, est *toujours* concurrentielle (au sens de la concurrence dynamique). En effet, en l'absence de barrière légale à l'entrée, un entrepreneur ne peut pas être protégé contre les activités entrepreneuriales d'autrui visant à proposer des occasions d'échange plus attractives que les siennes.

La conception hayékienne de Kirzner a fait l'objet de critiques au sein même de l'école autrichienne, dans le cadre du débat sur la « déshomogénéisation » de von Mises et de Hayek (voir § 8.3.8.)

3.2.6 *Entreprenariat et monopole.* Le monopole doit, d'une façon ou d'une autre, s'opposer à la concurrence. Mais si, comme l'affirme Kirzner, l'entreprenariat est « toujours » concurrentiel, alors le processus de marché semble ne laisser aucune place au monopole dès lors que la liberté d'entrée est respectée sur tous les marchés. Il reste néanmoins une possibilité, et une seule, qui empêcherait les entrepreneurs d'exploiter les occasions de profit qu'ils découvrent : la monopolisation d'une ressource. Seul ce contrôle monopolistique d'une ressource (ou un privilège étatique de monopole) permet à un producteur d'échapper en partie à la pression de la concurrence dynamique. Kirzner illustre très simplement cette éventualité en disant que « sans l'accès aux oranges, l'entrée dans la production de jus d'orange *est* bloquée » (1973, p. 103). La pression concurrentielle va bien sûr s'exercer entre le jus d'orange et les autres boissons, et même au-delà entre les boissons et les autres biens, mais le producteur qui a monopolisé les oranges bénéficie néanmoins d'une certaine protection dans son activité de fabrication de jus d'orange. Il n'y a pas de contradiction avec ce qui a été dit au paragraphe précédent, car ce n'est pas en tant que « pur » entrepreneur qu'il est ainsi protégé, mais en tant que propriétaire de ressource. Si l'on imagine un monopoleur – par exemple l'État – contrôlant successivement des ressources de plus en plus nombreuses, la place laissée au processus entrepreneurial concurrentiel se réduirait peu à peu jusqu'à disparaître lorsque l'économie serait entièrement collectivisée. Kirzner insiste aussi sur le fait que le monopole, au sens de la monopolisation d'une ressource, peut résulter d'une activité concurrentielle préalable (voir § 3.1.6).

3.2.7 *Différentiation des biens, coûts de vente, publicité.* Les conceptions standard de la concurrence ont tendance à considérer la différenciation des biens, les coûts de vente et la publicité comme des phénomènes de nature monopolistique. D'après la théorie de la concurrence entrepreneuriale de Kirzner, ces phénomènes apparaissent comme éminemment concurrentiels.

(1) La compétitivité, qui consiste à chercher à offrir à autrui de meilleures occasions d'échange, ne se limite évidemment pas à la concurrence sur les prix de types « donnés » de produits. La pression concurrentielle porte aussi, et peut-être même surtout, sur les types de biens et sur les qualités des biens proposés à la vente. La vigilance entrepreneuriale sert donc à repérer les nouveaux types ou qualités de biens que les clients seraient disposés à payer suffisamment cher pour que la différence entre le revenu et la dépense de production laisse apparaître un profit entrepreneurial. La *différentiation des biens et des qualités* n'est donc pas une échappatoire ou un subterfuge monopolistique, mais au contraire un aspect essentiel de la concurrence dynamique.

(2) Certains économistes (par exemple Chamberlin 1956 [1933]), ont cru pouvoir opérer une distinction entre les coûts de production (dépenses de fabrication et de transport du produit) et les *coûts de vente* (qui servent à augmenter la demande pour ce produit), puis ont dénoncé le risque monopolistique lié à l'extension de ces coûts de vente. Or, pour Kirzner, cette distinction est arbitraire car l'activité entrepreneuriale ne consiste pas à fabriquer un produit (effort de production) puis à essayer de le vendre (effort de vente). Elle vise *uniquement à vendre*. Les dépenses de l'entrepreneur forment une catégorie homogène, exclusivement constituée de coûts de vente. Comme l'avait déjà expliqué von Mises (1985 [1949], p. 340), on ne peut pas opérer de distinction scientifique, dans le processus dynamique du marché, entre une activité socialement utile de production et une activité stérile, voire anti-productive, de marketing.

(3) Parmi les efforts de vente, le plus connu et le plus débattu est la *publicité*. Dans un modèle d'équilibre, la publicité peut être considérée comme un service d'information distinct du produit lui-même, offert et demandé séparément (dans le modèle standard de concurrence parfaite, la publicité est inutile puisque les consommateurs sont supposés informés). Mais dans un processus dyna-

mique dans lequel l'entrepreneur cherche à faire découvrir aux consommateurs de nouveaux biens ou de nouvelles qualités, il n'a pas de sens de considérer la publicité comme un service distinct du bien lui-même puisque les deux sont irrémédiablement liés. Non seulement la publicité ne limite pas la concurrence, mais elle fait au contraire partie des outils indispensables à la rivalité compétitive. Si la publicité est souvent tapageuse et provocante, c'est parce que les entrepreneurs ne peuvent pas se contenter de présenter l'information sur leur produit aux consommateurs, ou de la mettre « sous leurs yeux », mais doivent leur faire *prendre conscience* – de façon récurrente – de l'existence de cette information.

3.2.8 *Critique de la « concurrence pure et parfaite »*. La concurrence sur le marché d'un bien homogène est dite « pure et parfaite » si chaque acteur est de taille suffisamment faible pour ne pas avoir d'influence perceptible sur la fixation du prix, s'il n'y a pas de barrière à l'entrée, et si les acteurs disposent d'une information parfaite. Sous ces hypothèses, le modèle démontre qu'à l'équilibre du marché le prix de vente du produit est unique et égal pour chaque entreprise à la fois à son coût marginal (c'est-à-dire à la dépense de production qui serait nécessaire à l'entreprise pour produire une unité supplémentaire du bien) et au minimum de son coût moyen (pour chaque entreprise, la courbe de coût moyen en fonction de la quantité produite est supposée décroissante jusqu'à un minimum, puis à nouveau croissante : voir figure 3.1 ci-dessous, et voir le manuel standard de Mankiw 1998, p. 356, pour une explication simple de cette forme de la fonction de coût moyen). Les économistes « autrichiens » rejettent totalement ce modèle.

(1) Von Mises (1988 [1944], p. 12) lui reproche d'être beaucoup trop réducteur. La concurrence ne se réduit pas à ce qui se passe sur un seul marché. Une entreprise ne rivalise pas seulement en baissant le prix par rapport à celles qui proposent le même type de bien, mais aussi en proposant des types de biens différents. Chaque chanteur ou chaque acteur connu délivre un type de service unique, et c'est justement cette spécificité qui le rend concurrentiel, non seulement vis-à-vis des autres chanteurs ou acteurs, mais aussi vis-à-vis de tous les autres types de biens sur lesquels les consommateurs dépensent leur argent (les livres, les vêtements,

etc.). Il remarque en outre que si une entreprise n'étend pas sa production jusqu'au point où le prix de vente est égal au coût marginal, cela ne signifie pas qu'elle pratique une politique monopolistique, mais plutôt que la situation d'équilibre final du système économique n'est pas atteinte (dans le modèle standard du monopole et dans celui de la concurrence monopolistique, à l'équilibre l'entreprise a intérêt à restreindre sa production de sorte que le prix de vente reste supérieur au coût marginal). Comme une telle situation d'équilibre ne peut pas survenir dans le monde réel, à cause de trop fréquents chocs dynamiques non anticipés, il n'est pas approprié de la considérer comme un point de référence normatif. L'écart positif entre prix unitaire et coût moyen (incluant l'intérêt) constitue un profit entrepreneurial qui n'est que temporaire et tendra à disparaître avec le développement de l'entreprise ou de la branche de production.

(2) Dans sa critique, Hayek (1946) insiste sur la question de l'acquisition de l'information. Le modèle de concurrence pure et parfaite suppose que les acteurs connaissent déjà les informations que seul le processus de concurrence dynamique a préalablement permis de découvrir, à savoir (a) quels sont les biens demandés par les consommateurs et (b) quels sont les moyens de produire ces biens au moindre coût. Il semble paradoxal à Hayek que la concurrence soit qualifiée de « parfaite » alors que les activités les plus compétitives – publicité, réduction des coûts, différenciation des produits – n'y ont plus le moindre rôle à jouer. Une autre difficulté de ce modèle provient de l'hypothèse d'homogénéité du produit : il n'existe pas deux entreprises qui produisent exactement le même bien. Enfin, considérer la situation de concurrence pure et parfaite comme un idéal normatif lui paraît inadéquat car c'est un critère trop exigeant. Ce que l'on peut attendre, concrètement et raisonnablement, de la concurrence n'est rien de plus – et rien de moins – qu'une tendance à la découverte de ce que veulent les consommateurs et à la satisfaction de leurs besoins à des coûts puis à des prix de plus en plus faibles.

(3) Rothbard (1962, p. 633) définit la concurrence « pure et parfaite » par l'horizontalité de la fonction de demande (demande parfaitement élastique), qui signifie que chaque entreprise est de taille négligeable par rapport au marché et ne peut avoir la moindre influence sur le prix de marché du bien. Il insiste sur l'irréalisme

de cet aspect du modèle : en réalité, la demande ne peut jamais être horizontale. Une augmentation, même très faible, de la quantité produite entraîne une baisse du prix, ce qui fait que toute entreprise a une influence sur le prix de vente.

(4) Reisman (1996) critique les implications normatives du modèle de la concurrence pure et parfaite. Si l'on considère la fixation du prix de vente au niveau du coût marginal comme un idéal d'efficacité à atteindre, alors de nombreuses entreprises devraient pratiquer un prix insuffisant pour couvrir leur coût unitaire moyen. En effet, le coût marginal de certaines activités peut être très faible, voire même nul, par exemple pour les entreprises dont les usines ne fonctionnent pas à pleine capacité, pour les compagnies ferroviaires ou aériennes (lorsqu'il reste des places libres), pour les salles de spectacle, etc. Il faudrait, pour respecter ce soi-disant critère d'efficacité, qu'elles fonctionnent à perte.

3.2.9 Critique de la « concurrence monopolistique ». Dans les années 1930, Chamberlin (1956 [1933]) et Robinson (1933) ont respectivement élaboré les modèles de la concurrence monopolistique et de la concurrence imparfaite, deux théories assez proches qui décrivent une structure de marché intermédiaire entre la concurrence pure et le monopole, censée représenter la situation la plus fréquente survenant dans une économie capitaliste développée.

Dans la conception de Chamberlin, la concurrence monopolistique sur le marché d'un type de biens se caractérise par la coexistence de nombreuses entreprises qui vendent des biens similaires mais légèrement différenciés les uns des autres : elles sont en « concurrence » puisqu'elles sont nombreuses et libres d'entrer sur ce marché, mais aussi en « monopole » puisque le bien vendu par chacune d'elles est spécifique. Cette différenciation des biens implique que la demande à laquelle fait face chaque entreprise est décroissante, alors qu'elle est horizontale en concurrence pure. La courbe des coûts moyens (en fonction de la quantité produite) est supposée avoir une forme en « U », décroissante puis croissante (Mankiw 1998, p. 356). À l'équilibre du marché, l'entreprise en concurrence pure maximise son profit en produisant la quantité qui correspond au minimum du coût moyen : les unités du bien sont donc produites de la façon la moins coûteuse possible. L'entreprise en concurrence monopolistique, du fait que sa demande est dé-

croissante, a intérêt à produire une quantité inférieure à celle qui correspond au minimum du coût moyen : elle ne se situe donc pas au point le plus économique en termes de coût. Par rapport à l'entreprise en concurrence pure, l'entreprise en concurrence monopolistique restreint sa production et augmente son prix de vente (voir figure 3.1), ce qui la rend moins satisfaisante du point de vue de l'efficacité économique.

Les économistes autrichiens sont tout aussi critiques vis-à-vis de cette théorie que vis-à-vis de celle de la concurrence pure et parfaite.

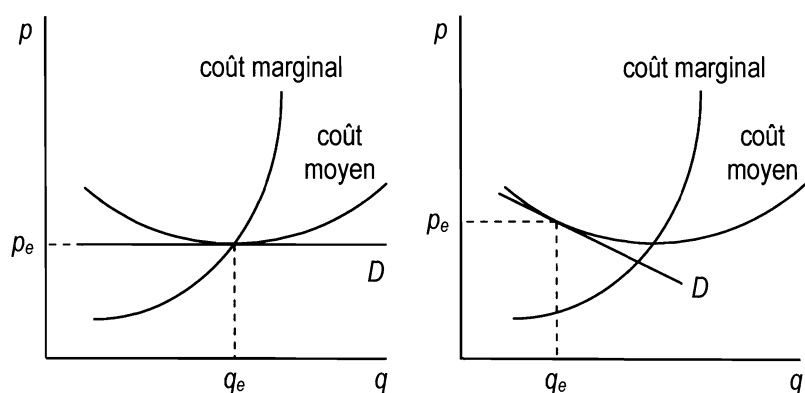


Figure 3.1. L'équilibre de concurrence parfaite (à gauche) et de concurrence monopolistique (à droite)

(1) Von Mises (1998 [1944]) définit la concurrence imparfaite comme une situation dans laquelle l'entreprise n'utilise pas ses équipements à pleine capacité, alors que si elle le faisait cela lui permettrait de diminuer son coût moyen. Il estime qu'une telle situation ne signifie pas nécessairement que l'entreprise applique une politique monopolistique. En effet, pour fonctionner à plein régime elle devrait se procurer des facteurs de production supplémentaires, qui seraient retirés d'autres branches de production. Et si la rentabilité de ces facteurs est plus élevée dans ces autres branches, cela implique que la demande pour les produits de ces autres branches est plus intense. En récupérant ces facteurs supplémentaires, l'entreprise réduirait son coût moyen, mais provo-

querait un gâchis du point de vue des consommateurs en détournant la production des voies qu'ils valorisent le plus.

(2) Hayek (1946, p. 94) ne propose pas d'analyse de la concurrence monopolistique proprement dite. Il se contente de dire que sa critique de la concurrence pure s'applique aussi à la concurrence monopolistique (et Kirzner montrera que c'est bien le cas : voir ci-après).

(3) Rothbard (1962) s'attaque d'abord à la distinction théorique entre la concurrence pure et la concurrence monopolistique. Selon lui, une fonction de demande est toujours décroissante, et la situation de concurrence pure ne peut pas survenir : il ne présente donc aucun intérêt de l'opposer à la concurrence monopolistique, ni de s'en servir comme situation de référence normative. Rothbard considère aussi que la différenciation des produits n'est pas un simple artifice de la part des producteurs, mais permet de mieux satisfaire les besoins variés des consommateurs.

(4) La théorie de la concurrence monopolistique est censée donner une image beaucoup plus réaliste du fonctionnement du système économique que la théorie de la concurrence pure et parfaite, grâce à la prise en compte de phénomènes importants comme ceux de la différenciation des biens et des efforts de vente. Mais pour Kirzner (1973), qui développe la perspective hayékienne, le modèle de concurrence monopolistique ne permet pas de combler les lacunes de celui de concurrence pure et parfaite, parce qu'il souffre exactement des mêmes défauts : les producteurs sont censés avoir *déjà identifié* leur fonction de coût et la demande qui s'adresse à eux. Or, cette hypothèse qui sous-tend les deux modèles montre que l'un et l'autre négligent complètement le processus de concurrence entrepreneuriale qui a préalablement permis de découvrir ces « données », et qu'ils éludent donc le problème central qui est celui de la découverte de ces informations cruciales. En ce qui concerne la différenciation entre les biens produits, la théorie de la concurrence monopolistique suppose qu'à l'équilibre final chaque producteur vend un bien spécifique. Mais si l'entrée est libre pourquoi ne produisent-ils pas le même bien, si c'est celui voulu par les consommateurs ? Comme aucune ressource n'est supposée monopolisée, rien ne les empêche de copier le producteur le plus efficace. Cette théorie se trouve alors dans l'incapacité d'expliquer le phénomène de différenciation qui constitue son

point de départ. Cette différenciation constitue selon Kirzner un choix des entrepreneurs qui déploient leur vigilance pour découvrir les occasions les plus attrayantes à offrir aux consommateurs. En résumé, la théorie de la concurrence monopolistique néglige complètement les aspects les plus fondamentaux de la concurrence dynamique, et commet une erreur en assimilant la différenciation des produits à une forme de monopole. Elle se trompe à la fois sur la nature de la concurrence et sur celle du monopole, ce qui conduit Kirzner à affirmer qu'elle constitue un « épisode regrettable » de la pensée économique moderne (1973, p. 114).

Chapitre 4

LA PRODUCTION ET SA STRUCTURE

Les fondements de l'analyse autrichienne de la production ont été posés par Böhm-Bawerk (1959 [1889]) avec sa théorie du « détour » de production, selon laquelle l'accumulation du capital consiste en un allongement, en un étirement de la structure de production. Hayek (1975 [1931]) a proposé une ingénieuse représentation graphique de cette structure sous la forme d'un triangle qui permet de visualiser la succession des étapes ainsi que l'accumulation du capital, et qui constitue le principal modèle de la macroéconomie autrichienne telle qu'elle a par la suite été développée par Rothbard (1962), Skousen (1990), Reisman (1996), Huerta de Soto (2006 [1998]) et Garrison (2001).

4.1. La production

4.1.1 *La production comme processus matériel.* Böhm-Bawerk (1959 [1889]) aborde la notion de production en reprenant une idée qui remonte aux économistes classiques, à savoir que la production d'un bien matériel n'est pas une véritable création mais un nouvel arrangement de la matière préexistante. Les processus de production des biens matériels consistent donc à contrôler les forces de la nature pour déplacer de la matière, soit par transfert simple (transport), soit par changement de forme (recomposition ou façonnage), soit par combinaison (assemblage). Fetter (1915) distingue ces processus selon quatre catégories : les changements de substance (par cuisson, fermentation, mûrissement, etc.), les changements de forme (tissage, menuiserie, etc.), les changements d'emplacement (déplacements d'objet) et les changements temporels (la date de disponibilité d'un bien peut être avancée, par exemple en agriculture grâce à des serres, ou au contraire retardée, grâce à des techniques de réfrigération ou de conservation sous vide).

4.1.2 *La production comme action.* Se dégageant complètement de la description physique des processus productifs, von Mises (1985

[1949], p. 147-149) considère que toute action est une production puisqu'elle consiste toujours à combiner des facteurs – travail, place au sol, etc. – en vue d'obtenir un bien. Dans cette perspective, il n'y a aucune raison de considérer que certains types d'actions ou certains types de facteurs sont par nature productifs alors que d'autres ne le sont pas. Les physiocrates français du XVIII^e siècle se trompaient lorsqu'ils croyaient que seules étaient productives les activités de l'agriculture, de la pêche, de la chasse et de l'extraction minière, et que les activités artisanales et commerciales étaient « stériles ». Les économistes classiques se trompaient lorsqu'ils qualifiaient d'improductives les prestations de services personnels, par opposition à la fabrication des biens matériels. Certains économistes contemporains répètent selon von Mises le même type d'erreur en dénigrant comme improductives les activités de marketing et de publicité. La « condition essentielle » de la production ne réside pas dans ses aspects matériels ou dans ses techniques mais dans la décision d'agir. Pour lui, et il se démarque explicitement ici de Marx, les forces productives ne sont pas matérielles mais spirituelles, intellectuelles et idéologiques, selon ses propres termes.

4.1.3 *La loi des rendements décroissants*. Découverte au début du XIX^e siècle par Ricardo et Malthus dans le cas de l'agriculture, cette loi était considérée par John Stuart Mill comme « la plus importante » de l'économie politique (1987 [1848], p. 177). Von Mises la présente en seconde position, après la loi de l'utilité marginale, parmi les lois de l'agir humain. Il en offre un exposé et une démonstration détaillés (1985 [1949], p. 134), tout comme Rothbard (1962, p. 30) et Reisman (1996, p. 67).

Loi des rendements : pour tout processus de production, si la quantité d'un facteur (facteur *variable*) augmente alors que les quantités des facteurs complémentaires restent constantes (facteurs *fixes*), le rendement du facteur variable (quantité totale produite divisée par quantité du facteur variable) va nécessairement finir par décroître, toutes choses égales par ailleurs. Si ce rendement finit par décroître, cela implique qu'il existe une quantité (ou plusieurs) du facteur variable, appelée *optimum*, pour laquelle le rendement atteint sa valeur maximale. En-deçà de cet optimum, les rendements vont être croissants, au moins dans certaines zones.

Von Mises démontre cette loi par l'absurde. Il suppose qu'il existe un processus de production auquel elle ne s'applique pas. Il en résulterait que l'on pourrait produire des quantités aussi grandes que l'on voudrait, même en réduisant à un niveau infime les quantités de facteurs complémentaires, à condition d'augmenter suffisamment la quantité du facteur variable (on pourrait par exemple nourrir la totalité de la population humaine à partir d'un lopin de terre aussi petit que l'on voudrait, à condition de lui appliquer des quantités suffisantes de travail et de capital). Mais cela signifierait que la productivité physique des facteurs complémentaires serait illimitée. Ces derniers ne seraient alors plus des biens, ce qui est absurde (puisqu'ils sont des biens par hypothèse). La loi ne peut pas être fausse, donc elle est vraie, et s'applique à tous les processus de production sans exception, pas seulement à l'agriculture.

4.1.4 Conséquences des rendements décroissants et forces en sens contraire. Les économistes classiques avaient déjà tiré de la loi des rendements décroissants l'une de ses conséquences les plus importantes, à savoir que face à une quantité donnée de terres agricoles (facteur fixe) l'augmentation de la population et donc du nombre de travailleurs (facteur variable) finit par entraîner, toutes choses égales par ailleurs, une baisse de la productivité moyenne du travail agricole et donc une baisse du niveau de vie moyen. Reisman (1996, p. 70) ajoute que compte tenu de l'épuisement des gisements (la nécessité, toutes choses égales par ailleurs, de creuser de plus en plus profondément pour en extraire les ressources), même si la population reste constante le niveau de vie va tendre à baisser sous l'effet des rendements décroissants de l'activité d'extraction.

De façon plus générale, la loi des rendements décroissants montre quelles sont les limites qui pèsent sur le progrès économique dès lors que le stock d'un facteur de production – que ce soit la terre ou un autre bien d'ordre supérieur – est fixé et ne peut pas être augmenté. Il existe cependant des forces qui contrebalancent l'effet des rendements décroissants :

(1) la découverte et l'application de nouvelles techniques de production plus efficaces (ou la découverte de nouveaux gisements riches en ressources) ;

(2) l'intensification de la division du travail ;

(3) l'accumulation du capital par augmentation de l'épargne.

4.1.5 *Le progrès de la connaissance technique.* Pour Menger (1976 [1871], p. 74), et contrairement à ce que pensait Adam Smith, la principale source du progrès économique n'est pas l'intensification de la division du travail, mais plutôt l'amélioration de la connaissance des phénomènes naturels. Dès lors que les relations de causalité entre ces phénomènes sont de mieux en mieux connues, il devient possible de mettre en œuvre des processus de production de plus en plus efficaces en préparant au préalable les matériaux et les outils ou machines requis. Le travail devient ainsi capable de produire des biens en plus grand nombre et de meilleure qualité. L'économie humaine passe de la simple récupération des biens de consommation dispersés dans l'environnement à des processus de production contrôlés capables d'atteindre les fins visées grâce à l'appui, qui peut être considérable, des forces de la nature. Menger ajoute que la mise en œuvre du progrès technique conduit à utiliser des biens d'ordre de plus en plus élevé, et donc à ajouter des étapes de production en amont.

Ces deux éléments – progrès technique et allongement de la structure de production – doivent cependant être distingués : d'une part, comme le reconnaîtra Böhm-Bawerk, certains progrès techniques peuvent raccourcir la structure de production ; et d'autre part, les acteurs économiques peuvent décider, pour économiser du temps, de réduire le nombre d'étapes, même s'ils connaissent des techniques plus productives mais qui exigeraient un délai de production plus long. En outre, la multiplication des étapes de production – qui servent par exemple à fabriquer au préalable des outils, des machines ou des bâtiments – est bien l'une des formes que peut prendre la division du travail, puisqu'il s'agit de la division *verticale* du travail (Hayek 1941, p. 71-73).

4.1.6 *Division du travail et société.* L'importance de la division du travail a bien sûr toujours été reconnue par les économistes de l'école autrichienne, mais von Mises (1981 [1922], p. 258-276) est le premier à lui avoir consacré des développements substantiels. Pour lui, la division du travail constitue un phénomène d'une immense importance puisqu'elle est à l'origine de la société humaine. Les individus se rendent compte qu'en se spécialisant dans les activités où ils sont les plus efficaces, ils peuvent augmenter la production par rapport à la situation où chacun essaierait de se procu-

rer par lui-même la totalité de ses moyens de subsistance. Ils entrent donc, en suivant leur intérêt bien compris, dans des relations de coopération fondées sur la division du travail. S'ils n'avaient pas conscience que la division du travail est productive, ou si par pure hypothèse elle ne permettait pas d'accroître la production, alors – pour von Mises – la société humaine n'existerait tout simplement pas. La productivité de la spécialisation des activités permet donc d'expliquer l'existence de la société humaine sans recourir à une explication circulaire comme par exemple un soi-disant instinct à se rassembler (les gens forment des sociétés parce qu'ils ont un instinct de rassemblement ; mais comment sait-on qu'ils ont cet instinct de rassemblement ? Parce qu'ils constituent des sociétés ! Ce raisonnement est circulaire et n'a rien de scientifique). La société n'est rien de plus qu'un moyen de coopération qui permet aux individus d'atteindre une meilleure satisfaction de leurs besoins : elle peut se développer ou au contraire régresser selon que la division du travail s'intensifie ou au contraire se délite. La décision d'un acteur de se spécialiser dans telle ou telle branche de production dépend des ressources naturelles et des biens du capital dont il peut disposer, de ses compétences personnelles, et de ses préférences entre les activités (Rothbard 1962, p. 80).

4.1.7 *Division du travail et production*. Von Mises (1985 [1949], p. 173) évoque trois conséquences de la division du travail, qui sont la répartition géographique des activités productives, la différenciation des facultés productives des individus en fonction de leur spécialisation, et l'utilisation des machines. Reisman (1996, p. 123) propose une analyse plus détaillée, en six points, des avantages productifs apportés par la division du travail :

(1) la multiplication des connaissances : les connaissances développées dans chacune des activités spécialisées s'ajoutent pour constituer une gigantesque somme de savoirs utilisés pour produire, alors que dans une société ayant un faible degré de division du travail chaque famille sait à peu près les mêmes choses que toutes les autres ;

(2) le bénéfice des génies : les grands scientifiques et les grands inventeurs peuvent consacrer la totalité de leur activité à améliorer les connaissances et les techniques ;

(3) la mise en œuvre des avantages individuels : chacun peut se

consacrer aux activités les mieux adaptées à ses capacités physiques ou intellectuelles, et la production peut ainsi profiter de ces avantages absolus ou relatifs ;

(4) la spécialisation géographique : les conditions locales (terre, climat, sous-sol, etc.) peuvent être particulièrement propices à certains types d'activités d'agriculture ou d'extraction ; la division du travail et l'échange des surplus permettent aux différentes régions de profiter des avantages de chacune des autres, les activités agricoles ou minières d'un territoire utilisant par exemple les ressources énergétiques tirées d'un autre territoire, et réciproquement ;

(5) l'économie d'apprentissage et de mouvement ; la répétition des mêmes gestes, l'application des mêmes techniques, rendent le producteur plus efficace et lui permettent de mieux rentabiliser son apprentissage (en minimisant le rapport entre le temps passé à apprendre ses techniques et le temps passé à les utiliser) ; la division du travail en usine permet de limiter au maximum les déplacements et les mouvements inutiles ;

(6) l'utilisation des machines ; la division du travail permet, non seulement d'utiliser des machines, mais aussi de les concevoir (grâce à des inventeurs spécialisés), de se procurer à travers le monde les matériaux nécessaires à leur fabrication, de les assembler efficacement en usine, et de les rentabiliser grâce à la production de masse.

Reisman (1996) montre aussi toute l'importance, pour l'instauration, l'intensification et la rationalisation de la division du travail, des institutions majeures du capitalisme : propriété privée des moyens de production, épargne et accumulation du capital, échange et monnaie, concurrence et inégalités de revenus, coordination par les prix de marché.

4.1.8 *La loi du « détour » de production.* L'élément central de la théorie de la production de Böhm-Bawerk est la loi du « détour » de production, déjà esquissée par Menger, qui analyse le phénomène de *l'accumulation du capital*. Des villageois veulent se procurer de l'eau à la rivière voisine. Ils peuvent aller la recueillir dans le creux de leurs mains : c'est la méthode la plus directe, mais qui va les conduire à effectuer de très nombreux allers-retours car ils ne pourront se procurer que de très petites quantités d'eau à

chaque fois. Ils peuvent aussi consacrer du temps à fabriquer des récipients, ce qui leur permettra de récupérer de plus grandes quantités d'eau à chaque voyage. La durée du processus de production a augmenté, puisqu'il faut au préalable prendre le temps de fabriquer les seaux, mais la productivité physique de leur travail s'est accrue. Enfin, ils peuvent accumuler encore davantage de capital en construisant un système d'adduction d'eau, ce qui leur prendra plus de temps encore avant de récupérer l'eau par les tuyaux, mais en mettra à leur disposition de beaucoup plus grandes quantités.

Böhm-Bawerk généralise cet exemple en affirmant que les méthodes plus détournées (*more roundabout*) de production sont « plus fructueuses » que les méthodes plus directes pour produire les biens de consommation (1959 [1889], p. 12). Une méthode de production est plus détournée, ou plus indirecte, si l'utilisation du facteur travail et des ressources naturelles est plus étalée dans le temps, de sorte que la période totale de production s'allonge : une plus longue durée de temps s'écoule – et de plus nombreuses étapes de production se déroulent – entre la mise en œuvre du processus et l'apparition des biens de consommation finaux. Cet allongement permet, soit de produire des biens de consommation en plus grande quantité, soit de les produire de meilleure qualité, soit de produire des types de biens qui ne peuvent tout simplement pas être produits en un plus court laps de temps.

La loi du détour de production caractérise donc l'aspect temporel des processus productifs. Elle signifie que la *même* quantité de travail, déployée sur une plus longue période de temps, permet de produire une plus grande quantité de biens de consommation par unité de temps, ou de produire des biens qui ne pourraient tout simplement pas être produits en un plus court laps de temps. Le détour permet ainsi d'accroître les quantités ou d'améliorer les qualités des produits, sans pour autant disposer d'une plus grande quantité de travail. Von Mises (1985 [1949], p. 506) critique néanmoins l'emploi du terme « détour », car si les gens adoptent ces processus indirects de production, c'est parce qu'ils constituent pour eux le plus bref chemin qui les conduit à leurs fins (leurs fins étant d'augmenter la quantité ou la qualité produite).

4.1.9 *La fonction de production intertemporelle.* Böhm-Bawerk applique au détour de production la loi des rendements décrois-

sants : l'accroissement de la productivité physique du processus sera de plus en plus faible au fur et à mesure que la période de production sera prolongée (bien qu'il ne précise pas ce point, nous pouvons considérer que ce sont les quantités de travail qui constituent le facteur fixe – ce ne peuvent être les ressources naturelles puisqu'elles sont récupérées en quantités de plus en plus grandes, comme expliqué au § 4.1.10). Wicksell (1954 [1893], p. 122) représente cette fonction de production intertemporelle par la courbe de la figure 4.1 : la quantité q/d de biens de consommation produits par unité de temps par un processus augmente avec la durée d de ce processus (loi du « détour »), mais de plus en plus lentement (loi des rendements).

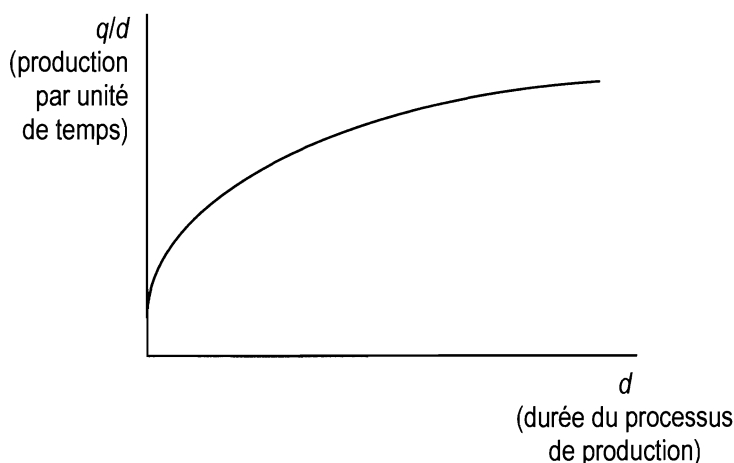


Figure 4.1. La fonction intertemporelle de production
(d'après Wicksell 1954 [1893], p. 122)

4.1.10 *L'explication de la productivité du « détour » de production.* La loi du « détour » n'affirme évidemment pas que tout allongement de la structure de production est nécessairement productif, mais seulement que *certain*s allongements, *judicieusement choisis*, sont productifs. Pour Böhm-Bawerk cette loi est « purement technique » et s'explique par le fait qu'un détour permet d'enrôler dans le processus de production des forces de la nature supplémentaires. L'utilisation d'un seau, par exemple, permet de

bénéficier des forces naturelles qui rendent ce récipient étanche et rigide pour transporter de plus grandes quantités d'eau. Hayek (1941, p. 60) explique lui aussi qu'un délai peut être indispensable pour bénéficier de l'apport de certaines forces et ressources naturelles. Plus le détour est important, plus ces forces seront puissantes et plus ces ressources seront nombreuses. Au fur et à mesure que le processus de production s'allonge, certains nouveaux types de ressources sont utilisés et deviennent des facteurs de production. Il précise que la productivité du « détour » de production vaut pour un niveau donné de connaissance technique (Hayek 1936). Si un progrès technique est découvert et mis en œuvre, alors la fonction de production de la figure 4.1 se déplace vers le haut : la production par unité de temps q/d s'élève pour les différentes durées d possibles du processus.

4.1.11 *L'épargne-investissement*. Les méthodes de production les plus directes (la cueillette par exemple) sont très peu productives. Les individus peuvent allonger le processus de production pour bénéficier d'une quantité ou d'une qualité accrue de produits. Mais allonger le détour implique de réallouer une partie de la quantité de travail vers les étapes hautes de la structure de production. Deux cas sont alors possibles. Soit la quantité de travail n'augmente pas et la production des biens de consommation va se trouver momentanément réduite puisqu'une partie du travail utilisé dans les étapes basses n'est plus disponible. Soit la quantité de travail augmente parce que les agents économiques renoncent à une partie de leur temps de loisir pour le consacrer à la production. Dans les deux cas, les individus consentent un sacrifice qui constitue leur épargne-investissement, et qui est destiné à accroître leur consommation future. Il leur faut donc décider si l'accroissement de la consommation future est suffisant pour justifier le sacrifice subi dans l'immédiat. Ce choix va dépendre, de façon très générale, de l'intensité de la préférence pour le présent, des stocks de biens du capital et de biens de consommation, et de l'efficacité relative des différentes méthodes de production connues.

4.1.12 *Mesurer la période de production ?* Böhm-Bawerk propose de mesurer, non pas la période de production totale d'un processus, mais plutôt sa période *moyenne* (1959 [1889], p. 86-87). En

effet, chaque produit fabriqué aujourd'hui est le résultat de l'enchaînement de toute une série de processus de production qui se sont succédé depuis des siècles et des siècles. La période de production totale ne présente donc pas un grand intérêt. Il veut tenir compte du fait que les quantités de travail effectuées très loin dans le passé n'ont qu'une influence négligeable sur le produit actuel, car elles ne représentent qu'une fraction infime de la totalité du travail accompli. Pour cela, il calcule la période moyenne de production qui donne une indication plus significative sur la durée du processus. Si un processus a utilisé une quantité q_3 de travail trois ans auparavant, une quantité q_2 de travail deux ans auparavant, q_1 un an auparavant et q_0 au dernier moment, alors d'après sa formule la période de production moyenne est égale à :

$$\frac{3q_3 + 2q_2 + 1q_1 + 0q_0}{q_3 + q_2 + q_1 + q_0}$$

Cette notion de période moyenne de production a été critiquée par von Mises (1985 [1949], p. 504) qui lui reproche sa nature rétroactive, déconnectée de l'action humaine qui est toujours tournée vers le futur. Hayek (1941, p. 141) souligne plutôt le caractère simpliste de cette mesure : dans les cas plus complexes et réalistes, il n'existe pas de moyen d'agréger les durées des différents processus qui composent le système économique en une mesure unique de la période de production du système (mais il est en revanche possible de savoir si la mise en œuvre de tel ou tel processus de production conduit à un allongement ou un raccourcissement de la structure).

4.2 La macroéconomie de la structure de production

4.2.1 *La structure de production hayékienne*. Böhm-Bawerk (1959 [1889], p. 106) offre une première illustration graphique de la structure de production globale, avec des cercles concentriques représentant les différentes « classes de maturité » des biens du capital. Il est inutile de s'attarder sur son schéma, car celui proposé par Hayek (1975 [1931]) et inspiré par Jevons (1965 [1871], p. 230) est beaucoup plus intéressant. Mais avant de présenter le

triangle hayékien, il est utile de revenir sur le concept de structure de production. La production des biens de consommation de l'année courante est le résultat d'un processus qui a commencé plusieurs années auparavant. La figure 4.2 illustre un système économique très simple dont le processus de production global s'effectue en *quatre* étapes successives qui durent *une année* chacune et qui sont numérotées de la plus proche à la plus éloignée de la consommation finale. Les facteurs originaires travail et terre sont peu à peu transformés, au cours de ces quatre étapes, en biens de consommation finaux.

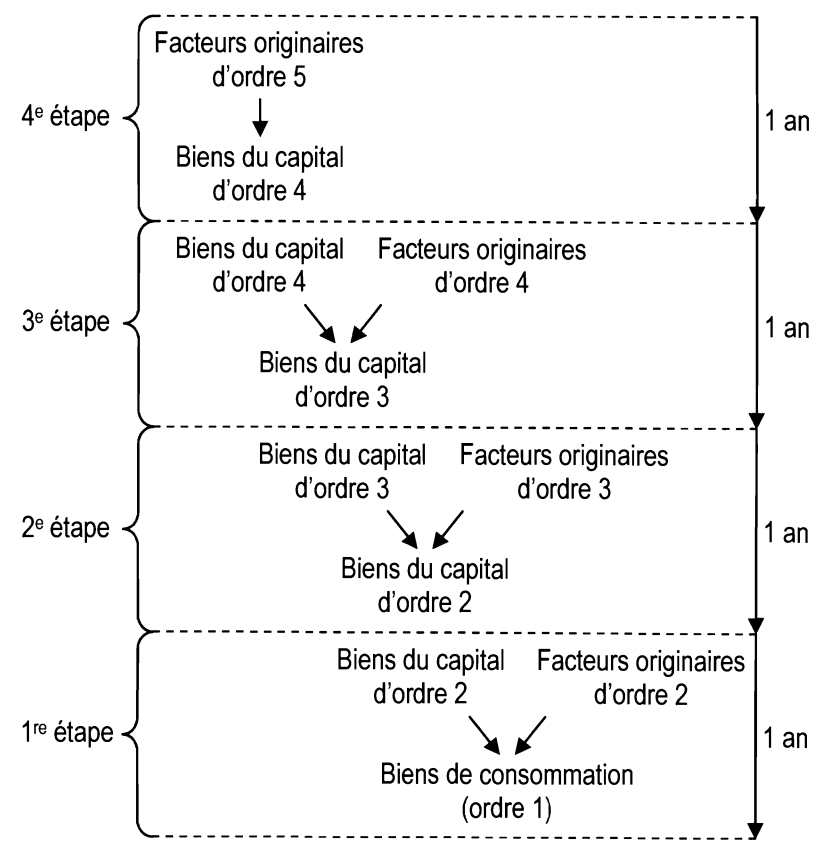


Figure 4.2. Une première illustration de la structure de production hayékienne

4.2.2 *Le triangle hayékien*. Le triangle de Hayek (1975 [1931]) va être présenté à l'aide d'une illustration plus détaillée et plus complète, qui a été réalisée ultérieurement par Rothbard (1962, p. 314). La structure est supposée composée de *six étapes* qui durent chacune *un an* (figure 4.3). Les biens du capital (facteurs de production produits) sont représentés par des rectangles grisés et les facteurs originaires avec lesquels ils sont combinés par des rectangles blancs. Rothbard fait aussi apparaître les paiements d'intérêt, et bien qu'il ne les intègre pas à son schéma il est très facile de les ajouter ainsi que cela a été fait à la figure 4.3 sur la partie gauche de la structure. Le taux d'intérêt annuel est supposé égal à 5 %. La structure est décrite ci-dessous dans l'ordre inverse du déroulement de la production, c'est-à-dire en commençant par la toute dernière étape, celle qui produit les biens de consommation :

(1) étape 1 : la production annuelle de biens de consommation est supposée valoir 100 unités monétaires ; comme le taux d'intérêt vaut 5 %, l'investissement I_1 réalisé au début de l'étape 1 (un an auparavant) est calculé grâce à l'équation $I_1(1 + 0,05) = 100$, soit $I_1 = 95$ (arrondi) ; le revenu d'intérêt à cette étape est donc $100 - 95 = 5$ unités monétaires ; on suppose que le prix agrégé des facteurs originaires utilisés à cette étape est 15 ;

(2) étape 2 : cette étape produit les biens du capital qui vont être utilisés lors de l'étape 1, pour un prix de 80 ($= I_1 - 15 = 95 - 15$) ; l'investissement I_2 du début de l'étape 2 est donné par l'équation $I_2(1 + 0,05) = 80$, soit $I_2 = 76$ (arrondi) ; le revenu d'intérêt est donc $80 - 76 = 4$; le prix des facteurs originaires utilisés à cette étape est *supposé* égal à 16 ;

(3) connaissant le taux d'intérêt et le nombre d'étapes, et en fixant arbitrairement un prix agrégé pour les facteurs originaires à chaque étape (15 pour la 1^{re} étape, 16 pour la 2^e, 12 pour la 3^e, 13 pour la 4^e et 7,5 pour la 5^e étape), des calculs similaires permettent de remonter d'étape en étape jusqu'à décrire la totalité de la structure.

Le processus de production complet dure six ans, et les consommateurs effectuent chaque année une dépense agrégée de 100 unités de monnaie qui représente la valeur ajoutée, ici répartie entre :

– le revenu d'intérêt annuel $16,5 = 5 + 4 + 3 + 2 + 1,5 + 1$ (partie gauche de la structure),

– et le revenu annuel des facteurs originaires travail et terre $83,5 = 15 + 16 + 12 + 13 + 7,5 + 20$ (partie droite de la structure).

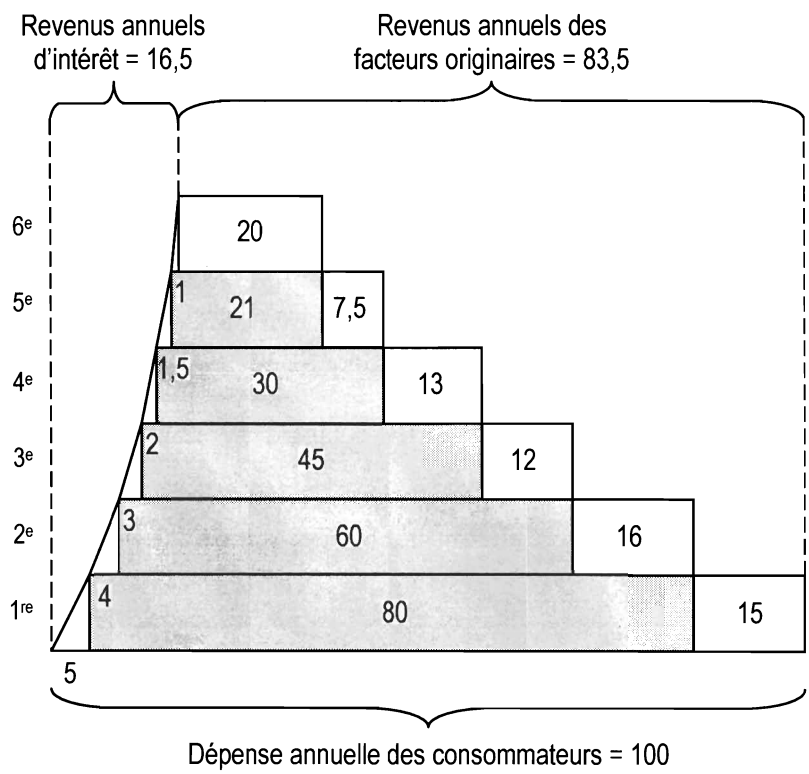


Figure 4.3. Le schéma de Rothbard
(adapté de Rothbard 1962, p. 314)

4.2.3 Structure synchronique et structure diachronique. Dans cet exemple, le processus de production dure six années pendant lesquelles les facteurs originaires travail et terre sont combinés pour produire des biens intermédiaires qui vont finalement être transformés en biens de consommation. Cela signifie-t-il qu'il faut attendre six ans avant de pouvoir consommer ? Non, car dans le système économique les processus de production sont « synchronisés », selon l'expression de Clark (1899). Six processus identiques se déroulent en parallèle avec une année d'écart, en sorte que *chaque année* l'un d'entre eux arrive à son terme et permet de bé-

néficier en continu du flux de biens de consommation finale.

Lors d'une année courante, les facteurs originaires disponibles sont répartis entre les différents processus. Dans l'exemple de Rothbard, des facteurs originaires pour une valeur 15 sont appliqués à l'étape 1, d'autres facteurs pour une valeur 16 à l'étape 2, d'autres encore pour une valeur 12 à l'étape 3, et ainsi de suite. Le schéma de la figure 4.3 peut donc être interprété de deux façons différentes :

(1) dans une perspective *diachronique*, comme la description du déroulement au cours du temps d'un processus de production de son début à son terme,

(2) dans une perspective *synchronique*, comme une vue « en coupe » ou une « photo » du système économique lors d'une année donnée, où l'on distingue tous les processus qui se déroulent simultanément, de celui qui parvient à son terme jusqu'à celui qui vient de débiter, en passant par les processus qui se trouvent à toutes les étapes intermédiaires.

La perspective synchronique sur le système économique est la plus intéressante car elle offre une représentation du système économique lors d'une année donnée.

4.2.4 Épargne et maintien de la structure. Cette illustration de la structure de production permet de comprendre la critique adressée par les économistes autrichiens (par exemple par Hayek 1936) à l'école américaine de Clark et Knight. Pour ces derniers, la production et la consommation sont, non seulement synchronisées, mais *simultanées*. Ils pensent en outre que le capital est un fonds *perpétuel* : dès lors que du capital a été formé, sa reproduction ne constitue plus un problème économique mais seulement un « détail technologique » (Knight 1934, p. 259). La figure 4.3 montre que ces deux assertions sont erronées.

(1) *La production et la consommation ne sont pas simultanées.* Le flux annuel de biens de consommation (de prix agrégé = 100) est le résultat d'un processus qui a débuté six années auparavant et qui arrive finalement à maturité ; les étapes de transformation intermédiaires (étapes 2, 3, 4, 5 et 6) ne permettent pas la consommation puisque pour chacune d'entre elles il faudra attendre que le processus concerné parvienne à son terme, dans un an pour l'étape 2, dans deux ans pour l'étape 3, et ainsi de suite. La synchronisa-

tion des six processus qui se déroulent en parallèle ne signifie pas du tout que la production et la consommation sont simultanées. Non seulement la synchronisation de la production ne fait pas disparaître la dimension temporelle de la production, mais c'est justement parce que la production prend du temps qu'elle doit être synchronisée pour engendrer un flux ininterrompu de biens de consommation.

(2) *Le capital n'est pas un fonds perpétuel.* La structure ne peut se maintenir que si les capitalistes décident, année après année, de renouveler leur épargne en sorte que les différentes étapes soient intégralement financées (Hayek 1975 [1931], p. 107). Chaque étape constitue ce que Rothbard (1962) appelle un « marché du temps », sur lequel les capitalistes investissent une certaine somme (par exemple $60 + 16$ à l'étape 2) et attendent une année avant de récupérer leur revenu brut (80) et leur revenu net ($80 - [60 + 16] = 4$) correspondant à l'intérêt sur le capital investi. Pour que la structure se reproduise, il faut que *tous* ces marchés du temps (il y en a 6 sur la figure 4.3) récoltent *chaque année* l'épargne-investissement nécessaire, pour un montant total de $20 + (21 + 7,5) + (30 + 13) + (45 + 12) + (60 + 16) + (80 + 15) = 319,5$ unités de monnaie. La reproduction de la structure dépend des décisions des acteurs, en l'occurrence de la décision des capitalistes de ne pas réduire leur épargne totale, et n'a donc rien d'automatique.

4.2.5 Le produit total : critique du PIB. L'agrégat statistique couramment utilisé pour mesurer la production totale d'un pays est le produit intérieur brut (PIB). Il est censé comptabiliser la richesse produite en une année. Or, à partir de sa représentation de la structure de production, Rothbard (1962, p. 343) conclut que le PIB est une mesure de la richesse *nette*, et non pas *brute*.

Le PIB (nominal) annuel du système économique représenté à la figure 4.3 est de 100 unités de monnaie. En effet, d'après la formule standard (Mankiw 1998, p. 614), $PIB = C + I + G + XN$, mais comme il n'y a pas ici d'investissement en biens durables ($I = 0$), comme les dépenses de l'État sont supposées nulles ou intégrées à la consommation C (voir § 8.2.1), et comme il n'y a pas de commerce extérieur ($XN = 0$), le PIB se réduit dans ce cas à la consommation annuelle : $PIB = C$. Ainsi, la statistique du PIB ne tiendrait ici compte que de la production de biens de consommation, et

laisserait entièrement de côté la production de tous les biens intermédiaires issus des étapes 2 à 6. Mesuré en unités de monnaie, le produit total se compose du produit de l'étape 6 (21 unités monétaires), plus du produit de l'étape 5 (30), plus du produit de l'étape 4 (45), etc., plus du produit de l'étape 1 (100). Le produit nominal total s'élève à $21 + 30 + 45 + 60 + 80 + 100 = 336$, et le PIB standard ne représente donc que 30 % du produit total puisque $100 \div 336 \approx 0,30$. Skousen (1990, p. 191) propose de nommer OIB (output intérieur brut, *gross domestic output* en anglais) ce produit total qui excède de très loin le PIB, mais qui reflète de façon beaucoup plus fidèle que ce dernier l'activité économique globale.

L'argument standard en faveur du PIB consiste à dire que si du blé a été utilisé pour produire de la farine, et que cette farine a été utilisée pour produire du pain, alors le prix du blé est déjà inclus dans celui de la farine, et le prix de la farine déjà inclus dans celui du pain : ajouter les trois prix reviendrait à comptabiliser deux fois la farine et trois fois le blé. La réponse de Reisman (1996, p. 674) est que le système économique a produit le blé, le système économique a produit la farine, et le système économique a produit le pain ; si l'on veut déterminer le produit *total*, on ne doit pas faire comme si le système économique n'avait produit *que* le pain, et il faut donc bien ajouter les trois prix.

4.2.6 Une structure de production avec biens durables. Les économistes autrichiens utilisent principalement le type de structure qui vient d'être représenté. En effet, ils préfèrent insister sur le rôle du capital circulant – les biens intermédiaires – que sur celui du capital fixe (Hayek 1941, p. 47). Dans ces structures, les ressources naturelles non produites (« terre ») sont peu à peu transformées par le travail en biens de consommation, et les biens du capital constituent les étapes intermédiaires de ce processus de transformation. La durée de vie des biens du capital n'excède pas la durée d'une étape (par exemple, une année), et il n'y a donc pas de biens durables. Hayek (1941, p. 131) illustre le cas inverse, qui est celui d'une structure exclusivement composée de biens durables, c'est-à-dire de biens qui fournissent des services de consommation pendant plusieurs périodes successives. Dans la figure 4.4, chaque année les facteurs originaires travail et terre sont consacrés à la production d'un bien durable qui va offrir des services

de consommation d'une valeur de 20 unités monétaires pendant cinq ans, puis deviendra entièrement inutilisable. La production est là aussi synchronisée : le flux annuel de biens de consommation reste constant grâce au fait que chaque année cinq biens durables fournissent simultanément leurs services, pour une valeur annuelle totale de $5 \times 20 = 100$ unités de monnaie. Ils sont produits avec

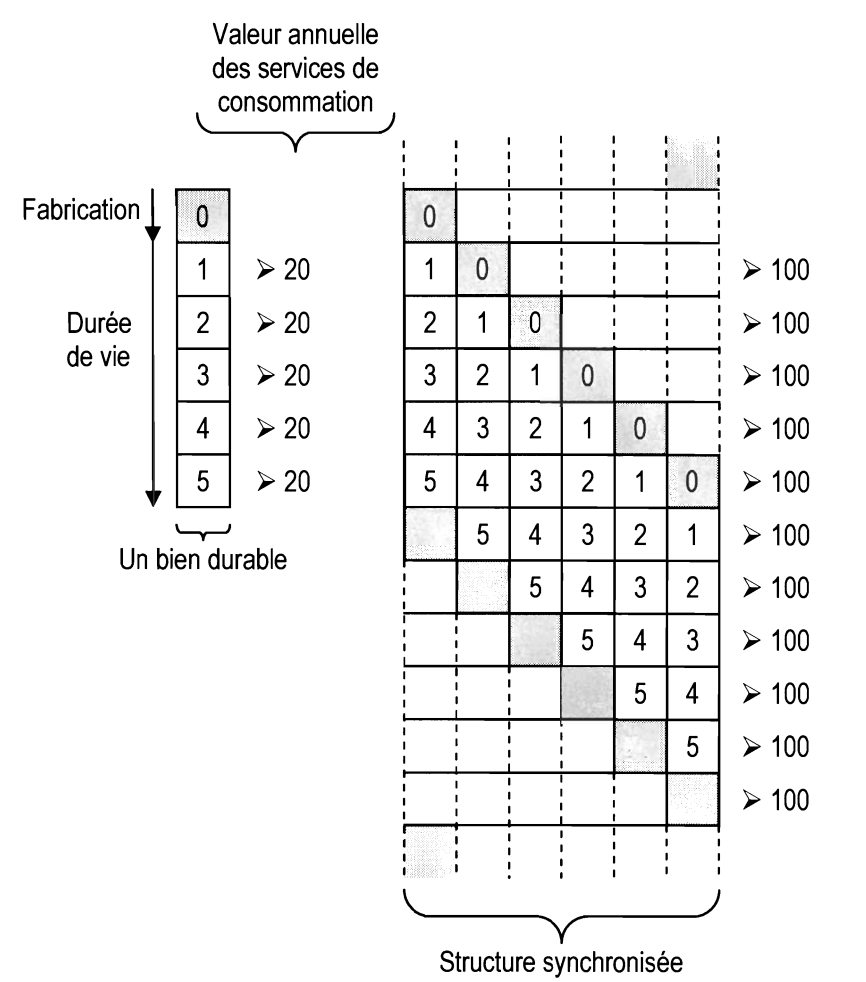


Figure 4.4. Structure de production avec biens durables (adapté de Hayek 1941, p. 131)
une année de décalage, et dès que l'un de ces biens durables atteint

sa limite d'âge et disparaît, une année de travail est consacrée à son remplacement.

Les systèmes économiques réels sont des combinaisons très complexes, avec une multitude d'étapes enchevêtrées, des deux grands types de structures, la structure avec biens intermédiaires non durables et celle avec biens durables.

4.2.7 *L'accumulation du capital*. Le triangle hayékien permet de décrire le processus d'accumulation du capital et d'illustrer du même coup la théorie du détour de production (Hayek 1975 [1931], p. 111). Supposons que les acteurs du système économique décident de restreindre leur consommation au profit de leur épargne-investissement (suite à une baisse de la préférence pour le présent). La structure va se rétrécir horizontalement et s'allonger verticalement. Plus précisément :

- la dépense se réduit dans les étapes basses (rétrécissement) puisque la dépense de consommation a baissé,
- la dépense augmente dans les étapes hautes (élargissement) puisque la dépense d'investissement a augmenté,
- de nouvelles étapes apparaissent au sommet de la structure (allongement).

L'accumulation du capital est donc très facile à visualiser sur ce type de schéma, puisqu'elle consiste en une déformation du triangle représentant la structure (voir figure 4.5). L'allongement de la structure va permettre, une fois achevée la réallocation des facteurs de production vers les étapes hautes, une *plus grande production finale* puisque le détour de production se sera accru (voir § 4.1.10). L'accumulation du capital s'accompagne donc bien d'une augmentation de la quantité de biens de consommation produits par période.

4.2.8 *Critique du « paradoxe de l'épargne »*. Le cas qui vient d'être examiné, celui d'une baisse de la dépense de consommation au profit de la dépense d'épargne-investissement, est souvent considéré, aujourd'hui encore, comme l'une des causes des crises économiques. En effet, d'après la *théorie de la sous-consommation*, les dépressions et le chômage de masse proviennent d'une baisse de la consommation. L'explication avancée est que la consommation devient insuffisante pour rentabiliser la production. Plus préci-

sément : si les ménages réduisent leurs dépenses de consommation, alors les prix des biens de consommation vont avoir tendance

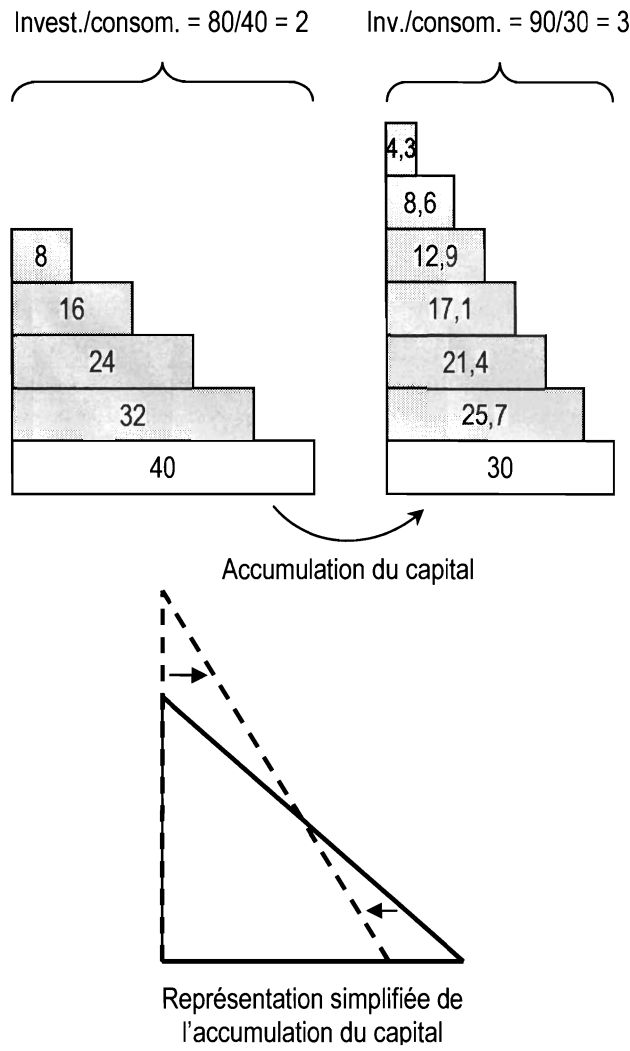


Figure 4.5. Illustration de l'accumulation du capital (d'après Hayek 1975 [1931, p. 111])

à baisser ; mais si, simultanément, les ménages accroissent leur

épargne-investissement, alors les prix des facteurs de production vont avoir tendance à augmenter ; du coup, les prix de vente des biens finaux vont s'avérer insuffisants pour couvrir leurs coûts, les entreprises vont subir des pertes, réduire leur production et licencier, d'où un chômage de masse et une récession. Il existerait donc un « paradoxe de l'épargne » : un individu qui épargne s'enrichit, mais si de nombreux individus tentent simultanément de s'enrichir par un surcroît d'épargne, alors ils ne feront qu'appauvrir la société en la plongeant dans une crise de sous-consommation.

Pour Hayek (1929, p. 224), ce raisonnement est erroné parce qu'il néglige l'existence de la structure de production, c'est-à-dire le fait que la production s'opère par étapes successives. Il est incorrect de raisonner en comparant l'évolution des prix moyens des biens de consommation et des facteurs de production, parce que ces prix varient différemment selon les étapes. Cette critique hayékienne du « paradoxe de l'épargne » repose sur une analyse du fonctionnement du système des prix pendant la phase d'accumulation du capital.

(1) Dans les étapes *basses*, proches de la consommation, la baisse de la demande de biens de consommation fait baisser leur prix, et cette baisse est imputée aux facteurs de production des étapes juste au-dessus, qui voient leurs prix diminuer aussi. Il en résulte une baisse de rentabilité dans les étapes basses, d'autant plus forte que l'étape est proche de la consommation finale (puisque l'effet de la baisse de la dépense de consommation est d'autant plus intense que l'étape est basse).

(2) L'augmentation de l'épargne-investissement « alimente » les étapes *hautes* où elle fait apparaître des profits. Les capitaux sont donc réalloués vers les étapes hautes, où ils contribuent à augmenter la demande et donc le prix de vente des facteurs de production, faisant apparaître des profits d'autant plus élevés que l'étape est proche du sommet de la structure (Hayek 1975 [1931], p. 137). Les facteurs de production convertibles tendent donc à être réalloués du bas vers le haut de la structure, ce qui augmente le « détour » de production et rend à terme le système économique plus productif.

En résumé, le changement initial des préférences intertemporelles – baisse de la préférence pour le présent conduisant à une réduction de la consommation et un accroissement de l'épargne-

investissement – diminue la rentabilité dans les étapes basses et l'augmente dans les étapes hautes. La réallocation des facteurs du bas vers le haut de la structure restaure l'équilibre en abaissant leur prix dans les étapes hautes et en le ré-augmentant dans les étapes basses : les taux de rentabilité de la structure tendent ainsi à nouveau vers l'égalité. Ce fonctionnement du système des prix est le même que celui présenté ci-dessus (§ 2.2.11 et § 2.2.13) : un choc fait diverger les taux de rentabilité entre différents secteurs du système économique, les entrepreneurs réallouent les facteurs des secteurs à faible rentabilité vers ceux à forte rentabilité, et cette réorganisation de la production ramène les taux de rentabilité à égalité. Finalement, le « paradoxe » de l'épargne s'évanouit : la baisse de la consommation et le surcroît d'épargne-investissement ne déclenchent pas de crise de surproduction globale, ils n'appauvrissent pas la société mais au contraire l'enrichissent en allongeant le détour de production.

4.2.9 La réévaluation de l'épargne-investissement. La théorie autrichienne de la structure de production réévalue grandement le rôle de l'épargne-investissement par rapport à la théorie macroéconomique keynésienne.

(1) D'un point de vue statique : la mesure standard du PIB lors d'une année donnée sous-estime de beaucoup l'importance et l'impact de l'investissement dans le fonctionnement du système économique. D'après une estimation de Skousen (1991, p. 45) sur les États-Unis pour l'année 1982, le produit total (output national brut = PNB + consommations intermédiaires) s'élève au double du PNB (produit national brut), et la consommation ne représente que 34 % de l'activité économique totale mesurée par l'ONB, alors qu'elle représente 65 % du PNB. En d'autres termes, par rapport aux évaluations des indicateurs standard, l'importance de l'investissement est double, et celle de la consommation moitié moindre.

(2) D'un point de vue dynamique : cette résolution du paradoxe de l'épargne permet de comprendre la critique que les économistes autrichiens adressent à l'idée keynésienne selon laquelle un surcroît d'épargne-investissement peut constituer une « fuite » néfaste pour le système économique. Pour eux, un accroissement de l'épargne n'a aucun effet paradoxal, et permet de produire davan-

tage de richesses et d'améliorer les niveaux de vie. Skousen (1990, p. 244) explique ainsi la forte corrélation observée entre taux d'épargne et croissance économique, et il y voit l'une des causes de la croissance élevée qu'ont connue, dans la seconde moitié du ^{xx}e siècle, les « dragons » du Sud-est asiatique (1991, p. 199).

Bien que Menger ait élaboré des théories du capital et de l'intérêt, c'est le grand traité de Böhm-Bawerk (1959 [1884], 1959 [1889]) qui constitue le point de départ des discussions approfondies qui ont eu lieu sur ces deux thèmes au sein de l'école autrichienne. Très vite, certaines de ses conceptions ont été contestées par Fetter (1900, 1902). Une fracture majeure s'est alors dessinée entre Böhm-Bawerk, Wicksell et Hayek d'une part, qui défendent une conception réelle du capital et une théorie productiviste de l'intérêt, et Fetter, von Mises et Rothbard d'autre part, qui adoptent une définition nominale du capital et une théorie subjectiviste de l'intérêt fondée sur la préférence pour le présent. Les divergences sont très profondes en ce qui concerne l'explication de l'intérêt, puisque cinq théories différentes – au moins – sont défendues par des économistes de l'école autrichienne : une théorie de l'usage, une théorie productiviste, une théorie de la préférence pour le présent, une théorie de l'échange, et la théorie du marché des fonds prêtables.

5.1 Facteurs de production et capital

5.1.1 *De la théorie classique de la distribution au principe d'imputation.* Ricardo (1951 [1817]) a développé la théorie classique de la distribution en distinguant trois grands types de revenus – rente, salaire et profit – respectivement attribués aux trois types de facteurs de production – terre, travail et capital. Chacun de ces types de revenus est expliqué par une théorie spécifique : la théorie de la rente pour le revenu de la terre, la théorie du minimum de subsistance pour le revenu du travail, et la théorie du revenu résiduel pour le « profit » (ce qu'il appelle « profit », et qui est l'intérêt sur le capital investi, est selon lui ce qui reste au capitaliste lorsque les dépenses de production ont été effectuées).

Avec sa théorie de l'imputation (voir § 1.3.1), Menger remet complètement en cause le paradigme ricardien. La terre et le travail, dans la mesure où ils sont des biens d'ordre supérieur, ont une

valeur subjective et un prix qui relèvent du *même type d'explication* : leur valeur pour un acteur correspond à l'importance des besoins qu'ils lui permettent de satisfaire, et leur prix se fixe par confrontation d'une offre et d'une demande déterminées par ces évaluations subjectives. Alors que Ricardo et ses successeurs proposent d'un côté une théorie générale des prix fondée sur le coût en quantité de travail (le prix des biens reproductibles à volonté est proportionnel à la quantité de travail qui a été nécessaire pour les produire), et d'un autre côté des théories *ad hoc* pour expliquer les prix des facteurs non produits (terre et travail), Menger réunit *tous les types de biens* dans un même cadre explicatif de la valeur et du prix. Quant au « profit » (revenu d'intérêt), Menger l'explique par l'usage du capital, mais cette question concerne la théorie de l'intérêt, et sera abordée plus bas.

Von Mises (1981 [1922], p. 131) remarque que le terme « distribution », bien qu'il ait été consacré par l'usage, n'est qu'une métaphore qui ne correspond pas à la réalité du fonctionnement d'une économie de marché. Dans ce système, les biens ne sont pas d'abord produits puis distribués : le processus de production est inextricablement lié aux échanges marchands qui déterminent les revenus. La question de la distribution, au vrai sens du terme, ne se pose que dans un régime interventionniste qui opère un prélèvement et une redistribution par l'impôt. Pour l'économie de marché, il vaut donc mieux parler d'une théorie des prix des biens d'ordre supérieur.

5.1.2 La conception « réelle » du capital. Les économistes autrichiens ont adopté deux conceptions très différentes du capital, l'une l'identifiant à des *biens matériels* (conception « réelle ») et l'autre à une *valeur monétaire* (conception « nominale »). Böhm-Bawerk adopte la première conception et définit le capital comme l'ensemble des facteurs de production produits (1959 [1889], p. 14). Le capital est ainsi distingué, d'une part de la « terre » (non produite : place au sol, ressources naturelles), d'autre part du travail (services immatériels), et enfin des biens de consommation (qui constituent la fin de la production, et non pas un moyen de produire). En ce sens, toute société humaine qui a dépassé le stade le plus élémentaire de la cueillette dispose de capital, ne serait-ce que sous forme d'outils.

Cette définition matérialiste, si l'on peut dire, a été reprise avec des nuances par Wicksell et Hayek. Wicksell (1954 [1893], p. 105) classe les facteurs matériels de production selon leur durabilité. Les facteurs les plus durables se retrouvent du côté de la « terre », même s'ils ont été produits, et il les appelle des « biens de rente » ; les facteurs qui sont rapidement consommés dans le processus de production se retrouvent du côté du capital, et il les nomme des « biens du capital » (*capital goods*). Hayek (1941, p. 55) reprend cette définition en remplaçant la durabilité par la permanence : pour lui, font partie du capital les facteurs de production non-permanents, c'est-à-dire ceux qui vont disparaître plus ou moins vite dans le processus de production, par opposition aux facteurs de production permanents (non-consommables) qui ne font pas partie du capital. Tout comme Wicksell, Hayek abandonne la distinction entre facteurs originaires et facteurs produits, qui lui apparaît comme un reliquat de la théorie de la valeur travail : la catégorie du capital, comme toute catégorie économique, doit être tournée vers le futur et non vers le passé. La question de savoir si un facteur est le résultat ou non d'un processus de production est une question historique et non pas économique. Ces trois partisans de la conception « réelle » du capital se retrouveront du côté de la théorie productiviste de l'intérêt, alors que les partisans de la conception « nominale » se retrouveront du côté de la théorie subjectiviste expliquant l'intérêt par la préférence pour le présent.

5.1.3 *Peut-on distinguer la « terre » et le « capital » ?* Böhm-Bawerk reprend la distinction classique entre « terre » (facteur de production matériel non produit) et capital (facteur de production matériel produit). Wicksell et Hayek, on l'a vu, retiennent une distinction similaire, fondée sur le critère de la durabilité ou permanence des facteurs matériels de production, excluant du capital les biens matériels permanents ou très durables, qu'ils aient ou non été produits.

Fetter (1900) conteste radicalement cette dichotomie classique puisqu'il estime qu'il faut inclure *tous* les facteurs matériels de production dans le capital. Plus précisément, il réfute les arguments avancés par Böhm-Bawerk en faveur d'une séparation terre/capital. Il montre d'une part que tous ces arguments sont fragiles, et d'autre part que Böhm-Bawerk se contredit parfois en

incluant dans le capital des biens qui devraient être inclus dans sa catégorie « terre ». Il aurait vraisemblablement aussi critiqué la distinction établie par Wicksell et Hayek, arguant du manque de précision du critère de durabilité (à partir de quelle durabilité un bien du capital devient-il un bien de rente, c'est-à-dire une « terre » ?), et arguant de l'inapplicabilité du critère de permanence (car, en toute rigueur, aucun bien matériel ne peut être permanent). Mais la raison profonde pour laquelle Fetter rejette la distinction entre terre et capital est qu'il défend une conception nominale, et non pas réelle, du capital.

Rothbard (1962) juge néanmoins la position de Fetter insatisfaisante, car il existe selon lui un critère important qui permet de distinguer terre et capital, celui de la *reproductibilité*. La terre n'est pas reproductible, et le loyer qu'elle rapporte est de ce fait un revenu net. Les biens du capital sont reproductibles et rapportent un revenu brut dont il faut retrancher les coûts de production pour déterminer le revenu net : ce dernier est, à l'équilibre, égal au revenu d'intérêt.

5.1.4 La conception « nominale » du capital. Dans la tradition autrichienne, la conception nominale du capital remonte à Menger, non pas dans son traité où il reprend la conception « réelle », mais dans un texte ultérieur consacré à la théorie du capital (Menger 1888). Il y définit le capital comme la valeur monétaire d'une propriété, c'est-à-dire qu'il reprend la signification utilisée dans le monde des affaires. Sous l'influence de Clark et de Fisher, Fetter (1900) adopte lui aussi cette conception nominale, et il montre qu'elle transcende les distinctions traditionnelles entre capital social et capital privé, entre « terre » et capital, et même entre facteurs de production et biens de consommation (puisque un bien de consommation, aussi longtemps qu'il peut être vendu ou revendu, possède une valeur en capital). Les choses matérielles ont des définitions différentes selon le point de vue où on les envisage :

(1) du point de vue de la satisfaction des besoins, elles sont des *biens* ;

(2) du point de vue légal (droit de contrôle), elles sont des *propriétés* ;

(3) et du point de vue de leur valeur (prix de marché), elles sont du *capital*.

Fetter (1977, p. 149) ajoute que la notion de capital ne se limite pas aux biens matériels puisqu'elle inclut selon lui les crédits, les clientèles, les franchises, les brevets, etc. Von Mises (1981 [1922], p. 106) s'inspire de Menger en définissant le capital comme un concept essentiellement lié au calcul économique et à la comptabilité. Il insistera aussi sur le fait que le concept de capital, avec son concept connexe de *revenu*, constitue un guide indispensable pour l'action en économie de marché, puisqu'il permet à l'acteur de déterminer la part de ses ressources qu'il peut consommer dans le présent sans pour autant sacrifier la satisfaction de ses besoins futurs (maintien du capital). Si la consommation excède le revenu, il y a consommation de capital, et si au contraire la consommation est inférieure au revenu il y a épargne et accumulation de capital. Le calcul monétaire est absolument indispensable pour mener à bien ces opérations, puisque dans un système économique complexe la description physique d'une multitude de facteurs de production hétérogènes ne donne aucune indication sur leur capacité à satisfaire les besoins humains. Von Mises (1985 [1949], p. 277) définit le capital comme « la somme de monnaie équivalente à tous les éléments d'actif, moins la somme de monnaie équivalente à toutes les obligations », ces actifs consistant en créances et en biens matériels de toutes sortes (terres, bâtiments, outils, biens de consommation, sommes de monnaie, etc.).

5.1.5 *La généralisation de la notion de rente.* Fetter a étendu la notion de capital à tous les facteurs matériels de production, et il généralise de la même façon la notion de rente. La rente est habituellement considérée comme le revenu de la terre, mais Fetter estime que la distinction entre terre et biens du capital est impossible à effectuer en pratique (1901, p. 320). Dès lors, la rente peut légitimement être définie comme le revenu – le *loyer par unité de temps* – de tous les facteurs matériels de production, et non pas seulement des facteurs « terre ».

La conception ricardienne de la rente doit cependant être corrigée. Ricardo pensait que la rente était un revenu de type résiduel, c'est-à-dire qu'elle constituait bien un revenu (obtenu grâce à l'efficacité productive d'un facteur terre) mais ne faisait pas partie des coûts de production. Fetter (1901) montre que Ricardo et Marshall, car ce dernier reprend sur ce point la conception ricar-

dienne, se trompent. L'idée classique selon laquelle la rente ne fait pas partie des coûts de production provient du fait que la terre est un cadeau gratuit de la nature : elle n'a pas de « coût réel » de production puisqu'aucun effort ou sacrifice pénible n'a été requis pour l'obtenir. Cette absence de coût réel est ensuite interprétée comme un coût monétaire nul, mais il y a là pour Fetter une erreur de raisonnement : la rente constitue bel et bien un coût de production (monétaire) dont le producteur – qu'il soit locataire ou propriétaire de la terre – doit tenir compte dans ses dépenses pour évaluer le revenu net de son activité. Bien que Fetter ne se réfère pas ici à Menger, on peut y voir une application directe du principe d'imputation : les coûts réels n'ont aucune influence sur la valeur ou le prix des biens ; le fait que le coût réel des facteurs « terre » soit nul ne leur confère donc aucune spécificité du point de vue de leur valeur ou de leur prix par rapport aux autres types de facteurs. La notion de « quasi-rente » de Marshall était déjà une première tentative de généralisation de la notion de rente, mais selon Fetter elle était encore incomplète et laissait subsister des contradictions dans le système conceptuel. Von Mises (1985 [1949], p. 668) rend hommage à la théorie classique de la rente, qui a anticipé le marginalisme, et il regrette que Ricardo ne se soit pas rendu compte que sa théorie fournissait une solution générale à la question du prix des facteurs, y compris des facteurs travail.

5.1.6 *Rente et actualisation.* Si les facteurs de production rapportent une rente – un loyer – à leurs propriétaires (et coûtent cette rente à leurs acheteurs), quelle est donc la nature du revenu d'intérêt ? Pour Fetter, la rente et l'intérêt ne doivent pas être considérés comme des revenus correspondant à des types de facteurs de production différents, mais plutôt comme les revenus des biens envisagés sous deux aspects différents :

- (1) la rente est le revenu des biens envisagés sous leur aspect de *richesse* ;
- (2) l'intérêt est le revenu des biens envisagés sous leur aspect de *capital*.

Plus exactement, la rente est le prix des biens pour l'unité de temps courante, alors que l'intérêt est un revenu d'actualisation correspondant à un écart entre prix présent et prix futur des biens.

La thèse de Fetter peut être illustrée de la façon suivante. Soit

un producteur qui souhaite se procurer maintenant (à l'instant t_0) un facteur F qui lui rapportera dans un an (à l'instant t_1) une rente R . Quelle somme sera-t-il prêt à déboursier (en t_0) pour l'acheter ? Il ne le paiera pas R , car en plaçant la somme R au taux d'intérêt annuel i en vigueur il obtiendrait en t_1 le revenu $R \times (1 + i)$ supérieur à R . Il achètera le facteur F en t_0 à sa valeur *actualisée* ou capitalisée $R \div (1 + i)$ (s'il parvient à l'acheter encore moins cher, il obtiendra en t_1 un profit entrepreneurial, et s'il le paye plus cher il subira une perte). Ainsi, l'utilisation du facteur F lui rapportera en t_1 à la fois :

(1) une rente R , qui est un revenu *synchronique* correspondant à la productivité en valeur monétaire du facteur F à l'instant t_1 ,

(2) et un intérêt $[R - R \div (1 + i)]$, qui est un revenu *diachronique* correspondant à la différence entre le revenu R de l'utilisation ou de la revente du facteur F à l'instant t_1 , et le prix d'achat $R \div (1 + i)$ de ce facteur à l'instant t_0 .

Ce cas se généralise aisément à celui des biens durables qui sont utilisés pendant plusieurs unités de temps successives. Soit un bien durable qui rapportera 5 000 € au bout d'un an, puis à nouveau 5 000 € au bout de deux ans, et ainsi de suite pendant 4 ans. Si le taux d'intérêt est supposé égal 3 %, alors le prix présent – *capitalisé* et *actualisé* – de ce bien est égal à :

$$\frac{5000}{1 + 0,03} + \frac{5000}{(1 + 0,03)^2} + \frac{5000}{(1 + 0,03)^3} + \frac{5000}{(1 + 0,03)^4}$$

Tout se passe comme si le producteur disposait d'une première unité de facteur achetée pour $5\,000 \div (1 + 0,03)$ et rapportant 5 000 € au bout d'un an (soit un taux d'intérêt de 3 %), d'une deuxième unité de facteur achetée $5\,000 \div (1 + 0,03)^2$ et rapportant 5 000 € au bout de deux ans (soit là aussi un taux d'intérêt annuel de 3 %), etc. Ainsi, le producteur obtient un taux d'intérêt de 3 % et une rente annuelle de 5 000 €.

Dans le cadre de la conception nominale du capital, les facteurs matériels de production rapportent à la fois une rente et un intérêt, ou, pour le dire de façon plus rigoureuse, ils rapportent, soit une rente, soit un intérêt, selon la perspective temporelle dans laquelle ce revenu est envisagé.

Le raisonnement présenté ci-dessus présuppose l'existence du

taux d'intérêt. Mais la *théorie de l'actualisation* (*capitalization theory*) de Fetter est plus profonde puisqu'elle explique – sur la base de la préférence pour le présent – l'existence même d'un taux d'intérêt : le phénomène de l'actualisation (décote des biens futurs) précède et explique celui du taux d'intérêt (voir § 5.2.6).

5.1.7 La théorie autrichienne des revenus : rente, intérêt, profit/perte. Chez Fetter, suivi par von Mises et Rothbard, les facteurs de production rapportent une rente (synchronique) et un intérêt (diachronique), sans qu'il y ait de différence de ce point de vue entre les facteurs originaires et les facteurs produits, ni entre les biens matériels et les services immatériels (dans un système esclavagiste, les services du travail asservi seraient eux aussi actualisés au taux d'intérêt : Rothbard 1977). La rente est le revenu qui reflète la productivité du facteur, c'est-à-dire sa capacité plus ou moins importante à contribuer à la satisfaction des besoins, évaluée d'un point de vue synchronique. L'intérêt reflète la différence de prix positive entre biens présents et biens futurs (voir ci-dessous, § 5.2.6).

Une distinction plus englobante encore est celle entre les *revenus statiques* (rente et intérêt), qui subsistent à l'équilibre, et les *revenus dynamiques* (profits et pertes entrepreneuriaux), qui ne peuvent exister qu'en déséquilibre. Au sein des facteurs de production, les distinctions théoriques importantes sont celles :

- entre facteurs d'ordre supérieur et d'ordre inférieur : les facteurs d'ordre supérieur reçoivent leur valeur et leur prix des biens d'ordre inférieur (principe d'imputation),

- entre facteurs convertibles et spécifiques : les prix des facteurs convertibles leurs sont imputés à partir de leur produit marginal ; le prix des facteurs originaires spécifiques est un prix résiduel, il n'est pas déterminé directement mais calculé comme ce qui reste lorsque les prix des autres types de facteurs du processus ont été soustraits du revenu de ce processus (Rothbard 1962),

- entre facteurs originaires et facteurs produits : les facteurs originaires reçoivent un revenu net, alors que le revenu brut des facteurs produits s'épuise entièrement dans leur coût de production et l'intérêt sur le capital investi pour les produire.

La théorie autrichienne de la distribution n'est évidemment plus une théorie de la répartition du revenu global entre des classes

sociales (travailleurs, propriétaires terriens et capitalistes), mais une théorie de la *distribution fonctionnelle* dans laquelle un individu concret peut être à la fois travailleur et capitaliste, ou propriétaire de terre et entrepreneur, etc., et reçoit donc la somme des revenus correspondant à ses différents types d'activités (Clark 1899).

5.1.8 *Le capitalisme.* Pour Böhm-Bawerk, qui identifie le capital aux facteurs de production produits (biens du capital), une économie est « capitaliste » dès lors qu'elle utilise des outils et des biens intermédiaires. Chez von Mises, qui adopte la conception nominale du capital, la notion de capitalisme est tout à fait différente. En effet, le capital, au sens d'un équivalent monétaire d'éléments d'actif, nécessite l'existence et le fonctionnement des institutions de la propriété privée, de l'échange et de la monnaie. Il n'est pas une catégorie de « l'agir total » mais de l'agir en économie de marché. Ce terme de « capitalisme » lui paraît donc bien adapté pour nommer le système dans lequel le calcul monétaire du capital constitue le principal outil de la rationalité productive.

5.2 Les théories de l'intérêt

5.2.1 *Le problème théorique de l'intérêt.* L'intérêt est le revenu du capital, et il présente selon Böhm-Bawerk (1959 [1884], p. 1) des caractéristiques notables :

- il est obtenu sans effort par le capitaliste,
- il ne dépend pas de la forme matérielle particulière des biens dans lesquels le capital est investi,
- il tend à être proportionnel au capital investi, c'est-à-dire à représenter un pourcentage uniforme du capital (taux d'intérêt),
- il peut en principe être obtenu indéfiniment.

Une théorie de l'intérêt doit donc expliquer d'où provient ce flux de richesses très particulier, qui s'acquiert indéfiniment et sans effort. Böhm-Bawerk distingue d'emblée le problème *théorique* de l'intérêt, qui s'interroge sur les causes de l'intérêt, du problème *moral* de l'intérêt, qui s'interroge sur la justice et le bien-fondé de ce revenu des capitalistes. Pour savoir si le revenu d'intérêt est moralement justifié, il lui paraît nécessaire de commencer par savoir d'où il provient.

5.2.2 *Intérêt originaire et intérêt contractuel.* Le revenu d'intérêt peut se présenter sous deux formes, l'une étant ce que Böhm-Bawerk appelle l'intérêt originaire et l'autre l'intérêt contractuel. Les producteurs vendent généralement leur produit plus cher que ce qu'ils ont dépensé pour le produire. Cet écart entre le prix de vente de l'unité produite et son coût monétaire de production constitue l'intérêt *originaire* (la distinction avec le profit entrepreneurial sera précisée juste après). Mais le revenu d'intérêt peut aussi être *contractuel*, s'il est fixé par contrat lors d'un prêt. La forme originaire de l'intérêt est la plus fondamentale des deux. En effet, s'il n'y avait pas d'intérêt originaire les épargnants capitalistes ne pourraient pas obtenir d'intérêt contractuel en prêtant des fonds aux entreprises. En d'autres termes, si les capitalistes sont en mesure de recevoir un intérêt contractuel pour les prêts qu'ils ont consentis aux entreprises, c'est parce que l'activité de production tend à rapporter un intérêt originaire. Fetter (1914, p. 234) explique lui aussi qu'il y a une priorité logique et historique du taux d'intérêt originaire sur le taux contractuel.

5.2.3 *Intérêt originaire et profit entrepreneurial.* L'écart entre le prix des produits et celui des facteurs de production n'est pas un tout homogène : il se décompose en une partie qui est à proprement parler l'intérêt, et une autre qui est le profit entrepreneurial. Böhm-Bawerk reconnaît qu'en pratique, il est difficile de distinguer les deux. Le critère qu'il propose est d'utiliser le taux d'intérêt habituel dans la branche à laquelle appartient l'entreprise. Si par exemple le capital investi dans l'entreprise est de 100 000 €, si le surplus lors d'une année donnée s'élève à 9 000 €, et si le taux d'intérêt habituel est de 5 %, alors le revenu d'intérêt est 5 000 € et le profit entrepreneurial 4 000 €. Von Mises (1985 [1949], p. 562) montre que le problème est en réalité plus délicat puisque les taux d'intérêt existant sur le marché intègrent tous une composante d'incertitude liée à l'imprévisibilité de l'avenir. Aucune donnée du marché ne peut donc refléter exactement le taux d'intérêt originaire, en tant que distinct du profit entrepreneurial.

Böhm-Bawerk explique le profit par un « effort personnel » de l'entrepreneur : une activité de supervision, une participation aux décisions stratégiques de l'entreprise, ou tout simplement l'acte de volonté consistant à décider d'investir dans cette entreprise plutôt

que dans une autre. Il considère donc le profit comme la rémunération d'un travail. Von Mises (1985 [1949], p. 561) effectue une distinction tripartite entre l'intérêt originaire proprement dit, le salaire managérial (rémunération d'un travail), et le profit/perte entrepreneurial (dû à l'incertitude de l'avenir). À ces trois composantes peut s'ajouter une composante monétaire due aux variations de la valeur de la monnaie (voir § 6.2.14).

5.2.4 *L'intérêt comme phénomène d'évaluation intertemporelle (Böhm-Bawerk)*. Böhm-Bawerk (1959 [1889], p. 259) fonde sa théorie de l'intérêt sur le principe selon lequel les biens présents ont en général davantage de valeur que les biens futurs (voir § 1.2.11). Selon lui, ce principe permet de comprendre les trois formes majeures sous lesquelles apparaît le phénomène de l'intérêt.

(1) *L'intérêt sur les prêts de monnaie (intérêt contractuel)*. Un prêt d'argent est un échange entre un individu *A* qui cède un bien présent (une somme de monnaie présente) et un individu *B* qui cède un bien futur en remboursement (somme future). Comme une somme égale à la somme présente a moins de valeur dès lors qu'elle n'est disponible que dans le futur, le prêteur demandera à recevoir une somme plus élevée en remboursement, et l'emprunteur devra accéder à cette demande sous peine de ne pas obtenir de prêt (sauf s'il trouve un prêteur disposé à lui faire un cadeau, par exemple un membre de sa famille). Ainsi, le créateur recevra davantage d'euros futurs qu'il n'en a prêtés dans le présent, et cette différence constitue l'intérêt.

(2) *L'intérêt sur investissement (intérêt originaire)*. Pour Böhm-Bawerk, les facteurs de production sont des biens futurs, puisqu'il faut du temps pour que se déroule le processus de production qui donnera naissance aux biens de consommation susceptibles de satisfaire des besoins. Les facteurs subissent donc la décote de valeur et de prix appliquée aux biens futurs par rapport aux biens présents. Par exemple, les facteurs utilisés aujourd'hui pour produire 10 tonnes de blé dans un an, valent moins que les 10 tonnes de blé aujourd'hui. Plus généralement, les facteurs de production (biens futurs) valent moins que leurs produits une fois ces derniers devenus disponibles (biens présents), et l'écart de prix constitue là aussi l'intérêt.

(3) *L'intérêt sur les biens durables*. Un bien durable fournit des services au cours des périodes de temps successives de sa durée de vie. Les services d'une période future, lorsqu'ils sont évalués aujourd'hui, sont des biens futurs et ont donc moins de valeur que s'ils étaient immédiatement disponibles. Si un bien durable rend des services qui seront vendus 1 000 € par an pendant 3 ans, alors sa valeur en capital, c'est-à-dire le prix de la totalité de ses services, ne sera pas égale à 3 000 € mais inférieure à ce prix pour tenir compte de la décote de valeur des services futurs, ou en d'autres termes de *l'actualisation* (voir § 5.1.6). Si cette décote est de 5 % par an par exemple, le premier service vaudra bien 1 000 €, mais le suivant vaudra $1\,000 \div (1 + 5\%) = 952$ €, et le dernier vaudra $1\,000 \div (1 + 5\%)^2 = 907$ €, pour un prix total de $1\,000 + 952 + 907 = 2\,859$ € (valeur en capital). Cette actualisation ou décote de valeur permet aussi de comprendre pourquoi la place au sol, qui rend des services pendant un nombre infini de périodes futures, a un prix fini et non pas infini (Böhm-Bawerk 1959 [1889], p. 334).

5.2.5 *La théorie de la productivité du capital (Böhm-Bawerk, Wicksell, Hayek)*. Pour répondre à la question de la détermination du taux d'intérêt originaire, Böhm-Bawerk développe une théorie *productiviste* qui explique l'intérêt par la productivité du capital (dans le paradigme autrichien, cette productivité est liée à la longueur de la période de production : voir § 4.1.8). Sa théorie est assez technique et nous ne pourrions en présenter qu'un bref résumé, en signalant que Wicksell (1954 [1893]) en propose une élégante illustration graphique (enrichie par Dorfman 1959).

Böhm-Bawerk (1959 [1889]) considère un système économique pour lequel la fonction de production (figure 4.1) est fixée. Il suppose qu'il n'existe qu'une seule qualité de travail, et dans un premier temps que le salaire mensuel est lui aussi fixé. Il montre alors que les capitalistes peuvent maximiser le taux d'intérêt en choisissant une certaine période totale de production. Il existe donc, pour chaque valeur du salaire mensuel w , une longueur « optimale » l^* de la structure qui maximise le taux d'intérêt (et l^* croît avec w). Böhm-Bawerk tient ensuite compte de la contrainte de capital, c'est-à-dire du capital total disponible pour financer la structure de production. Si ce capital est insuffisant pour financer

une structure de longueur $l^*(w)$, alors la concurrence entre les travailleurs fait baisser le salaire w et la structure se raccourcit jusqu'à ce que soit atteint le salaire d'équilibre w_e compatible avec le capital disponible. Si au contraire le capital disponible est excessif, alors la concurrence entre les capitalistes pour embaucher les travailleurs fait monter le salaire jusqu'à w_e et allonge la structure jusqu'à $l^*(w_e)$. Le taux d'intérêt est ainsi déterminé par le choix maximisateur des capitalistes, compte tenu de la fonction de production globale et de la quantité totale de capital disponible.

Hayek (1941) ne présente pas une théorie détaillée de la détermination de l'intérêt, mais il adhère explicitement lui aussi à la théorie productiviste selon laquelle l'intérêt est la productivité marginale du capital, c'est-à-dire le surcroît de quantité produite lorsque la période de production s'allonge d'une unité de temps. Il rejette néanmoins les hypothèses simplificatrices sur lesquelles s'appuie Böhm-Bawerk, à savoir qu'il existe une période moyenne de production l et une quantité définie de capital K . En outre, il tient compte des préférences intertemporelles des agents comme facteur supplémentaire, mais selon lui d'importance secondaire dans la détermination du taux d'intérêt (1941, p. 222).

5.2.6 La théorie de la préférence pour le présent ou théorie de l'actualisation (Fetter, von Mises). Pour Fetter et von Mises, Böhm-Bawerk avait définitivement réfuté les théories productivistes de l'intérêt (voir § 5.2.9). Selon eux, il se contredit donc, et se trompe, en développant une théorie productiviste. Car l'intérêt originaire ne peut pas s'expliquer par un raisonnement portant sur un accroissement des quantités ou des valeurs produites : si des facteurs de production permettent de produire des biens qui ont davantage de valeur, alors cette valeur sera imputée aux facteurs et l'écart entre prix et coût restera inexpliqué. Selon Fetter et von Mises, l'intérêt originaire est une émanation des choix intertemporels des acteurs, une conséquence agrégée de leurs actions. Tous les biens sont datés, c'est-à-dire disponibles à une certaine date. Les décisions des acteurs en matière d'offre et de demande ont toutes une dimension temporelle puisqu'elles concernent des biens datés. Compte tenu de la préférence pour le présent, ces décisions font apparaître une prime en faveur des biens présents, ou, ce qui revient au même, une décote pour les biens futurs.

Ces écarts de prix ont tendance à converger les uns vers les autres en pourcentage – en d’autres termes, le taux d’intérêt tend à être uniforme – parce que les entrepreneurs recherchent le pourcentage d’écart le plus élevé possible. Lorsque sur un marché le taux de rendement est supérieur à la moyenne, ils surenchérissent pour acheter les facteurs et accroissent la production, ce qui fait à la fois augmenter le coût unitaire et baisser le prix de vente : le taux d’écart entre prix et coût tend ainsi à se réduire et à revenir vers la moyenne ; symétriquement, si l’écart en pourcentage est inférieur à la moyenne il a tendance à remonter vers elle (voir § 2.2.11). Ainsi, l’économie de marché est entièrement parcourue par le phénomène de l’intérêt originaire qui est inextricablement mêlé aux prix des biens disponibles à des dates différentes. Si les services des biens futurs sont actualisés (par exemple si un bien qui rapportera 1 000 € dans un an s’achète aujourd’hui 950 €), ce n’est pas parce que le taux d’intérêt existe. C’est au contraire parce que les services futurs subissent une décote de valeur – parce qu’ils sont actualisés – qu’il va exister un taux d’intérêt, d’où le nom de *théorie de l’actualisation* que Fetter (1914, p. 230) donne à sa théorie de l’intérêt. L’actualisation – décote de valeur et donc de prix des biens futurs – précède logiquement et historiquement le phénomène du taux d’intérêt, et en fournit l’explication.

Les variations de l’intérêt originaire s’expliquent, dans ce cadre, par des changements de la préférence pour le présent. Si elle baisse, c’est-à-dire si les gens accordent un surcroît de valeur à la consommation future par rapport à la consommation présente, alors la valeur et le prix des biens futurs vont augmenter, la valeur et le prix des biens présents vont diminuer, ce qui va réduire l’écart entre ces prix et donc constituer une diminution du taux d’intérêt originaire. Inversement, si la préférence pour le présent augmente, alors c’est la demande de biens présents qui s’accroît et la demande de biens futurs qui se réduit, d’où un écart accru entre les prix des biens présents et des biens futurs, c’est-à-dire une augmentation du taux d’intérêt originaire. Von Mises (1985 [1949], p. 553) prend l’exemple d’une population qui croit que la fin du monde approche : la préférence pour le présent va considérablement augmenter, la prime aux biens présents va devenir très forte (la décote des biens futurs très importante) et l’intérêt originaire va fortement s’élever.

5.2.7 *La théorie de l'échange (Rothbard)*. Böhm-Bawerk (1959 [1889], p. 287) présente brièvement une théorie de l'échange, que Rothbard (1962, p. 319) développe en détail. D'après cette théorie, le taux d'intérêt se fixe sur le marché d'échange entre biens présents et biens futurs, que Rothbard appelle le *marché du temps*. Ce marché ne se limite pas aux prêts à la consommation, ni même aux prêts et emprunts en monnaie. Il s'étend à tous les processus de production. En effet, l'achat de facteurs de production consiste à céder un bien présent (la monnaie, qui est selon Rothbard le bien présent par excellence) pour obtenir en échange un bien futur (le revenu monétaire obtenu au terme de la production). Rothbard étudie ici la détermination du taux d'intérêt « pur », c'est-à-dire en l'absence de composante d'incertitude.

Chaque acteur est caractérisé par son profil de préférences intertemporelles (avec une préférence plus ou moins forte pour le présent), et par les richesses dont il dispose dans le présent et dont il s'attend à disposer dans le futur. Pour les différents taux d'intérêt possibles entre la période courante et la période future (par exemple : 1 %, puis 2 %, puis 3 %, etc.), chacun envisage, soit d'emprunter jusqu'à la période future parce que le taux est suffisamment faible, soit de prêter parce que le taux est suffisamment rémunérateur :

- la somme des portions de courbes correspondant à des prêts constitue la courbe d'offre de biens présents (contre biens futurs obtenus en remboursement lors de la période suivante) ; cette courbe est croissante : plus le taux d'intérêt est élevé, plus la quantité de monnaie présente cédée est importante puisque la rémunération est de plus en plus forte ;

- la somme des portions correspondant à des emprunts constitue la courbe de demande de biens présents (contre biens futurs) ; cette courbe est décroissante : plus le taux d'intérêt est élevé, plus le remboursement est onéreux, plus la quantité de monnaie empruntée est faible.

La confrontation entre ces deux courbes fait apparaître un taux d'intérêt d'équilibre pour lequel la quantité prêtée est égale à la quantité empruntée (voir figure 5.1). En effet, si le taux est supérieur au taux d'équilibre, alors les offreurs de biens présents ne trouvent pas tous des demandeurs et décident de baisser leur « prix », c'est-à-dire leur taux d'intérêt, pour trouver preneur (et

inversement : si le taux est inférieur au taux d'équilibre, alors il va avoir tendance à remonter). Le taux d'intérêt pur se détermine ici de la même façon que le prix sur un marché d'enchères.

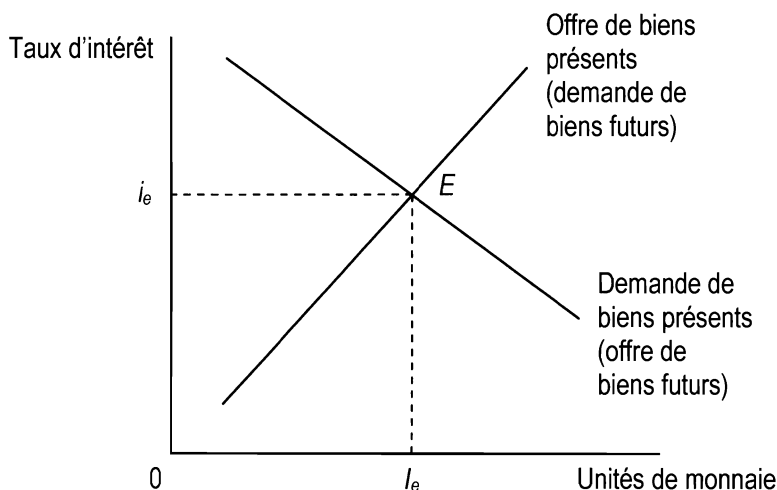


Figure 5.1. La détermination du taux d'intérêt pur d'équilibre (Rothbard 1962, p. 332)

5.2.8 La théorie des fonds prêtables. Böhm-Bawerk attribue très clairement et dès le début de son traité un objectif primordial à la théorie de l'intérêt : rendre compte du phénomène de l'intérêt *originnaire*. Beaucoup de membres de l'école autrichienne pensent, comme lui, que l'existence de l'intérêt originnaire précède logiquement celle de l'intérêt contractuel (Fetter 1927, p. 275, von Mises 1985 [1949], p. 553, Hayek 1941, p. 266, Rothbard 1962, p. 364). L'intérêt contractuel qui se forme sur le marché des fonds prêtables est donc pour eux un phénomène secondaire et dérivé, qui ne peut pas constituer une explication fondamentale de l'intérêt. Cependant, des auteurs plus récents de l'école autrichienne, comme Skousen (1990), Horwitz (2000) et Garrison (2001), adoptent la théorie du marché des fonds prêtables pour expliquer le niveau du taux d'intérêt.

Cette théorie est une application du modèle de l'enchère, comme la théorie de l'échange de Rothbard, mais limitée ici au

marché des fonds prêtables (voir figure 5.2). L'épargne monétaire (croissante en fonction du taux d'intérêt) constitue l'offre de fonds. Les entrepreneurs sont disposés à emprunter des fonds d'autant plus importants que le taux d'intérêt auquel ils devront rembourser est bas, ce qui constitue la demande de fonds (décroissante avec le taux d'intérêt). Le taux d'intérêt a tendance à se fixer au niveau qui égalise la quantité prêtée et la quantité empruntée, c'est-à-dire au taux d'équilibre. Si le taux est supérieur au taux d'équilibre, alors la quantité de fonds offerts est excédentaire, les épargnants sont en concurrence pour placer leurs fonds, ils réduisent leurs exigences pour ne pas manquer une occasion d'échange profitable, et le taux a tendance à baisser vers le taux d'équilibre (et inversement : si le taux est inférieur au taux d'équilibre, il a tendance à remonter). Cette théorie des fonds prêtables n'est pas spécifiquement « autrichienne » puisqu'elle est aussi la théorie standard, présentée dans les ouvrages d'introduction à la macroéconomie (Mankiw 1998, p. 691).

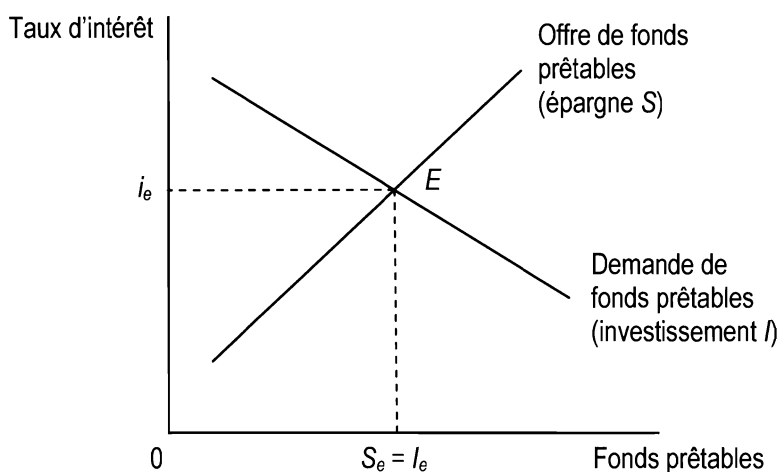


Figure 5.2. Le marché des fonds prêtables
(d'après Skousen 1990, p. 204)

5.2.9 Critique des théories classiques de l'intérêt : productivité, abstinence, rémunération, exploitation. Dans la période « classique » située entre la publication du traité d'Adam Smith (1776)

et la révolution néo-classique des années 1870, plusieurs théories importantes de l'intérêt ont été élaborées. Ces théories ont fait l'objet d'une remarquable analyse critique de la part de Böhm-Bawerk (1959 [1884]).

(1) *Les théories de la productivité du capital* (Say, Lauderdale, Malthus, von Thünen, etc.). D'après ces théories, le « capital » (c'est-à-dire ici les outils, machines, etc.) est « productif » en ce qu'il permet de produire des biens en plus grandes quantités ou de meilleure qualité. Ces quantités ou qualités accrues rapportent donc un revenu plus élevé que celles qui seraient produites sans ce capital. L'intérêt serait constitué par cet accroissement de revenu et proviendrait ainsi de la « productivité » du capital. Or, même si l'on admet que la production accrue obtenue grâce au « capital » rapporte bien ce supplément de revenu (ce qui n'a rien d'évident), cette théorie n'explique *pas* l'existence de l'intérêt. Ce dernier est un *écart* de prix entre un produit et les facteurs qui ont servi à le produire, et à aucun moment cette théorie ne rend compte de l'existence de cet écart. Elle échoue donc à expliquer l'intérêt, tout simplement parce qu'elle ne s'attaque pas au problème posé.

(2) *La théorie de l'abstinence* (Senior, Bastiat, Cairnes). Cette théorie considère l'intérêt comme la récompense obtenue par le capitaliste en échange de son abstinence, c'est-à-dire du sacrifice de consommation présente que constitue son épargne-investissement. Pour Senior (1836), cette abstinence du capitaliste est un coût qui doit être intégré en tant que tel dans le prix de vente. Comme l'écrit Böhm-Bawerk, l'intérêt est alors en quelque sorte le « salaire de l'abstinence », qui s'ajoute aux salaires des travailleurs pour constituer le prix de vente. Menger (1976 [1871], p. 156) propose une brève critique de cette théorie, en disant que les coûts de production quels qu'ils soient – en travail, en abstinence, etc. – ne peuvent *jamaïs* expliquer pourquoi les biens ont de la valeur, et ne peuvent donc pas rendre compte de l'intérêt. Böhm-Bawerk met en évidence une difficulté supplémentaire dans le raisonnement de Senior. Un capitaliste paye aujourd'hui une certaine somme, mettons 100 €, pour le salaire d'un travail qui portera ses fruits dans un an. Pour Senior, ce capitaliste subit deux sacrifices, l'un en salaire (100 €) et l'autre en abstinence (puisque'il doit attendre un an avant de toucher le revenu de son investissement). Mais Böhm-Bawerk montre qu'en réalité il n'a subi *qu'un*

seul sacrifice, celui des 100 €, et aucun autre. Ainsi, même si l'on acceptait une explication « classique » de l'intérêt en termes de coûts de production, la théorie de l'abstinence proposée par Senior ne saurait être satisfaisante.

(3) *La théorie de la rémunération (Courcelle-Seneuil, les socialistes de la chaire)*. L'intérêt est ici tout simplement conçu comme une rémunération des services de travail du capitaliste. Ce travail consisterait par exemple en la constitution et la préservation de l'épargne, c'est-à-dire en un effort de volonté et d'intelligence, et recevrait un salaire sous la forme de l'intérêt. Cette théorie bute cependant sur un fait crucial, qui est qu'un gros capital rapporte un revenu d'intérêt élevé, alors qu'un petit capital ne rapporte qu'un faible revenu d'intérêt. Or, à supposer que les capitalistes aient accompli un effort, il n'a pas été plus grand dans le premier cas ni plus petit dans le second. L'écart de revenu entre le gros et le petit capital reste donc inexpliqué par la théorie de la rémunération. Böhm-Bawerk résume sa critique en disant que l'intérêt n'est pas un revenu du travail mais un revenu de la propriété.

(4) *La théorie de l'exploitation (les socialistes : Rodbertus, Proudhon, Lassalle, Marx, etc.)*. Selon cette théorie, seul le travail est producteur de richesse. L'intérêt est donc une richesse produite par les travailleurs, mais que les capitalistes leur extorquent en profitant de la domination que leur confère la propriété privée du capital. En effet, les travailleurs sont menacés de famine s'ils n'ont pas d'emploi, et ils sont donc contraints à accepter un salaire inférieur à leur véritable contribution (il s'agit là de la version de Rodbertus ; d'après celle de Marx, les travailleurs sont payés à leur salaire de reproduction, c'est-à-dire au minimum de subsistance, et le surplus est confisqué sous forme d'intérêt). Menger (1976 [1871], p. 168) avait déjà brièvement critiqué l'un des piliers de cette théorie en disant que la richesse n'est pas le produit exclusif du travail. Böhm-Bawerk reprend cette critique, puis développe toute une série d'autres arguments dont le plus intéressant est le suivant.

Considérons un processus de production qui n'utilise que du travail (les biens du capital sont produits en cours de route et la « terre » est supposée gratuite). Supposons que ce processus dure un an et produise un bien vendu 10 500 €. Le travail, s'il est rémunéré à la valeur pleine et entière de son produit, vaut 10 500 €.

Mais à *quel moment* vaut-il cette somme ? Les partisans de la théorie de l'exploitation négligent de répondre à cette question cruciale. Si un capitaliste fait une avance pour financer ce processus, combien va-t-il être disposé à payer ? En supposant que le taux d'intérêt annuel soit de 5 %, le capitaliste ne financera ce projet que si le travailleur accepte d'être payé 10 000 € (ou moins). S'il doit payer davantage que 10 000 €, le capitaliste sera perdant puisqu'en recevant 10 500 € un an plus tard il obtiendra un taux d'intérêt inférieur au taux en vigueur qui est de 5 %. Ainsi, la valeur du travail payée *au début* du processus n'est que de 10 000 €, alors que si elle est payée *à la fin* du processus elle s'élève à 10 500 €. Et l'on comprend que le travailleur qui est payé 10 000 € au début de ce processus reçoit la valeur pleine et entière de son travail, puisque s'il plaçait cette somme pendant un an au taux en vigueur de 5 % il récupérerait bien un an plus tard les 10 500 € qui constituent le prix de vente du produit. Même dans un processus qui n'utilise *que* du travail, il n'y a donc *aucune incompatibilité* entre le fait que le capitaliste touche un revenu d'intérêt (ici, 500 € *à la fin* du processus) et le fait que le travailleur récupère la *totalité* de la valeur de son travail (ici, 10 000 € *au début* du processus). Si le capitaliste paye le travailleur moins de 10 000 €, alors il réalise un profit entrepreneurial en plus de son revenu d'intérêt, mais ceci suppose que le système est en déséquilibre et que le salaire va augmenter sous l'effet des forces concurrentielles. À l'équilibre, le capitaliste touche l'intérêt et le travailleur n'est pas « exploité » au sens de Marx puisqu'il récupère bien l'intégralité de la valeur de son travail.

5.2.10 *La théorie de l'usage (Menger)*. Cette théorie est née dans la période classique comme un développement des théories productivistes. Elle est reprise et améliorée par Menger (1976 [1871], p. 157), pour qui l'intérêt ne provient ni de l'abstinence des capitalistes ni de l'exploitation des travailleurs. Il constitue selon lui le prix payé pour l'usage des services du capital. En effet, la production prend du temps, et elle nécessite de pouvoir disposer de ces services du capital pendant une certaine durée. Disposer de ces services a donc une valeur, et un prix s'il y a échange contre de la monnaie : ce prix est l'intérêt sur le capital investi, et il est d'autant plus élevé que la durée de mise à disposition de ces ser-

vices est longue. Menger explique ainsi l'existence de l'écart entre le prix des facteurs de production (capital) et le prix du produit, c'est-à-dire l'existence de l'intérêt originaire.

Böhm-Bawerk, de son côté, présente une critique convaincante de cette théorie. Il a montré, dans sa théorie des biens, qu'une relation entre un acteur et un bien n'était pas elle-même un bien (voir § 1.1.8). Or, le pouvoir de disposer d'un bien est une relation entre l'acteur et ce bien, et ce pouvoir n'est donc pas un bien. Il montre, en d'autres termes, qu'il n'existe pas un second bien (le pouvoir de disposer des facteurs de production) en plus de ces facteurs eux-mêmes. Le prix des services des facteurs de production (pendant toute leur durée d'utilisation) est déjà inclus dans leur prix. La théorie de l'usage n'est pas en mesure d'expliquer l'écart entre le prix des produits et celui des facteurs, et elle n'explique donc pas l'existence de l'intérêt.

Au terme de cette présentation des théories autrichiennes de l'intérêt, force est de constater que les divergences entre les auteurs sont profondes et irrémédiables. Aucune synthèse ne paraît possible entre ces différents points de vue.

Chapitre 6

LA MONNAIE ET SON POUVOIR D'ACHAT

Après des travaux précurseurs de Menger (1976 [1871]) et de Wieser (1910), l'intégration de la monnaie au paradigme autrichien a été effectuée par von Mises (1980 [1912]). Son traité sur les questions monétaires reste l'ouvrage de référence, que l'on peut aujourd'hui compléter par le traité très complet et beaucoup plus récent de Huerta de Soto (2006 [1998]). La théorie autrichienne de la monnaie est d'abord une théorie des *prix monétaires*, en ce sens qu'elle prend pleinement en compte la nature monétaire des prix. Les prix monétaires sont des taux d'échange entre les biens et la monnaie, et ils dépendent nécessairement de l'évaluation subjective par les acteurs économiques des unités de monnaie dont ils disposent. L'utilisation d'une monnaie a donc une influence directe sur la détermination de *tous* les prix. Selon le terme consacré, la monnaie n'est pas un simple « voile » posé sur des échanges de troc – elle n'est *jamais neutre*. Quant à son pouvoir d'achat, il s'explique par la confrontation de l'offre de monnaie et de la demande d'encaisse des individus, selon le schéma classique de la *théorie quantitative*, revue et corrigée à la lumière des principes du subjectivisme et du marginalisme. Von Mises rejette néanmoins vigoureusement la fameuse *équation des échanges* censée offrir une formalisation mathématique de la théorie quantitative.

6.1. La notion de monnaie

6.1.1 *La fonction de la monnaie.* La monnaie n'a pas été inventée, elle n'a pas été créée ni imposée par un pouvoir politique : elle est apparue comme un développement de l'économie de marché permettant aux individus d'accroître la satisfaction de leurs besoins en surmontant les limites du troc et en intensifiant la division du travail (voir § 2.1.9). Von Mises se place dans cette perspective menagérienne et conclut que la fonction de la monnaie n'est autre que de faciliter le fonctionnement des marchés en servant *d'intermédiaire commun entre les échanges*. Les autres fonctions

qui lui sont parfois attribuées ne sont pas distinctes de cette fonction première puisqu'elles s'en déduisent, comme par exemple les fonctions de réserve de valeur, de facilitation des transactions de crédit, ou de moyen de paiement (1980 [1912], p. 47).

La monnaie suppose l'échange et donc la propriété privée, non seulement des biens de consommation, mais surtout des facteurs de production. Elle n'aurait évidemment aucune utilité dans les systèmes économiques qui ne reposent pas sur l'échange, à savoir les économies domestiques isolées d'une part, et les économies collectivistes de l'autre.

6.1.2 *La monnaie en tant que bien*. Pour von Mises, la monnaie est un bien, mais elle n'est ni un bien de consommation ni un facteur de production. Elle appartient à une troisième catégorie de biens, celle des *moyens d'échange*. Le bien qui sert de monnaie peut aussi être, par ailleurs, un bien de consommation ou un facteur de production. Il suffit de prendre l'exemple de l'or qui, lorsqu'il servait de monnaie, était utilisé dans des bijoux et comme composant dans des appareillages électriques. Mais lorsque ce bien est utilisé comme intermédiaire entre les échanges, il n'est pas un bien de consommation (puisque'il est transmis et non pas consommé) et pas non plus un facteur de production (puisque'une augmentation de sa quantité ne permet pas d'accroître le produit global du système économique).

6.1.3 *Les types de monnaies*. Von Mises rejette la distinction traditionnelle entre monnaie métallique et monnaie de papier, qu'il considère comme superficielle, et distingue trois types de monnaies (1980 [1912], p. 73, Hülsmann 2008a, chap. 1) :

(1) La *monnaie marchandise* : le bien utilisé comme monnaie est aussi un bien que les individus peuvent souhaiter se procurer pour son usage direct ; il en a existé dans l'histoire de multiples exemples : le sel, les fourrures, les métaux précieux, etc. ;

(2) La *monnaie décrétée* ou *estampillée* (*Zeichengeld* dans l'original allemand, *fiat money* en anglais) : le bien utilisé comme monnaie doit son caractère de monnaie, non pas à des spécificités physiques, mais à une qualification légale ; si le gouvernement décide que seules peuvent servir de monnaie les pièces d'or fabriquées et estampillées par une certaine entreprise, alors il s'agit là

d'une monnaie décrétée ; les monnaies scripturales étatiques actuelles comme l'euro, le dollar, etc., sont bien sûr elles aussi des monnaies décrétées ;

(3) La *monnaie crédit* : des titres de créance sur une personne physique ou légale sont parfois utilisés comme intermédiaires entre les échanges ; si ces titres sont convertibles en monnaie sur demande et considérés comme parfaitement sûrs, alors les acteurs ne les distinguent pas de la quantité de monnaie qu'ils représentent : ils sont pour von Mises un simple « substitut de monnaie » (voir § suivant) ; mais si en revanche ils ne sont pas parfaitement sûrs, alors ils font l'objet d'une évaluation spécifique et constituent donc une monnaie distincte qui est une monnaie crédit.

6.1.4 *Les substituts de monnaie*. Bien souvent, nous dit von Mises, l'échange monétaire s'effectue, non pas avec de la monnaie proprement dite, mais avec des titres de créance sur une somme de monnaie équivalente. Dans la mesure où ces titres sont parfaitement sûrs et convertibles sur demande en monnaie proprement dite, ils peuvent même circuler comme intermédiaires entre les échanges sans jamais que leurs détenteurs successifs ne demandent leur conversion en monnaie : ce sont des substituts de monnaie. L'exemple le plus évident est celui des billets de banque dans un système de monnaie or. Les gens utilisent les billets émis par les banques, ces billets sont convertibles à vue en or au guichet des banques, et dans la mesure où ils sont considérés comme parfaitement sûrs ils se substituent à la monnaie, c'est-à-dire à l'or, comme moyen d'échange. Mais dès que le moindre doute surgit concernant leur convertibilité en or, ils ne sont plus des substituts de monnaie puisque les gens font alors une différence entre ces titres et la monnaie proprement dite, et ils subissent une dévalorisation plus ou moins importante par rapport à la monnaie. Dans certains pays, la petite monnaie a elle aussi pu faire partie des substituts de monnaie puisque sa valeur en métal était inférieure à sa valeur faciale, mais qu'elle était néanmoins convertible en or sur demande à la Banque centrale.

Dans sa théorie du crédit et des crises, von Mises opérera une distinction fondamentale entre deux types de substituts de monnaie :

(1) d'une part les *certificats de monnaie* qui sont couverts à

100 % par de la monnaie (or ou argent) détenue par la banque qui a émis ces substituts,

(2) et d'autre part la *monnaie fiduciaire* dont la couverture n'est pas totale mais seulement partielle ou fractionnaire (la banque ne conserve pas la totalité mais seulement une fraction de l'or ou argent correspondant aux billets émis).

La typologie de von Mises est représentée sur la figure 6.1.

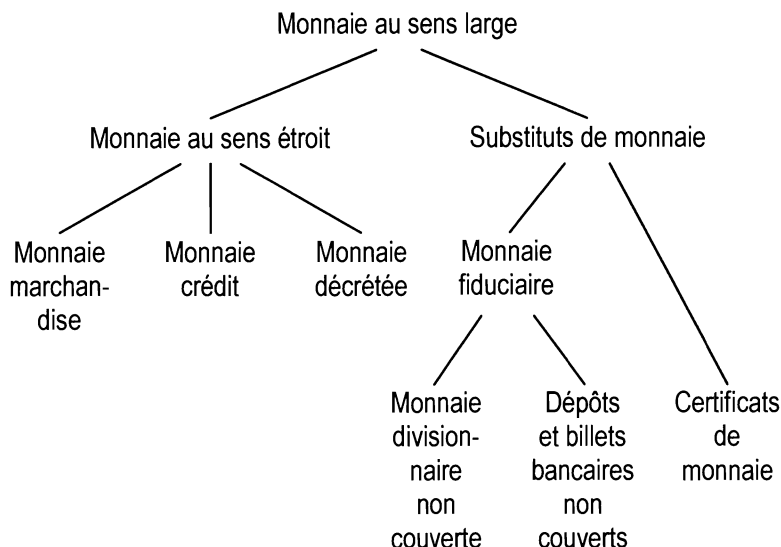


Figure 6.1. La typologie de la monnaie de von Mises
(von Mises 1980 [1912], p. 526)

6.1.5 La tendance à l'unification monétaire. La façon dont l'institution de la monnaie est apparue permet d'expliquer la très grande diversité des biens qui ont servi de monnaie au cours de l'histoire. Chaque société formant un système économique suffisamment développé a secrété ses propres monnaies, en fonction de la nature de ses activités économiques (chasse, élevage, agriculture, etc.), de son degré de sédentarisation, et des biens qu'elle produisait ou trouvait dans son environnement (les coquillages dans les sociétés polynésiennes, les graines de cacao dans les sociétés sud-américaines, etc.).

Lorsque deux systèmes économiques monétaires entrent en contact, la monnaie est indispensable pour procéder aux échanges, et ce sont à nouveau les biens les plus « échangeables » qui vont être sélectionnés comme intermédiaires d'échange, en l'occurrence les biens les plus échangeables parmi les monnaies existantes dans chacun des deux systèmes. Il y a donc une tendance à l'unification des systèmes monétaires au fur et à mesure que le commerce s'étend au-delà des limites d'une société particulière.

Le développement économique et l'internationalisation des échanges ont conduit au cours de cette évolution historique à sélectionner les métaux précieux or et argent comme monnaies, car ce sont les deux biens qui possèdent au plus haut point les qualités d'échangeabilité (forte valeur par unité de volume, malléabilité, divisibilité, durabilité, facilité de transport). Cependant, au cours du ^{xx}e siècle, les États ont les uns après les autres *interdit* à leurs citoyens d'utiliser les métaux précieux comme monnaies, et les ont remplacés par des monnaies scripturales étatiques ayant cours forcé comme l'euro et le dollar actuels (en devenant ainsi les producteurs de la monnaie de leur pays, les gouvernements étaient désormais en mesure de créer *ex nihilo* du pouvoir d'achat pour financer leurs dépenses, et aussi de mener à bien des politiques monétaires, par exemple de relance de la consommation : voir chap. 7). Von Mises (1980 [1912], p. 46) semble penser que compte tenu des désavantages de la coexistence de plusieurs monnaies, une seule – soit l'or, soit l'argent – aurait fini par s'imposer, mais il est tout à fait possible que l'or, l'argent et d'autres monnaies utilisées pour le petit commerce aient continué à circuler côte à côte dans un régime de *concurrence monétaire* qui présente lui aussi une série d'avantages (par exemple la possibilité de délaisser une monnaie dont la qualité se dégrade ; voir Hülsmann 2008a, p. 46, p. 76).

6.1.6 Monnaie et calcul économique. Dans ses trois grands traités, successivement sur la monnaie, le collectivisme et l'action humaine, von Mises accorde une place de plus en plus grande à la question du calcul économique. La thèse qu'il défend est que la monnaie, et en amont l'échange et les droits de propriété privée sur les facteurs de production, sont indispensables pour faire fonctionner efficacement un système économique développé car ils permet-

tent aux acteurs de procéder aux calculs nécessaires.

Dans une économie domestique isolée, l'acteur (« Robinson Crusoé ») n'a pas de difficulté à évaluer l'importance relative des facteurs de production dont il dispose. Les processus qu'il met en œuvre sont courts et peu nombreux, et il lui est donc possible d'imputer une valeur subjective aux unités des différents facteurs, en particulier à son travail, à partir de la valeur qu'il attribue aux biens de consommation produits. Si par exemple il souhaite consommer davantage de viande, il saura quelle quantité supplémentaire de travail consacrer à son activité de chasse au détriment de ses activités de cueillette, de pêche et de loisir.

La situation est totalement différente dans un système économique développé, où les processus de production sont nombreux, longs et entrelacés, et où la plupart des multiples facteurs de production sont convertibles. Les producteurs ont dans ce cas besoin d'un instrument de calcul leur permettant d'estimer *ex ante* et de vérifier *ex post* si les facteurs sont bien alloués de façon à satisfaire au mieux les besoins que les consommateurs considèrent comme les plus importants. Or, pour pouvoir calculer les revenus et les coûts, le capital, les profits et les pertes, etc., il faut une unité de mesure. Les valeurs subjectives ne possèdent pas d'unité de mesure puisqu'elles sont ordinales et non pas cardinales (voir § 1.2.8). Seuls les prix en monnaie permettent de calculer des revenus et des dépenses, et donc d'estimer si l'entreprise projetée gaspillera ou non – du point de vue de la satisfaction des consommateurs – les facteurs convertibles qui pourraient être utilisés dans d'autres processus. Les valeurs subjectives ordinales que les acteurs attribuent aux biens de consommation ne peuvent être imputées à l'ensemble des facteurs de production convertibles sans un système de prix monétaires. L'échange et donc la propriété privée des facteurs de production sont indispensables pour effectuer le calcul économique dans un système économique développé. Cette thèse conduira von Mises à sa célèbre critique du collectivisme (voir § 8.3.4).

6.1.7 Les limites du calcul économique. Une première limite du calcul économique provient des variations de valeur de la monnaie. Ces variations sont inévitables, mais von Mises explique que si la monnaie est « solide », c'est-à-dire si sa valeur ne subit que de

faibles variations au cours du temps (voir § 7.1.10), alors sur le court ou moyen terme le calcul économique n'est pas gravement faussé.

La seconde limite est plus profonde : le calcul économique ne s'applique qu'aux biens qui s'échangent contre de la monnaie. Or, il existe aussi de nombreux biens qui ne sont pas échangés : ils n'ont pas de prix monétaire et ne peuvent donc pas être intégrés au calcul économique. Est-il souhaitable ou non de construire un barrage hydroélectrique à cet endroit ? Si la construction gâche le paysage, alors le calcul économique n'apporte qu'une réponse partielle à la question posée puisque cet enlaidissement, lui, ne peut pas être pris en compte dans le calcul monétaire des revenus et des coûts. Il est cependant tout à fait possible d'en tenir compte et de renoncer, pour des raisons esthétiques, à construire un barrage à cet endroit. Plus généralement, le choix entre les biens matériels et les idéaux (idéaux moraux, spirituels ou esthétiques), entre « le pain et l'honneur » comme le résume von Mises, ne peut être réglé par le calcul économique. Il est réglé par un arbitrage direct entre les valeurs subjectives de biens du premier ordre (von Mises 1935 [1920], p. 98-101).

6.2 Le pouvoir d'achat de la monnaie

6.2.1 La valeur subjective de la monnaie. La valeur subjective d'un bien est le degré d'importance que l'acteur accorde à ce bien relativement aux autres. Une somme de monnaie possède pour l'acteur, comme tous les autres biens, une valeur subjective. La pensée de von Mises sur la source de cette valeur a évolué entre son traité monétaire de 1912 et son traité de synthèse de 1949 (Hülsmann 2007, p. 786).

Dans le premier, il explique que si l'acteur attribue une valeur subjective à une somme de monnaie, c'est parce que cette dernière lui permet de se procurer d'autres biens par l'échange. Cette valeur subjective repose donc sur la quantité *objective* de biens que la somme permet d'obtenir. Elle repose, en d'autres termes, sur la valeur objective d'échange de la monnaie, ou pour le dire plus simplement sur le *pouvoir d'achat* de la monnaie (von Mises 1980 [1912], p. 118).

Plus tard, dans son traité de synthèse, il centre son analyse sur la notion d'*encaisse* : une encaisse monétaire (un stock de monnaie) a pour l'acteur une valeur subjective qui n'est pas une valeur dérivée de celle des biens qu'elle permet d'acheter ; elle rend un service spécifique qui est celui de la monnaie et qui consiste à fournir un pouvoir d'achat immédiatement disponible face à un avenir incertain. Dans cette perspective, la valeur subjective de la monnaie *dépend* de son pouvoir d'achat, en ce sens qu'une unité de monnaie a davantage de valeur subjective si son pouvoir d'achat est plus fort, toutes choses égales par ailleurs ; mais elle *repose* sur le service spécifique que fournit une encaisse monétaire, à savoir une réserve de pouvoir d'achat utilisable à tout instant dans un univers d'incertitude radicale. Si l'avenir était certain, la monnaie deviendrait inutile et personne ne voudrait détenir la moindre encaisse liquide (voir § 6.2.5). Quoi qu'il en soit, la tâche principale de la théorie de la monnaie est de découvrir comment se détermine le pouvoir d'achat de la monnaie à partir de la valeur subjective que lui confèrent les acteurs économiques.

6.2.2 La composante historique de la valeur de la monnaie (théorème de la régression). Le premier point souligné par von Mises est que le pouvoir d'achat de la monnaie, à la différence de celui de tous les autres biens, présente une composante historique (1980 [1912], p. 134). Le prix de la bière aujourd'hui, c'est l'exemple qu'il prend, ne dépend en rien de celui d'hier. Si les prix des biens ont en général tendance à être stables, c'est parce que la situation économique change peu d'une période à l'autre. Mais dans le cas de la monnaie, comme l'avait déjà expliqué Wieser (1910), il existe une chaîne causale qui relie son pouvoir d'achat d'une période à l'autre. En effet, les acteurs forment leurs préférences entre les biens et la monnaie en s'appuyant sur les prix monétaires qu'ils connaissent déjà, c'est-à-dire sur les prix monétaires de la période précédente. Ainsi, les prix monétaires qui apparaissent lors de la période courante – et qui définissent le pouvoir d'achat courant de la monnaie – sont dépendants des prix monétaires et donc du pouvoir d'achat de la monnaie de la période antérieure. Le pouvoir d'achat – la valeur objective d'échange – de la monnaie aujourd'hui dépend de son pouvoir d'achat hier, son pouvoir d'achat hier dépend de son pouvoir d'achat avant-hier, et ainsi de suite.

Von Mises nomme cet enchaînement causal le *théorème de la régression* (1985 [1949], p. 429).

Logiquement, cette régression monétaire remonte jusqu'à la période où le bien qui sert de monnaie n'était pas encore utilisé comme intermédiaire entre les échanges, et n'était valorisé que pour sa capacité à satisfaire des besoins comme bien de consommation ou comme facteur de production : arrivée à ce point, la chaîne causale s'arrête. Toute monnaie marchandise peut ainsi être ramenée jusqu'au moment qui précède son utilisation comme moyen d'échange. De façon similaire, les billets de banque qui circulaient comme substituts de monnaie conservent une valeur, et peuvent donc continuer à être utilisés comme moyen d'échange commun, lorsque leur convertibilité en métal précieux or ou argent est temporairement ou définitivement suspendue par décision gouvernementale.

6.2.3 *Le fondement de la théorie subjectiviste des prix monétaires.*

Le prix en monnaie d'un bien se détermine sur un marché d'échange entre ce bien et la monnaie. Chaque acheteur et chaque vendeur procède à une évaluation relative des unités de monnaie et des unités de biens pour former sa demande ou son offre individuelle, et la confrontation de la somme de ces demandes et de ces offres détermine le prix du bien en monnaie. Mais pour effectuer leurs évaluations, les acheteurs et les vendeurs doivent connaître les prix monétaires des biens. Un *cercle vicieux* semble alors apparaître dans la théorie des prix monétaires : les prix monétaires sont déterminés en tenant compte du pouvoir d'achat de la monnaie, mais ce pouvoir d'achat dépend lui-même des prix monétaires.

Von Mises peut être considéré comme le fondateur de la théorie subjectiviste de la monnaie puisqu'il est le premier à avoir résolu cette difficulté et donc à avoir pleinement intégré la monnaie au paradigme subjectiviste et marginaliste issu de Menger. Il montre que la contradiction n'est qu'apparente (1980 [1912], p. 142). Elle disparaît dès que l'on tient compte de la composante historique des prix en monnaie. Les acteurs déterminent les prix monétaires courants en s'appuyant sur les prix monétaires passés. Ils peuvent ainsi, sans contradiction, former leur offre et leur demande sur le marché d'un bien en comparant la valeur subjective de ce bien et celle de la monnaie. Le prix monétaire déterminé de cette manière

pourra ensuite à son tour jouer un rôle dans la formation des prix monétaires de la période suivante, et ainsi de suite. Les prix monétaires sont ainsi expliqués dans le cadre de la théorie autrichienne de la valeur, et leur explication ne souffre ni d'une circularité ni d'une régression à l'infini (von Mises 1985 [1949], p. 429-430, Rothbard 1962, p. 235).

6.2.4 *La théorie quantitative de la monnaie.* La théorie dite « quantitative » remonte très loin dans l'histoire de la pensée économique, bien avant la période classique du XIX^e siècle, mais elle n'en est pas moins, pour von Mises, la théorie fondamentale permettant d'expliquer le pouvoir d'achat de la monnaie. Il s'oppose aux versions mécanistes et holistes de cette théorie (voir ci-dessous § 6.2.12), mais il en retient l'idée essentielle issue de la loi de l'offre et de la demande : le pouvoir d'achat de la monnaie dépend de la relation entre l'offre et la demande de monnaie. Si l'offre de monnaie devient plus abondante alors que sa demande n'a pas changé, son pouvoir d'achat va tendre à baisser (et inversement). Si sa demande augmente alors que son offre reste la même, son pouvoir d'achat va tendre à s'accroître (et inversement).

6.2.5 *La demande de monnaie.* Un acteur économique demande de la monnaie pour effectuer toutes les dépenses qu'il envisage de faire à plus ou moins brève échéance en biens de consommation, et éventuellement en facteurs de production ou en impôts. La demande individuelle de monnaie est donc pour von Mises une *demande d'encaisses*. La demande de marché de la période courante est, comme pour tous les autres biens, la somme des demandes individuelles (von Mises 1980 [1912], p. 154). La quantité de monnaie que l'acteur détient à un certain moment n'est donc pas un simple reliquat, un simple résidu non dépensé ; elle est la conséquence d'un choix conscient et volontaire, d'un arbitrage effectué entre dépenser et conserver de la monnaie (von Mises 1985 [1949], p. 422, Salin 1990, p. 136).

La source ultime de la demande individuelle d'encaisses est *l'incertitude de l'avenir*. En l'absence de toute incertitude, les acteurs économiques n'auraient pas besoin d'encaisses liquides puisque tous leurs échanges seraient parfaitement prévus à l'avance. Et si personne ne possédait d'encaisses monétaires, cela

signifierait qu'il n'y aurait pas de monnaie dans le système économique. Dans l'hypothèse (irréaliste) où l'avenir serait certain, la monnaie ne servirait à rien et n'existerait donc pas (von Mises 1985 [1949], p. 438).

6.2.6 *L'offre de monnaie*. L'offre de monnaie est constituée par la totalité du stock de monnaie disponible dans le système économique. Ce stock correspond à la quantité de monnaie « au sens large » de von Mises (voir figure 6.1) :

- si les réserves bancaires sont non fractionnaires, c'est-à-dire si les banques émettent des certificats de monnaie (substituts de monnaie couverts à 100 %), alors le stock de monnaie correspond au stock de l'or utilisé à des fins monétaires,

- si les réserves bancaires sont fractionnaires, alors il faut faire la somme du stock de monnaie au sens strict (quantité d'or) et du montant des substituts non couverts (monnaie fiduciaire),

- dans le système de monnaie décrétée étatique qui est le nôtre, le stock de monnaie est constitué par les billets de banque centrale et par le montant des comptes de dépôts des clients des banques (ce qui correspond à l'agrégat monétaire standard M1).

Dans un article consacré à la notion d'offre de monnaie, Rothbard (1997 [1978], p. 346) estime qu'il faut élargir le concept d'offre de monnaie à tous les instruments convertibles en monnaie sur demande à un taux fixé, y compris les comptes d'épargne et les obligations (à leur prix de rachat). Mais dans un texte ultérieur, il définit à nouveau la quantité de monnaie *M* comme la somme des billets de banque centrale et des montants des comptes de dépôt (2008 [1983], p. 132).

6.2.7 *La convergence vers l'équilibre monétaire*. Von Mises ne propose pas à proprement parler de théorie de la convergence du pouvoir d'achat de la monnaie vers sa valeur d'équilibre. Il se contente de dire que le pouvoir d'achat de la monnaie pour la période courante est déterminé par la relation monétaire, c'est-à-dire par la confrontation entre l'offre et la demande de monnaie, cette demande étant elle-même en partie déterminée par le pouvoir d'achat de la monnaie de la période précédente compte tenu du théorème de la régression (1985 [1949], p. 432).

Rothbard (1962, p. 662), de son côté, présente un modèle dé-

taillé de la convergence du pouvoir d'achat de la monnaie vers son niveau d'équilibre, dans le cadre d'une analyse en termes de *demande totale-stock* (il n'est pas certain que la théorie de Rothbard soit compatible avec celle de von Mises, mais cette question ne sera pas abordée ici). Soit un bien quelconque A qui n'est pas la monnaie. Sa « demande totale » pour la période courante est la somme de deux demandes :

- une *demande d'échange* qui est le concept habituel de demande utilisé dans le modèle de l'enchère (voir § 2.2.2) : les quantités du bien A que les acteurs souhaitent acheter, lors de la période courante, respectivement pour les différents prix monétaires possibles de ce bien (et toutes choses égales par ailleurs) ;

- une *demande de réserve* : les quantités du bien A que ses possesseurs veulent conserver selon les différents prix possibles ; soit un acteur qui possède 12 unités de A : si le prix unitaire est 1 € il souhaite en conserver 8 sur les 12 (et vendre les 4 autres), si le prix est 2 € il souhaite en conserver 5 sur les 12 (et vendre les 7 autres), etc.

Ces concepts de demande sont applicables à la monnaie. La demande d'échange correspond dans ce cas à la quantité de monnaie que les acteurs économiques souhaitent obtenir, pour les différents pouvoirs d'achat possibles de cette monnaie, en vendant leurs biens (par exemple leur travail). La demande totale de monnaie pour la période courante est la somme de sa demande d'échange et de sa demande de réserve (auxquelles s'ajoute, pour une monnaie marchandise, la demande de monnaie pour usages non monétaires).

La demande d'échange de la monnaie est *décroissante* par rapport au pouvoir d'achat de la monnaie : quand le pouvoir d'achat de la monnaie augmente, la quantité de monnaie demandée par les acteurs en échange de leurs biens tend à diminuer. La demande de réserve est elle aussi *décroissante* : plus le pouvoir d'achat de la monnaie est élevé, plus est faible la quantité de monnaie que les acteurs souhaitent conserver comme provision pour faire face à leurs dépenses courantes. La demande totale de monnaie est donc décroissante, comme somme de deux fonctions décroissantes. Le stock de monnaie est la quantité totale de monnaie dans le système économique. Il ne dépend pas du pouvoir d'achat de la monnaie et il est donc représenté par une droite *verticale*.

La figure 6.2 représente l'offre et la demande de monnaie en fonction du pouvoir d'achat de la monnaie. Ce dernier est mesuré sur l'axe vertical comme s'il était une simple variable numérique, ce qui constitue bien sûr une grossière simplification. Si le pouvoir d'achat de la monnaie est au-dessus de son point d'équilibre, alors la demande totale de monnaie est inférieure à son stock, les gens dépensent davantage de monnaie, ce qui accroît la demande – et par suite le prix – des biens, et fait donc baisser le pouvoir d'achat de la monnaie qui est ramené vers sa valeur d'équilibre. Si le pouvoir d'achat de la monnaie est au-dessous de sa valeur d'équilibre, alors il a tendance à remonter vers elle selon un processus symétrique au précédent.

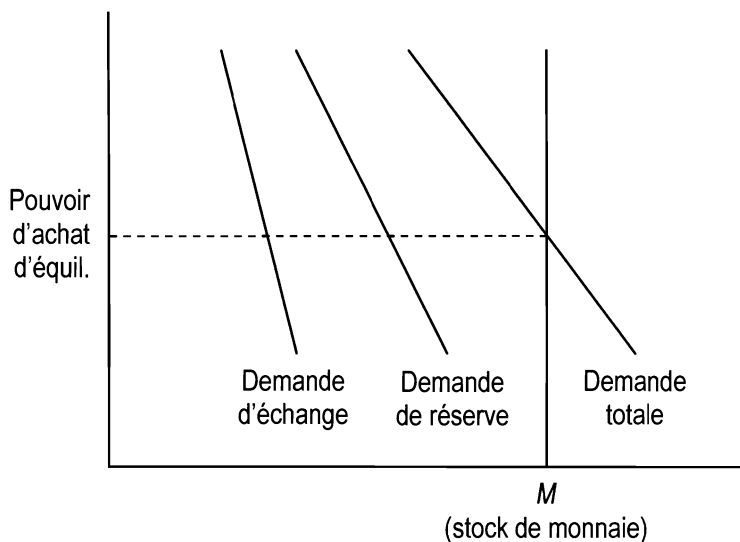


Figure 6.2. L'équilibre du pouvoir d'achat de la monnaie
(d'après Rothbard 1962, p. 667)

6.2.8 Les effets des changements d'offre ou de demande de monnaie sur son pouvoir d'achat. Si l'offre de monnaie augmente, alors les acteurs disposent de stocks accrus de monnaie relativement à leurs stocks des autres biens. L'utilité marginale de la monnaie de ces acteurs a donc tendance à baisser. Les dépenses de monnaie augmentent, ce qui accroît la demande monétaire des

biens, ce qui fait monter les prix monétaires des biens, et donc diminue le pouvoir d'achat de la monnaie. Si c'est la demande de monnaie qui s'accroît (thésaurisation accrue), alors les acteurs conservent par devers eux de plus grandes sommes de monnaie, les dépenses monétaires diminuent, les demandes des autres biens baissent, les prix monétaires ont donc tendance à baisser et le pouvoir d'achat de la monnaie augmente. On retrouve bien les effets classiques de la loi de l'offre et de la demande, en l'occurrence de la théorie quantitative de la monnaie, mais fondés ici sur un raisonnement subjectiviste et marginaliste (von Mises 1980 [1912], p. 161, p. 175).

Ces résultats peuvent aussi être visualisés sur le schéma de Rothbard (figure 6.2) : une augmentation de la quantité de monnaie M conduit à une baisse du pouvoir d'achat d'équilibre, alors qu'une hausse de la demande totale de monnaie élève ce pouvoir d'achat.

Les changements importants de l'offre de monnaie peuvent avoir un impact sur sa demande, dans la mesure où les gens anticipent des variations significatives du pouvoir d'achat de la monnaie. Si par exemple l'offre de monnaie augmente assez rapidement, les prix ont tendance à augmenter eux aussi assez vite. Si les gens prennent conscience de cette montée des prix et s'attendent à ce qu'elle se poursuive, ils réduisent leur demande de monnaie. En effet, plus longtemps ils conservent une unité de monnaie, plus sa valeur se réduit (plus la quantité de biens que cette unité de monnaie peut acheter diminue). L'accroissement de l'offre de monnaie peut ainsi entraîner une réduction de la demande de monnaie, ce qui renforce le mouvement à la baisse du pouvoir d'achat de la monnaie (von Mises 1985 [1949], p. 448).

6.2.9 L'influence de la sphère réelle. Les variations du pouvoir d'achat de la monnaie peuvent provenir de la sphère monétaire, comme on vient de le voir, mais elles peuvent aussi trouver leur origine dans la sphère réelle, du côté de la production. Si l'offre globale de biens s'accroît, par exemple grâce au progrès technique, alors ces stocks accrus de biens vont être mis en vente contre de la monnaie, et toutes choses égales par ailleurs leur prix unitaire va donc avoir tendance à baisser – ce qui signifie que le pouvoir d'achat de la monnaie va tendre à augmenter. Inversement, si la

production et l'offre globale de biens se contractent, alors les prix tendent à monter et le pouvoir d'achat de la monnaie diminue (voir § 7.1.3).

6.2.10 *L'effet redistributif d'une augmentation de la quantité de monnaie (effet Cantillon)*. Supposons que la quantité de monnaie augmente, parce qu'elle a été produite (monnaie marchandise) ou émise (monnaie décrétée) en plus grande quantité. Le pouvoir d'achat de la monnaie va avoir tendance à diminuer, toutes choses égales par ailleurs. Mais cette diminution ne sera pas instantanée. Elle s'opérera tout au long du processus de diffusion de la nouvelle monnaie dans le système économique.

La nouvelle monnaie entre dans le système en arrivant entre les mains de certains acteurs économiques. Ces derniers disposent d'un stock accru de monnaie, leur utilité marginale de la monnaie tend à diminuer, et ils vont donc dépenser au moins une partie de cette monnaie supplémentaire pour se procurer davantage de certains biens. Les demandes de ces biens augmentent, ainsi que leurs prix. Les vendeurs de ces biens récupèrent la nouvelle monnaie, qu'ils vont à leur tour dépenser pour acheter d'autres biens, dont la demande et le prix vont eux aussi augmenter, et ainsi de suite. L'effet est le plus fort aux points d'entrée de la nouvelle monnaie, et il s'atténue au fur et à mesure que des acteurs de plus en plus nombreux reçoivent puis dépensent la nouvelle monnaie.

Ce processus de diffusion avantage clairement les premiers détenteurs de monnaie, puisqu'ils peuvent se procurer les biens qu'ils demandent avant que les prix n'aient augmenté sous l'effet de l'accroissement de la quantité de monnaie. En revanche, ceux qui reçoivent en dernier la nouvelle monnaie vont devoir payer des prix plus élevés avant de bénéficier d'un surcroît de revenu monétaire, et subissent donc une baisse de leur niveau de vie le temps que s'effectue la diffusion (von Mises 1980 [1912], p. 161-162). Hayek (1975 [1931], p. 66) montre que cette théorie remonte à Cantillon et que le type de phénomène qu'elle analyse peut donc être appelé l'effet Cantillon.

6.2.11 *La monnaie n'est jamais « neutre »*. Non seulement les effets des variations de l'offre (ou de la demande) de monnaie ne sont pas instantanés, mais en outre ils ne sont pas – et ne peuvent

jamais être – uniformes. Une monnaie *neutre* serait une monnaie dont les changements de pouvoir d'achat entraîneraient des changements proportionnels de tous les prix, mais sans affecter les activités de production et de consommation qui continueraient à se dérouler comme auparavant. Par exemple, un doublement de la quantité de monnaie multiplierait tous les prix par deux en laissant inchangés les aspects « réels » du système. Von Mises (1985 [1949], p. 437) affirme qu'une monnaie ne peut jamais être neutre. Chaque individu réagit différemment des autres à une variation de son stock de monnaie. En fonction de sa situation particulière et de ses préférences subjectives, il modifie dans un sens ou dans un autre, avec une intensité plus ou moins grande, ses dépenses monétaires. Ces réactions différenciées ont nécessairement un impact dans le domaine de la production, puisqu'elles favorisent certains biens et leurs producteurs au détriment de certains autres. Et il en est de même des variations individuelles de la demande de monnaie. Les changements issus de la sphère monétaire débordent *nécessairement* sur la sphère réelle et rendent impossible la neutralité de la monnaie.

6.2.12 *Critique de l'équation des échanges*. L'équation des échanges est une version mathématisée de la théorie quantitative de la monnaie, popularisée au début du XX^e siècle par Fisher (1911, p. 26-27) sous la forme suivante :

$$MV = PT$$

Cette équation est nécessairement vraie (tautologique) compte tenu de la définition même des termes employés : M est la quantité de monnaie, V la vélocité de la monnaie (le nombre moyen de fois qu'une unité de monnaie est échangée sur la période), P le niveau général des prix, et T le nombre de transactions effectuées pendant la période (pendant une année, par exemple). Fisher utilise l'équation des échanges en supposant que V et T sont constants, et conclut que dans ce cas un doublement de la quantité de monnaie M entraîne un doublement du niveau des prix P (c'est-à-dire une division par deux du pouvoir d'achat de la monnaie).

Von Mises rejette totalement cette équation qui lui paraît viciée à la base par sa conception holiste des phénomènes économiques.

Les taux d'échange se fixent à partir des évaluations subjectives des acteurs. Une relation globale comme l'équation des échanges est donc dans l'incapacité d'expliquer les variations de pouvoir d'achat de la monnaie. Quant aux conséquences mécaniques de cette équation – un doublement de la quantité de monnaie divise par deux le pouvoir d'achat de la monnaie –, elles sont en toute rigueur erronées puisque la monnaie n'est jamais neutre (von Mises 1980 [1912], p. 166-168).

6.2.13 *Les influences de long terme sur la demande et l'offre de monnaie.* Sur le très long terme, la demande de monnaie augmente au fur et à mesure que des individus de plus en plus nombreux s'intègrent à l'ordre économique monétaire. Lorsque la division du travail et l'échange monétaire remplacent l'autarcie, ou tout simplement lorsque la population augmente, des acteurs plus nombreux ont besoin de monnaie et la demande totale d'encaisses tend donc à s'accroître (von Mises 1980 [1912], p. 174). Cette influence à la hausse est contrebalancée par le développement du système des chambres de compensation : les mouvements de monnaie entre les comptes des clients de banques différentes ne sont pas effectués pour chaque transaction mais au terme d'une période ; de ce fait, la plupart de ces mouvements se compensent et seuls les soldes sont échangés, ce qui nécessite beaucoup moins de monnaie.

Du côté de l'offre, pour les monnaies marchandises les variations sont très limitées. La quantité d'or, par exemple, n'augmente chaque année que d'un faible pourcentage (environ 2 %), ce qui ne pourrait avoir qu'une légère influence à la baisse sur le pouvoir d'achat d'une monnaie or (Skousen 1996 [1977], p. 86). Avec le remplacement des monnaies or et argent par des monnaies scripturales étatiques, l'augmentation de la quantité de monnaie n'a plus aucune limite physique. Les Banques centrales ont pu créer de très grandes quantités de monnaie de réserve, qui ont ensuite été démultipliées par les banques grâce au fractionnement des réserves (on démontre que si les banques ne sont tenues de conserver qu'une fraction x de monnaie en réserve face aux dépôts à vue, elles vont pouvoir multiplier la quantité de monnaie par $1/x$; au taux de réserve actuel de 10 %, $x = 0,10$, les banques multiplient par dix la quantité de monnaie de réserve). Il en a résulté, au cours

du XX^e siècle et dans tous les pays, des baisses considérables du pouvoir d'achat des monnaies. En utilisant l'indice des prix à la consommation, Hazlitt (1978, p. 20) montre que le dollar – qui est pourtant l'une des monnaies les plus stables (ou les moins instables) – a perdu les 3/4 de son pouvoir d'achat entre 1940 et 1976. Hayek (1990 [1976], p. 136-137) fait état de ce qu'il appelle la « destruction de la monnaie de papier » en remarquant qu'entre 1950 et 1975 les monnaies des différents pays ont perdu entre 40 % et 99 % de leur valeur : par exemple 57 % pour les États-Unis et la Suisse, 75 % pour la France, 78 % pour le Royaume-Uni, et plus de 99 % pour des pays d'Amérique latine comme le Chili, l'Argentine et le Brésil.

6.2.14 *Changement du pouvoir d'achat de la monnaie et taux d'intérêt.* Seuls les effets à long terme des changements du pouvoir d'achat de la monnaie sur le taux d'intérêt originaire sont analysés ici. Les effets à court terme dus à une baisse rapide du pouvoir d'achat de la monnaie – baisse résultant elle-même d'une expansion inflationniste du crédit – relèvent de la théorie du cycle (voir § 7.2.3).

Pour von Mises (1980 [1912]), les changements de valeur de la monnaie n'ont aucun effet direct sur le taux d'intérêt réel. Ils peuvent en revanche avoir sur lui un effet *indirect*. Lorsque le pouvoir d'achat de la monnaie varie, que ce soit sous l'effet d'un changement de l'offre ou de la demande de monnaie, il s'opère une redistribution (effet Cantillon). Si cette dernière s'effectue au profit des capitalistes et au détriment des autres types d'acteurs économiques, alors l'épargne va avoir tendance à augmenter et le taux d'intérêt à baisser. Si, inversement, la redistribution s'opère au détriment des capitalistes, alors l'épargne va avoir tendance à diminuer et le taux d'intérêt à augmenter. Mais il n'y a aucune raison de privilégier la première ou la seconde de ces deux possibilités, et il faut en conclure que les variations du pouvoir d'achat de la monnaie n'ont pas d'effet systématique dans un sens ou dans un autre sur le taux d'intérêt originaire réel de long terme.

Le taux d'intérêt originaire nominal, en revanche, dépend des variations du pouvoir d'achat de la monnaie. Si les prix augmentent en moyenne de 3 % par an, alors les prix de vente se trouvent accrus de 3 % (environ) par rapport aux coûts de production et le

taux originaire nominal annuel intègre cette hausse. Si par exemple le taux d'intérêt réel est de 4 %, alors le taux d'intérêt nominal va s'élever à approximativement $4 + 3 = 7$ % (Reisman 1996, p. 763).

Hülsmann (2009) nuance la conception misésienne en distinguant le cas d'une monnaie marchandise de celui d'une monnaie scripturale. Il montre que dans le premier cas (monnaie marchandise), les variations de la demande de monnaie ont bel et bien un effet systématique sur le taux d'intérêt : elles entraînent une variation dans le même sens du taux d'intérêt d'équilibre. En effet, si par exemple la monnaie est l'or, et si la demande de monnaie augmente, alors les prix monétaires ont tendance à baisser. La rentabilité de la production d'or augmente, puisque le prix de l'or en les autres biens s'est élevé – une même quantité d'or permet d'acheter davantage de biens – alors que ses coûts de production ont baissé (puisque tous les prix monétaires, en particulier ceux des facteurs de production de l'or, tendent à baisser). Dans le même temps, la rentabilité moyenne de la production des autres biens (intérêt originaire) n'a pas changé. Des capitaux vont être réalloués vers la production d'or où le taux de rentabilité va diminuer, pendant que le taux de rentabilité des autres branches de production va augmenter sous l'effet du reflux des capitaux. Le nouveau taux d'intérêt d'équilibre est donc nécessairement plus élevé que le taux initial. Ainsi, dans le cas d'une monnaie métallique, le taux d'intérêt varie dans le même sens que la demande de monnaie. Dans le cas d'une monnaie scripturale, cet effet ne se manifeste pas puisqu'elle est produite sans coût.

6.2.15 *La théorie de la parité du pouvoir d'achat.* Considérons deux monnaies distinctes, la monnaie *A* et la monnaie *B*, respectivement utilisées dans deux pays différents (ou côte à côte sur le même territoire). Des producteurs (exportateurs et importateurs), des investisseurs, ou tout simplement des voyageurs échangent chacune de ces monnaies en l'autre, ce qui fait apparaître un *taux de change* entre ces deux monnaies. Ce taux de change va tendre à s'aligner sur leur « parité de pouvoir d'achat », selon l'expression forgée par Cassel (von Mises 1980 [1912], p. 207). À cette parité, les gens peuvent se procurer les mêmes quantités de biens, soit avec une somme de monnaie *A*, soit avec la somme de monnaie *B* qu'ils obtiennent en échange contre cette somme de *A*. En d'autres

termes, le taux d'échange entre les monnaies correspond au rapport de leurs pouvoirs d'achat respectifs. Si par exemple une unité de *A* permet d'acheter deux unités d'un bien ou d'un panier de biens, et si une unité de *B* permet d'acheter quatre unités de ce bien ou de ce panier, alors le pouvoir d'achat de *B* est double de celui de *A*, et à la parité une unité de *B* s'échangera contre deux unités *A*.

Si le taux de change s'écarte de la parité, les gens opèrent des arbitrages qui vont ramener ce taux à égalité avec le rapport des pouvoirs d'achat. Dire que le taux de change ne correspond pas à la parité implique que : (1) l'une des deux monnaies est *surévaluée* par rapport à l'autre (par exemple : son pouvoir d'achat est double de l'autre, mais son taux de change est triple en ce sens qu'une de ses unités s'échange contre trois de l'autre monnaie), et (2) l'autre monnaie est *sous-évaluée* (par exemple : son pouvoir d'achat est la moitié de celui de l'autre monnaie, mais son taux de change est trois fois moindre en ce sens qu'il faut trois de ses unités pour se procurer une seule unité de l'autre monnaie). Dans ces conditions, les gens ont intérêt à échanger leurs biens contre la monnaie surévaluée, puis à échanger cette monnaie contre la monnaie sous-évaluée, et enfin à acheter des biens avec la quantité de monnaie ainsi obtenue : grâce à cette succession d'échanges, ils vont en effet pouvoir se procurer davantage de biens qu'ils n'en possédaient initialement. Au cours de ce processus, la monnaie surévaluée est échangée contre la monnaie sous-évaluée, ce qui conduit à réduire leur différence de valeur jusqu'à ce qu'elle finisse par correspondre à la parité de pouvoir d'achat (ce processus est décrit plus en détail, et sous d'autres aspects, par Rothbard 1962, p. 725).

Ainsi, le taux de change entre deux monnaies tend à égaler la parité de leurs pouvoirs d'achat respectifs. Lorsqu'une monnaie perd de la valeur – se dévalue – par rapport à une autre, cela est dû à la baisse relative de son pouvoir d'achat, et cette baisse provient en général d'une politique inflationniste : dans le pays *A*, l'offre de monnaie *A* augmente sous l'effet d'une création monétaire, et son pouvoir d'achat baisse par rapport à celui de la monnaie étrangère *B* ; le taux de change finit par refléter cette nouvelle parité en indiquant une dégradation de la valeur de *A* par rapport à celle de *B*. Von Mises s'oppose totalement aux thèses selon lesquelles la perte

de valeur d'une monnaie nationale par rapport aux autres proviendrait d'une balance des paiements dite « défavorable » ou de l'activité des spéculateurs (1980 [1912], p. 282, p. 286).

Les travaux monétaires de von Mises (1980 [1912]) sont d'emblée très hostiles à l'inflation et aux politiques étatiques de création monétaire, qui vont selon lui dérégler le système des prix, provoquer un appauvrissement économique, et creuser des déséquilibres qui finiront par se solder par des crises. Cette hostilité restera une caractéristique majeure de l'école autrichienne, qui sera amenée à prôner un remède contre l'inflation qui peut aujourd'hui paraître radical, à savoir le retour à une monnaie marchandise.

7.1 L'inflation

7.1.1 Les définitions autrichiennes. Dans le paradigme autrichien, l'inflation et la déflation sont *par définition* des phénomènes monétaires, correspondant respectivement à une création et à une destruction de monnaie. Von Mises en propose les définitions théoriques suivantes (1980 [1912], p. 272) :

- l'inflation est un *accroissement de la quantité de monnaie* (monnaie au sens large) qui n'est pas contrebalancé par une augmentation de la demande de monnaie, et qui a donc pour conséquence une baisse du pouvoir d'achat de la monnaie (une hausse des prix) ;

- la déflation est une *contraction de la quantité de monnaie* qui n'est pas contrebalancée par une réduction de la demande de monnaie, et qui a donc pour conséquence une hausse du pouvoir d'achat de la monnaie (une baisse des prix).

Les définitions de Rothbard sont différentes (1962, p. 940, p. 942). Pour lui, l'inflation est la situation dans laquelle la quantité de monnaie *augmente plus vite que la quantité d'or* (ou plus vite que la quantité de la monnaie marchandise qui serait utilisée sur un marché libre). Symétriquement, il y a déflation lorsque la quantité de monnaie diminue ou augmente moins vite que la quantité d'or (une déflation n'est possible que s'il y a eu au préalable une inflation).

Pour ne pas osciller entre plusieurs conceptions différentes, la

définition qui sera utilisée dans ce chapitre sera *celle de Rothbard*, qui a aussi été adoptée par Hazlitt (1978, p. 11), Reisman (1996, p. 895) et Hülsmann (2008a, p. 85). (Cette définition ne fait pourtant pas l'unanimité puisque Horwitz 2000, p. 104, p. 142, utilise les définitions misésiennes de l'inflation et de la déflation comme modifications de l'équilibre monétaire.)

Au chapitre précédent, les raisonnements étaient présentés en termes de pouvoir d'achat de la monnaie. Ils vont plutôt être présentés ici en termes de niveau des prix. Faut-il le préciser, plus les prix monétaires dans leur ensemble sont élevés, plus le pouvoir d'achat de la monnaie est faible, et inversement.

7.1.2 Inflation et hausse des prix. L'inflation est habituellement définie, aussi bien dans les médias que dans les manuels d'économie, comme une hausse générale des prix. Il faut donc répéter et souligner que les économistes autrichiens ne définissent pas l'inflation comme une hausse des prix. D'après la définition de Rothbard, l'inflation entraîne une hausse des prix, *toutes choses égales par ailleurs*. Mais dans des situations historiques particulières, où la condition « toutes choses égales par ailleurs » n'est pas respectée, il peut y avoir *inflation sans hausse des prix*, et *hausse des prix sans inflation* – le terme « inflation » étant bien employé ici dans le sens que lui donne Rothbard, à savoir que la quantité de monnaie augmente plus vite que la quantité d'or.

(1) Inflation sans hausse des prix : si la production augmente assez fortement, par exemple grâce au progrès technique, et si l'inflation est modérée, alors la hausse des prix due à l'inflation peut être plus que contrebalancée par la baisse due à l'augmentation de l'offre de biens. Cette situation s'est produite aux États-Unis au cours des années 1920 : le niveau des prix de gros et des prix à la consommation est resté à peu près stable parce que l'effet à la hausse dû à l'inflation était compensé par l'effet à la baisse dû à la forte croissance économique de cette période (Rothbard 1983 [1963], p. 154).

(2) Dans des conjonctures historiques très rares, il pourrait y avoir hausse des prix sans inflation. Si les prix avaient tendance à monter alors que la quantité de monnaie n'augmentait pas plus vite que la quantité d'or, cela pourrait être dû, soit à une diminution globale de la production, soit à une baisse de la demande de mon-

naie. Cependant, ces deux facteurs ne peuvent survenir que de façon ponctuelle et limitée. Aucun d'eux ne peut expliquer une hausse des prix forte et continue.

7.1.3 *Une « inflation » par les coûts ?* L'explication la plus naïve de la hausse générale des prix est celle de l'avidité des producteurs capitalistes qui feraient monter les prix pour accroître leurs profits (*profit-push inflation*). Or, les producteurs fixent d'emblée le prix qui maximise leur revenu, en sorte que toute augmentation de ce prix leur serait défavorable puisqu'il impliquerait une baisse de leur revenu (Rothbard 1977 [1970], p. 89). Il est donc impossible d'expliquer de cette façon des hausses de prix généralisées.

Peu satisfaisantes elles aussi sont les explications courantes « d'inflation par les coûts » (*cost-push inflation*), selon lesquelles la hausse des prix serait due à l'augmentation des coûts (Reisman 1996, p. 907-915). Ces théories reviennent à expliquer la hausse des prix par une réduction de l'offre globale de biens. La variante la plus répandue est l'« inflation par les salaires » (*wage-push inflation*). En effet, si les syndicats parviennent à obtenir dans toutes les branches des hausses de salaires, les producteurs devront licencier une partie des travailleurs à cause de la hausse du coût du travail, la quantité globale produite aura tendance à diminuer, et les prix auront tendance à monter, toutes choses égales par ailleurs. L'autre variante est celle des situations de crise (*crisis-push inflation*), comme par exemple la raréfaction brutale d'une ressource originaire (le pétrole lors de la crise de 1973) ou de mauvaises récoltes. Dans ces conjonctures la production globale diminue, ce qui conduit bien, toutes choses égales par ailleurs, à une hausse du niveau des prix.

Cependant, compte tenu de leur caractère exceptionnel, intermittent, aucune de ces formes d'« inflation » ne pourrait avoir un effet durable à la hausse sur les prix. Seules des augmentations substantielles et répétées de la quantité de monnaie peuvent rendre compte d'une dévalorisation forte et continue de la monnaie.

7.1.4 *Les causes de l'inflation.* Avec la définition de Rothbard, l'inflation n'est pas causée par une production de monnaie marchandise (extraction d'or), mais par une *création* de monnaie supplémentaire.

(1) Dans un système où l'or (par exemple) est la monnaie au sens étroit du terme, les banques provoquent une inflation en émettant – légalement ou non – davantage de monnaie fiduciaire (non couverte), c'est-à-dire en réduisant leur taux de réserve (leur taux de couverture en or des comptes de dépôt).

(2) Dans un système de monnaie scripturale étatique (décrétée) comme le nôtre, ce sont les Banques centrales qui déclenchent l'inflation (Hazlitt 1978, p. 118-120) :

- en imprimant davantage de billets de Banque centrale,
- en autorisant les banques à réduire leur taux de réserves obligatoires, ce qui leur permet de créer davantage de monnaie sur la base de leurs réserves existantes,
- en diminuant le taux d'escompte ou taux directeur, qui est le taux d'intérêt auquel les banques peuvent emprunter à la Banque centrale,
- en créant *ex nihilo* de la monnaie puis en l'injectant dans l'économie en rachetant des Bons du Trésor (cette politique de marché ouvert – politique d'*open market* ou de monétisation de la dette – est utilisée par les Banques centrales des États-Unis et d'Angleterre ; la Banque centrale européenne doit opérer d'une façon plus détournée en prêtant aux banques commerciales qui achètent des titres des Trésors nationaux).

Les moyens les plus utilisés aujourd'hui par les Banques centrales sont les deux derniers. Ils permettent d'accroître la quantité de réserves dont disposent les banques, puis de démultiplier cette quantité additionnelle grâce au fractionnement des réserves.

7.1.5 *Les conséquences de l'inflation.* Von Mises énumère, en plus de la baisse du pouvoir d'achat de la monnaie, trois conséquences (1980 [1912], p. 226-243).

(1) Sur la répartition des richesses : les richesses sont redistribuées en faveur des premiers acteurs qui obtiennent la monnaie nouvellement créée et au détriment des derniers qui la reçoivent (effet Cantillon, voir § 6.2.10).

(2) Sur les relations entre créditeurs et débiteurs : si la baisse de la valeur de la monnaie n'a pas été prévue ou a été sous-estimée, alors les débiteurs sont avantagés au détriment des créditeurs puisque ces derniers sont remboursés avec une monnaie dont la valeur a baissé. Mais si la baisse a été correctement anticipée, alors

les prêteurs demanderont et obtiendront un taux d'intérêt plus élevé compensant la perte de pouvoir d'achat de la monnaie : dans ce cas, il n'y aura pas de redistribution en faveur des emprunteurs.

(3) Sur le calcul économique du profit : lorsque la valeur de la monnaie baisse, le calcul du profit, qui consiste à retrancher des coûts passés aux revenus courants, se trouve falsifié. En effet, les revenus sont comptabilisés avec une monnaie dont la valeur a baissé par rapport à la période où les coûts ont été encourus. Le profit calculé en monnaie tend donc à être surévalué, avec un risque de consommation du capital qui sera évoqué ci-après.

7.1.6 *Les méfaits économiques de l'inflation.* Von Mises et ses héritiers (Hayek 1975 [1931], Rothbard 1990 [1963], Sennholz 1977, Hazlitt 1978, Reisman 1996, Huerta de Soto 2006 [1998], Hülsmann 2008a, etc.) sont très critiques vis-à-vis de l'inflation, pour deux raisons. La première est qu'ils considèrent que l'inflation est *la cause des crises économiques* récurrentes subies par les économies capitalistes : ce point fera l'objet de la section suivante. La seconde raison est que selon eux l'inflation *entrave l'accumulation du capital*, et donc restreint le progrès économique :

- par une consommation du capital (surconsommation) : si les producteurs ne tiennent pas compte du fait que leur calcul économique est faussé par l'inflation, ils vont surévaluer leur profit et risquent de consommer une partie de leur capital en croyant le maintenir intact ; cette surconsommation n'est que temporaire puisqu'elle s'opère au détriment de la production future (von Mises 1980 [1912], p. 235) ;

- par une surtaxe pesant sur le capital : si les bénéfices des entreprises sont taxés par l'État à un taux fixe, par exemple 50 %, alors la surévaluation de ces bénéfices due à l'inflation conduit à les taxer en réalité au-delà des 50 % ; si les investisseurs ne réduisent pas leur consommation en proportion, alors la taxation du capital va réduire l'investissement productif au profit de la consommation étatique (von Mises 1980 [1912], p. 236, Reisman 1996, p. 931) ;

- en défavorisant les créiteurs par rapport aux débiteurs lorsque l'inflation est mal anticipée et sous-estimée ;

- en transformant les placements les plus sûrs, ceux dont les

intérêts sont des sommes de montant fixé par contrat (obligations, assurances-vie, bons du Trésor), en placements les moins sûrs ; dès lors qu'un placement rapporte des intérêts fixes (non indexés sur la hausse des prix), les baisses non anticipées du pouvoir d'achat de la monnaie constituent des atteintes directes à la valeur des remboursements (Reisman 1996, p. 930).

En outre, les institutions qui rendent l'inflation possible, à savoir la monnaie scripturale et la Banque centrale en tant que prêteur de dernier ressort, engendrent de forts aléas moraux (*moral hazards*). D'un côté, les autorités politiques risquent d'utiliser à leur avantage, c'est-à-dire risquent d'abuser de, leur capacité à créer autant de monnaie qu'elles le veulent (avec pour seule véritable limite l'hyperinflation et la destruction de la monnaie). D'un autre côté, les banques (et même les autres entreprises), sachant que le gouvernement a les moyens de les renflouer instantanément par création monétaire, vont avoir tendance à prendre des risques excessifs. Ainsi, aussi bien du côté des producteurs que des utilisateurs de monnaie, les institutions de l'inflation accroissent le risque que les acteurs prennent des décisions téméraires et irresponsables (Hülsmann 2006).

7.1.7 Autres méfaits de l'inflation. Pour Hülsmann (2008a, chap. 13), les institutions et les politiques inflationnistes – monnaie scripturale, Banque centrale, politique de « marché ouvert », etc. – ont des conséquences néfastes qui vont bien au-delà de ces aspects économiques puisqu'ils affectent aussi les sphères morale, politique, culturelle, et même spirituelle.

(1) Les monnaies de papier sont selon lui immorales parce qu'elles proviennent toujours d'une violation massive des droits de propriété privée perpétrée par un gouvernement à l'encontre de ses citoyens : ces derniers se voient imposer l'usage d'une monnaie scripturale qu'ils rejetteraient s'ils avaient le droit d'utiliser à la place des monnaies marchandises.

(2) La redistribution opérée par l'inflation (effet Cantillon) est arbitraire et donc injuste. En particulier, les politiques de « marché ouvert » pratiquées par les Banques centrales profitent indûment aux banques commerciales qui vont pouvoir s'enrichir en touchant des intérêts sur une monnaie qu'elles créent *ex nihilo* sous forme de crédit.

(3) Lorsqu'un gouvernement utilise pour régler ses dépenses une monnaie qui vient d'être créée *ex nihilo*, il entre en concurrence avec les autres acheteurs pour se procurer les biens. Tout se passe comme s'il augmentait les impôts puisqu'il peut désormais obtenir davantage de biens au détriment des résidents du pays. Mais cet impôt est *caché* aux citoyens, il échappe à leur contrôle démocratique et constitue donc une confiscation abusive (von Mises 1980 [1912], p. 468).

(4) Cet impôt d'inflation a été utilisé pendant les guerres (par exemple au cours des deux guerres mondiales) pour dissimuler aux citoyens le coût réel de ces conflits. Il contribue aussi à prolonger les affrontements militaires en donnant aux gouvernements des moyens supplémentaires que les citoyens ne leur auraient peut-être pas démocratiquement consentis.

(5) Le système inflationniste confère aux institutions financières comme les banques une importance disproportionnée, non seulement parce qu'elles s'enrichissent en empochant une part des gains de la création monétaire, mais aussi parce que les autres acteurs se trouvent fortement incités à recourir à leurs services. En régime inflationniste, les ménages n'ont pas du tout intérêt à épargner en thésaurisant de la monnaie puisque celle-ci se dévalue plus ou moins rapidement : il leur faut donc placer leur épargne en faisant appel aux services du secteur financier. Les gens sont ainsi amenés à se préoccuper davantage de leur argent et tendent à devenir plus matérialistes. En outre, l'endettement, et avec lui la dépendance financière et le risque de surendettement, tendent à se substituer à l'autonomie et la responsabilité personnelle.

7.1.8 *L'« épargne forcée »*. Il existe un cas dans lequel l'inflation peut favoriser l'accumulation du capital, celui de l'épargne forcée, qui était déjà bien connu des économistes classiques (Hayek 1939 [1932], p. 194). Lorsque la monnaie nouvellement créée est prêtée aux entreprises, une redistribution s'opère de la consommation vers l'investissement. Cette redistribution est favorable à l'accumulation du capital et donc à la production.

Le phénomène de l'épargne forcée constitue pour les économistes autrichiens le fonds de vérité qui sous-tend les politiques gouvernementales inflationnistes. Cependant, pour eux, même si l'épargne forcée pourrait en théorie compenser les effets négatifs

de l'inflation, cette éventualité est en réalité très peu probable. Dans le meilleur des cas, l'épargne forcée ne ferait qu'atténuer quelque peu les méfaits de l'inflation (von Mises 1980 [1912], p. 252-253, Reisman 1996, p. 937, Huerta de Soto 2006 [1998], p. 409-413).

7.1.9 Les fausses solutions contre l'inflation. Compte tenu de leur profonde hostilité à l'encontre de l'inflation, les économistes autrichiens se sont beaucoup intéressés depuis von Mises aux procédés qui permettraient de la limiter ou même de l'éviter. Un dispositif à l'efficacité totalement illusoire mais à la nocivité bien réelle doit d'emblée être écarté : le *contrôle des prix*. Ce dernier ne s'attaque pas aux causes de la hausse des prix mais à ses conséquences, et il ajoute ses propres déséquilibres (voir § 8.1.3) aux dérèglements déjà créés par l'inflation.

Que penser d'une politique visant à *stabiliser la valeur de la monnaie* ? Pour von Mises, elle se heurterait à des obstacles insurmontables. Il faudrait tout d'abord mesurer le pouvoir d'achat de la monnaie. Or toute mesure de ce type comporte une part inévitable d'arbitraire. Et il faudrait ensuite disposer d'instruments permettant de contrebalancer rapidement et avec juste l'intensité requise toute variation du pouvoir d'achat de la monnaie. Une telle politique est impraticable, et elle présente un inconvénient supplémentaire : un pouvoir gouvernemental discrétionnaire sur la valeur de la monnaie peut très facilement passer de l'objectif de stabilisation à un objectif inflationniste. En outre, pour stabiliser la valeur de la monnaie dans une période de forte hausse de la production, la Banque centrale doit pratiquer une vigoureuse expansion du crédit, ce qui a pour conséquence de déséquilibrer le système économique et de le conduire vers la crise (voir section 7.2).

Les variations de la valeur de la monnaie sont inévitables. À défaut de pouvoir être supprimées, elles doivent être limitées grâce à des institutions adéquates, en particulier grâce à l'utilisation de monnaies métalliques (von Mises 1980 [1912], p. 269).

7.1.10 Des institutions monétaires pour une monnaie « solide ». Les économistes autrichiens rejettent totalement les institutions monétaires actuelles, fondées sur une Banque centrale monopolisant la production d'une monnaie de réserve purement scripturale.

Dans un tel système, le risque et les dangers de manipulations inflationnistes par le gouvernement leur paraissent beaucoup trop grands. Pour favoriser le progrès économique et éviter les crises, il faut une monnaie « solide » (*sound*). La solution est simple puisqu'elle réside dans l'utilisation des *monnaies marchandises* au premier rang desquelles se trouvent les métaux précieux or et argent.

Cependant, même si l'or et l'argent constituent les monnaies au sens étroit du terme et s'il n'existe pas de Banque centrale, deux types d'institutions sont envisageables selon que les banques sont autorisées à créer de la monnaie fiduciaire (« banques libres ») ou non (auquel cas les banques de dépôt sont de simples « entrepôts de monnaie »). Dans les deux cas, les banques sont en concurrence pour recueillir les dépôts à vue des acteurs économiques. Elles émettent des substituts de monnaie tels que les billets de banque, et effectuent les ordres de virement sous forme de chèques (et aujourd'hui de cartes bancaires). Dans le système de banque libre, les banques sont tenues par obligation contractuelle vis-à-vis de leurs clients de *convertir à vue en monnaie marchandise* (or, argent ou autre) les montants correspondant à ces substituts de monnaie : le propriétaire d'un billet ou d'un chèque de la banque *B* peut s'il le souhaite l'échanger au guichet de cette banque contre la quantité de monnaie marchandise qu'il représente. Dans le système de banque entrepôt de monnaie, les banques sont en outre tenues de conserver en réserve la *totalité* des dépôts de leurs clients. Si un client a déposé 3 onces d'or et 8 onces d'argent à sa banque, des quantités égales doivent rester dans les coffres de la banque jusqu'à ce qu'il décide d'en retirer tout ou partie.

Hayek (1990 [1976]) propose un schéma institutionnel original, celui de la production concurrentielle de monnaies *scripturales* par des établissements financiers privés. Comme les gens souhaitent disposer de monnaies stables, ces établissements seraient incités à maintenir le pouvoir d'achat de leur monnaie, exprimé comme la capacité à se procurer un panier de ressources naturelles, sous peine de perdre leurs clients – de perdre les utilisateurs de leur monnaie. Ce projet ne sera pas davantage évoqué ici.

7.1.11 *La banque libre (réserves fractionnaires)*. Dans ce système, les banques ont l'autorisation légale de créer puis de prêter de la

monnaie fiduciaire, c'est-à-dire de fractionner leurs réserves de monnaies marchandises au taux qu'elles choisissent. Compte tenu de leurs obligations contractuelles de convertibilité à vue, elles subissent néanmoins une contrainte sur la création monétaire : plus une banque crée de monnaie fiduciaire (alors que les autres n'en créent pas), plus elle reçoit de demandes de retrait de monnaie marchandise de la part des clients des autres banques (auxquels ses propres clients effectuent des paiements), plus son stock de cette monnaie marchandise se réduit, plus elle fragilise sa capacité à convertir à vue ses substituts de monnaie, et plus elle risque de faire faillite. Cette *loi des compensations adverses* est analysée en détail par von Mises (1985 [1949], p. 459) et par Selgin (1991 [1988], p. 66).

Il existe une profonde divergence d'interprétation sur le fonctionnement du système de la banque libre (Rothbard 1988, p. 234). Pour von Mises (1985 [1949], p. 469), la prudence incite les banques à limiter strictement leur création monétaire, sous peine d'avoir des difficultés à assurer leurs obligations contractuelles de convertibilité, de mettre en danger leur réputation et de faire faillite suite à une demande massive de retrait (retrait d'or par exemple) de la part de leurs clients. Si, historiquement, les banques libres ont adopté des taux de réserve assez bas, c'est parce qu'elles s'attendaient à ce que l'État les relève de leurs obligations contractuelles de convertibilité en cas de problème (aléa moral). Mais si elles savent qu'elles ne peuvent pas compter sur une aide gouvernementale, alors elles adopteront par précaution un taux de couverture qui, sans nécessairement atteindre les 100 %, sera très élevé. Le système de banque libre sera alors proche d'un système avec réserves non fractionnaires.

Pour un défenseur de la banque libre comme Selgin (1991 [1988]), l'avantage de ce système est au contraire qu'il permet de s'affranchir de la contrainte trop stricte des 100 % de réserves, et d'adapter aisément la quantité de monnaie aux besoins monétaires. Supposons que la demande de monnaie augmente. Dans un système de pur étalon or (monnaie or avec réserves bancaires de 100 %), la valeur de la monnaie augmenterait et l'extraction minière de l'or se développerait au détriment des autres activités économiques. Un système de banque libre permettrait d'éviter ce coût. En effet, la hausse de la demande de monnaie conduirait à une

diminution du volume des compensations interbancaires, ce qui augmenterait les réserves des banques qui seraient alors en mesure de créer davantage de monnaie (1991 [1988], p. 112). Inversement, une baisse de la demande de monnaie conduirait à une contraction de la masse monétaire. Par ailleurs, dans ses illustrations numériques Selgin utilise un taux de réserves de 10 %, ce qui est très éloigné du pur étalon or (1991 [1988], p. 72).

En résumé, von Mises défend le système de banque libre parce qu'il estime qu'il se rapproche du pur étalon or, alors que Selgin le défend parce qu'il s'en éloigne et permet des ajustements qui seraient impossibles avec une couverture or de 100 %.

7.1.12 La banque entrepôt (100 % de réserves). Ce système diffère du précédent en ce qu'il y est *interdit* aux banques de créer de la monnaie fiduciaire. Elles sont tenues, par obligation légale, de couvrir la totalité du montant des comptes de dépôt par des réserves en monnaie marchandise (dépôts couverts à 100 %). Les banques de dépôt à vue sont dans ce cas de simples *entrepôts de monnaie* puisqu'elles conservent l'intégralité des monnaies marchandises déposées à vue par leurs clients. Dans de telles institutions monétaires, d'après les définitions autrichiennes il ne peut y avoir ni inflation, ni déflation. Les revenus des banques – pour leur activité de dépôt à vue – proviennent alors essentiellement de la rémunération des services d'entreposage et de moyens de paiement qu'elles fournissent à leurs clients. À la suite de Rothbard (2005 [1962]), de nombreux économistes de l'école autrichienne se sont prononcés en faveur de réserves bancaires non fractionnaires (par exemple Skousen 1996 [1977], Reisman 1996, p. 954, Hoppe *et al.* 1998, Huerta de Soto 2006 [1998], Hülsmann 2008a).

Pourquoi s'opposer au système de banque libre ? D'abord parce que le fractionnement des réserves fait peser une menace constante sur la stabilité du système monétaire. Au moindre doute sur la solvabilité d'une banque, les déposants vont se précipiter pour récupérer leur monnaie marchandise, mais ils n'en récupéreront qu'une fraction d'autant plus faible que le taux de réserve était bas. L'inquiétude risque alors de se propager aux autres banques et de déclencher par contagion une panique bancaire générale qui ruinera la totalité des banques et provoquera une grave déflation au sens d'une forte contraction de la quantité de monnaie. Selgin et

White (1996, p. 91) répliquent que les banques à réserves fractionnaires pourraient faire signer à leurs clients des clauses de sauvegarde leur donnant un certain délai en cas de panique bancaire pour rassembler l'or requis, avec compensation pour les clients lésés par cette attente supplémentaire : l'argument paraît assez peu convaincant compte tenu des difficultés insurmontables auxquelles les banques seraient confrontées en cas de panique.

Un autre argument, très différent, est de nature morale et juridique : la création de monnaie fiduciaire est pour Rothbard (2008 [1983], p. 97-99) une activité frauduleuse parce qu'elle consiste pour les banques à prêter de la monnaie qui ne leur appartient pas puisque c'est celle de leurs déposants. La seule différence entre l'activité de création de monnaie fiduciaire et celle de faux monnayeur, nous dit Rothbard, est que la première n'est pas interdite ! Selgin et White (1996, p. 86-92) contestent cet argument en disant que si les clients sont au courant que les banques prêtent leurs dépôts, alors le transfert de propriété du titre sur la monnaie vers les banques est, non seulement légal comme c'est le cas aujourd'hui, mais aussi légitime puisqu'il s'effectue par consentement mutuel. Huerta de Soto (2006 [1998], p. 19) critique le raisonnement de Selgin et White, en soulignant que des différences majeures séparent contrat de prêt et contrat de dépôt, aussi bien du point de vue économique (paiement ou non d'un intérêt, etc.) que du point de vue légal (transfert ou garde, etc.) ; le système des réserves fractionnaires repose selon lui sur une confusion inadmissible entre ces deux types de contrats puisque le client d'une banque à réserves fractionnaires est censé à la fois déposer et prêter son argent à la banque – il est donc censé effectuer simultanément deux activités incompatibles (voir aussi Hoppe *et al.* 1998, Gentier 2003).

7.1.13 *Une baisse tendancielle mais non déflationniste des prix.*

Dans le cadre des institutions monétaires prônées par les économistes autrichiens, le niveau des prix des biens de consommation aurait tendance à baisser sur le long terme (alors que dans notre système actuel il monte plus ou moins fortement). Cette lente baisse refléterait les progrès de la productivité du travail – un équivalent, en beaucoup plus général, de la baisse que nous connaissons sur les prix de produits tels que les ordinateurs, les écrans plats, les téléphones portables, etc. Cette baisse généralisée des

prix des biens de consommation ne constituerait pourtant pas une déflation puisque la quantité de monnaie (or, par exemple) continuerait à augmenter (lentement) d'année en année : il n'y aurait donc pas de difficulté accrue pour les débiteurs à rembourser leurs dettes, ni de baisse des taux de rentabilité des entreprises puisque la baisse des prix unitaires des produits serait plus que compensée par l'accroissement de la production (Reisman 1996, p. 955, Selgin 1997).

7.2 La théorie du cycle

7.2.1 La théorie du crédit de circulation. La théorie autrichienne du cycle – appelée théorie « monétaire » du cycle ou de façon plus précise « théorie du crédit de circulation » – a été conçue par von Mises au début des années 1910 (1980 [1912]). Il raisonnait alors dans le cadre d'un système économique à monnaie métallique (or ou argent) avec émission de monnaie fiduciaire par les banques. Sa théorie explique l'enchaînement d'une phase ascendante (boom) et d'une phase descendante (crise). Elle n'est donc pas une simple théorie de la panique bancaire, même si une panique bancaire due au fractionnement des réserves peut bien sûr venir s'ajouter à la crise. Cette théorie a été appliquée aux deux plus grandes crises économiques du ^{xx}e siècle, la Grande Dépression des années 1930 d'une part, et la stagflation des années 1970 d'autre part, ainsi qu'à toute une série d'autres crises, y compris la crise actuelle dite des « sous-primes ».

7.2.2 Le crédit de circulation. Von Mises part de la distinction entre deux types d'activités bancaires du point de vue des opérations de crédit :

- la *négociation de crédit* qui consiste pour les banques à emprunter puis investir l'épargne de leurs clients, et qui leur rapporte la différence entre le taux d'intérêt qu'elles obtiennent sur ces placements et celui auquel elles empruntent aux épargnants (capitalistes) ;

- l'*émission de crédit* qui consiste pour les banques à fournir du crédit, non pas à partir d'une épargne préalablement constituée par des acteurs économiques, mais par création de monnaie fidu-

ciaire à partir des comptes de dépôt de leurs clients (les banques prêtent l'argent déposé à vue par leurs clients) ; leur revenu est alors constitué de la totalité de l'intérêt qu'elles touchent sur ces prêts.

Il appelle le premier type de crédit le *crédit marchandise* et le second le *crédit de circulation*. Lorsqu'un prêt consenti à partir d'un crédit de circulation est remboursé, la quantité de monnaie qui avait été créée à cette occasion disparaît, mais la banque peut immédiatement recréer la même quantité et consentir un nouveau prêt.

7.2.3 *Expansion du crédit de circulation et taux d'intérêt monétaire*. Wicksell (1936 [1898], p. 102-106) a forgé une distinction bien connue entre le *taux d'intérêt naturel*, qui est le taux d'intérêt qui prévaudrait dans une économie sans monnaie (et qui dépend donc de la situation économique générale), et le *taux d'intérêt monétaire*, qui est le taux d'intérêt sur les prêts en monnaie, en particulier sur les prêts consentis par les banques. Il étudiait principalement l'effet d'une divergence entre ces deux taux sur le niveau des prix (montrant que si le taux monétaire est inférieur au taux naturel, les producteurs bénéficient de cet écart, donc accroissent leur demande de biens et services, d'où une hausse des prix, et inversement si le taux monétaire est supérieur au taux naturel).

Von Mises (1980 [1912]) reprend cette distinction, mais se pose une question différente. Il part de l'hypothèse que les banques accroissent simultanément leurs crédits de circulation (une banque seule n'est pas en mesure d'accroître fortement son crédit de circulation, car elle subirait alors un prélèvement trop important de ses réserves, selon la loi des compensations adverses : voir § 7.1.11). Pour placer ces prêts supplémentaires, elles doivent bien sûr réduire le taux d'intérêt monétaire, qui passe alors au-dessous du taux naturel. Si l'émission de crédit de circulation est suffisamment large, le taux monétaire peut s'approcher de zéro, mais les banques ne descendront pas au-dessous du taux (très faible) qui leur permet de couvrir les frais de fonctionnement de leur activité d'émission de crédit de circulation. Von Mises analyse ce processus en se demandant s'il existe des forces de rappel qui vont ramener le taux monétaire vers le taux naturel (1980 [1912], p. 398-404, 1928, p. 118-130).

7.2.4 *Le boom.* La première phase du processus est celle de l'expansion du crédit (de circulation) émis par les banques en direction du marché des capitaux. Comme on vient de le voir, le taux d'intérêt monétaire baisse. Certaines entreprises qui auparavant n'auraient pas trouvé de financement parce que leur rentabilité était insuffisante comparée au taux d'intérêt, non seulement deviennent rentables, mais parviennent à obtenir des crédits puisque ceux-ci sont désormais plus abondants. Une entreprise qui rapporterait 4 % quand le taux d'intérêt (naturel) est à 5 % ne sera pas financée ; mais si le taux passe à 3 % sous l'effet de l'expansion du crédit, alors elle pourra être lancée.

La conséquence est la même que si le crédit marchandise avait augmenté, à savoir que la structure de production tend à s'allonger, comme sous l'effet d'une véritable accumulation du capital. La conjoncture économique *semble* alors très favorable : les entreprises disposent de fonds supplémentaires ; leur demande de facteurs productifs, et en particulier de travail, augmente, ce qui tend à accroître les prix de ces facteurs. Les prix des biens de consommation vont aussi augmenter, mais avec un certain retard par rapport aux facteurs de production, puisqu'il faut attendre que les propriétaires de ces facteurs (et surtout les travailleurs) dépensent leurs revenus monétaires accrus sur les biens de consommation.

7.2.5 *La crise.* Lorsque l'expansion du crédit s'arrête ou même seulement ralentit (nous verrons pourquoi elle ne peut durer au § 7.2.6), les effets qui viennent d'être décrits *s'inversent* et une crise économique se déclenche.

Tout d'abord, le taux d'intérêt remonte puisqu'il n'est plus artificiellement abaissé par l'expansion du crédit : la rentabilité de toute une série de projets, entrepris en profitant de la baisse des taux et de l'abondance du crédit, s'effondre subitement. Mais c'est surtout dans les projets lancés dans la partie « haute » de la structure que se manifeste la crise. En effet, l'expansion du crédit a provoqué un allongement de la structure de production, c'est-à-dire une réallocation des facteurs vers les étapes hautes. Lorsque l'expansion du crédit prend fin, les forces conduisant à un *racourcissement* reprennent le dessus. Dans la phase de boom les prix des facteurs augmentaient avant ceux des biens de consommation, ce qui faisait apparaître des profits dans la partie haute et y

attirait les facteurs de production. Mais avec la fin de l'expansion du crédit qui alimentait d'abord la demande de facteurs, la demande de consommation des travailleurs joue à nouveau pleinement son rôle et fait monter les prix des biens de consommation par rapport à ceux des facteurs : les profits apparaissent désormais dans les étapes basses (proches de la consommation finale).

La baisse de rentabilité qui se manifeste dans les étapes hautes en même temps que sa hausse dans les étapes basses marque le début de la crise. En un bref laps de temps, de nombreuses entreprises situées en amont de la structure de production doivent réduire leur activité ou même fermer leurs portes. La réallocation massive des facteurs de production, et en particulier du travail, vers les étapes basses constitue la crise proprement dite. En d'autres termes, l'allongement de la structure qui s'est produit sous l'effet de l'émission de crédit de circulation était artificiel : il ne correspondait pas aux souhaits intertemporels des acteurs économiques. La crise est le moment où la véritable situation économique se dévoile, et où s'effectue le pénible ajustement – raccourcissement de la structure – qui ramène le système vers son équilibre intertemporel entre consommation et investissement.

Il est important de noter que s'il existe une symétrie entre le boom (allongement de la structure) et la crise (raccourcissement), leur temporalité est très différente. Le boom dure en général pendant plusieurs années et creuse lentement, subrepticement, le déséquilibre intertemporel. La crise, en revanche, révèle d'un coup toute l'ampleur du désajustement. Une partie des facteurs de production spécifiques fabriqués pendant le boom perdent aussitôt leur valeur et entraînent de lourdes pertes en capital, et d'autres facteurs spécifiques, adaptés à l'équilibre intertemporel, vont devoir être produits en vue d'être combinés avec les unités de facteur travail qui sont réallouées vers le bas de la structure.

7.2.6 Le boom peut-il durer indéfiniment ? Les banques pourraient-elles, en accélérant l'émission de crédit de circulation, reporter indéfiniment le déclenchement de la crise et le rééquilibrage du système ? Pourraient-elles prolonger le boom sans jamais devoir y mettre un terme ? Von Mises et ses successeurs répondent par la négative. L'accélération de l'expansion du crédit conduit nécessairement à l'hyperinflation et donc à la destruction de la

monnaie. Il n'y a que deux issues possibles :

- soit la *crise de réajustement*, qui se déclenche dès lors que les banques stoppent ou même seulement ralentissent l'expansion de crédit ; les banques peuvent en effet être amenées à réfréner leur expansion de crédit parce que leur taux de réserves métalliques devient trop faible (pour leurs opérations de convertibilité à vue ou bien par rapport à la réglementation gouvernementale) ;

- soit l'*hyperinflation* et la destruction de la monnaie si les banques poursuivent au delà de toute mesure leur expansion du crédit.

Finalement, il existe bien une force de rappel qui ramène le taux d'intérêt vers sa valeur naturelle après que l'expansion du crédit de circulation l'en a éloigné. Même la fuite en avant vers l'hyperinflation ne permettra pas, une fois la monnaie détruite, d'échapper au retour vers le taux naturel.

7.2.7 Illustration du cycle par les triangles hayékiens. Les analyses de Hayek (1975 [1931]) et de Rothbard (1962, p. 850-877) ne présentent pas de différences de fond avec celle de von Mises. Elles partent de la distinction entre les simples fluctuations économiques et le cycle des affaires proprement dit. Les chocs économiques usuels donnent lieu à de simples fluctuations : ils déclenchent un processus qui conduit le système vers un nouvel équilibre, comme dans le cas des chocs de demande, des chocs techniques et des chocs de ressource (§ 2.2.13), ainsi que dans le cas des chocs de préférences intertemporelles (qui conduisent à une accumulation ou une consommation du capital : § 4.2.7 et figure 4.5).

L'expansion du crédit, en revanche, donne naissance au cycle avec une phase de boom qui *éloigne* le système de l'équilibre : l'injection de nouvelles liquidités sur le marché des capitaux provoque un *mal-investissement* caractérisé par un surinvestissement dans les étapes hautes et un sous-investissement dans les étapes basses, ce qui entraîne un déplacement des facteurs vers les étapes hautes (éloignées de la consommation finale). L'équilibre entre consommation et investissement, voulu par les acteurs économiques compte tenu de leur préférence intertemporelle, est bousculé par l'expansion du crédit qui augmente la dépense d'investissement plus vite que la dépense de consommation. La structure de production s'étire (allongement de la période moyenne

de production), mais elle s'élargit aussi suivant la dimension horizontale (monétaire) sous l'effet de l'augmentation de la quantité de monnaie, comme le montre la figure 7.1.

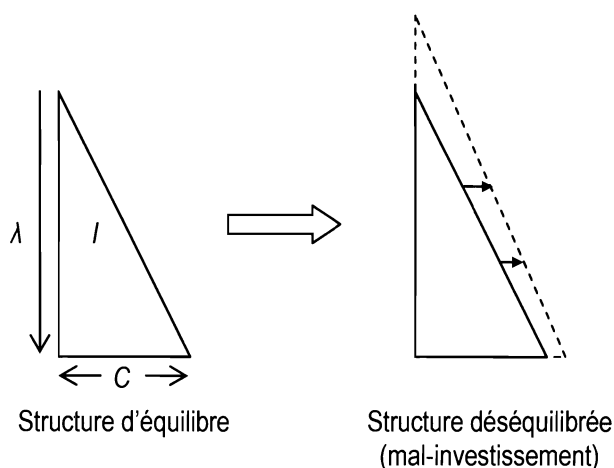


Figure 7.1. L'expansion du crédit provoque un malinvestissement (la base d'un triangle représente la dépense monétaire annuelle C en biens de consommation, la hauteur λ représente la durée totale de production, et l'aire représente la dépense monétaire totale d'investissement I : voir figures 4.3 et 4.5)

Lorsque les agents économiques (propriétaires de facteurs et capitalistes) récupèrent leurs revenus, ils répartissent la part qu'ils dépensent entre consommation et épargne-investissement conformément à leur préférence temporelle. Mais une *nouvelle injection de liquidités*, si elle est suffisamment importante, peut à nouveau court-circuiter le rééquilibrage entre consommation et investissement, et maintenir voire accroître la longueur artificielle de la structure, et ainsi de suite avec une injection supplémentaire, puis une autre, etc. (figure 7.2).

Cependant, à moins qu'une inflation exponentielle ne finisse par détruire la monnaie, il vient un moment où l'expansion du crédit ralentit et ne suffit donc plus à supplanter les choix intertemporels des agents économiques. Ces derniers réinstaurent alors l'équilibre qu'ils souhaitent entre consommation et investissement,

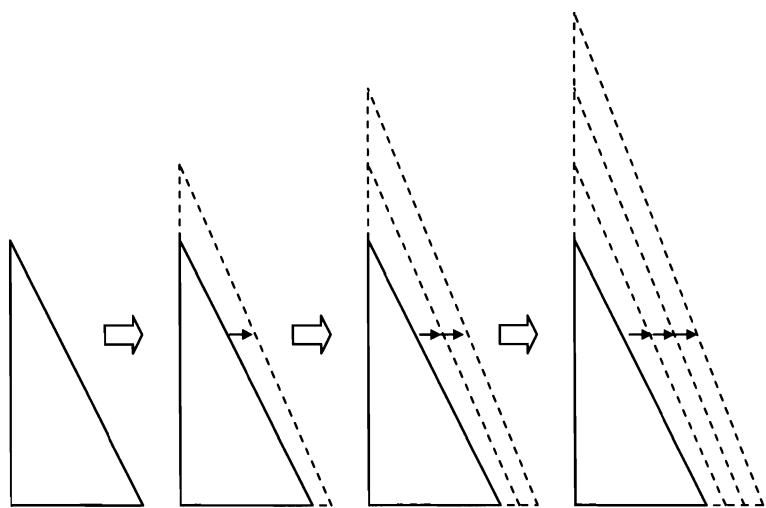


Figure 7.2. Le boom : des injections de crédit successives déforment et déséquilibrent de plus en plus la structure de production (d'après Fillieule 2005, p. 15)

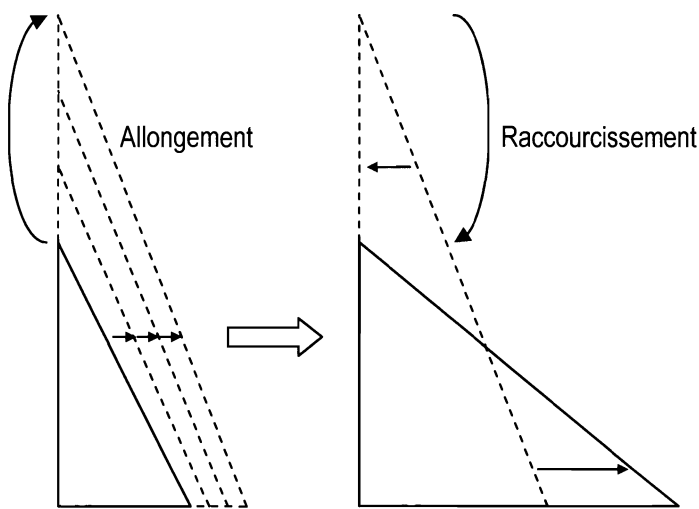


Figure 7.3. La crise : lorsque l'expansion de crédit s'arrête ou ralentit, la structure de production revient à l'équilibre intertemporel entre consommation et investissement

et le taux d'intérêt remonte vers sa valeur d'équilibre (qui peut avoir changé compte tenu de la redistribution opérée par l'inflation). Ce *réajustement* de la structure de production aux préférences intertemporelles constitue la crise, qui est d'autant plus grave que la phase de boom a été plus longue et a creusé plus profondément le déséquilibre entre consommation et investissement (figure 7.3).

7.2.8 Deux métaphores. Von Mises illustre sa théorie du cycle par la métaphore du chantier de construction (1946, p. 223). Dans la phase de boom, tout se passe comme si un chantier trop ambitieux était lancé. Les ressources disponibles sont employées pour creuser les fondations et mettre en place les premiers murs et piliers. Mais en cours de route les bâtisseurs se rendent compte qu'ils ont été présomptueux et que leurs ressources s'avèrent insuffisantes pour achever la construction : c'est la phase de crise. Certaines des ressources investies dans la construction initiale sont gâchées et il faudra finalement se contenter d'un édifice – d'une structure de production – plus modeste.

Rothbard (1962, p. 861) utilise une autre métaphore. Il compare la succession du boom et de la crise à une ivresse (l'euphorie du boom) suivie d'un dégrisement (la « gueule de bois » de la crise), ou à une toxicomanie suivie d'une pénible cure de désintoxication. Cette métaphore permet de comprendre que l'interprétation autrichienne du cycle prend le sens commun complètement à rebours : la phase de boom, bien qu'elle paraisse favorable et bénéfique, est en réalité celle où le système tombe malade, et la phase de crise est la période de guérison ou de réadaptation.

7.2.9 Cycle d'affaires et niveau des prix. Dans la théorie de von Mises, le cycle s'accompagne d'une hausse des prix dans la phase de boom, hausse plus rapide pour les facteurs de production que pour les biens de consommation. En effet, toutes choses égales par ailleurs, l'émission de crédit de circulation supplémentaire augmente la quantité de monnaie (au sens large) et tend donc à accroître le niveau des prix. Mais dans une conjoncture historique donnée, cette tendance à la hausse des prix peut être contrebalancée par des forces qui réduisent le niveau des prix, par exemple une croissance de la production (Hayek 1966 [1929], Rothbard

1962, p. 862). Ainsi, le déséquilibre qui finira par provoquer la crise peut se creuser alors même que le niveau des prix reste stable, comme cela a été le cas lors des années qui ont précédé la crise de 1929.

Au cours de la phase de crise, les prix des biens de consommation augmentent plus vite que ceux des facteurs de production. Mais il peut aussi arriver que les prix baissent dans le cas où une *déflation* (une contraction de la quantité de monnaie) s'ajoute à la crise, par exemple sous l'effet d'une panique bancaire qui va conduire certaines banques à la faillite et détruire une partie de la monnaie fiduciaire. Il s'agit là d'un effet particulier lié au fractionnement des réserves dans le cadre d'une monnaie métallique. Dans le cas d'une monnaie scripturale étatique, la création monétaire peut se poursuivre pendant la crise et donner naissance à une *stagflation* comme dans la période des années 1970, où la dépression s'accompagne d'une forte hausse des prix.

Le point important à noter est que la théorie autrichienne du cycle ne fait jouer *aucun rôle causal à la variation du niveau des prix* puisqu'elle se focalise sur les prix relatifs des biens de consommation et des facteurs de production (Hayek 1966 [1929], p. 103-105).

7.2.10 Combattre la crise ? Pour les économistes autrichiens, la crise n'est pas une maladie mais une phase de guérison qui ramène le système économique vers l'équilibre intertemporel voulu par les acteurs économiques. Le gouvernement devrait donc *ne rien faire qui ralentisse ou bloque ce processus de réajustement* : la crise ne doit pas être combattue. Les options politiques sont alors très réduites. En taxant davantage les travailleurs au profit des capitalistes, le gouvernement pourrait en principe augmenter l'épargne-investissement et amortir le choc en justifiant *a posteriori* certains des investissements qui avaient au départ artificiellement allongé la structure. Cependant, aucun économiste autrichien ne propose une telle politique.

La principale erreur consisterait pour le gouvernement à accélérer l'expansion du crédit et abaisser encore le taux d'intérêt : cela ne ferait que maintenir un déséquilibre intertemporel de la structure de production et prolonger la dépression. Cette politique inflationniste, rejetée par les économistes autrichiens, est précisément

celle qui a été appliquée aux États-Unis pour combattre la dépression issue de la crise des « sous-prime » de 2007. L'opposition avec les remèdes keynésiens est ici frontale : ce qui, du point de vue keynésien, constitue une voie pour sortir de la crise, est bien au contraire pour les économistes autrichiens un facteur d'aggravation.

7.2.11 *Éviter le cycle*. Le meilleur moyen de combattre la crise est, à plus long terme, d'éviter le déclenchement du cycle d'affaires. Pour Rothbard (1962, p. 862), une monnaie or avec réserves bancaires non fractionnaires prémunirait complètement contre le cycle, alors qu'une monnaie scripturale étatique (décrétée) est en revanche la plus propice à l'expansion du crédit et donc au cycle. Selon von Mises, même dans un système de banque libre (monnaie métallique avec réserves fractionnaires), l'expérience des crises passées conduirait les banques soucieuses de la réputation de leur monnaie fiduciaire à beaucoup de prudence vis-à-vis du crédit de circulation. Mais les interférences des gouvernements – création d'une Banque centrale privilégiée, suspension en cas de crise de l'obligation pour les banques de convertir à vue en monnaie métallique les billets et dépôts (et l'on peut ajouter aujourd'hui : garantie des dépôts) – ont toujours eu pour but de promouvoir et de faciliter l'expansion du crédit en déresponsabilisant les banques. Le problème « ultime » lui paraît donc *idéologique* : les cycles existeront aussi longtemps que les gens croiront que la baisse des taux d'intérêt par expansion du crédit rend possible une prospérité générale et durable (von Mises 1928, p. 139-141).

7.2.12 *La Grande Dépression (1929-1941)*. Les économistes autrichiens donnent de cette crise majeure une interprétation aux antipodes de celle qui est couramment admise (von Mises 1931, Rothbard 1983 [1963], Murphy 2009).

D'après la narration habituelle, le système capitaliste a subi en 1929 une de ses crises récurrentes dues à son instabilité intrinsèque, elle-même résultant de l'exubérance irrationnelle des entrepreneurs et des investisseurs. Le Président américain Herbert Hoover n'a mené qu'une politique timide, trop peu interventionniste, et le système s'est enfoncé dans la dépression, avec un taux de chômage atteignant un niveau inconnu jusqu'alors, culminant à 25 %

en 1933. Heureusement, le Président Franklin Roosevelt élu en 1932 s'est au contraire appliqué à mener une politique économique très active, déterminée et soutenue, qui a fini par avoir raison de la crise et par « sauver le capitalisme de lui-même », selon l'expression consacrée.

L'interprétation autrichienne s'oppose point par point à ce récit « canonique ». L'expansion du crédit des années 1920, voulue et encouragée par la Banque fédérale, a subrepticement déséquilibré le système économique. Ce déséquilibre n'a pas été soupçonné parce que le niveau des prix est resté stable au cours de cette période (la tendance à la hausse des prix due à l'expansion du crédit était contrebalancée par la tendance à la baisse due à une forte croissance). Lorsque la crise s'est déclenchée avec le krach boursier d'octobre 1929, le Président Hoover a conduit une politique d'interventions tous azimuts, qui était déjà celle du futur *New Deal* : maintien du niveau des salaires, expansion du crédit, développement des travaux publics, renforcement du protectionnisme, creusement du déficit public, restriction de l'immigration, augmentation des impôts. Cependant, le maintien des salaires nominaux face à la très forte déflation des années 1931-1933 s'est traduit par une *hausse* des salaires réels, ce qui a fait exploser le chômage. Après trois années d'un activisme politique sans précédent dans l'histoire américaine pour combattre une crise économique, celle-ci était plus profonde que jamais. Et pour les économistes autrichiens, c'est bien *cet activisme politique* qui a empêché le rééquilibrage et qui a prolongé et aggravé la dépression.

Franklin D. Roosevelt est élu président des États-Unis en 1932 sur un programme promettant une diminution des immixtions de l'État dans l'économie. Entré en fonction en mars 1933, il va au contraire poursuivre et amplifier avec son *New Deal* la politique de son prédécesseur, et avec aussi peu de succès puisqu'il n'est pas parvenu à mettre fin à la crise : le taux de chômage était encore supérieur à 17 % en 1939.

Quant à la disparition du chômage au cours de la seconde guerre mondiale, selon les économistes autrichiens elle n'est pas due aux dépenses publiques massives effectuées pour l'effort de guerre, mais plutôt à la très forte augmentation de l'épargne privée (Skousen 1989) et au déclin relatif des salaires réels (les salaires réels augmentant moins vite que la productivité du travail : Vedder

et Gallaway 1997 [1993], p. 152-157). L'interprétation autrichienne de l'origine, de la nature, du déroulement et de la sortie de la Grande Dépression s'oppose radicalement à l'interprétation courante.

7.2.13 *La stagflation des années 1970*. Autant les sources du déséquilibre économique des années 1920 étaient restées cachées derrière la stabilité du niveau des prix, autant celles de la stagflation – combinaison de stagnation et d'inflation – des années 1970 ont été bien visibles. Dans les deux décennies qui ont suivi la seconde guerre mondiale, les gouvernements ont appliqué des politiques inflationnistes de « relance par la demande » pour combattre les hausses du chômage. Dès 1950, Hayek prévenait que ces politiques keynésiennes risquaient de conduire à des expansions monétaires de plus en plus fortes, et surtout de moins en moins efficaces contre le chômage (1950, p. 271). En effet, lorsqu'un gouvernement se lance dans ce type de politique, il peut sauvegarder des emplois qui sans cela auraient été détruits, mais ce sauvetage n'est dû qu'à un artifice monétaire qui leur confère une rentabilité. Le maintien de ces emplois exige des doses d'inflation de plus en plus grandes pour empêcher le réajustement qui ferait disparaître ces emplois (voir § 7.2.7). Il est alors très difficile de faire machine arrière car arrêter ou même seulement ralentir la création monétaire conduirait à des destructions massives d'emplois. Cependant, l'inflation affecte négativement l'efficacité productive du système économique (voir § 7.1.6), comme cela s'est manifesté au cours des années 1970.

Lorsque les niveaux élevés d'inflation de cette période n'ont plus suffi à combattre la hausse du chômage, alors même qu'ils déréglaient de plus en plus le fonctionnement du système économique, le gouvernement américain a pris en 1979 la décision de sortir de la stagflation en refermant les vannes de la création monétaire. Le taux annuel d'inflation a culminé à 13,5 % en 1980, puis est revenu aux alentours de 3 à 4 % à partir de 1983. Mais deux récessions successives en ont résulté, de juin à juillet 1980 puis de juillet 1981 à novembre 1982 (source : *National Bureau of Economic Research*), et le taux de chômage a atteint son niveau le plus élevé depuis la Grande Dépression, dépassant les 10 % de septembre 1982 à juin 1983 (source : *US Department of Labor*). Si les

politiques « keynésiennes » de relance inflationniste de l'après-guerre ont bien permis dans un premier temps de combattre les hausses du chômage, elles se sont finalement soldées à plus long terme par le marasme des années 1970 (Hayek 1978, p. 191-231), dont les États-Unis ne sont sortis qu'au prix d'un taux de chômage record au début des années 1980.

7.2.14 *La crise des subprimes*. La crise actuelle dite des « sous-primes » (*subprime crisis*), qui s'est déclenchée aux États-Unis en 2007, constitue pour les Autrichiens un cas d'école d'application de leur théorie.

De 2001 à 2003, la Banque fédérale a conduit une politique monétaire très active pour sortir les États-Unis de la récession ayant suivi l'explosion en mars 2000 de la bulle internet (ainsi que les attentats du 11 septembre 2001). Elle a fait passer le taux d'intérêt interbancaire (*federal fund rate*) de plus de 6 % début 2001 à 1 % en juin 2003 (Thornton 2008). Pour cela, elle a pratiqué une politique de « marché ouvert » en créant puis injectant dans le système économique de nouvelles quantités de monnaie de réserve, démultipliées par les banques (Murphy 2007). Les grandes quantités de monnaie ainsi créées ont été placées sur le marché des capitaux (expansion de crédit) à des taux d'intérêt monétaires de plus en plus faibles, ou vers des emprunteurs de moins en moins sûrs – l'abaissement artificiel du taux monétaire étant bien sûr à la base de l'explication autrichienne du cycle. Cette baisse du taux a encouragé les emprunts à long terme, et en particulier les emprunts immobiliers dont le taux à 30 ans a atteint un plus bas historique au cours du 1^{er} semestre 2003 (Thornton 2008).

À partir de 2004 et jusqu'à la mi-2006, par crainte d'un retour de l'inflation, la Banque fédérale a fait remonter les taux d'intérêt. Mais en resserrant ainsi les vannes du crédit, elle a privé le marché immobilier de fonds qui alimentaient la demande et faisaient monter le prix de vente des logements. Leur demande a diminué, ainsi que leur prix. Non seulement certains emprunteurs à taux variable se sont retrouvés dans l'impossibilité de payer leurs traites à cause de la hausse du taux d'intérêt, mais les institutions de prêts n'avaient plus en garantie que des logements d'une valeur insuffisante. Cette masse d'emprunts immobiliers non remboursables a constitué un profond mal-investissement, une grande quantité de

capital gâchée lors de la phase d'expansion de crédit de 2001-2003 : la profitabilité du marché immobilier était artificielle, entretenue par la création monétaire et disparaissant avec elle. La correction de ce gâchis a conduit à la faillite de nombreux établissements financiers et aurait mené toutes choses égales par ailleurs à une déflation (lorsque la monnaie fiduciaire émise par les banques en faillite aurait disparu). Mais compte tenu des politiques inflationnistes de « relance » mises en œuvre par les diverses Banques centrales, il faut plutôt s'attendre pour l'avenir à un retour de l'inflation (Reisman 2007).

Bien que la crise soit récente, de nombreux commentaires en ligne et quelques publications sont d'ores et déjà disponibles sur ses différents aspects. La bulle immobilière a été rapidement détectée (Shostak 2003) et analysée (Thornton 2006). Huerta de Soto (2008), Rallo (2009) et Horwitz (2009) ont montré la pertinence de la théorie autrichienne pour analyser cette crise. La responsabilité des autorités monétaires dans le déclenchement du cycle a été soulignée par Reisman (2008). Hülsmann (2008b) a insisté sur l'aléa moral que fait peser la banque centrale sur le système bancaire : les organismes prêteurs sont incités à prendre des risques inconsidérés parce qu'ils savent que la banque centrale leur viendra en aide en cas de problème majeur. Enfin, Woods (2009) a rédigé le premier livre qui analyse dans la perspective autrichienne les différents aspects de la crise et critique les politiques mises en œuvre pour la combattre (voir aussi Norberg 2009).

Au cours des années 1920, von Mises (1981 [1922], chap. 34, 1977 [1929]) développe une théorie de l'interventionnisme étatique qui est à la fois simple et sans complaisance à l'égard des politiques gouvernementales : dès lors que l'État intervient dans le marché, soit en contrôlant des prix, soit en réglementant la production, soit en prélevant puis en redistribuant des richesses, soit en créant et en injectant de la monnaie dans le système économique par l'intermédiaire de la Banque centrale (politique inflationniste), il *appauvrit* la société en ce sens qu'il provoque une diminution du niveau de vie moyen. Ces interventions peuvent éventuellement se justifier par des jugements de valeur favorables à la redistribution de la majorité vers une minorité ciblée, mais jamais par l'objectif d'enrichir la société dans son ensemble. Sous l'influence des arguments anti-interventionnistes de von Mises, la plupart des économistes ultérieurs de l'école autrichienne – si ce n'est tous – seront des libéraux. Certains de ses successeurs iront jusqu'à l'anarcho-capitalisme (Rothbard 1991 [1982]), alors que d'autres, moins intransigeants ou moins cohérents, seront favorables à certaines interventions limitées de l'État dans le marché (Hayek 1985 [1944]).

8.1 L'interventionnisme

8.1.1 *Capitalisme, collectivisme, interventionnisme.* Von Mises (1998 [1940]) distingue trois types majeurs de systèmes économiques :

(1) le capitalisme ou économie de marché « non entravée » (*unhampered*) est défini par la propriété privée des facteurs de production ; le rôle de l'État se limite à la protection de « la vie, la santé et la propriété privée de ses citoyens contre la force et la fraude » (1998 [1940], p. 8),

(2) le collectivisme ou socialisme est défini par la propriété publique des facteurs matériels de production,

(3) l'interventionnisme est un système qui se distingue d'une

part du collectivisme en ce qu'il ne cherche pas à abolir la propriété privée des facteurs de production, et d'autre part du capitalisme en ce qu'il vise à restreindre, à limiter le champ d'action des propriétaires de facteurs productifs.

Le gouvernement recourt à l'interventionnisme dans le but d'améliorer la situation des gouvernés. La question fondamentale qui se pose alors, et que von Mises (1926, p. 17) appelle le *problème de l'interventionnisme*, est la suivante : une politique interventionniste peut-elle atteindre l'objectif qu'elle se fixe, à savoir améliorer sur le long terme la situation des gouvernés dans leur ensemble ? À cette question, von Mises répond *non* : les mesures interventionnistes finissent par conduire à une baisse du niveau de vie moyen, toutes choses égales par ailleurs, même si la redistribution qu'elles opèrent peut bénéficier – surtout à court terme – à certaines parties de la population.

8.1.2 *La notion d'intervention.* L'autorité gouvernementale effectue une intervention dans le marché lorsqu'elle « force » les propriétaires de moyens de production à employer ces moyens « d'une autre façon qu'ils ne l'auraient fait eux-mêmes » (von Mises 1926, p. 20). Les deux types d'interventions sont selon lui :

- le *contrôle des prix* (salaire minimum imposé, etc.),
- le *contrôle de la production* (par exemple l'interdiction légale de produire un certain bien).

Dans cette perspective misésienne, les mesures gouvernementales destinées à protéger la propriété privée ne sont évidemment pas des interventions. Les nationalisations et la redistribution étatique n'en sont pas non plus, puisqu'elles laissent par ailleurs le processus de marché se dérouler sans interférence. Rothbard (1977 [1970]) et Lavoie (1982) défendent une conception plus large et sans doute plus satisfaisante de l'intervention, qui inclut les entreprises nationalisées ainsi que les prélèvements et dépenses du gouvernement.

8.1.3 *Le contrôle des prix.* Le contrôle des prix, c'est-à-dire l'interdiction légale faite aux vendeurs de vendre un certain bien au-dessus ou au-dessous d'un certain prix, constitue pour von Mises (1923) l'exemple paradigmatique de l'intervention gouvernementale dans le marché.

Supposons qu'un gouvernement souhaite améliorer le niveau de vie de ses populations les moins favorisées, et décide dans ce but de réduire le prix d'un bien de première nécessité en imposant un prix plafond – prix maximum légal (comme l'ont fait par exemple les révolutionnaires français en 1793 avec le prix du grain). La baisse administrative du prix entraîne un excès de la quantité demandée par rapport à la quantité offerte : les acheteurs désireux de se procurer le bien pour ce prix réduit deviennent plus nombreux, alors que les vendeurs deviennent au contraire moins nombreux (si par exemple ces derniers préfèrent garder pour leur consommation personnelle une partie de leur stock). Le rationnement est d'autant plus marqué que les courbes d'offre et de demande de court terme sont plus élastiques.

Les acheteurs ne pourront pas tous être servis. Les vendeurs vont réagir, soit selon le principe du « premier arrivé premier servi » (d'où des files d'attente et pertes de temps pour les acheteurs), soit en revendant au marché noir ou contre des dessous-de-table (avec de fortes hausses du prix pour couvrir le risque encouru en violant la loi), soit en discriminant entre les acheteurs (en vendant de préférence à leur famille, à leurs amis, à leurs meilleurs clients, aux membres de leur ethnie, etc.). Le gouvernement peut alors aggraver les sanctions pour dissuader les ventes au marché noir, et imposer une politique de rationnement pour que chacun puisse être servi, au moins en partie.

Mais des difficultés plus graves apparaissent lors de la phase suivante du processus. Car la baisse imposée au prix de vente a réduit le taux de rentabilité de la production de ce bien. Les producteurs n'ont tout simplement plus intérêt à produire ce bien en aussi grande quantité qu'auparavant. Le gouvernement, s'il tient à maintenir en vigueur le contrôle du prix, va devoir rendre à nouveau la production rentable en réduisant les prix des facteurs, c'est-à-dire des ressources naturelles, des biens du capital intermédiaires et même du travail. Le contrôle des prix va ainsi avoir tendance à se généraliser à l'ensemble du système économique, et doit logiquement déboucher sur la socialisation complète de l'économie (von Mises 1923, p. 146).

8.1.4 Du contrôle sélectif au contrôle universel des prix. Lorsque le contrôle des prix est « sélectif », c'est-à-dire limité à un petit

nombre de marchandises, la dépense des acheteurs sur ces marchandises diminue puisque ces dernières sont désormais vendues moins cher et en moins grande quantité. Les acheteurs reportent alors leurs dépenses sur des substituts des biens contrôlés (Rothbard 1977 [1970], p. 26). Les prix de ces substituts s'élèvent sous l'effet de cette augmentation de leur demande, la rentabilité de leur production s'accroît et leurs branches se développent sous l'afflux de capitaux. Il en résulte une distorsion de la structure de production, qui produit de plus faibles quantités des biens contrôlés et de plus grandes quantités de leurs substituts que ne le souhaitent les consommateurs. Cette distorsion n'est pas non plus voulue par le gouvernement, qui peut alors :

- soit lever le contrôle des prix, ce qui inverse le processus et supprime la distorsion en restaurant la structure de production voulue par les consommateurs,

- soit au contraire chercher à circonvenir les effets pervers du contrôle des prix initial en l'étendant aux prix des facteurs de production (de façon à restaurer la rentabilité des branches des biens contrôlés), puis aux facteurs de ces facteurs et ainsi de suite jusqu'à ce que *tous* les prix soient contrôlés.

Dans ce second cas, le contrôle des prix devient « universel ». Les acheteurs ne peuvent plus reporter leurs dépenses sur des substituts et une partie de leur monnaie devient comme « anesthésiée », selon le terme de Rothbard, puisqu'elle ne peut plus être dépensée. Les rationnements et les files d'attente vont alors se multiplier, les marchés noirs et le favoritisme se généraliser.

8.1.5 *Le salaire minimum*. Alors qu'un prix maximum (prix plafond) entraîne un rationnement sous une forme ou une autre, un prix minimum (prix-plancher) provoque au contraire des invendus puisque la quantité demandée pour ce prix devient inférieure à celle que les vendeurs souhaitent écouler. L'exemple le plus typique est celui du salaire minimum (von Mises 1923, p. 148-149).

Si le salaire minimum n'est appliqué que dans certaines branches, alors l'augmentation des coûts de production y réduit le taux de rentabilité puisqu'un travailleur coûte davantage (salaire-plancher imposé) que ce qu'il rapporte (productivité marginale nominale). Les capitaux ont donc tendance à refluer vers les branches non contrôlées où le taux de rentabilité est resté plus éle-

vé. La baisse de l'investissement dans les branches contrôlées conduit à une réduction de l'offre et une hausse du prix de vente (qui finit par ramener le taux de rentabilité au taux d'équilibre), mais aussi à des licenciements, toutes choses égales par ailleurs. Les travailleurs excédentaires vont devoir chercher un emploi dans les branches non contrôlées où le salaire va avoir tendance à baisser sous l'effet de cette augmentation de l'offre de travail (la baisse du salaire est en partie atténuée par l'accroissement du capital dont bénéficient les branches non contrôlées). La hausse du salaire dans les branches contrôlées se traduit donc finalement par une baisse du salaire dans les branches non contrôlées et par une distorsion de la structure de production par rapport à celle initialement choisie par les consommateurs. C'est ainsi, par exemple, que les travailleurs des branches syndiquées peuvent obtenir des hausses de salaires au détriment des travailleurs des branches non syndiquées. Si le gouvernement souhaite éviter les licenciements dans les branches contrôlées, il doit effectuer des interventions supplémentaires.

Dans le cas où le salaire minimum est appliqué dans l'ensemble des branches, la quantité de travail offerte pour ce prix dépasse la quantité demandée. Les travailleurs surnuméraires ne peuvent pas trouver d'emplois dans d'autres branches puisqu'elles sont toutes contrôlées, et il apparaît un « chômage institutionnel » qui peut être « massif » et « permanent » (von Mises 1985 [1949], p. 809-810). La diminution de la quantité de travail fait baisser la production et donc aussi le niveau de vie moyen, toutes choses égales par ailleurs. Le salaire minimum affecte surtout le travail peu qualifié, et introduit là encore une distorsion de la production puisque les branches qui utilisent une plus grande proportion de cette qualité de travail subiront une plus forte réduction relative de leur activité.

Si le gouvernement, non seulement impose un salaire minimum à l'ensemble de l'économie, mais en outre interdit les licenciements, alors le salaire nominal va s'élever au détriment des revenus d'intérêt. Les capitalistes seront dissuadés d'épargner et d'investir par la baisse de leur revenu : ils consommeront leur capital, d'où un affaiblissement des capacités productives et finalement là encore une baisse du niveau de vie moyen. (Il faudrait aussi tenir compte ici des effets dynamiques d'appauvrissement :

l'augmentation administrative des salaires réduit les profits entrepreneurs, ce qui ralentit voire bloque les ajustements productifs.)

Il existe néanmoins une solution pour le gouvernement, qui consiste à dissoudre la hausse du salaire réel dans l'inflation : les effets néfastes sur l'emploi d'un salaire minimum peuvent être annulés grâce à un accroissement suffisant de la quantité de monnaie et donc des prix, ce qui réduit la valeur réelle du salaire plancher nominal et fait disparaître le chômage involontaire. C'est d'ailleurs là selon Hayek (1959, p. 282) l'une des explications des politiques inflationnistes menées au cours du xx^e siècle : lorsque les syndicats parvenaient d'une façon ou d'une autre à obtenir des augmentations généralisées de salaire, les gouvernements créaient de la monnaie en quantité suffisante pour compenser cette hausse et éviter l'apparition d'un chômage de masse.

Pour les Autrichiens, l'intervention gouvernementale ou syndicale ne peut pas augmenter durablement le niveau moyen des salaires réels : seuls en sont capables l'accumulation du capital, le progrès technique et l'intensification de la division du travail.

8.1.6 *La prohibition.* La première forme de contrôle étatique de la production est la prohibition, c'est-à-dire l'interdiction pure et simple de produire un certain bien ou d'utiliser une certaine méthode de production. L'exemple historique le plus célèbre est celui de la prohibition de l'alcool aux États-Unis entre 1920 et 1933. Rothbard (1977 [1970], p. 34) analyse la prohibition des biens sans prendre position du point de vue moral sur les biens visés par l'interdiction. Il conclut que ce type de contrôle défavorise surtout les consommateurs, qui ne peuvent plus se procurer un bien qu'ils désirent, sauf au marché noir. Mais le prix est alors beaucoup plus élevé à cause de la prime de risque du producteur et à cause de la perte d'efficacité productive due à l'impossibilité de recourir à la production et la distribution de masse (voir aussi Thornton 1991). Selon Rothbard, les seuls vrais gagnants sont les agents gouvernementaux rémunérés pour faire respecter l'interdiction.

8.1.7 *Le privilège de monopole ou de quasi-monopole.* La seconde forme de contrôle de la production est l'octroi d'un privilège monopolistique d'État. Ce privilège autorise la production mais limite

d'une façon ou d'une autre le nombre d'acteurs qui ont le droit de l'effectuer. Si ce nombre est réduit à un seul producteur, il s'agit d'un monopole, et sinon d'un « quasi-monopole ». Les conséquences sont similaires à celles de la prohibition, mais moins radicales. Les consommateurs sont perdants puisque la production du bien monopolisé est plus faible qu'ils ne le souhaiteraient. Les producteurs qui auraient décidé d'entrer sur ce marché mais ne le peuvent pas à cause de la restriction de concurrence, sont eux aussi perdants puisqu'ils doivent se diriger vers d'autres marchés et se contenter d'une plus faible rémunération. La structure de production s'en trouve altérée, comme dans le cas de la prohibition et du contrôle des prix. Les gagnants sont les producteurs privilégiés ainsi que les agents gouvernementaux chargés d'élaborer la réglementation et de contrôler sa mise en œuvre.

Si la demande à laquelle fait face le monopole est inélastique, alors il peut fixer un prix de monopole (voir § 3.1.4) supérieur au prix concurrentiel. Cet écart de prix constitue le gain de monopole, imputé au privilège monopolistique (Rothbard 1977 [1970], p. 37-39). Et si le gouvernement veut limiter ce gain en contrôlant le taux de rentabilité du monopoleur, il est très facile pour ce dernier d'accroître ses coûts au lieu de baisser son prix. Dans le cas où le quasi-monopole est conféré à un facteur originaire (travail ou ressource naturelle), l'offre de ce facteur va tendre à décroître, et ses propriétaires vont bénéficier d'un prix « restrictionniste » (issu de la politique gouvernementale de restriction) supérieur au prix pratiqué sur un marché non entravé.

8.1.8 *Les types de monopoles et quasi-monopoles d'État.* Von Mises (1998 [1944]) énumère les types suivants :

- les licences (permis spéciaux pour vendre certains biens : services de taxi, médicaments, etc.),
- les lois sur la propriété intellectuelle (marques déposées, copyrights, brevets : voir plus bas),
- les cartels nationaux protégés par des tarifs douaniers,
- les cartels internationaux (l'exemple typique est aujourd'hui celui de l'OPEP),
- les restrictions à la concurrence (visant à protéger les petits commerçants contre les grands magasins, les vendeurs locaux contre les entreprises de vente à distance, etc.).

Rothbard (1977 [1970], p. 41-80) ajoute, entre autres, à cette liste les types suivants :

- les normes de qualité et de sécurité ; pour von Mises, les certifications étatiques (comme par exemple les diplômes de médecine) sont des protections légitimes instaurées par l'État pour défendre le public contre les charlatans ; pour Rothbard, ces certifications constituent des entraves monopolistiques à la concurrence, qui interdisent aux consommateurs de bénéficier de services médicaux (par exemple) moins chers pour des maux moins graves, et leur interdisent aussi de recourir à un autre type de médecine que celui approuvée par l'État,
- les lois de restriction à l'immigration (quasi-monopole conféré aux travailleurs à l'intérieur des frontières d'un pays),
- les lois sur le travail des enfants ou sur l'obligation de scolarité jusqu'à un certain âge (quasi-monopole conféré aux travailleurs adultes ou à ceux dépassant l'âge de la scolarité obligatoire),
- la conscription (qui soustrait du marché du travail une partie des travailleurs adultes),
- le salaire minimum et les allocations chômage (qui maintiennent hors du marché du travail une partie des travailleurs parmi ceux dont la productivité est la plus faible),
- les lois anti-trust (voir paragraphe suivant),
- les lois de conservation (qui restreignent l'exploitation des ressources naturelles non renouvelables et confèrent un quasi-monopole aux propriétaires de celles qui sont exploitées ; lorsque, comme c'est le cas aux États-Unis, le gouvernement délimite de vastes parcs naturels publics, il raréfie la place au sol et fait monter le prix de la terre restant disponible pour produire et bâtir).

8.1.9 *Les lois anti-trust*. Les lois anti-trust sont censées protéger les consommateurs contre les abus qui pourraient être commis par les entreprises de très grande taille. Von Mises (1981 [1922], p. 326) remarque que la question des trusts est souvent abordée de façon émotionnelle, même par les économistes, dans la crainte d'un pouvoir monopolistique croissant exercé par ces grandes organisations. Il admet l'existence des prix de monopole, et donc la possibilité pour les monopoles ou les cartels faisant face à une demande inélastique de restreindre leur production et d'élever leur prix au détriment des consommateurs. Mais il ajoute que cette si-

tuation ne peut survenir que dans le cas rare de la monopolisation ou cartellisation d'une ressource naturelle, comme le diamant à son époque ou le pétrole aujourd'hui (1985 [1949], p. 390). La crainte d'un pouvoir excessif des grandes entreprises manufacturières se dissipe dès lors que l'on comprend que dans une économie de marché il n'y a pas de tendance à la multiplication, mais au contraire à la raréfaction des prix de monopole (voir § 3.1.8). La toute première justification des lois anti-trust, à savoir la restriction du commerce (*restraint of trade*) du *Sherman Act* de 1890, s'avère donc peu convaincante. Et elle l'est d'autant moins que l'histoire économique montre que lors de l'une des condamnations les plus symboliques de l'anti-trust, celle de la *Standard Oil* en 1909, l'entreprise produisait des quantités de pétrole de plus en plus grandes à des prix de plus en plus faibles, et ne menait donc absolument pas une politique restrictive (Armentano 1996 [1982], p. 67).

Rothbard (1977 [1970], p. 60) conteste d'autres arguments qui sous-tendent les législations anti-trust.

1^{er} argument : lorsque les entreprises sont trop grandes ou trop peu nombreuses dans un secteur, cela « réduit la concurrence ». Cet argument suppose que l'on peut mesurer un degré de concurrence. Or, la concurrence n'est pas une quantité mais un processus dans le déroulement duquel les critères de la taille des entreprises ou de la concentration du marché sont dénués de pertinence : un secteur peut être concurrentiel alors qu'il ne se compose que d'une seule ou de quelques grosses entreprises (c'est l'exemple de la *Standard Oil* au début du XX^e siècle, et celui des producteurs de processeurs de micro-ordinateurs ou d'écrans plats aujourd'hui). Reisman (1996, p. 376) ajoute que la grande taille des entreprises dans certains secteurs ne doit pas être considérée comme une barrière à l'entrée, mais au contraire comme un signe que le processus concurrentiel est bien à l'œuvre : cette taille indique que de très grandes quantités de capital doivent être rassemblées afin de réaliser une baisse des coûts unitaires suffisamment forte pour pouvoir rivaliser avec les entreprises en place.

2^e argument : les ententes, collusions ou cartellisations sont « anticoncurrentielles ». Or, ces formes d'organisation consistent à rassembler des ressources productives sous une autorité et une volonté uniques. De ce point de vue, elles sont similaires à la cons-

titution d'une entreprise. Toute législation s'opposant aux ententes, collusions et cartellisations en tant que telles devrait logiquement aussi se retourner contre la constitution d'entreprises.

3^e argument : une entreprise suffisamment puissante pourrait pratiquer une concurrence « sauvage » (*dumping*) en augmentant sa production et en abaissant ses prix le temps de ruiner ses concurrentes de plus petite taille, puis, une fois débarrassée de ces dernières, exploiterait sa clientèle en restreignant sa production et en augmentant son prix de vente bien au-delà du prix initial. Ce scénario est très peu vraisemblable. Non seulement il suppose que la demande soit inélastique, mais surtout il coûterait très cher à l'entreprise pendant la phase où elle ferait du *dumping* : les entreprises concurrentes pourraient attendre, en minimisant leurs coûts de fonctionnement, que la grande entreprise finisse par s'essouffler. Et même en supposant qu'une entreprise parvienne à monopoliser ainsi un marché, la menace d'entrants potentiels la retiendra de proposer un prix trop élevé : son intérêt est plutôt de dissuader les concurrents potentiels en proposant un prix qui ne lui laisse qu'une marge raisonnable, insuffisante pour attirer des rivaux (Rothbard 1962, p. 600).

Du point de vue autrichien, les lois anti-trust s'appuient sur des modèles profondément insatisfaisants de la concurrence, et en tirent donc des conclusions inefficaces – voire nocives – en matière de politique de la concurrence (Armentano 1996 [1982]). Le modèle de concurrence parfaite (critiqué ci-dessus, § 3.2.8) et l'approche structuraliste en théorie de l'organisation industrielle (structure-conduite-performance) reposent sur une conception *statique* et *restreinte* de la concurrence. Bien qu'ils sous-tendent plus ou moins explicitement toutes ces législations, ces modèles ne permettent pas d'évaluer une réalité économique qui est celle d'un processus concurrentiel dynamique et généralisé. Tous les types de phénomènes qui révèlent une inefficacité productive du point de vue statique peuvent être interprétés, lorsqu'on les envisage d'un point de vue dynamique, comme des signes d'un vigoureux processus concurrentiel. Si par exemple on observe une corrélation entre la concentration et le taux de profitabilité sur les marchés, cela signifie que les entreprises plus profitables sont en train de gagner des parts au détriment de celles qui le sont moins. En cherchant à corriger ces soi-disant inefficacités, les lois anti-trust, et

plus généralement les politiques de la concurrence, vont pénaliser les entreprises qui réussissent le mieux. Rothbard et Armentano en concluent, de façon assez surprenante, que les lois anti-trust ont en réalité des effets *anticoncurrentiels* : elles protègent les producteurs moins efficaces contre leurs concurrents plus efficaces, ou les plus nombreux moins efficaces contre les moins nombreux plus efficaces, et constituent donc un quasi-monopole gouvernemental en faveur des premiers et à l'encontre des seconds.

8.1.10 *La propriété intellectuelle*. Les lois sur la propriété intellectuelle entrent dans quatre catégories différentes : les marques déposées®, les copyrights©, les brevets et les secrets de fabrication (Kinsella 2008 [2001]). Seules les trois premières vont être évoquées ici. La question qui se pose est la suivante : ces lois constituent-elles des interventions de l'État dans le marché, ou bien font-elles intégralement partie du système des droits de propriété privée qui caractérise une économie de marché libre (non entravée) ? Les réponses apportées par les auteurs de l'école autrichienne sont très diverses. À un extrême, von Mises (1985 [1949], p. 388) considère la propriété intellectuelle sous ses différentes formes comme un privilège monopolistique restrictif (une licence). À l'autre extrême, Reisman (1996, p. 389) la considère comme faisant partie intégrante d'un système purement capitaliste : les brevets n'ont selon lui pas le caractère d'un monopole gouvernemental restrictif, bien au contraire, puisqu'ils incitent à l'invention productive qui conduit à la multiplication de la production et donc à la réduction des prix.

Von Mises considère les lois sur la propriété intellectuelle – marques déposées, copyrights, brevets – comme des licences monopolistiques, mais il refuse de se prononcer sur leur bien-fondé. En effet, d'un côté ces lois favorisent certaines inventions et certaines créations artistiques, mais de l'autre elles permettent aux inventeurs ou aux propriétaires intellectuels de bénéficier d'un prix de monopole. Il lui paraît donc impossible de pouvoir les justifier ou les condamner d'un point de vue strictement scientifique. Il se contente de dire que ces institutions posent un problème de définition et d'extension des droits de propriété (1985 [1949], p. 697).

Rothbard (1962, p. 592) exclut d'emblée les marques déposées de la catégorie des monopoles étatiques. Pour lui, le nom d'un in-

dividu ou d'une entreprise définit son identité et constitue donc sa propriété privée. En protégeant l'utilisation d'un nom ou d'une marque, l'État n'entrave pas la concurrence, il ne fait qu'exercer son activité de défense des droits de propriété et de protection contre les faux, en l'occurrence contre les falsifications de leur identité dont les producteurs ou leurs acheteurs pourraient être victimes. De même, les copyrights, lui paraissent constituer un droit de propriété tout à fait légitime dans une économie de marché libre : l'acheteur d'une œuvre sous copyright (un livre, par exemple) s'engage *par contrat* à ne pas la reproduire ; s'il viole les termes de ce contrat, il commet selon Rothbard un vol implicite. L'institution du brevet lui semble en revanche contrevenir aux règles du marché libre pour la raison suivante : si une invention est brevetée et si quelqu'un d'autre fait, indépendamment, la même découverte, alors le second inventeur n'aura pas le droit d'utiliser sa propre invention. Il y a donc là un privilège monopolistique accordé au premier inventeur, qui n'existe pas dans le cas du copyright. L'institution du brevet ne peut être légitimée, dans un ordre capitaliste, que si elle s'inspire de celle du copyright. En outre, ce copyright ne devrait avoir aucune limite dans le temps : il devrait être *perpétuel* puisque toute durée finie – par exemple jusqu'à 70 ans après la mort du créateur ou de l'inventeur – serait purement arbitraire (Rothbard 1962, p. 654-657).

Beaucoup plus récemment, Kinsella (2008 [2001]) est revenu à la conception de von Mises selon laquelle le brevet et le copyright sont tous deux des types de privilèges monopolistiques gouvernementaux. Mais alors que ce dernier refusait de se prononcer pour ou contre ces deux institutions, Kinsella élabore une critique approfondie des différents types d'arguments qui sont communément avancés en leur faveur, entre autres l'argument utilitariste (employé par Reisman) et l'argument du droit contractuel (de Rothbard). Dans cette perspective, la propriété intellectuelle de type copyright et brevet est bien une intervention de l'État dans le marché, une intervention à la fois inefficace et injuste (injuste, entre autres, parce qu'elle contredit le principe libéral fondamental d'appropriation des ressources rares qui est celui de « l'appropriation originelle » : Hoppe 1989, p. 17).

8.1.11 *Les externalités*. Une externalité est une situation dans la-

quelle l'action d'un agent a des conséquences non voulues – appelées « effets externes » – sur un autre agent. Si mon voisin fait pousser de beaux parterres de fleurs, alors le plaisir que me procure la vue de son jardin est un effet externe, à condition bien sûr que cette ornementation n'ait pas été faite dans le but de me plaire. Lorsqu'une externalité est favorable, comme dans cet exemple, elle est dite « positive », et « négative » dans le cas contraire.

L'existence des externalités est souvent utilisée pour justifier l'intervention de l'État, sur la base d'une argumentation du type suivant : dans le cas où un effet externe bénéficie à la population, il est souhaitable que l'État mette en place un système d'incitations qui favorise les actions donnant naissance à ces externalités positives ; ces incitations peuvent, entre autres, consister en une redistribution par l'impôt au profit des actions qui génèrent ces bienfaits. Rothbard (1962, p. 889) adresse une critique de principe à ce type de politique. Cet interventionnisme consiste à obliger les gens à payer pour des services qu'ils n'achèteraient pas s'ils étaient libres de leurs choix. Je prends plaisir à regarder le beau jardin de mon voisin, mais cela ne suffit pas à légitimer une taxe que l'État prélèverait sur moi et lui reverserait pour qu'il entretienne ses massifs de fleurs. Deux autres difficultés de principe peuvent être indiquées (Block 1996, p. 347) : tout d'abord, les externalités positives sont tellement nombreuses que si l'État devait intervenir pour chacune d'elles il ouvrirait une véritable boîte de Pandore (Block 1983, p. 2, prend l'exemple amusant du port des chaussettes) ; et ensuite, la subjectivité des préférences fait que ce qui est une externalité positive pour quelqu'un peut être une externalité négative pour quelqu'un d'autre.

La principale intervention de l'État censée se justifier par des externalités positives est celle du financement public de l'éducation. L'éducation est en effet supposée diffuser des bénéfices sociaux qui vont bien au-delà de l'instruction acquise par l'individu pour lui-même : réduction de la délinquance, promotion de la cohésion sociale, participation à la citoyenneté, amélioration de la vie démocratique, accélération de la croissance économique. Un financement par l'impôt des établissements scolaires et universitaires – et donc un accès « gratuit » – est-il pour autant justifié ? West (1994 [1965]) apporte une réponse négative et offre une critique détaillée des divers arguments visant à légitimer par des ex-

ternalités positives le financement public du système éducatif. Très brièvement, deux de ses principaux contre-arguments sont : (1) que l'éducation ne se réduit pas, loin s'en faut, à l'instruction formelle prodiguée dans les établissements scolaires, et (2) que nous n'avons pas de raison de penser que les familles sous-investiraient dans l'instruction scolaire de leur progéniture en l'absence d'un service public d'éducation.

8.1.12 *La critique misésienne de l'interventionnisme.* L'analyse misésienne du contrôle des prix et de la production conduit à trois caractéristiques générales de l'interventionnisme gouvernemental :

- il est *appauvrissant* puisqu'il tend à réduire le niveau de vie moyen ; certains acteurs peuvent bien sûr prospérer et s'enrichir grâce à l'interventionnisme, mais c'est au détriment de la majorité, et au prix d'une baisse de la production globale ou d'une inadaptation des offres aux demandes de biens ;

- il est *cumulatif* en ce sens qu'il se nourrit de lui-même ; chaque intervention génère des effets pervers qui requièrent pour être corrigés des interventions supplémentaires, qui requièrent à leur tour d'autres interventions, et ainsi de suite ; le gouvernement peut décider de renoncer à l'intervention initiale, mais s'il souhaite absolument la maintenir, il va lui falloir multiplier les interventions, ce qui peut mener – comme le montre l'exemple du contrôle des prix – à un contrôle total (collectiviste) de la production ;

- et surtout, l'interventionnisme est *illogique* puisqu'il conduit à des conséquences qui sont exactement contraires aux buts visés ; il fait disparaître les biens de première nécessité qu'il voulait au contraire rendre plus accessibles, il diminue le niveau de vie alors que son objectif était de l'élever, il provoque des récessions suite aux expansions du crédit alors que son intention était de dynamiser la production et l'emploi, etc.

Von Mises insiste beaucoup sur ce dernier aspect, l'illogisme, parce qu'il lui permet d'adresser à l'interventionnisme une critique purement scientifique, indépendante de tout jugement de valeur. Il conclut (1977 [1929], p. 97) qu'il n'y a pas de troisième voie entre capitalisme et collectivisme (*tercium non datur*, selon l'expression latine) : soit les interventions s'intensifient et se multiplient, et le système économique se dirige vers le collectivisme, soit elles sont levées et le système retourne vers le capitalisme de laissez-faire

(Ikeda 1997 propose une élaboration de cette thèse).

8.1.13 *Libéralisme classique, néo-libéralisme, anarcho-capitalisme.* La démarche de von Mises et ses conclusions se rattachent au libéralisme classique de laissez-faire (État minimum ou « minarchie »), mais on trouve aussi dans l'école autrichienne des conceptions qui vont jusqu'à l'anarchisme et au refus de tout État, par exemple chez Rothbard (1977 [1970]) et Hoppe (1989). Dans la direction opposée, la conception néo-libérale accepte certaines interventions étatiques censées améliorer le fonctionnement du marché ou remédier à certaines de ses carences. Hayek (1985 [1944], p. 34, p. 90), par exemple, prône toute une série de régulations s'étendant à de nombreux domaines comme ceux du travail (limitation légale du nombre d'heures travaillées), de l'aide sociale (un minimum de subsistance garanti à tous), de l'assurance (assistance publique contre les risques de catastrophes naturelles), de la politique monétaire, ou encore de la politique de la concurrence. Hoppe (1994) et Block (1996) offrent des présentations détaillées – et critiques – de cet interventionnisme hayékien. Von Mises occupe une position intermédiaire entre les anarchistes (pro-capitalistes) d'un côté, qui sont hostiles à l'existence même de l'État, et les néo-libéraux de l'autre, qui pensent qu'un certain nombre d'interventions gouvernementales sont utiles pour faire fonctionner le marché plus efficacement (Hülsmann 2007, p. 708, p. 857).

8.2 Impôts et dépenses de l'État

8.2.1 *Nature des revenus et des dépenses de l'État.* L'État peut tirer ses revenus de quatre sources qui sont l'impôt, l'emprunt, la production et l'inflation (via la Banque centrale). Seul l'impôt, qui est en général la source la plus importante, sera évoqué ici. Il consiste en un *prélèvement forcé* de richesses sur la population (dans de très rares cas, le paiement de l'impôt peut être volontaire). Le produit de cet impôt est ensuite dépensé par l'État, soit sous forme de dons, soit sous forme d'échanges. Les dons sont des allocations ou des subventions. Les échanges servent à rémunérer des services de travail (par exemple, à payer les traitements des fonctionnaires)

ou à acheter des biens produits par des entreprises privées (par exemple, des avions militaires).

Dans tous ces cas, les dépenses de l'État sont des *dépenses de consommation* et non pas des dépenses d'investissement. En effet, une dépense monétaire constitue un investissement lorsqu'elle est faite dans le but de produire un bien dont la vente couvrira cette dépense initiale (plus les intérêts). Elle constitue une consommation dans le cas contraire, c'est-à-dire lorsqu'elle est effectuée pour acheter des biens qui ne seront pas revendus (ou dont la vente n'est pas censée couvrir les coûts de production). Ainsi, les dépenses publiques consacrées à rémunérer des fonctionnaires ou à construire des édifices publics (routes, écoles, hôpitaux, etc.), sont des dépenses de consommation. Dans cette perspective, l'État n'est pas un investisseur mais un consommateur (Rothbard 1977 [1970], p. 86, Reisman 1996, p. 446, p. 454).

Les économistes autrichiens s'opposent clairement ici à la tradition keynésienne qui prône « l'investissement public », et se rattachent à la tradition classique. Cette dernière distinguait entre la « consommation productive » (investissement) qui a pour but de reproduire avec un profit les richesses détruites dans le processus de production, et la « consommation improductive » qui implique une destruction irrémédiable et sans reproduction des richesses utilisées. James Mill (1826 [1821], p. 246) expliquait très clairement que le gouvernement est un consommateur improductif, non pas au sens où ses activités – par exemple de police et de justice – sont inutiles, mais au sens où ses dépenses ne contribuent pas à remplacer les richesses qu'elles détruisent, et où il ne peut continuer ses activités qu'en récupérant par l'impôt des richesses produites par d'autres agents économiques.

8.2.2 Impôt et redistribution. Ces prélèvements et versements effectués par l'État ont bien sûr des conséquences sur la structure de production. En général, l'impôt est prélevé en de multiples points de la structure, puis dépensé sur des postes très divers. Il est néanmoins intéressant, pour comprendre les effets du processus d'imposition, de traiter le cas simple où le prélèvement et la dépense sont étroitement localisés.

Supposons que le gouvernement prélève un nouvel impôt sur une certaine branche de production d'un bien de consommation *A*

et le dépense en achetant des biens produits par une autre branche *B* (Rothbard 1977 [1970], p. 85). L'impôt accroît les coûts de production des entreprises de la branche *A* et réduit leur taux de rentabilité. Inversement, les revenus supplémentaires dont bénéficient les entreprises de la branche *B* augmentent la demande qui s'adresse à elles, leur prix de vente, et donc aussi leur taux de rentabilité. Les capitaux tendent à être réalloués de la branche *A* vers la branche *B*. L'investissement et donc la production se réduisent du côté de *A* (les entreprises les plus fragiles font faillite, les autres se contractent), et ils s'accroissent au contraire du côté de *B* (les entreprises se développent, et de nouvelles entreprises apparaissent). Les effets se diffusent ensuite vers le haut de la structure de production : les entreprises et branches en amont de *A* tendent à se contracter, et celles en amont de *B* tendent à se développer. Les facteurs originaires spécifiques subissent des pertes de revenus dans la filière de *A* et des gains dans la filière de *B*. Les facteurs non spécifiques, et en particulier le travail, tendent à être réalloués de la filière de *A* vers la filière de *B*. Au final, la branche *A* et toute sa filière subissent l'impôt et se contractent, alors que la branche *B* et toute sa filière bénéficient de l'impôt et s'agrandissent. Ce processus est similaire à celui qui résulte d'un changement des préférences des consommateurs, qui a déjà été décrit (§ 2.2.13), et il donne une raison supplémentaire de considérer l'État comme un consommateur. Des conséquences analogues s'ensuivent si les acteurs sur lesquels l'impôt est prélevé ont des profils de dépense différents de ceux auxquels l'impôt est reversé.

Deux caractéristiques essentielles de l'impôt sont ici bien mises en lumière.

(1) L'impôt opère une *redistribution* des richesses qui peut être complexe mais qui conduit néanmoins à distinguer les *payeurs de l'impôt*, qui subissent à court ou à long terme une perte de revenu, et les *bénéficiaires de l'impôt*, qui obtiennent au contraire un accroissement de revenu.

(2) L'impôt n'est *jamais neutre* puisqu'il opère toujours une distorsion de la structure de production en développant certaines branches ou entreprises au détriment de certaines autres.

8.2.3 *Les types d'impôts.* Rothbard (1977 [1970]) analyse en détail les différents types d'impôts et leurs conséquences. Dans une éco-

nomie monétaire, l'impôt peut être calculé :

- soit à partir des revenus monétaires,
- soit à partir des dépenses monétaires,
- soit à partir des capitaux (c'est-à-dire des valeurs en capital calculées en monnaie).

Les principaux types d'impôts sont ceux qui portent sur les revenus bruts des ventes de détail, sur les différentes sortes de revenus nets (salaires, loyers, intérêts, profits/pertes) et sur les capitaux individuels.

8.2.4 *L'impôt sur les ventes de détail.* Supposons que le gouvernement instaure un nouvel impôt qui prélève 5 % du revenu de l'ensemble des ventes de détail, et qu'il s'en serve par exemple pour construire et entretenir des bâtiments publics. Les entreprises vont-elles répercuter cette taxe de 5 % sur leur prix de vente et la faire payer sous forme de hausse de prix par les consommateurs ? La réponse est non. Les entreprises ont fixé leur prix de vente au niveau qui maximise leur revenu net, et tout mouvement à la hausse ou à la baisse conduirait à une diminution de leurs recettes nettes. Elles n'ont donc pas intérêt à changer leur prix de vente, et il est du reste évident que si elles avaient eu intérêt à augmenter leur prix elles l'auraient déjà fait, sans attendre d'être taxées. Les entreprises qui subissent cet impôt sont dans l'incapacité de le faire payer par les consommateurs. Elles vont devoir le répercuter sur les revenus qu'elles versent aux différents types d'agents économiques, à savoir aux travailleurs, propriétaires de « terre », capitalistes et entrepreneurs de l'étape 1, ainsi qu'aux capitalistes situés en amont dans la structure de production (auxquels elles achètent leurs biens du capital). Cet exemple illustre la *règle de l'incidence* énoncée par Rothbard (1977 [1970], p. 88) : un impôt ne peut jamais être transféré vers l'aval de la structure de production.

Lors de la phase de transition pendant laquelle le nouvel impôt commence à faire sentir ses effets, des capitalistes subissent des pertes en capital (et des gains dans les branches qui bénéficient de la redistribution). Mais si l'on passe sur ce phénomène transitoire qui concerne les profits/pertes, et si l'on suppose pour simplifier qu'aucun effet ne se manifeste sur le taux d'intérêt (une consommation étatique remplace une consommation privée sans modifier

la préférence pour le présent ni la répartition entre consommation et investissement), alors cet impôt sera finalement payé par les propriétaires des facteurs originaires travail et terre, puisque les baisses de prix des biens du capital finissent par être imputées aux salaires et aux loyers. L'impôt sur les ventes de détail est en réalité un *impôt sur le revenu* du travail et de la terre. Les prix des biens de consommation n'ont pas changé (en moyenne), mais les salaires et les loyers ont globalement baissé du montant de l'impôt.

Il faut maintenant tenir compte du fait que le montant taxé est dépensé par l'État, en l'occurrence pour la construction de bâtiments publics. Cette dépense de consommation se substitue à celle des travailleurs et des propriétaires de terre imposés. Elle entraîne dans la branche des travaux publics une hausse des prix qui se diffuse dans la filière en amont, et qui est imputée là aussi aux revenus des facteurs originaires travail et terre. La baisse des salaires et loyers des branches taxées et la hausse du côté de la filière bénéficiaire entraînent une réallocation des facteurs convertibles des premières vers les secondes. Les facteurs originaires spécifiques, en revanche, sont par définition immobiles et subissent une hausse de revenu dans la branche bénéficiaire et sa filière, et une baisse dans la branche taxée et sa filière. Ainsi, l'impôt de 5 % sur les ventes de détail ne va pas du tout se traduire par une réduction de 5 % des loyers et des salaires. Certains revenus originaires vont être affectés plus que proportionnellement, d'autres moins que proportionnellement, et d'autres encore augmenteront sous l'effet de la recomposition de la structure de production. Les gouvernements mettent souvent en place l'impôt sur les ventes de détail pour taxer la consommation, mais en fait ils taxent – et subventionnent – les revenus des propriétaires de facteurs, à des taux variés qui peuvent être très différents du taux initial de 5 %.

8.2.5 *L'impôt sur les revenus nets*. Dans une économie de marché, les quatre types de revenus nets sont les salaires, les loyers, l'intérêt et les profits. Une taxe sur ces types de revenus ne peut être transférée vers l'amont de la structure de production, et elle est donc payée respectivement par les travailleurs, les propriétaires de terre, les capitalistes et les entrepreneurs. En réduisant la rémunération attachée à l'activité correspondante, l'impôt sur le revenu tend à décourager, du côté des payeurs de l'impôt, respectivement

le travail, la location de la « terre », l'épargne et l'entrepreneuriat. Dans le cas (rare) où l'offre de travail de l'individu taxé est décroissante, alors il réagit à la baisse de son revenu en travaillant davantage, c'est-à-dire en sacrifiant du loisir. Deux autres effets généraux peuvent être évoqués. L'impôt sur le revenu, d'une part défavorise les activités payées en monnaie au profit de celles payées en nature (dans la mesure où seuls sont taxés les revenus monétaires), et d'autre part incite les acteurs économiques à ne pas déclarer officiellement leur activité (dans la mesure où seuls les revenus déclarés à l'État sont taxés).

8.2.6 *L'impôt sur le capital.* Supposons que le gouvernement impose une taxe annuelle de 1 % sur le capital. Son effet va être de réduire l'intérêt perçu sur le capital investi, ce qui va pénaliser l'épargne, réduire l'investissement, et conduire à une baisse de la production globale. La satisfaction des consommateurs va être affectée négativement, non seulement à cause de cette baisse de la production, mais aussi parce que ce type d'impôt altère la composition de la production. En effet, plus une branche est capitaliste, c'est-à-dire plus son taux d'équipement en capital est élevé, plus l'investissement dans cette branche sera taxé et donc découragé (Rothbard 1977 [1970], p. 116).

8.2.7 *Niveau et progressivité de l'impôt.* L'impôt entraîne une redistribution des richesses et une recomposition de la structure de production qui sont d'autant plus marquées que le niveau de l'impôt est plus élevé. L'autre question qui se pose est celle du poids de l'impôt entre les contribuables des différentes tranches de revenu. De ce point de vue, le débat classique oppose l'impôt proportionnel à l'impôt progressif.

Lorsque l'impôt est fortement progressif, c'est-à-dire atteint des taux très élevés (par exemple 75 %) sur la tranche la plus haute du revenu, son objectif est de tendre à égaliser les revenus. En effet, les enquêtes statistiques ont toujours montré que, compte tenu du très petit nombre de contribuables concernés, les recettes fiscales obtenues au moyen d'une forte progressivité étaient presque négligeables en pourcentage de l'impôt total (Hayek 1960, p. 312). Un impôt résolument progressif est destiné à mettre en œuvre une politique égalitaire, et non pas à augmenter les ressources de

l'État. Il présente en outre toute une série d'effets pervers énumérés par Hayek (1960, p. 317, p. 320) :

- un effet désincitatif : dans une économie de marché concurrentielle, le revenu monétaire est un indicateur de la valorisation de l'activité correspondante par les consommateurs (principe d'imputation) ; une taxe très disproportionnée sur les hauts revenus dissuade les individus de mener à bien les activités économiques les plus valorisées ;

- un détournement des ressources : si l'impôt est très fortement progressif, non seulement il réduit l'activité globale, mais il détourne les acteurs des activités les plus productives – où ils sont beaucoup plus taxés – pour les inciter à se consacrer plutôt à des activités moins productives et beaucoup moins taxées en pourcentage ;

- une diminution de certains types d'investissements : les investissements de long terme et susceptibles de procurer des gains très importants en un bref laps de temps vont se trouver taxés de façon disproportionnée par rapport à ceux dont le flux de revenus est plus étalé dans le temps (Hayek cite le cas des inventeurs et des artistes qui peuvent subir des années de vaches maigres avant de rencontrer un succès momentané) ; le même problème va se poser pour les investissements particulièrement risqués qui ne se justifient que parce qu'ils peuvent rapporter un revenu exceptionnellement élevé, par rapport aux investissements moins risqués, moins rentables et donc moins taxés ;

- un délitement de la division du travail : supposons qu'un travailleur perçoive un très haut revenu sur lequel il est fortement taxé en pourcentage ; il n'aura plus les moyens d'embaucher des travailleurs pour faire à sa place un certain nombre de tâches plus simples ; à moins que l'écart de revenu soit considérable, il devra accomplir ces tâches lui-même au lieu de se consacrer entièrement à son activité principale hautement productive ;

- une restriction de la formation de nouveaux capitaux et de la concurrence : le progrès économique provient avant tout de la mise en œuvre d'innovations techniques ou organisationnelles ; ces innovations peuvent rapporter dans un premier temps des revenus très élevés ; les nouveaux capitaux requis pour les mettre en place sont pour partie obtenus en réinvestissant ces forts profits initiaux ; or, si ces revenus sont amplement taxés, il sera plus long et plus

difficile d'obtenir les capitaux nécessaires au développement de ces nouvelles techniques ou ces nouveaux biens ; non seulement l'innovation productive sera ralentie, mais les entreprises plus anciennes et déjà en place bénéficieront alors d'une protection quasi-monopolistique contre les nouveaux arrivants plus productifs (von Mises 1985 [1949], p. 851).

8.2.8 *Les subventions étatiques*. Les subventions sont le premier type de dépense étatique. Comme toutes les autres dépenses gouvernementales, elles divisent la société en deux groupes, l'un de bénéficiaires et l'autre de payeurs : une partie des individus va pouvoir se procurer des ressources grâce au *moyen politique*, à savoir la confiscation, plutôt que grâce au *moyen économique*, à savoir la production et l'échange volontaires (Oppenheimer 1914). Comme les non-producteurs sont favorisés aux dépens des producteurs, le groupe des premiers a tendance à s'étendre au détriment du groupe des seconds, toutes choses égales par ailleurs, ce que Rothbard (1977 [1970], p. 172) illustre par les exemples des allocations sociales (subventions à la pauvreté), des allocations chômage (subventions à l'inactivité) et des allocations familiales (subventions à la procréation). Il en résulte une réduction de la production globale, qui survient même lorsque les subventions sont destinées à aider des entreprises privées. En effet, ces transferts bénéficient alors en général à des entreprises en difficulté, et ils ont pour conséquence de maintenir en place des facteurs de production qui sans cela seraient réalloués vers des lignes plus productives.

8.2.9 *Les services étatiques*. Le second type de dépense étatique vise à financer en partie ou en totalité des organismes publics qui fournissent à la population des services divers, de protection (police, armée, pompiers), d'éducation (écoles publiques, etc.), de soins médicaux (hôpitaux publics, etc.), de voirie, de transport collectif, d'acheminement du courrier, d'eau, d'électricité, et ainsi de suite. Les organismes publics se caractérisent souvent par leur inefficacité et par leurs déficits chroniques. Ces problèmes sont la conséquence inévitable de leur nature étatique (von Mises 1944, p. 45 et suivantes, Rothbard 1977 [1970], p. 173 et suivantes).

Mais pourquoi des « entreprises » étatiques seraient-elles

moins efficaces que des entreprises privées ? Ne pourrait-on pas tout simplement s'inspirer de la gestion de ces dernières pour diriger efficacement les organismes étatiques ? Non, car cette solution est inapplicable. Même lorsque le service étatique est payant, l'« entreprise » publique peut toujours recourir à des fonds gouvernementaux prélevés par l'impôt, alors que l'entreprise privée doit être en mesure d'offrir un taux de rentabilité suffisant à des investisseurs capitalistes à la recherche de gains. La possibilité d'obtenir des fonds prélevés par la force plutôt que par l'échange volontaire a toute une série de conséquences négatives en termes d'efficacité : elle réduit l'incitation à l'effort et à l'innovation, et fait passer au second plan le souci des déficits budgétaires et de la satisfaction des usagers (d'où une offre de services uniformisés qui s'adaptent mal à la diversité des préférences individuelles). Les entreprises étatiques jouissent souvent d'un monopole légal qui renforce encore ces différentes tendances, mais même si ce n'est pas le cas elles peuvent se permettre de fixer des prix inférieurs aux coûts, ce qui élimine une grande partie de la concurrence par des entreprises privées et leur confère de fait une situation de privilège monopolistique. En outre, avec la fixation de prix faibles, voire nuls (dans ce dernier cas, les services sont gratuits en ce sens qu'ils ne sont pas directement payés par leurs utilisateurs), la quantité demandée tend à excéder la quantité offerte, d'où des problèmes de file d'attente et de rationnement caractéristiques des marchés sur lesquels les prix sont plafonnés (voir § 8.1.3).

Les entreprises étatiques rencontrent une difficulté plus profonde, qui tient à leur incapacité à procéder au calcul économique et donc à l'impossibilité d'allouer rationnellement les ressources. Une entreprise privée survit si sa rentabilité est suffisante pour attirer des investisseurs, ou plus précisément si le prix que les clients consentent à payer permet de couvrir les coûts de production tout en laissant une marge suffisante pour rémunérer le capital investi. À chaque instant, le calcul économique permet de déterminer le profit réalisé et de calculer le profit futur estimé, pour savoir si – compte tenu des désidératas des acheteurs et en dernière analyse des consommateurs – telle ou telle branche de l'entreprise devrait être agrandie, réduite, ou même supprimée. Or, la situation d'un organisme étatique est tout à fait différente. Dans le cas d'un commissariat de police, par exemple, comment déterminer le

nombre de policiers nécessaires pour faire respecter les lois à cet endroit ? Comment faire en sorte qu'il n'y en ait, ni trop (pour éviter un gâchis de ressources), ni trop peu (pour éviter une dégradation de la qualité du service) ? Et au niveau au-dessus, quelles sommes allouer respectivement au commissariat, aux écoles et à l'entretien des routes de la commune ? Comme le prix de ces services est nul (« gratuité ») et que leurs coûts sont financés par l'impôt, aucun calcul économique n'est possible et l'État ne peut pas déterminer rationnellement – c'est-à-dire du point de vue de la satisfaction des besoins des gens – la répartition des dépenses publiques entre ces différents postes, ni même le montant total à prélever. Dans ces conditions, la décision publique ne peut être que très approximative. Si les fonctionnaires de chaque service pouvaient décider eux-mêmes du montant de leurs propres ressources, ils en demanderaient toujours davantage. C'est pour cette raison que les arbitrages publics se font au plus haut niveau gouvernemental, et même à ce niveau-là constituent des décisions politiques qui sont l'équivalent pour un État de ce que sont pour un individu les décisions de *consommation* – et non pas l'équivalent des décisions d'investissement fondées sur le calcul économique du capital et du revenu net.

8.2.10 *Minimiser la dépense publique.* La théorie autrichienne de l'interventionnisme concluait à un effet *appauvrissant* des interventions de l'État dans le marché. La théorie de l'impôt et de la dépense publique aboutit au même résultat : ces activités étatiques distordent la structure de production, entravent le processus concurrentiel, et dissuadent l'activité productive tout en politisant la société (les talents de l'activisme politique se développent au détriment des capacités productives : Hoppe 1989, p. 55). D'un point de vue macroéconomique, les prélèvements et dépenses de l'État tendent à :

- diminuer le produit total (dans le cas où les dépenses sont consacrées à la redistribution ou au financement de services qui pourraient être fournis par le marché),
- modifier les proportions de biens et services produits par rapport à celles voulues par les consommateurs,
- amoindrir l'épargne-investissement et augmenter les taux d'intérêt (du fait des emprunts publics et des taxes sur le capital et

sur les revenus d'intérêt),

– et réduire la valeur de la monnaie (dans la mesure où les dépenses étatiques sont financées par la création monétaire).

Selon le principe de *neutralité axiologique*, il est impossible de déduire logiquement une préconisation à partir d'un raisonnement strictement factuel ou scientifique. Il est néanmoins assez facile de comprendre l'hostilité générale des économistes de l'école autrichienne vis-à-vis de tout excès de dépense étatique. Bien entendu, la définition de ce qu'est un « excès » dépend des positions des uns et des autres en matière de philosophie politique : pour les anarcho-capitalistes, favorables à la suppression de l'État, la moindre dépense est déjà excessive ; pour les minarchistes, seules se justifient les dépenses nécessaires pour entretenir un État minimal, c'est-à-dire les dépenses de police, de justice et d'armée ; pour les néolibéraux enfin, certaines dépenses étatiques supplémentaires sont jugées souhaitables (voir § 8.1.13).

8.3 Le collectivisme

8.3.1 *L'État collectiviste*. Alors que le capitalisme se définit par la propriété privée des facteurs de production, le collectivisme se caractérise au contraire par la propriété collective de ces facteurs (terres agricoles, gisements, usines, entrepôts, etc.). L'appropriation collective ne signifie évidemment pas que chaque individu peut s'approprier les facteurs qu'il veut pour les utiliser comme il l'entend : il en résulterait des conflits incessants et insurmontables. Non, elle signifie que l'utilisation de l'ensemble des facteurs de production est confiée à une institution unique, une autorité économique suprême que l'on peut appeler « État » et qui est chargée d'organiser – de planifier – la production au service du « bien commun ». Les questions concernant la nature de ce « bien commun » et la constitution de cette autorité suprême (élection démocratique ou dictature) relèvent de la morale et de la politique, et n'entrent pas dans le champ de l'investigation économique (von Mises 1981 [1922], p. 112). La question de savoir si les travailleurs pourraient choisir leur occupation ou seraient assignés de force à un poste de travail, sera elle aussi laissée de côté.

Pour les partisans d'un régime collectiviste, le capitalisme

souffre de deux défauts majeurs : tout d'abord, il constitue un système « anarchique » dans lequel les producteurs prennent leurs décisions indépendamment les uns des autres et se trouvent engagés dans une concurrence destructrice, et ensuite, il ne fait qu'exprimer et même exacerber les égoïsmes individuels. Le collectivisme peut alors apparaître comme une bien meilleure solution, puisqu'il remplace l'anarchie de la production par un plan gouvernemental unique et cohérent, et substitue aux égoïsmes individuels la poursuite des objectifs définis par un État supposé soucieux de l'intérêt général. Un régime collectiviste présente donc a priori un indéniable attrait. Mais peut-il fonctionner ? Von Mises (1935 [1920]) développe une théorie du calcul économique qui le conduit à apporter à cette question une réponse négative.

8.3.2 *Le problème économique.* Un système économique développé se compose d'une multitude de facteurs de production, d'un très grand nombre de techniques plus ou moins indirectes permettant de transformer ces facteurs en biens de consommation finale, et d'une foule d'individus dont les préférences sont extrêmement diverses. Le problème économique fondamental est le suivant : comment utiliser ces facteurs de production pour faire en sorte de satisfaire au mieux les besoins des gens ? En d'autres termes : comment les facteurs devraient-ils être répartis entre les différentes lignes de production pour permettre aux gens d'atteindre les fins qu'ils considèrent comme les plus importantes ? Ou encore : comment éviter autant que possible le gâchis de facteurs consistant à produire des biens jugés moins importants, alors que d'autres jugés plus importants par les individus concernés auraient pu être produits à la place ? En bref, le problème économique est celui de l'allocation des ressources rares à usages alternatifs en vue de fins préalablement choisies (Robbins 1947 [1932], p. 30).

Il est essentiel de noter que ce problème n'est pas susceptible d'une solution purement technique, sauf dans certains cas très simples ou très irréalistes. Une solution technique serait par exemple envisageable si les facteurs étaient tous purement spécifiques ou s'ils étaient tous parfaitement substituables les uns aux autres (si, en d'autres termes, il n'y avait au fond qu'un seul facteur de production). Comme, en réalité, la plupart des facteurs sont convertibles et les possibilités de substitution plus ou moins limi-

tées, un *instrument de calcul économique* est indispensable pour effectuer une répartition rationnelle des ressources productives.

8.3.3 *Calcul économique et prix de marché.* Dans un régime capitaliste, le calcul économique est rendu possible par le système des prix de marché exprimés en monnaie (voir § 6.1.6). En effet, ces prix permettent à chaque producteur d'évaluer numériquement la profitabilité anticipée de son projet. Si cette rentabilité est insuffisante, cela signifie que le prix de vente prévu ne couvre pas les coûts unitaires (compte tenu de l'intérêt sur le capital investi). Pourquoi les coûts sont-ils trop élevés ? Parce que d'autres producteurs surenchérisent pour se procurer les facteurs, et en font monter le prix : ils anticipent une demande suffisamment forte de la part de leurs clients, et en fin de compte de la part des consommateurs finaux qui sont à l'origine de l'ensemble de ce processus d'imputation. Tous ces calculs reposent sur des prix et des demandes futurs anticipés face à l'incertitude radicale de l'avenir. Les producteurs vont parfois s'apercevoir après coup qu'ils se sont trompés. Ceci ne remet pas du tout en cause l'utilité du système des prix de marché, puisque c'est uniquement grâce à lui que les producteurs, d'une part sont en mesure de prendre leur décision initiale, et d'autre part se rendent compte de leurs erreurs s'ils en commettent et peuvent les rectifier dans le cours du processus de marché.

8.3.4 *L'impossibilité du calcul économique en régime collectiviste.* Dans un système économique collectiviste, par définition, l'État est le seul propriétaire des facteurs de production. Le système des échanges et l'utilisation d'une monnaie, à supposer qu'il y ait l'un et l'autre, sont strictement limités aux biens de consommation. Compte tenu de la diversité des préférences, il est en effet souhaitable de laisser les individus échanger ce type de biens, de façon à ce qu'ils puissent accroître leur satisfaction en procédant à des transactions volontaires. Les facteurs de production, en revanche, sont la propriété d'un seul acteur, l'État, et ils ne peuvent donc pas faire l'objet d'échanges – puisque cela supposerait l'existence d'au moins *deux* propriétaires distincts. Or, en l'absence d'échange, il ne peut évidemment pas apparaître le moindre prix pour les facteurs de production. Le calcul économique est donc *impossible*.

L'objection adressée par von Mises aux partisans du collectivisme est au fond très simple : le collectivisme ne peut pas fonctionner parce qu'il ne dispose d'aucun moyen de calcul économique. Les dirigeants d'un tel système se trouvent dans l'incapacité de résoudre le problème économique fondamental qui est celui de l'allocation rationnelle des facteurs de production en vue de satisfaire au mieux les fins visées.

Supposons que l'État collectiviste souhaite bâtir une usine d'automobiles. Tout d'abord, comment savoir si les facteurs de production qui vont être consacrés à la construction de cette usine puis à la fabrication des véhicules ne devraient pas plutôt être employés pour satisfaire d'autres besoins, que les gens estiment plus importants ? Pourquoi placer l'usine ici plutôt qu'ailleurs ? Quelle taille et quelle durabilité lui attribuer ? Pourquoi la construire avec tels matériaux plutôt que tels autres ? Pourquoi implémenter telle méthode de fabrication des voitures plutôt que telle autre ? Quel type de métal ou d'une autre matière utiliser pour fabriquer telle ou telle pièce des véhicules ? Alors que les prix de marché apportent une réponse rationnelle à ces questions et de nombreuses autres, l'État collectiviste, étant privé d'instrument de calcul, en sera réduit à prendre des décisions de production arbitraires. Von Mises en conclut qu'un système collectivisé ne pourra jamais approcher, même de très loin, l'efficacité productive d'un système capitaliste, et pourrait même conduire à la désintégration de la division du travail et à un appauvrissement considérable (1981 [1922], p. 105, p. 117).

8.3.5 Critique de l'argument historique. La première objection, et la plus évidente, qui peut être adressée à von Mises est que des régimes collectivistes ont existé dans l'histoire – et existent encore, en très petit nombre, aujourd'hui –, l'exemple le plus emblématique étant celui de l'ex-Union soviétique. Pour exister, il a bien fallu qu'ils fonctionnent d'une façon ou d'une autre. La réponse à cette objection est que l'argument misésien sur le calcul économique doit en toute rigueur être appliqué à un régime collectiviste *isolé*. Or, historiquement, les pays communistes ont coexisté avec des pays à base capitaliste. Ils disposaient ainsi d'informations précieuses sur les prix de marché des divers biens, et en particulier ceux des ressources naturelles affichés sur les

marchés mondiaux. Ils pouvaient aussi se procurer (légalement ou non) les techniques de production utilisées et les nouveaux types de biens produits dans les pays capitalistes. Malgré cette situation avantageuse qui leur permettait dans une certaine mesure de résoudre le problème du calcul économique, il est bien connu que ces régimes étaient affectés de graves dysfonctionnements. Leur production était chaotique et surtout insuffisante, incapable de pourvoir convenablement aux besoins de la population (en dépit d'une dotation considérable de ressources naturelles dans le cas de l'URSS). Ces graves difficultés s'expliquent aisément grâce à la théorie du calcul économique (les prix des pays capitalistes ne pouvaient jamais correspondre exactement à ce qu'ils auraient été dans les pays collectivisés) et grâce à la théorie de l'interventionnisme (un contrôle universel des prix et une vaste redistribution des richesses conduisent à multiplier les rationnements et à démotiver les producteurs).

8.3.6 *Critique de la solution marxiste.* Marx et Engels adhéraient à la théorie de la valeur travail formulée par les Classiques, et ils ont laissé entendre dans certains de leurs écrits qu'il serait possible pour une économie collectivisée de recourir à un calcul économique fondé sur le temps de travail (von Mises 1935 [1920], p. 112, Lavoie 1985, p. 68). L'État pourrait ainsi, connaissant d'une part la quantité totale de travail disponible, et d'autre part les quantités requises pour produire les différents types de biens, affecter les travailleurs aux différents postes de production. Pour von Mises, cette solution marxiste se heurte à deux difficultés insurmontables. D'abord, comme elle s'appuie exclusivement sur le travail, elle ne prend pas en compte l'utilisation des ressources naturelles (facteurs originaires « terre »). Soit deux processus de production qui exigent le même temps de travail, mais utilisent, l'un peu et l'autre beaucoup de ressources naturelles : cette différence très importante entre les deux processus ne sera pas prise en compte dans un calcul qui se limite aux heures de travail. La seconde difficulté de la solution marxiste est qu'elle repose sur des taux de substitution arbitraires entre les différentes qualités de travail. Une heure d'un certain type de travail qualifié vaut davantage qu'une heure de travail non qualifié. Mais quelle est la relation quantitative précise entre les deux ? En l'absence de marchés du

travail, les relations entre les multiples qualités de travail ne sont pas connues : le planificateur ne peut pas correctement intégrer les différentes sortes de travail qualifié dans ses calculs (sauf si l'on suppose que des pays capitalistes avoisinants lui fournissent cette information ; mais les partisans du collectivisme veulent supprimer le capitalisme, et non pas le laisser subsister pour s'en servir comme d'une sorte de base de données).

8.3.7 *Critique de la solution mathématique.* Dans la tradition de l'économie mathématique (Walras 1874), l'équilibre du système économique – et l'allocation optimale des ressources – est conçu comme la solution d'un système d'équations dont les inconnues sont les prix et dont les paramètres sont les préférences individuelles, les stocks de biens disponibles, et les techniques de production des entreprises. Pourquoi ne pas envisager que le planificateur d'un régime collectiviste organise l'économie en résolvant le système des équations walrasiennes de l'équilibre général ? Pour Hayek (1935, p. 208), une telle solution est totalement inapplicable. Elle requiert une connaissance centralisée par l'organisme planificateur : (1) de la totalité des facteurs de production disponibles en les distinguant bien selon leur localisation, leur état d'usure, et toute autre dimension ayant un rapport avec leur capacité productive (la qualité et la quantité des biens complémentaires, par exemple), (2) de l'ensemble des techniques de production dans les domaines les plus variés, et (3) des préférences individuelles entre tous les biens de consommation concevables. Cette masse d'information est, non seulement colossale, mais subit à chaque instant des modifications imprévisibles en fonction des conditions locales de production et de consommation. Il est tout à fait inconcevable de pouvoir résoudre un gigantesque système d'équations dont les paramètres sont extrêmement nombreux et varient dans un délai beaucoup plus court que celui requis pour les transmettre au centre de calcul. La solution mathématique est impraticable : elle surestime de beaucoup les capacités de calcul et d'acquisition d'information d'un centre de planification. Cette critique permet de comprendre, en creux, que le système capitaliste mobilise une quantité phénoménale de connaissances techniques et d'informations qui vont des plus générales aux plus localisées, et dont les changements sont rapidement pris en compte par les ac-

teurs concernés (voir § 2.2.14).

L'argument de von Mises (1985 [1949], p. 747) à l'encontre de la solution mathématique est différent de celui de Hayek. Plutôt que de s'attacher comme ce dernier à la complexité du problème à résoudre, il évoque la difficulté pour le planificateur de découvrir par quel *chemin* le système économique va parvenir, à partir de sa situation présente, jusqu'à la situation d'équilibre général. À l'instant initial t_0 , le système économique se caractérise par une certaine dotation de facteurs de production répartis entre les différentes entreprises. Supposons que le planificateur dispose de toutes les informations nécessaires et résolve les équations de l'équilibre général walrasien. Il sait donc quel est l'état d'équilibre final à atteindre à l'instant futur t_1 . Mais il ignore en revanche totalement comment atteindre cet équilibre final. En effet, la solution du système des équations de l'équilibre général ne lui apprend rien sur le processus qui pourrait permettre de faire passer le système de sa situation initiale à la situation visée. La résolution du problème mathématique, même en supposant qu'elle soit possible, n'indique que l'objectif à atteindre, pas la succession d'étapes qui vont effectivement permettre de l'atteindre. Or, en l'occurrence, il est vain de savoir où l'on va si l'on ignore totalement comment y aller.

8.3.8 « *Déshomogénéiser* » von Mises et Hayek ? Pour von Mises, contrairement à Hayek, le problème fondamental du collectivisme n'est pas un problème d'information, mais bien un problème de calcul économique, c'est-à-dire au fond un problème de *droits de propriété* puisque ce sont les droits de propriété qui donnent naissance à l'échange, qui donne à son tour naissance aux prix, qui rendent possible le calcul économique ; en l'absence de droits de propriété privée des facteurs de production, le calcul économique est donc impossible. On vient de le voir, von Mises affirme que même si les informations détenues par le planificateur central sont supposées parfaites – c'est-à-dire s'il connaît les préférences, les techniques et les facteurs disponibles –, en l'absence d'un moyen de calcul il lui sera impossible de mettre en œuvre une rationalité instrumentale : il ne sera pas en mesure d'appliquer efficacement les moyens dont il dispose aux fins qu'il vise. Pour Hayek, la difficulté majeure qui se pose en régime collectiviste est celle de la *dissémination* des connaissances et des informations à travers le

système économique : si le planificateur pouvait centraliser les informations ainsi dispersées, il pourrait élaborer et appliquer un plan efficace ; mais cette tâche de centralisation est impraticable et un régime collectiviste échoue à mobiliser la masse de connaissances et d'informations sur laquelle peut en revanche s'appuyer un régime capitaliste.

Selon la thèse de la « déshomogénéisation », von Mises et Hayek ont élaboré deux critiques tout à fait distinctes l'une de l'autre, qui soulèvent des problèmes de nature différente, un problème de droits de propriété pour le premier et un problème d'information pour le second. Cette thèse est cependant controversée : elle est défendue par Salerno (1990, 1993), Hoppe (1996) et Hülsmann (1997) entre autres, mais critiquée par Kirzner (1996) et Yeager (1997) qui estiment au contraire que la critique hayékienne se situe dans le prolongement de celle de von Mises.

8.3.9 Critique du système de planification décentralisée. Le calcul en nature est impossible (on ne peut ajouter des pommes et des oranges), la solution marxiste qui consiste à recourir à un calcul en heures de travail donne des résultats faux, la résolution du système des équations de l'équilibre général est à la fois impraticable et vaine. Comment donc venir à bout de ce problème profond et difficile qui est celui du calcul économique ?

À partir des années 1930, les économistes néo-classiques partisans du socialisme imaginent un nouveau type de collectivisme se rapprochant par sa forme du fonctionnement d'un régime capitaliste (Lange 1936, 1937). Leur idée est de recourir à une planification dite « décentralisée ». Le Centre envoie aux directeurs des entreprises un ensemble de prix sur lesquels ils doivent baser leur calcul économique. Chacun calcule son plan de production, compte tenu de ces prix, en faisant en sorte de maximiser son profit. Il en résulte des quantités offertes et des quantités demandées pour chaque bien. Les nombres de ces quantités offertes et demandées sont adressés au Centre qui les ajoute puis calcule le déséquilibre sur chaque marché : si la quantité offerte d'un bien excède sa quantité demandée il baisse son prix, et il l'augmente dans le cas contraire. Il peut ensuite envoyer aux directeurs d'entreprises une liste de prix rectifiés, à partir desquels ils vont devoir réaliser leurs nouveaux plans de production ; la confrontation des quantités of-

fertes et demandées va faire apparaître de nouveaux déséquilibres qui vont à leur tour être corrigés par des changements de prix annoncés par le Centre, et ainsi de suite. Les directeurs d'entreprises ne sont évidemment pas propriétaires de leurs établissements et ne peuvent pas s'approprier les profits : ces derniers sont récupérés puis distribués par l'État.

Cette tentative de solution est paradoxale (von Mises 1985 [1949], p. 743). Le collectivisme et la planification avaient justement pour but à l'origine de supprimer les marchés et la concurrence capitalistes. Or, il est désormais admis – par les économistes néo-classiques – que des sortes de prix de marché sont indispensables, même dans un système qui se veut pleinement socialiste.

La principale critique que Hayek (1940) adresse à cette méthode de planification décentralisée est que, si elle permet *en principe* de résoudre le problème du calcul économique, elle est très loin d'être satisfaisante *en pratique*. Elle pourrait fonctionner si l'économie restait proche d'une situation d'équilibre statique, et donc si le Centre pouvait se contenter d'effectuer des ajustements mineurs. Dans ces conditions, cette procédure par essais et erreurs serait en effet suffisante. Mais en réalité, des chocs dynamiques de toutes sortes et d'ampleur plus ou moins grande surviennent constamment à travers l'ensemble du système. Si une procédure de planification décentralisée est adaptée à un univers statique, elle ne l'est pas du tout à l'univers mouvant, changeant, incertain, qui est celui d'une économie développée et complexe. C'est vraisemblablement pour cette raison que ces procédures n'ont jamais été appliquées en pratique, en dépit des nombreux travaux théoriques dont elles ont fait l'objet de la part des économistes néo-classiques standard.

Von Mises (1981 [1922], p. 105) avait par avance critiqué cette solution en précisant que le problème du calcul économique ne se pose que dans une économie dynamique, soumise dans tous ses secteurs à de fréquents changements non anticipés. Si l'on se place dans l'hypothèse irréaliste d'une économie statique, alors en effet un régime collectiviste efficace est tout à fait concevable. Sous des hypothèses réalistes, en revanche, un tel régime est très inefficace, quel que soit son type de planification – centralisée ou décentralisée.

8.3.10 *La perspective autrichienne sur le débat*. L'article initial de von Mises (1935 [1920]) a déclenché un débat intense et prolifique sur la possibilité du collectivisme. Boettke (2000) en rassemble les principaux textes dans un coffret de neuf volumes ! Mais l'histoire de ce débat est racontée de façon tout à fait différente selon que l'on se place dans la perspective néo-classique standard ou dans la perspective autrichienne (Lavoie 1985).

Dans la perspective néo-classique standard, les économistes autrichiens ont rapidement perdu ce débat face aux arguments néo-classiques. Von Mises prétendait que le calcul économique était « impossible » en régime collectiviste. Or, grâce à un système de planification décentralisée, il s'avère que ce type de calcul est en principe possible : fin du débat, défaite des économistes autrichiens. L'interprétation autrichienne n'est pas du tout la même. Les économistes néo-classiques standard n'ont pas compris la véritable nature de l'argument de von Mises, parce qu'ils raisonnent dans un cadre trop marqué par le concept d'équilibre. Or, la réalité économique n'est jamais celle de l'équilibre final – ce concept n'est qu'un outil théorique qui permet de comprendre par antithèse que le processus de marché se caractérise par de multiples déséquilibres, par des profits et des pertes d'ampleur variable en divers points du système, par l'importance de la fonction entrepreneuriale, etc. Le problème n'est donc pas du tout de savoir si, en principe, le calcul économique est possible ou non en régime collectiviste. Il est de savoir si, concrètement, le planificateur collectiviste disposerait oui ou non d'un outil de calcul économique fiable et précis. La méthode de planification décentralisée repose sur une représentation statique – et donc inadéquate – de l'ordre économique. Elle ne constitue pas une réponse satisfaisante à l'argument misésien. Les économistes autrichiens n'ont pas été défaits, mais plutôt incompris, par les néo-classiques standard.

La critique du collectivisme parachève ainsi la critique autrichienne de l'interventionnisme : non seulement les interventions étatiques diminuent le niveau de vie moyen, mais en se multipliant elles rapprochent peu à peu le système économique d'un régime collectiviste incapable de fonctionner. Ce chapitre et le précédent nous permettent de comprendre que le libéralisme des économistes autrichiens est d'origine, non pas idéologique, mais théorique : leurs *théories scientifiques* conduisent à conclure que

c'est l'État, et non pas le marché libre, qui est à l'origine de l'inflation, des crises, du chômage de masse et de la plupart des monopoles restrictifs.

L'une des caractéristiques les plus marquantes de l'école autrichienne est qu'elle constitue un système *intégré* d'analyse économique : sa théorie de la valeur débouche successivement sur une théorie des prix, sur une théorie de l'intérêt (théorie de l'actualisation), et sur une théorie de la monnaie ; ces dernières donnent à leur tour naissance à une théorie de la structure de production, une théorie du progrès économique, puis enfin à une théorie du cycle. Dans cet édifice successivement érigé par Menger, Böhm-Bawerk et Fetter, puis von Mises et Hayek, chaque nouvel apport s'appuie sur le précédent sans remettre en cause la cohérence de l'ensemble – même si le désaccord sur l'origine de l'intérêt, subjectiviste chez Fetter et von Mises, productiviste chez Böhm-Bawerk et Hayek, peut donner lieu à deux interprétations différentes de cette construction théorique. La figure C.1 montre comment s'articulent ces différents éléments, en suivant plutôt ici l'interprétation issue de Fetter et von Mises.

Le contraste avec l'école néo-classique standard est tout à fait frappant. Car cette dernière, telle qu'elle est par exemple présentée dans le manuel de Mankiw (1998), donne à voir une science économique *éclatée*. Éclatée d'abord entre une microéconomie centrée sur les modèles de concurrence parfaite et monopolistique, et une macroéconomie principalement centrée sur la question de la croissance : il n'y a pas de pont, pas de relation théorique entre ces deux aspects de l'analyse économique standard. Du côté autrichien, cette séparation étanche entre la théorie de la concurrence et la théorie de la croissance n'a pas lieu d'être. La concurrence entrepreneuriale consiste à retirer les facteurs de leurs emplois relativement moins productifs pour les allouer à des usages relativement plus productifs, et la croissance n'est autre que la conséquence au niveau agrégé de toutes ces actions entreprises en parallèle sur les différents marchés. Il est bien sûr toujours possible de distinguer un point de vue « micro » centré sur la concurrence et un point de vue « macro » centré sur la croissance – ou sur le « progrès économique » comme l'appelle von Mises –, mais ces deux aspects sont ici étroitement liés, au point que l'on ne peut pas les concevoir convenablement l'un sans l'autre.

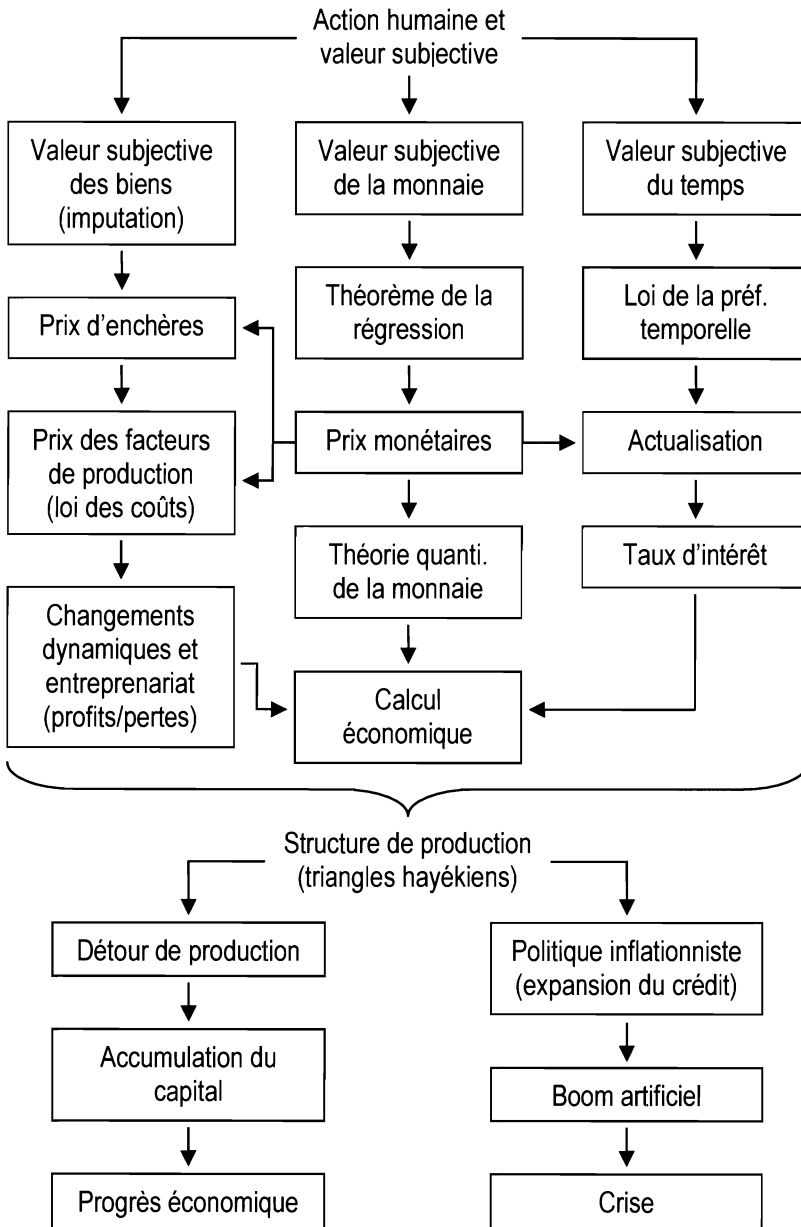


Figure C.1. Un système intégré d'analyse économique

Au sein des deux domaines étanches de la micro- et de la macroéconomie standard, d'autres fractures apparaissent : la microéconomie orthodoxe est elle-même scindée entre les modèles d'équilibre partiel et ceux d'équilibre général. Cette distinction, non seulement n'existe pas dans l'école autrichienne, mais constitue dans ce cadre une grave erreur théorique.

Raisonner en équilibre partiel revient d'une part à négliger l'influence sur l'extérieur de ce qui se passe sur le marché étudié, et d'autre part à considérer les forces extérieures comme des données ou des paramètres dont la nature et l'origine sont dénuées d'importance. Il s'agit là, dans la perspective autrichienne, non pas d'une simplification légitime du point de vue scientifique, mais d'une double erreur de raisonnement. Dès sa fondation par Menger, le paradigme autrichien met l'accent sur le phénomène de l'imputation, qui est la façon dont la valeur et les prix anticipés des produits déterminent la valeur et les prix des facteurs servant à les produire. La loi des coûts, réinterprétée par Böhm-Bawerk, illustre ces relations verticales ascendantes totalement négligées dans les modèles d'équilibre partiel, et montre la nécessité de raisonner sur des marchés interdépendants. L'étude de l'adaptation du système économique aux chocs dynamiques montre elle aussi comment ces chocs – de demande, de ressource, etc. – font apparaître des profits sur certains marchés, compensés par des pertes sur d'autres marchés, et comment ces profits et pertes se propagent vers les marchés placés en amont ou en aval, affectant l'ensemble de l'économie et conduisant en fin de compte à une réallocation des facteurs originaires de production. Raisonner en équilibre partiel rompt le fil de cette analyse causale et empêche de comprendre les aspects les plus importants du fonctionnement du système des prix et de l'économie de marché.

Les modèles d'équilibre général peuvent de ce point de vue sembler plus satisfaisants, puisqu'ils tiennent compte de l'interdépendance généralisée entre les marchés. Mais ils souffrent eux aussi de graves défauts. D'abord, ils se contentent d'expliquer des prix d'équilibre. Or, les prix de marché, tels qu'ils apparaissent dans le cours du processus de concurrence entrepreneuriale du monde réel, ne sont pas des prix d'équilibre : ils s'accompagnent, sur toute une série de marchés, de déséquilibres qui se manifestent par des profits, des pertes, des ruptures de stock ou des invendus.

Ensuite, l'explication de la détermination des prix d'équilibre offerte par les modèles d'équilibre général est très discutable : ces prix sont censés constituer les inconnues d'un système d'équations mathématiques, et être tous déterminés *simultanément* comme les solutions de ce système. Tout se passe comme si le processus entrepreneurial de convergence vers l'équilibre final n'avait aucune influence sur son résultat, sur l'état-limite auquel il conduit. L'analyse typiquement autrichienne de l'enchaînement causal consécutif aux chocs dynamiques permet à la fois d'étudier le principe de détermination des prix réels de marché et d'éviter la métaphore inappropriée du système d'équations : dans le monde réel, il n'y a pas de détermination simultanée de tous les prix d'équilibre. Enfin, les modèles d'équilibre général excluent la monnaie : ils reposent sur un système de troc – alors que grâce au théorème misésien de la régression, il devient possible d'étudier l'interdépendance entre les marchés en tenant compte des prix monétaires. Ainsi, non seulement la théorisation autrichienne ne rencontre pas les difficultés de chacun des deux types de modèles partiels et généraux, mais elle ne scinde pas la théorie des prix en ces deux compartiments étanches, et offre là encore une théorisation unifiée.

Quant à la macroéconomie standard, elle est segmentée entre une analyse de court terme plutôt keynésienne (non-neutralité de la monnaie, préférence pour la liquidité) et une analyse de long terme plutôt classique (monnaie neutre, théorie quantitative de la monnaie). Ce découplage entre le court et le long terme est dénoncé par Garrison (2001) comme l'une des principales faiblesses de cette macroéconomie. Et de fait, il y a là un problème : à quel « moment » le court terme se transforme-t-il en long terme, à quel « moment » faut-il abandonner les principes keynésiens pour adopter les principes classiques ? Ces questions apparaissent dans la perspective autrichienne comme dénuées de sens : à partir d'un instant donné, la période qui suit est nécessairement du court terme, et les mêmes principes d'explication doivent être mobilisés que ceux utilisés lors de la période de court terme précédente. La macroéconomie standard se heurte donc à une double difficulté : ses principes de raisonnement sont différents et même contradictoires selon les circonstances (keynésiens *versus* classiques), et son articulation entre court et long terme est insatisfaisante.

Ces problèmes n'existent pas dans le paradigme autrichien.

Chez von Mises, par exemple, le long terme n'est pas du tout conçu comme une période qui va survenir dans le monde réel après un délai suffisamment long : il est une construction théorique purement imaginaire, une situation limite dans laquelle on suppose que le système économique s'est pleinement adapté à un choc initialement subi, et sans que d'autres chocs se produisent pendant la phase d'adaptation (condition « toutes choses égales par ailleurs »). Le long terme est donc ici un pur outil de raisonnement permettant d'analyser les effets ultimes d'un changement dynamique, et non pas une réalité qui advient ou adviendra dans le monde réel. Dans ce cadre, il n'a pas de sens de dire que la monnaie n'est pas neutre à court terme et neutre à long terme, puisqu'à long terme le système économique se trouve dans un équilibre statique où la monnaie ne joue plus aucun rôle. Il n'a pas non plus de sens de changer de principes de raisonnement entre le court et le long terme, puisque le long terme n'est autre que l'issue d'une succession de périodes de court terme dans chacune desquelles les *mêmes* principes d'analyse vont être appliqués, à savoir le principe de la préférence pour le présent (ou de la productivité du capital) pour expliquer le taux d'intérêt, et la théorie quantitative de la monnaie pour expliquer le pouvoir d'achat de la monnaie.

Alors que le paradigme néo-classique standard est fragmenté en sous-champs qui sont au mieux disjoints et au pire contradictoires, l'école autrichienne est parvenue au fil des générations successives qui la composent, et avec une mention spéciale à von Mises pour sa synthèse de *L'action humaine*, à élaborer un *système* global et cohérent de connaissances économiques qui se déduisent d'un petit nombre de principes. Nous espérons que ce bref ouvrage aura fait apprécier cette voie de réflexion fondée sur la synthèse et la logique plutôt que sur le découpage en spécialités et la formalisation mathématique.

La date de publication de la première édition dans la langue d'origine est indiquée entre crochets.

Le symbole ☐ indique les textes de l'école autrichienne qui sont disponibles en libre accès sur internet, principalement sur le site de l'Institut von Mises (<http://mises.org>).

- Aimar, Thierry (2005) *Les apports de l'école autrichienne d'économie. Subjectivisme, ignorance et coordination*, Paris, Vuibert.
- Armentano, Dominick T. (1996 [1982]) *Antitrust and Monopoly. Anatomy of a Policy Failure*, 2^e éd. 1990, Oakland, Cal., The Independent Institute.
- ☐ Block, Walter (1983) « Public Goods and Externalities: The Case of Roads », *Journal of Libertarian Studies*, vol. 7, n° 1, p. 1-34.
- ☐ Block, Walter (1996) « Hayek's Road to Serfdom », *Journal of Libertarian Studies*, vol. 12, n° 2, p. 339-365.
- Boettke, Peter (dir.) (2000) *Socialism and the Market: The Socialist Calculation Debate Revisited*, 9 volumes, Londres, Routledge.
- ☐ Böhm-Bawerk, Eugen von (1959 [1884]) *History and Critique of Interest Theories*, 4^e éd. 1921, South Holland, Ill., Libertarian Press.
- ☐ Böhm-Bawerk, Eugen von (1959 [1889]) *Positive Theory of Capital*, 4^e éd. 1921, South Holland, Ill., Libertarian Press.
- Böhm-Bawerk, Eugen von (1962 [1881]) « Whether Legal Rights and Relationships Are Economic Goods », in E. von Böhm-Bawerk (1962), p. 25-138.
- ☐ Böhm-Bawerk, Eugen von (1962 [1894]) « The Ultimate Standard of Value », in E. von Böhm-Bawerk (1962), p. 303-370.
- Böhm-Bawerk, Eugen von (1962) *Shorter Classics of Böhm-Bawerk*, vol. 1, South Holland, Ill., Libertarian Press.
- Chamberlin, Edwin H. (1956 [1933]) *The Theory of Monopolistic Competition*, 7^e éd., Cambridge, Harvard University Press.
- ☐ Clark, John Bates (1899) *The Distribution of Wealth*, New York, N.Y., Macmillan.

- Čuhel, Franz (1907) *Zur Lehre von den Bedürfnissen*, Innsbruck, Wagner.
- Dorfman, Robert (1959) « A Graphical Exposition of Böhm-Bawerk's Interest Theory », *The Review of Economic Studies*, February, p. 153-158.
- 📖 Fetter, Frank (1900) « Recent Discussion of the Capital Concept », in F. Fetter (1977), p. 33-73.
- 📖 Fetter, Frank (1901) « The Passing of the Old Rent Concept », in F. Fetter (1977), p. 318-355.
- 📖 Fetter, Frank (1902) « The “Roundabout Process” in the Interest Theory », in F. Fetter (1977), p. 172-191.
- 📖 Fetter, Frank (1914) « Interest Theories, Old and New », in F. Fetter (1977), p. 226-255.
- 📖 Fetter, Frank (1915) *Economic Principles*, New York, N.Y., Century.
- 📖 Fetter, Frank (1927) « Interest Theory and Price Movements », in F. Fetter (1977), p. 260-316.
- 📖 Fetter, Frank (1977) *Capital, Interest, and Rent. Essays in the Theory of Distribution*, Kansas City, Sheed Andrews and McMeel.
- 📖 Fillieule, Renaud (2005) « The “Values-Riches” Model: An Alternative to Garrison's Model in Austrian Macroeconomics of Growth and Cycle », *Quarterly Journal of Austrian Economics*, vol. 8, n° 2, p. 3-19.
- Fisher, Irving (1911) *The Purchasing Power of Money. Its Determination and Relation to Credit Interest and Crises*, 2^e éd. 1913, New York, N.Y., Macmillan.
- Garrison, Roger (2001) *Time and Money. The Macroeconomics of the Capital Structure*, Londres, Routledge.
- Gentier, Antoine (2003) *Économie bancaire. Essai sur les effets de la concurrence et de la réglementation sur le financement du crédit*, Paris, Éditions Publibook Université.
- 📖 Hayek, Friedrich A. (1929) « The “Paradox” of Saving », in F. A. Hayek (1939), p. 199-263.
- 📖 Hayek, Friedrich A. (1932) « A Note on the Development of the Doctrine of “Forced Saving” », in F. A. Hayek (1939), p. 183-197.

- 📖 Hayek, Friedrich A. (1935) *Collectivist Economic Planning. Critical Studies on the Possibilities of Socialism*, Londres, George Routledge and Sons.
- 📖 Hayek, Friedrich A. (1936) « The Mythology of Capital », *Quarterly Journal of Economics*, vol. 50, n° 2, p. 199-228.
- 📖 Hayek, Friedrich A. (1939) *Profits, Interest and Investment, and Other Essays on the Theory of Industrial Fluctuations*, Londres, George Routledge & Sons.
- 📖 Hayek, Friedrich A. (1940) « The Competitive “Solution” », in F. A. Hayek (1948), p. 181-208.
- 📖 Hayek, Friedrich A. (1941) *The Pure Theory of Capital*, Chicago, Ill., The University of Chicago Press.
- 📖 Hayek, Friedrich A. (1946) « The Meaning of Competition », in F. A. Hayek (1948), p. 92-106.
- 📖 Hayek, Friedrich A. (1948) *Individualism and Economic Order*, Chicago, Ill., The University of Chicago Press.
- Hayek, Friedrich A. (1952) *The Counter-Revolution of Science. Studies on the Abuse of Reason*, Glencoe, Ill., The Free Press.
- Hayek, Friedrich A. (1960) *The Constitution of Liberty*, Londres, Routledge & Kegan Paul.
- Hayek, Friedrich A. (1964) « The Theory of Complex Phenomena », in F. A. Hayek (1967), p. 22-42.
- 📖 Hayek, Friedrich A. (1966 [1929]) *Monetary Theory and the Trade Cycle*, New York, N.Y., Augustus M. Kelley.
- Hayek, Friedrich A. (1950) « Full Employment, Planning and Inflation », in F. A. Hayek (1967), p. 270-279.
- Hayek, Friedrich A. (1959) « Unions, Inflation and Profits », in F. A. Hayek (1967), p. 280-294.
- Hayek, Friedrich A. (1967) *Studies in Philosophy, Politics and Economics*, Londres, Routledge & Kegan Paul.
- 📖 Hayek, Friedrich A. (1974) « The Pretence of Knowledge », in F. A. Hayek (1978), p. 23-34.
- 📖 Hayek, Friedrich A. (1975 [1931]) *Prix et production*, 2^e éd. 1935, Paris, Calmann-Levy.

- Hayek, Friedrich A. (1978) *New Studies in Philosophy, Politics, Economics and the History of Ideas*, Londres, Routledge & Kegan Paul.
- Hayek, Friedrich A. (1985 [1944]) *La route de la servitude*, Paris, PUF.
- 📖 Hayek, Friedrich A. (1990 [1976]) *Denationalisation of Money—The Argument Refined. An Analysis of the Theory and Practice of Concurrent Currencies*, 3^e éd., Londres, The Institute of Economic Affairs.
- 📖 Hazlitt, Henry (1978) *The Inflation Crisis, and How to Resolve It*, New Rochelle, N.Y., Arlington House.
- 📖 Hazlitt, Henry (2006 [1946]) *L'économie politique en une leçon*, 2^e éd. 1978, Paris, Éditions Charles Coquelin.
- 📖 Hoppe, Hans-Hermann (1989) *A Theory of Socialism and Capitalism: Economics, Politics, and Ethics*, Boston, Kluwer.
- 📖 Hoppe, Hans-Hermann (1994) « F. A. Hayek on Government and Social Evolution: A Critique », *The Review of Austrian Economics*, vol. 7, n° 1, p. 67-93.
- 📖 Hoppe, Hans-Hermann (1996) « Socialism: A Property or Knowledge Problem? », *The Review of Austrian Economics*, vol. 9, n° 1, p. 143-149.
- 📖 Hoppe, Hans-Hermann, Hülsmann, Jörg Guido et Block, Walter (1998) « Against Fiduciary Media », *Quarterly Journal of Austrian Economics*, vol. 1, n° 1, p. 19-50.
- Horwitz, Steven (2000) *Microfoundations and Macroeconomics. An Austrian Perspective*, Londres, Routledge.
- 📖 Horwitz, Steven (2009) « the Microeconomic Foundations of Macroeconomic Disorder: An Austrian Perspective on the Great Recession of 2008 », *Working Paper*, Mercatus Center, George Mason University, juin 2009, 23 p.
- 📖 Huerta de Soto, Jesús (2006 [1998]) *Money, Bank Credit, and Economic Cycles*, 2^e éd., Auburn, Ala., Ludwig von Mises Institute.
- 📖 Huerta de Soto, Jesús (2008) « Financial Crisis and Recession », *Mises Daily*, 6 octobre 2008, <http://mises.org>.
- 📖 Hülsmann, Jörg Guido (1997) « Knowledge, Judgment, and the Use of Property », *The Review of Austrian Economics*, vol. 10, n° 1, p. 23-48.

- ▣ Hülsmann, Jörg Guido (2006) « The Political Economy of Moral Hazard », *Politická Ekonomie*, vol. 1, p. 35-47.
- ▣ Hülsmann, Jörg Guido (2007) *Mises: The Last Knight of Liberalism*, Auburn, Ala., Ludwig von Mises Institute.
- ▣ Hülsmann, Jörg Guido (2008a) *The Ethics of Money Production*, Auburn, Ala., Ludwig von Mises Institute.
- ▣ Hülsmann, Jörg Guido (2008b) « Beware the Moral Hazard Trivializers », *Mises Daily*, 3 juin 2008, <http://mises.org>.
- ▣ Hülsmann, Jörg Guido (2009) « The Demand for Money and the Time-Structure of Production », in J. G. Hülsmann et S. Kinsella (dir.), *Property, Freedom, Society. Essays in Honor of Hans-Hermann Hoppe*, Auburn, Ala., Ludwig von Mises Institute, p. 309-324.
- Ikeda, Sanford (1997) *Dynamics of the Mixed Economy: Toward a Theory of Interventionism*, Routledge.
- Jevons, W. Stanley (1965 [1871]) *The Theory of Political Economy*, 5^e éd., New York, N.Y., Augustus M. Kelley.
- ▣ Kinsella, Stephan N. (2008 [2001]) *Against Intellectual Property*, Auburn, Ala., Ludwig von Mises Institute.
- Kirzner, Israel M. (1973) *Competition and Entrepreneurship*, Chicago, Ill., The University of Chicago press.
- ▣ Kirzner, Israel M. (1996) « Reflections on the Misesian Legacy in Economics », *Review of Austrian Economics*, vol. 9, n° 2, p. 143-154.
- Knight, Frank H. (1921) *Risk, Uncertainty, and Profit*, Houghton Mifflin Company.
- Knight, Frank H. (1934) « Capital, Time, and the Interest Rate », *Economica*, vol. 1, n° 3, p. 257-286.
- Lange, Oskar (1936) « On the Economic Theory of Socialism », *Review of Economic Studies*, vol. 4, p. 53-71.
- Lange, Oskar (1937) « On the Economic Theory of Socialism », *Review of Economic Studies*, vol. 5, p. 132-142.
- ▣ Lavoie, Don (1982) « The Development of the Misesian Theory of Interventionism », in I. M. Kirzner (éd.), *Method, Process, and Austrian Economics. Essays in Honor of Ludwig von Mises*, Lexington, Mass., Lexington Books.

- Lavoie, Don (1985) *Rivalry and Central Planning. The Socialist Calculation Debate Reconsidered*, Cambridge, Cambridge University Press.
- Mankiw, Gregory N. (1998) *Principes de l'économie*, Paris, Economica.
- Marshall, Alfred (1920 [1890]) *Principles of Economics. An Introductory Volume*, 8^e éd., Londres, Macmillan.
- Menger, Carl (1888) « Zur Theorie des Kapitals », *Conrad's Jahrbücher für Nationalökonomie und Statistik*, vol. 17, Jena, Gustav Fischer Verlag.
- 📖 Menger, Carl (1976 [1871]) *Principles of Economics*, New York, N.Y., New York University Press.
- 📖 Menger, Carl (1985 [1883]) *Investigations Into the Method of the Social Sciences*, New York, N.Y., New York University Press.
- Mill, James (1826 [1821]) *Elements of Political Economy*, 3^e éd., Londres, Baldwin and Cradock.
- Mill, John Stuart (1987 [1848]) *Principles of Political Economy with Some of their Applications to Social Philosophy*, 7^e éd. 1871, Fairfield, N.J., Augustus M. Kelley.
- 📖 Mises, Ludwig von (1923) « Theory of Price Control », in L. von Mises (1977 [1929]), p. 139-151.
- 📖 Mises, Ludwig von (1926) « Interventionism », in L. von Mises (1977 [1929]), p. 15-55.
- 📖 Mises, Ludwig von (1928) « Monetary Stabilization and Cyclical Policy », in L. von Mises (1978), p. 57-171.
- 📖 Mises, Ludwig von (1931) « The Causes of the Economic Crisis », in L. von Mises (1978), p. 173-203.
- 📖 Mises, Ludwig von (1935 [1920]) « Economic Calculation in the Socialist Commonwealth », in F. A. Hayek (1935), p. 87-130.
- 📖 Mises, Ludwig von (1944) *Bureaucracy*, New Haven, Yale University Press.
- 📖 Mises, Ludwig von (1946) « The Trade Cycle and Credit Expansion: The Economic Consequences of Cheap Money », in L. von Mises (1978), p. 215-226.
- 📖 Mises, Ludwig von (1962) *The Ultimate Foundation of Economic Science. An Essay on Method*, Princeton, N.J., Van Nostrand.

- 📖 Mises, Ludwig von (1977 [1929]) *A Critique of Interventionism*, New Rochelle, N.Y., Arlington House.
- 📖 Mises, Ludwig von (1978) *On the Manipulation of Money and Credit*, Dobbs Ferry, N.Y., Free Market Books.
- 📖 Mises, Ludwig von (1980 [1912]) *The Theory of Money and Credit*, 2^e éd. 1924, Indianapolis, Ind., LibertyClassics.
- 📖 Mises, Ludwig von (1981 [1922]) *Socialism. An Economic and Sociological Analysis*, 2^e éd. 1932, Indianapolis, Ind., LibertyClassics.
- 📖 Mises, Ludwig von (1981 [1933]) *Epistemological Problems of Economics*, New York, New York University Press.
- 📖 Mises, Ludwig von (1985 [1949]) *L'action humaine. Traité d'économie*, 3^e éd. 1963, Paris, PUF.
- 📖 Mises, Ludwig von (1998 [1940]) *Interventionism: An Economic Analysis*, Irvington-on-Hudson, N.Y., The Foundation for Economic Education.
- 📖 Mises, Ludwig von (1998 [1944]) « Monopoly Prices », *The Quarterly Journal of Austrian Economics*, vol. 1, n° 2, p. 1-28.
- 📖 Mises, Ludwig von (2005 [1927]) *Liberalism. The Classical Tradition*, Indianapolis, Ind., Liberty Fund.
- 📖 Mises, Ludwig von (2007 [1957]) *Theory and History. An Interpretation of Social and Economic Evolution*, Auburn, Ala., Ludwig von Mises Institute.
- 📖 Murphy, Robert P. (2007) « The Worst Recession in 25 years? », *Mises Daily*, 1^{er} octobre 2007, <http://mises.org>.
- Murphy, Robert P. (2009) *The Politically Incorrect Guide to the Great Depression and the New Deal*, Washington, D. C., Regnery.
- Norberg, Johan (2009) *Financial Fiasco: How America's Infatuation with Homeownership and Easy Money Created the Economic Crisis*, Washington, D. C., Cato Institute.
- Oppenheimer, Franz (1914) *The State*, New York, N.Y., Vanguard Press.
- O'Driscoll, Gerald P. et Rizzo, Mario (1985) *The Economics of Time and Ignorance*, Oxford, Basil Blackwell.
- Pareto, Vilfredo (1981 [1909]) *Manuel d'économie politique*, Genève, Droz.

- ☞ Rallo, Juan Ramón (2009) « Economic Crisis and Paradigm Shift », *Mises Daily*, 6 janvier 2009, <http://mises.org>.
- ☞ Reisman, George (1996) *Capitalism. A Treatise on Economics*, Ottawa, Ill., Jameson Books.
- ☞ Reisman, George (2007) « The Housing Bubble and the Credit Crunch », 10 août 2007, www.capitalism.net.
- ☞ Reisman, George (2008) « The Myth that Laissez Faire is Responsible for Our Present Crisis », *Mises Daily*, 23 octobre 2008, <http://mises.org>.
- Ricardo, David (1951 [1817]) *On the Principles of Political Economy and Taxation*, 3^e éd. 1821, Cambridge, Cambridge University Press.
- Robbins, Lionel (1947 [1932]) *Essai sur la nature et la signification de la science économique*, 2^e éd. 1935, Paris, Librairie de Médicis.
- Robinson, Joan (1933) *The Economics of Imperfect Competition*, Londres.
- ☞ Rothbard, Murray N. (1962) *Man, Economy and State. A Treatise on Economic Principles*, Princeton, N.J., D. Van Nostrand.
- ☞ Rothbard, Murray N. (1977 [1970]) *Power and Market. Government and the Economy*, 2^e éd., Kansas City, Sheed Andrews and McMeel.
- ☞ Rothbard, Murray N. (1977) « Introduction », in F. Fetter (1977), p. 1-24.
- ☞ Rothbard, Murray N. (1983 [1963]) *America's Great Depression*, 4^e éd., New York, N.Y. Richardson & Snyder.
- ☞ Rothbard, Murray N. (1988) « The Myth of Free Banking in Scotland », *Review of Austrian Economics*, vol. 2, n^o 1, p. 229-245.
- ☞ Rothbard, Murray N. (1990 [1963]) *What Has Government Done to our Money?*, Auburn, Ala., Ludwig von Mises Institute.
- Rothbard, Murray N. (1991 [1954]) « Vers une reconstruction de la théorie de l'utilité et du bien-être », in *Économistes et charlatans*, Paris, Les Belles Lettres, p. 105-161.
- ☞ Rothbard, Murray N. (1991 [1982]) *L'éthique de la liberté*, Paris, Les Belles Lettres.
- Rothbard, Murray N. (1997 [1978]) « Austrian Definitions of the Supply of Money », in *The Logic of Action I: Money, Method and the Austrian School*, Cheltenham, Edward Elgar, p. 337-349.

- 📖 Rothbard, Murray N. (2005 [1962]) *The Case for a 100 Percent Gold Dollar*, Auburn, Ala., Ludwig von Mises Institute.
- 📖 Rothbard, Murray N. (2008 [1983]) *The Mystery of Banking*, 2^e éd., Auburn, Ala., Ludwig von Mises Institute.
- 📖 Salerno, Joseph T (1990) « Ludwig von Mises as Social Rationalist », *The Review of Austrian Economics*, vol. 4, p. 26-54.
- 📖 Salerno, Joseph T (1993) « Mises and Hayek Dehomogenized », *The Review of Austrian Economics*, vol. 6, n° 2, p. 113-146.
- Salin, Pascal (1990) *La vérité sur la monnaie*, Paris, Odile Jacob.
- Samuelson, Paul A. (1954) « The Pure Theory of Public Expenditure », *Review of Economics and Statistics*, vol. 36, n° 4, p. 387-389.
- Schumpeter, Joseph A. (1939) *Business Cycles. A Theoretical, Historical, and Statistical Analysis of the Capitalist Process*, New York, McGraw-Hill.
- Schumpeter, Joseph A. (1951 [1942]) *Capitalisme, socialisme et démocratie*, Paris, Payot.
- Schumpeter, Joseph A. (1999 [1911]) *Théorie de l'évolution économique. Recherches sur le profit, le crédit, l'intérêt et le cycle de la conjoncture*, 2^e éd. 1926, Paris, Dalloz.
- Selgin, George (1991 [1988]) *La théorie de la banque libre*, Paris, Les Belles Lettres.
- 📖 Selgin, George et White, Lawrence H. (1996) « In Defense of Fiduciary Media—or, We are Not Devo(lutionists), We are Misesians! », *Review of Austrian Economics*, vol. 9, n° 2, p. 83-107.
- Selgin, George (1997) *Less than Zero. The Case for a Falling Price Level in a Growing Economy*, Londres, Institute of Economic Affairs.
- Senior, Nassau (1836) *An Outline of the Science of Political Economy*, Londres.
- Sennholz, Hans F. (1977) *Age of Inflation*, Belmont, Mass., Western Islands.
- 📖 Shostak, Frank (2003) « Housing Bubble: Myth or Reality? », *Mises Daily*, 4 mars 2003, <http://mises.org>.
- 📖 Skousen, Mark (1989) « Saving the Depression: A New Look at World War II », *The Review of Austrian Economics*, vol. 2, n° 1, p. 211-226.

- Skousen, Mark (1990) *The Structure of Production*, New York, N.Y., New York University Press.
- Skousen, Mark (1991) *Economics on Trial. Lies, Myths, and Realities*, Burr Ridge, Ill., Irwin.
- Skousen, Mark (1996 [1977]) *Economics of a Pure Gold Standard*, 3^e éd., Irvington-on-Hudson, N.Y., The Foundation for Economic Education.
- Smith, Adam (1976 [1776]) *An Inquiry into the Nature and Causes of the Wealth of Nations*, Chicago, Ill., The University of Chicago Press.
- Thirlby, G. F. (1981 [1946]) « The Subjective Theory of Value and Accounting "Cost" », in J. Buchanan et G. F. Thirlby (éd.), *L.S.E. Essays on Cost*, New York, N.Y., New York University Press, p. 137-161.
- ☞ Thornton, Mark (1991) *The Economics of Prohibition*, Salt Lake City, Utah, University of Utah Press.
- ☞ Thornton, Mark (2006) « The Economics of Housing Bubbles », *Mises Institute Working Paper*, 12 juin 2006, <http://mises.org>.
- ☞ Thornton, Mark (2008) « The Housing Bubble in Four Easy Steps », *Mises Daily*, 27 septembre 2008, <http://mises.org>.
- Vedder, Richard K. et Gallaway, Lowell E. (1997 [1993]) *Out of Work. Unemployment and Government in Twentieth-Century America*, New York, N.Y., New York University Press.
- Walras, Léon (1874) *Éléments d'économie politique pure*, Corbaz et C^{ie}, Lausanne.
- West, Edwin G. (1994 [1965]) *Education and the State. A Study in Political Economy*, 3^e éd., Indianapolis, In., Liberty Fund.
- ☞ Wicksell, Knut (1936 [1898]) *Interest and Prices. A Study of the Causes Regulating the Value of Money*, Londres.
- ☞ Wicksell, Knut (1954 [1893]) *Value, Capital and Rent*, Londres, George Allen and Unwin.
- Wieser, Friedrich von (1884) *Über den Ursprung und die Hauptgesetze des Wirthschaftlichen Werthes*.
- ☞ Wieser, Friedrich von (1893 [1889]) *Natural Value*, Londres, Macmillan.

- Wieser, Friedrich von (1910) « Der Geldwert und seine geschichtliche Veränderungen », *Zeitschrift für Volkswirtschaft, Sozialpolitik und Verwaltung*, Leipzig.
- Woods, Thomas E. Jr. (2009) *Meltdown. A Free-Market Look at Why the Stock Market Collapsed, the Economy Tanked, and Government Bailouts Will Make Things Worse*, Washington, Regnery.
- 📖 Yeager, Leland B. (1997) « Calculation and Knowledge: Let's Write *Finis* », *Review of Austrian Economics*, vol. 10, n° 1, p. 133-136.

Table des matières détaillée

Remerciements	7
Sommaire.....	9
Préface des éditeurs	11
Introduction	15
Chapitre 1 : biens et valeur	21
1.1 La théorie des biens.....	21
1.1.1 Les quatre pré-requis d'un bien. 1.1.2 Biens et subjectivité. 1.1.3 Les services. 1.1.4 Biens d'ordre supérieur et étapes de production. 1.1.5 Biens complémentaires ou substituables, convertibles ou spécifiques. 1.1.6 Les biens « non économiques ». 1.1.7 Les biens « imaginaires ». 1.1.8 Les droits de propriété et les « relations » sont-ils des biens ? 1.1.9 Le temps. 1.1.10 Critique de la notion de « bien public ».	
1.2 La valeur subjective.....	28
1.2.1 Définition et origine de la valeur. 1.2.2 Valeur et subjectivité. 1.2.3 L'échelle des valeurs. 1.2.4 La résolution du « paradoxe de la valeur ». 1.2.5 Unité et stock. 1.2.6 La loi de l'utilité marginale. 1.2.7 Variation de stock et utilité marginale. 1.2.8 La mesure des valeurs. 1.2.9 Valeur et indifférence. 1.2.10 La comparaison interpersonnelle des valeurs. 1.2.11 La préférence pour le présent. 1.2.12 La loi de la préférence temporelle.	
1.3 La valeur des facteurs de production.....	37
1.3.1 Le principe d'imputation. 1.3.2 La valeur des facteurs complémentaires. 1.3.3 Proportions fixes ou variables des facteurs. 1.3.4 Imputation et facteurs convertibles. 1.3.5 Le coût d'opportunité.	
Chapitre 2 : échange et prix	45
2.1 La théorie de l'échange	45
2.1.1 Les formes d'interaction sociale. 2.1.2 L'explication de l'échange. 2.1.3 L'échange à prix fixé. 2.1.4 Échange et maximisation de la valeur. 2.1.5 Productivité et coûts de l'échange. 2.1.6 Valeur d'usage et valeur d'échange. 2.1.7 Échange et division du travail. 2.1.8 Marchandises et	

« échangeabilité ». 2.1.9 Du troc à l'échange monétaire.

2.2 Le système des prix de marché52

2.2.1 La notion de prix. 2.2.2 Le marché d'enchères. 2.2.3 La version classique de la loi des coûts. 2.2.4 L'inversion de la causalité classique : la détermination des coûts par les prix. 2.2.5 La convergence vers l'équilibre par résorption des profits et pertes entrepreneuriaux. 2.2.6 Un cas apparent de détermination des prix par les coûts. 2.2.7 Le produit marginal monétaire. 2.2.8 L'escompte. 2.2.9 Des « obstacles frictionnels » aux « changements dynamiques ». 2.2.10 Profits et pertes monétaires. 2.2.11 La tendance à l'égalité des taux de rentabilité. 2.2.12 L'économie en rotation uniforme. 2.2.13 Les processus d'adaptation aux grands types de changements dynamiques. 2.2.14 La rationalité de l'économie de marché. 2.2.15 Les entrepreneurs face à la souveraineté des consommateurs.

Chapitre 3 : monopole et concurrence69

3.1 La théorie du prix de monopole69

3.1.1 Monopole économique et monopole politique. 3.1.2 La relativité du monopole. 3.1.3 La restriction monopolistique. 3.1.4 Le prix de monopole. 3.1.5 Les types de monopoles économiques chez von Mises. 3.1.6 La nature du revenu de monopole. 3.1.7 Les conséquences systémiques d'un prix de monopole. 3.1.8 Les prix de monopole ont-ils tendance à remplacer les prix concurrentiels ? 3.1.9 Prix de monopole ou alignement sur les coûts de production ? 3.1.10 Le prix de monopole : une illusion ?

3.2 La concurrence entrepreneuriale76

3.2.1 La notion de concurrence. 3.2.2 Une conception large de la concurrence. 3.2.3 La concurrence comme processus dynamique de découverte. 3.2.4 L'élément entrepreneurial de l'action humaine. 3.2.5 La concurrence entrepreneuriale. 3.2.6 Entrepreneuriat et monopole. 3.2.7 Différentiation des biens, coûts de vente, publicité. 3.2.8 Critique de la « concurrence pure et parfaite ». 3.2.9 Critique de la « concurrence monopolistique ».

Chapitre 4 : la production et sa structure.....89

4.1 La production89

4.1.1 La production comme processus matériel. 4.1.2 La production comme action. 4.1.3 La loi des rendements décroissants. 4.1.4 Conséquences des rendements décroissants et forces en sens contraire. 4.1.5 Le progrès de la connaissance technique. 4.1.6 Division du travail et société. 4.1.7 Division du travail et production. 4.1.8 La loi du « détour » de production. 4.1.9 La fonction de production intertemporelle. 4.1.10 L'explication de la producti-

vité du « détour » de production. 4.1.11 L'épargne-investissement. 4.1.12 Mesurer la période de production ?

4.2 La macroéconomie de la structure de production 98

4.2.1 La structure de production hayékienne. 4.2.2 Le triangle hayékien. 4.2.3 Structure synchronique et structure diachronique. 4.2.4 Épargne et maintien de la structure. 4.2.5 Le produit total : critique du PIB. 4.2.6 Une structure de production avec biens durables. 4.2.7 L'accumulation du capital. 4.2.8 Critique du « paradoxe de l'épargne ». 4.2.9 La réévaluation de l'épargne-investissement.

Chapitre 5 : capital et intérêt 111

5.1 Facteurs de production et capital 111

5.1.1 De la théorie classique de la distribution au principe d'imputation. 5.1.2 La conception « réelle » du capital. 5.1.3 Peut-on distinguer la « terre » et le « capital » ? 5.1.4 La conception « nominale » du capital. 5.1.5 La généralisation de la notion de rente. 5.1.6 Rente et actualisation. 5.1.7 La théorie autrichienne des revenus : rente, intérêt, profit/perte. 5.1.8 Le capitalisme.

5.2 Les théories de l'intérêt 119

5.2.1 Le problème théorique de l'intérêt. 5.2.2 Intérêt originaire et intérêt contractuel. 5.2.3 Intérêt originaire et profit entrepreneurial. 5.2.4 L'intérêt comme phénomène d'évaluation intertemporelle (Böhm-Bawerk). 5.2.5 La théorie productiviste (Böhm-Bawerk, Wicksell, Hayek). 5.2.6 La théorie de la préférence pour le présent ou théorie de l'actualisation (Fetter, von Mises). 5.2.7 La théorie de l'échange (Rothbard). 5.2.8 La théorie des fonds prêtables (O'Driscoll, Skousen, Garrison). 5.2.9 Critique des théories classiques de l'intérêt : productivité, abstinence, rémunération, exploitation. 5.2.10 La théorie de l'usage (Menger).

Chapitre 6 : la monnaie et son pouvoir d'achat 133

6.1 La notion de monnaie 133

6.1.1 La fonction de la monnaie. 6.1.2 La monnaie en tant que bien. 6.1.3 Les types de monnaies. 6.1.4 Les substituts de monnaie. 6.1.5 La tendance à l'unification monétaire. 6.1.6 Monnaie et calcul économique. 6.1.7 Les limites du calcul économique.

6.2 Le pouvoir d'achat de la monnaie 139

6.2.1 La valeur subjective de la monnaie. 6.2.2 La composante historique de

la valeur de la monnaie (théorème de la régression). 6.2.3 Le fondement de la théorie subjectiviste des prix monétaires. 6.2.4 La théorie quantitative de la monnaie. 6.2.5 La demande de monnaie. 6.2.6 L'offre de monnaie. 6.2.7 La convergence vers l'équilibre monétaire. 6.2.8 Les effets des changements d'offre ou de demande de monnaie sur son pouvoir d'achat. 6.2.9 L'influence de la sphère réelle. 6.2.10 L'effet redistributif d'une augmentation de la quantité de monnaie (effet Cantillon). 6.2.11 La monnaie n'est jamais « neutre ». 6.2.12 Critique de l'équation des échanges. 6.2.13 Les influences de long terme sur la demande et l'offre de monnaie. 6.2.14 Changement du pouvoir d'achat de la monnaie et taux d'intérêt. 6.2.15 La théorie de la parité du pouvoir d'achat.

Chapitre 7 : inflation et crise155

7.1 L'inflation155

7.1.1 Les définitions autrichiennes. 7.1.2 Inflation et hausse des prix. 7.1.3 Une « inflation » par les coûts ? 7.1.4 Les causes de l'inflation. 7.1.5 Les conséquences de l'inflation. 7.1.6 Les méfaits économiques de l'inflation. 7.1.7 Autres méfaits de l'inflation. 7.1.8 L'« épargne forcée ». 7.1.9 Les fausses solutions contre l'inflation. 7.1.10 Des institutions monétaires pour une monnaie « solide ». 7.1.11 La banque libre (réserves fractionnaires). 7.1.12 La banque entrepôt (100 % de réserves). 7.1.13 Une baisse tendancielle mais non déflationniste des prix.

7.2 La théorie du cycle167

7.2.1 La théorie du crédit de circulation. 7.2.2 Le crédit de circulation. 7.2.3 Expansion du crédit de circulation et taux d'intérêt monétaire. 7.2.4 Le boom. 7.2.5 La crise. 7.2.6 Le boom peut-il durer indéfiniment ? 7.2.7 Illustration du cycle par les triangles hayékiens. 7.2.8 Deux métaphores. 7.2.9 Cycle d'affaires et niveau des prix. 7.2.10 Combattre la crise ? 7.2.11 Éviter le cycle. 7.2.12 La Grande Dépression (1929-1941). 7.2.13 La stagflation des années 1970. 7.2.14 La crise des *subprimes*.

Chapitre 8 : État et marché181

8.1 L'interventionnisme181

8.1.1 Capitalisme, collectivisme, interventionnisme. 8.1.2 La notion d'intervention. 8.1.3 Le contrôle des prix. 8.1.4 Du contrôle sélectif au contrôle universel des prix. 8.1.5 Le salaire minimum. 8.1.6 La prohibition. 8.1.7 Le privilège de monopole ou de quasi-monopole. 8.1.8 Les types de monopoles et quasi-monopoles d'État. 8.1.9 Les lois anti-trust. 8.1.10 La propriété intellectuelle. 8.1.11 Les externalités. 8.1.12 La critique scientifique de l'interventionnisme. 8.1.13 Libéralisme classique, néo-libéralisme, anarcho-capitalisme.

8.2 Impôts et dépenses de l'État.....	195
8.2.1 Nature des revenus et des dépenses de l'État. 8.2.2 Impôt et redistribution. 8.2.3 Les types d'impôts. 8.2.4 L'impôt sur les ventes de détail. 8.2.5 L'impôt sur les revenus nets. 8.2.6 L'impôt sur le capital. 8.2.7 Niveau et progressivité de l'impôt. 8.2.8 Les subventions étatiques. 8.2.9 Les services étatiques. 8.2.10. Minimiser la dépense publique.	
8.3 Le collectivisme	205
8.3.1 L'État collectiviste. 8.3.2 Le problème économique. 8.3.3 Calcul économique et prix de marché. 8.3.4 L'impossibilité du calcul économique en régime collectiviste. 8.3.5 Critique de l'argument historique. 8.3.6 Critique de la solution marxiste. 8.3.7 Critique de la solution mathématique. 8.3.8 « Déshomogénéiser » von Mises et Hayek ? 8.3.9 Critique du système de planification décentralisée. 8.3.10 La perspective autrichienne sur le débat.	
Conclusion.....	217
Bibliographie.....	223
Table des matières détaillée	235

Ouvrage finalisé par
Yvon Bruant
Achevé d'imprimer - septembre 2010
Imprimerie de l'Université Charles-de-Gaulle – Lille 3
Dépôt légal - octobre 2010
1 232^e volume édité par les
Presses Universitaires du Septentrion
Villeneuve d'Ascq – France

L'école autrichienne d'économie

Une autre hétérodoxie

Cet ouvrage offre une présentation synthétique des concepts et des théories de l'une des principales écoles de pensée hétérodoxes en économie, l'école autrichienne. Fondée à la fin du XIX^e siècle lors de la révolution néo-classique et développée par certains des plus grands théoriciens du XX^e siècle, elle connaît aujourd'hui un regain d'intérêt de par son explication monétaire de la crise des *subprimes* et sa critique des politiques gouvernementales qui sont appliquées pour la combattre.

L'école autrichienne constitue un paradigme à part entière, qui aborde avec ses propres concepts et dans un cadre théorique spécifique tous les grands thèmes de l'analyse économique, depuis la notion de valeur jusqu'aux effets des interventions de l'État, en passant par la formation des prix de marché, la nature du processus concurrentiel, les lois de la production, les phénomènes monétaires et les cycles d'affaires. Les théoriciens de l'école autrichienne présentent aussi l'une des défenses les plus cohérentes et les plus solides du libéralisme économique.

Le contenu de cet ouvrage intéressera principalement les étudiants en économie, et plus généralement en sciences sociales, qui cherchent une présentation approfondie, rigoureuse et néanmoins non mathématisée des fondements de l'analyse économique.



Renaud Fillieule

est professeur à l'Université de Lille 1 et membre du laboratoire CLERSÉ. Ses travaux de recherche actuels portent sur l'école autrichienne d'économie, dans les domaines de la théorie des prix et de la macroéconomie de la structure de production.

18 €



F 112321
ISBN : 978-2-7574-0163-7
ISSN : 1773-8814



9 782757 401637