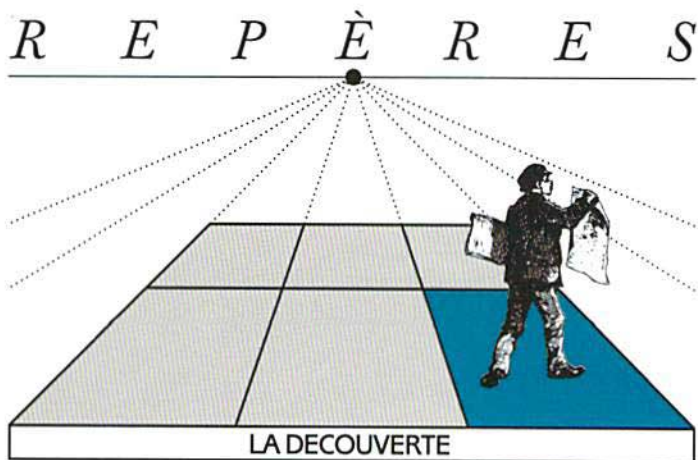


Nathalie Moureau
Dorothee Rivaud-Danset

L'incertitude dans les théories économiques



68860

Nathalie Moureau
et Dorothée Rivaud-Danset

L'incertitude dans les théories économiques

Éditions La Découverte
9 *bis*, rue Abel-Hovelacque
75013 Paris

Remerciements

Nous tenons à remercier Jean-Marc Tallon pour ses conseils qui nous ont été précieux ainsi que François Facchini, Olivier Gergaud, Fabrice Philippe et Fabien Rivaud pour leur lecture attentive de certains chapitres.

Le logo qui figure au dos de la couverture de ce livre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, tout particulièrement dans le domaine des sciences humaines et sociales, le développement massif du photocopillage.

Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or cette pratique s'est généralisée dans les établissements d'enseignement supérieur, provoquant une baisse brutale des achats de livres, au point que la possibilité même pour les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée.

Nous rappelons donc qu'en application des articles L. 122-10 à L. 122-12 du Code de la propriété intellectuelle, toute reproduction à usage collectif par photocopie, intégralement ou partiellement, du présent ouvrage est interdite sans autorisation du Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris). Toute autre forme de reproduction, intégrale ou partielle, est également interdite sans autorisation de l'éditeur.

Si vous désirez être tenu régulièrement informé de nos parutions, il vous suffit d'envoyer vos nom et adresse aux Éditions La Découverte, 9 bis, rue Abel-Hovelacque, 75013 Paris. Vous recevrez gratuitement notre bulletin trimestriel **À la Découverte**. Vous pouvez également retrouver l'ensemble de notre catalogue et nous contacter sur notre site www.editionsladecouverte.fr.

Introduction

L'incertitude est indissociable de la vie économique. Keynes, Knight et Hayek l'avaient souligné en leur temps, montrant que l'économie ne peut pas être conçue sur le registre de la mécanique. Les néoclassiques soucieux d'ordre ont, en revanche, tardé à en reconnaître l'importance. L'introduction, dans les années 1960, de l'incertitude dans l'analyse économique néoclassique a provoqué une « révolution » [Stiglitz, 2002]. De nombreux résultats tenus pour acquis sont tombés. Dans son article pionnier sur le marché des voitures d'occasion, Akerlof [1970] a démontré que l'introduction de l'incertitude sur la qualité d'un produit peut conduire à la disparition du marché. Le concept d'équilibre, la vision traditionnelle de la coordination du marché par le système des prix ne sont plus évidents. L'incertitude est le talon d'Achille de l'analyse économique du marché concurrentiel. Jusqu'à quel point l'édifice théorique est-il ébranlé ? Le principe de l'optimisation, central pour les néoclassiques, peut-il être conservé ?

Bien entendu, ces questions ne sont pas restées sans réponse, des concepts et des courants nouveaux ont émergé. L'attribution du prix Nobel à des spécialistes de l'incertitude témoigne de la vitalité des recherches dans ce domaine (Georges Akerlof, Michael Spence et Joseph Stiglitz en 2001, Daniel Kahneman en 2002). La façon dont les économistes contemporains se saisissent des questions liées à l'incertitude est à l'origine d'un clivage profond, qui renvoie à des conceptions distinctes de la rationalité.

Pour les néoclassiques, les individus utilisent les probabilités pour traiter l'incertitude, ont des capacités de prévision et de calcul telles que l'optimisation est toujours possible. La rationalité demeure « parfaite ».

Pour les autres courants de pensée, l'omniprésence de l'incertitude dans la vie économique fait obstacle à cette conception de la

rationalité. Ceci légitime la position critique qu'adoptent de nombreux économistes vis-à-vis de l'utilitarisme et de la maximisation. Les questions se posent, alors, en des termes bien différents. Comment expliquer que les personnes réussissent à se coordonner, alors qu'elles ne sont pas dotées de capacités de calcul et d'anticipation exceptionnelles ? Comment éviter que l'économie ne tourne au chaos ?

Le thème de l'incertitude ouvre aussi sur d'autres registres de questions. Des comportements spécifiques, comme l'imitation ou l'attentisme, sont-ils efficaces ? Peut-on dire quelque chose sur le comportement à adopter en incertitude radicale ? Telle est la question posée aux économistes par le principe de précaution.

PREMIÈRE PARTIE

LES MONDES DE L'INCERTITUDE

Les premiers économistes à s'intéresser à l'incertitude ont privilégié une conception très extensive conforme à leur vision du monde. L'univers ne peut pas être représenté comme un ensemble fini de situations possibles et la connaissance des agents est limitée. Comment les agents se comportent-ils ? Selon les cas, ils utilisent ou non les probabilités, le qualificatif de risque étant réservé à l'incertitude probabilisable [Knight, 1921] (chapitre 1). Aujourd'hui, prédomine une acception étroite de l'incertitude. Les individus appréhendent leur environnement, qu'il soit présent ou futur, comme étant composé d'un ensemble fini de situations possibles dont la réalisation est aléatoire. Dans les modèles, l'incertitude a une origine précise — par exemple, l'information sur une variable cruciale fait défaut au décideur — et peut être traitée comme un risque. En introduisant l'incertitude au sein de l'analyse économique, même dans son acception étroite, les économistes ont montré l'extrême fragilité des résultats obtenus dans le cadre de la concurrence parfaite (chapitre 11).

I / L'incertitude souveraine, une vision partagée par des courants hétérogènes

Les visions philosophiques de l'incertitude développées en parallèle par Knight, Keynes et Hayek présentent de nombreux points communs. Tous récusent l'hypothèse d'omniscience dont sont dotés les agents par la théorie néoclassique. Ils mettent en évidence les difficultés de la connaissance et/ou de la prévision que rencontrent les individus face à la complexité, à la nouveauté ou, tout simplement, à l'existence d'aléas. Ils admettent qu'en incertitude le calcul ne peut suffire à guider l'action. Dans un contexte incertain, les êtres humains sont pris en tension entre deux pôles. Le premier est celui du calcul et de l'action expliquée rationnellement par les objectifs. Le second pôle expliquant les comportements dépend de la conception que les auteurs se font des êtres humains. Ce sera l'intuition pour Knight, les règles de la vie en société pour Hayek, l'humeur et le hasard pour Keynes. Aujourd'hui, plusieurs courants hétérogènes rassemblés sous la bannière de l'incertitude radicale prolongent ces réflexions en les radicalisant.

1. L'incertitude « épistémique » : l'apport d'économistes-philosophes

Par leur réflexion à la fois de philosophe et d'économiste, Knight et Keynes ouvrent au début du xx^e siècle un débat essentiel. Est-il possible de rendre compte scientifiquement d'un objet incertain ? Autrement dit, peut-on représenter l'incertitude dans la vie économique ordinaire, comme le font pour les jeux de hasard les statisticiens qui, à partir des propriétés du jeu, définissent des lois de probabilité ? Peut-on rendre compte de la façon dont les êtres humains traitent l'incertitude ? Hayek, en s'interrogeant sur les limites de la connaissance, contribue également à ce débat. Knight occupe une

place de choix, en raison de sa célèbre distinction entre le risque et l'incertitude. Bien que celle-ci soit devenue une référence incontournable, sa réflexion reste largement méconnue.

Le risque et l'incertitude selon Knight

Dans *Risk, Uncertainty and Profit* [1921], Knight effectue une distinction entre risque et incertitude qui fait autorité depuis. « La différence pratique entre les deux catégories, le risque et l'incertitude, est que, s'agissant de la première, la distribution du résultat parmi un ensemble de cas est connue (soit par le calcul *a priori*, soit par des statistiques fondées sur les fréquences observées), tandis que ceci n'est pas vrai de l'incertitude en raison de l'impossibilité de regrouper les cas, parce que la situation à traiter présente un degré élevé de singularité » (p. 233).

Une situation est risquée quand la prévision peut se faire à partir de probabilités mathématiques ou de probabilités fréquentistes. Les probabilités mathématiques (nombre de cas favorables/nombre de cas total) sont calculées *a priori*, comme dans les jeux de hasard où les chances sont égales. Les probabilités fréquentistes (induites de l'expérience) sont calculées à partir d'un grand nombre d'observations d'un événement qui se répète avec une certaine fréquence, comme le nombre de jours de pluie dans une année. Elles sont toutes deux qualifiées d'objectives.

Une situation incertaine est considérée comme unique et n'est pas réductible à un groupe de cas similaires, elle n'est donc pas probabilisable. La prévision repose alors sur deux exercices séparés de jugement :

- le premier consiste à former une *estimation* ou un jugement personnel. Cette conjecture s'appuie sur l'expérience personnelle ou peut relever de la pure intuition ;

- le second mesure la validité du jugement effectué, il dépend de la confiance que l'individu a dans son estimation. Ce degré de confiance ou de croyance est appelé probabilité épistémique, du grec *épistémé* qui signifie connaissance [Hacking, 1975].

Cette distinction entre les deux moments du jugement a été obscurcie au ^{xx} siècle avec le développement des probabilités subjectives (de Finetti, Ramsey, Savage). Celles-ci ont conduit à confondre les deux moments que constituent la formation d'une estimation et l'évaluation du degré de confiance dans cette estimation, toute anticipation devenant alors mesurable.

La possibilité d'associer ou non des probabilités à des situations incertaines a de grandes répercussions sur la nature de la prévision. Utiliser les probabilités permet d'étendre à l'univers incertain le

calcul et donc l'optimisation. Pour Knight, les situations singulières ne permettent pas l'usage des probabilités. Il s'agit de situations inclassables, soit parce qu'elles sont inédites, en raison d'éléments nouveaux dont on ignore les propriétés, soit parce qu'elles sont complexes ou mal structurées. Il ne semble plus possible de modéliser les comportements en incertitude, ce qui explique l'importance pour les économistes de la distinction entre risque et incertitude. Pour autant, cette distinction n'est pas de nature dichotomique car « rien dans le monde de l'expérience, n'est absolument unique de même que deux choses ne sont pas absolument semblables » (p. 227). La singularité ne fait pas obstacle à la prévision mais exclut qu'elle s'appuie seulement sur des modèles probabilistes. Dans une telle situation, les personnes sont prises en tension entre 1) le besoin de se référer à des catégories générales, ce qui conduit à simplifier, en privilégiant des caractéristiques connues permettant une description et une prévision objectives ; 2) l'intérêt de saisir la singularité, ce qui requiert de percevoir des traits nouveaux et/ou de privilégier la complexité, en pensant ensemble des données hétérogènes.

• *Un exemple de complémentarité entre probabilité objective et jugement.* — Dans l'espace qui s'organise entre la prévision parfaite des jeux de hasard et l'ignorance totale du futur, les deux modalités de traitement de l'incertitude que sont les probabilités fréquentistes et le jugement, ont leur place. Le domaine du crédit bancaire en fournit une illustration. Pour qualifier une demande de crédit concernant une entreprise de petite taille (x), les banquiers recourent conjointement à l'« analyse objective » et à l'« analyse subjective » du risque. D'une part, le calcul de ratios financiers standard permet de classer la demande en référence à une catégorie donnée d'entreprises dont la probabilité de défaillance a été estimée (« x est un risque de niveau z ayant telle probabilité de faillite »), l'expert utilise une méthode formalisée fondée sur le regroupement de cas similaires en une classe homogène de risque. D'autre part, l'expert s'attache à saisir les spécificités de la demande de l'entreprise irréductibles à une catégorie donnée, ce qui mobilise l'intuition et l'expérience accumulée. Il formule une estimation personnelle de la qualité du demandeur qui complète le résultat obtenu par la méthode précédente. L'expert lui associe une mesure de la croyance dans son jugement qui peut prendre, ou non, la forme d'un ratio (« la demande de crédit est formulée par Dupont que je connais bien, son entreprise est viable et les résultats attendus permettront de rembourser le prêt, j'en ai l'intime conviction ou j'ai 90 % de chances d'avoir bien prédit »). Ce savoir-faire tacite ne conduit pas à une

approche quantifiée du risque et le caractère subjectif du jugement limite sa possibilité d'être communiqué à autrui [Rivaud-Danset, 1998].

• *Le profit, à mi-chemin entre le hasard et la bonne prévision.*

— Dans son ouvrage de 1921, Knight voulait expliquer le profit, un point aveugle de la théorie néoclassique qui le définit comme la simple rémunération du capital dans une économie concurrentielle. C'est en approfondissant la question de l'action en incertitude qu'il recherche les fondements du profit. Pour Knight, l'entrepreneur, même en l'absence de progrès technique, doit affronter une incertitude qui a une double origine : 1) le déroulement du processus de production et la productivité des travailleurs ; 2) les préférences des consommateurs et l'évolution de la demande. « Le profit a pour origine l'imprévisibilité inhérente, absolue des choses, le fait brut et abrupt que les conséquences de l'activité humaine ne peuvent pas être anticipées, dans la mesure où même un calcul de probabilité les concernant est impossible et n'a pas de sens » [Knight, 1921, p. 311]. Le profit est le résultat d'une action entreprise dans un monde où tout n'est pas planifié et où l'on n'agit pas en conformité parfaite avec son plan d'action. Le profit d'une firme, ayant pour origine la capacité de prévision de l'entrepreneur, ne peut pas être estimé *ex ante* et incorporé dans le prix des biens et services qu'offre cette firme, à la différence du salaire. Symétriquement, le résultat malheureux d'une activité engendre une perte, dès lors que la dépense ou le manque à gagner étant imprévisibles ne peuvent pas être traités comme un coût fixe. En revanche, un dommage aléatoire peut être transformé en un risque. Par exemple, les pertes dues à l'explosion des bouteilles de champagne se produisent avec une régularité telle qu'elles peuvent être évaluées en référence à la loi des grands nombres et imputées dans les coûts fixes et, donc, dans le prix [Knight, 1921, p. 213].

Pour Knight le profit suppose ainsi l'absence de procédure objective de prévision sur laquelle s'accorderaient les acteurs. C'est un revenu résiduel qui, n'étant pas susceptible d'une mesure *ex ante*, ne peut pas être stipulé dans un contrat.

« *Keynes sans l'incertitude, ce serait comme Hamlet sans le prince* ¹ »

Comme Knight, Keynes considère que l'incertitude ne peut pas être réduite à une affaire de probabilité. Pour preuve, sa distinction

1. Minsky [1975, p. 57].

entre l'incertain et l'improbable. L'absence de fondement scientifique sur lequel construire le calcul probabiliste du prix du cuivre dans vingt ans nous place dans une situation d'ignorance et d'incertitude. En revanche, pouvoir calculer la chance de gagner à la roulette range cette éventualité dans la catégorie du probable et non de l'incertain [Keynes, 1937, 113-114]. On reconnaît ici la distinction opérée par Knight entre l'incertitude non probabilisable et le risque. On retrouve également chez Keynes l'idée d'un continuum entre risque et incertitude. Dans de nombreuses situations, la prévision se forme, en partie, à partir de l'observation passée de régularités susceptibles d'être captées par des probabilités. C'est ainsi que le temps qu'il fera dans un avenir proche n'est, selon Keynes, que modérément incertain. Dans un tel contexte d'incertitude, formuler des prévisions relève d'une double logique : « Ce que nous voulons simplement rappeler, c'est que les décisions humaines engageant l'avenir sur le plan personnel, politique ou économique ne peuvent être inspirées par une stricte prévision mathématique, puisque la base d'une telle prévision n'existe pas ; c'est que notre besoin inné d'activité constitue le véritable moteur des affaires, notre intelligence choisissant de son mieux entre les solutions possibles, calculant chaque fois qu'elle le peut, mais se trouvant souvent désarmée devant le caprice, le sentiment ou la chance » [1936 (1942), p. 178].

L'histoire a retenu cette conception de l'incertitude que l'on trouve dans la *Théorie générale* [1936] et dans le célèbre article « The General Theory of Employment » [1937]. Est mise en avant la difficulté, voire l'impossibilité, de formuler des prévisions, en raison du manque de connaissances. Dans le cas extrême, l'incertitude dite radicale s'identifie à l'ignorance, à l'absence de base scientifique sur laquelle construire le calcul probabiliste, ce que résume la formule : « Tout simplement, nous ne savons pas » [Keynes, 1937, p. 114]. Suivant cette conception, c'est le futur et, en particulier, le long terme qui sont incertains.

Cette conception de l'incertitude n'a pas toujours dominé. Initialement, pour Keynes, c'est le conflit entre des corps de connaissances différents qui place les personnes en incertitude, c'est-à-dire dans une situation où il n'est pas possible de trancher. Dans le *Traité des probabilités* [1921], l'incertitude s'identifie au dilemme qui rend le choix difficile et/ou à l'incommensurabilité qui trouve son origine dans la pluralité des arguments possibles et des jugements de valeur qui ne peuvent pas se hiérarchiser. Suivant cette conception, les difficultés à traiter l'incertitude n'ont pas pour origine le caractère imprévisible du futur mais la complexité de la situation présente.

• *Le poids des anticipations.* — L'incertitude participe chez Keynes d'une vision novatrice des questions macroéconomiques dont rend compte l'expression originale d'« économie monétaire de production » [Barrère 1985 ; Combemale 2003]. « Une économie monétaire est essentiellement, comme nous le verrons, une économie où la variation des vues sur l'avenir peut influencer sur le volume actuel de l'emploi, et non sur sa seule orientation » [Keynes, 1936, préface de l'édition anglaise (1942), p. 15]. Sont ainsi pointés le rôle majeur des anticipations et, avec elles, la capacité du futur à influencer le niveau actuel de l'activité économique et l'emploi.

Pourquoi le qualificatif de monétaire est-il associé à l'incertitude ? Parce que la monnaie relie le présent et le futur [Keynes, 1936, p. 309]. Elle permet d'attendre, de reporter des décisions et sera d'autant plus désirée que la situation est troublée. Elle apporte de la sécurité et aide à supporter l'incertitude. « La possession de monnaie apaise notre inquiétude et la prime que nous exigeons pour nous dessaisir de la monnaie est la mesure du degré de notre inquiétude » [Keynes, 1937, p. 116]. D'où une conception originale du taux d'intérêt qui récompense non pas l'épargne mais la renonciation à la liquidité. La prime de liquidité, c'est-à-dire la rémunération exigée pour acheter des actifs moins liquides que la monnaie, dépend du degré de confiance que chaque individu place dans ses conjectures. Avec la prime de liquidité, on retrouve la décomposition en deux moments du jugement en incertitude : d'une part l'anticipation, d'autre part la confiance du décideur dans cette anticipation. Cette décomposition, explicite dans le *Traité des probabilités* de Keynes [1921], rejoint celle proposée par Knight, la même année.

Pourquoi privilégier la production et non l'échange ? Au moins deux arguments permettent de répondre : 1) la préférence pour la liquidité réduit les fonds prêtables et nuit en conséquence à l'investissement physique ; conserver la richesse sous forme de monnaie — thésauriser — entre en compétition avec la décision d'investir dans la production ; 2) la décision d'investissement dépend — *via* l'efficacité du capital — des anticipations concernant la demande selon le principe de la demande effective. Or l'investissement, par un effet multiplicateur, joue un rôle moteur dans la détermination du niveau de l'activité économique et donc de l'emploi. Cette construction confère un rôle clef aux prévisions et à l'état de la confiance des entrepreneurs dans le futur [Lavoie, 1985]. L'économie monétaire de production fait dépendre les variations de l'emploi des anticipations des agents. Dans une économie d'échange, le temps peut être aboli. En effet, la question essentielle est celle de la répartition des biens et les anticipations des agents sur le futur n'interviennent pas. Même si l'on se place dans la situation où les agents ont à effectuer

des échanges dans le futur, on suppose que toutes les éventualités sont connues ainsi que les prix des biens futurs (système complet de marché), le temps se trouve dès lors artificiellement « écrasé ». Dans une économie de production, ce sont les décisions qui engagent l'avenir (l'investissement) et qui dépendent de la vision du futur qui déterminent le présent. L'incertitude et les anticipations sont ainsi placées au centre de la théorie macroéconomique.

• *Anticipations et conventions.* — Quelle logique gouverne un système économique où les décisions majeures reposent sur des anticipations ? L'incertitude ne risque-t-elle pas d'inhiber toute action ? Répondre en se référant « aux esprits animaux » [Keynes, 1937], c'est-à-dire à l'esprit d'entreprise qui pousse les individus à agir, ne conduit-il pas à une autre impasse où l'économie ressemblerait à un casino, où tout serait le fruit du hasard ? Plus concrètement, la conception keynésienne apparemment très subjective du taux d'intérêt conduit-elle à ce qu'il y ait autant de taux que d'individus ? Keynes nous rassure, chaque individu n'est pas plongé dans une introspection permanente. Le taux observé n'est pas une simple moyenne calculée à partir d'opinions très dispersées, il reflète une opinion dominante. « Peut-être, au lieu de dire que le taux d'intérêt est au plus haut degré un phénomène psychologique, serait-il plus exact de dire qu'il est au plus haut degré un phénomène conventionnel, car sa valeur effective dépend dans une large mesure de sa valeur future telle que l'opinion dominante estime qu'on la prévoit. Un taux d'intérêt *quelconque* que l'on accepte avec une foi suffisante en ses *chances* de durée *durera* effectivement... » [Keynes, 1936, p. 219]. Le taux d'intérêt établi par convention stabilise les anticipations. Sans cette vision du futur largement partagée, la coordination serait très difficile et la décision très fragile. L'ancrage conventionnel permet que l'incertitude ne soit pas nécessairement préjudiciable au fonctionnement du système économique. Le courant français de l'économie des conventions cherchera à développer cette intuition (voir chapitre v).

Hayek : quand l'incertitude et l'ordre spontané sont indissociables

Comme Keynes et Knight, Hayek, qui reçut le prix Nobel d'économie en 1974, s'oppose à une double idée : 1) l'action en économie est la simple résultante d'une décision fondée sur le calcul ; 2) la décision est prise par des agents omniscients, c'est-à-dire des agents qui connaissent l'ensemble des données pertinentes. Cette opposition le conduit à se démarquer, lui aussi, de la théorie

Keynes et les probabilités épistémiques*

L'incertitude keynésienne désigne fréquemment une forme d'incertitude qui ne peut être saisie ni par les probabilités des jeux de hasard ni par les probabilités fréquentistes. Pour autant, les individus ne sont pas démunis pour fonder des anticipations. La *Théorie générale* insiste sur l'importance des croyances partagées par un groupe (conventions). Dans le *Traité des probabilités* [1921], Keynes avait suivi une autre voie en privilégiant la méthode inductive. Celle-ci déduit les propositions sur le futur de la connaissance des faits [Bateman, 1996]. En dehors des sciences exactes, une proposition obtenue par induction affirme, non pas qu'une chose est ainsi, mais qu'elle est probablement ainsi, relativement aux observations passées. Comme Keynes, prenons l'exemple des cygnes, classique en philosophie. Supposons que je n'ai observé dans ma vie que des cygnes blancs, si je généralise ces observations par une relation universelle invariable, du type « tous les cygnes sont blancs », celle-ci s'effondre à la première exception, par exemple l'existence de cygnes noirs. Le raisonnement par induction n'aboutit pas à une conclusion universelle mais à un argument non conclusif, comme dans la formule « la plupart des cygnes sont blancs » ou « la probabilité qu'un cygne soit blanc est de tel ordre ». Autrement dit, la méthode inductive conduit à formuler des jugements ou des opinions qui sont de l'ordre du probable et non pas du certain [1921 (1973), p. 244-245].

On en déduit que, dans toute généralisation à partir de l'observation des faits, un élément d'incertitude est plus ou

moins présent. Multiplier les observations n'y changera rien. Ce que nous prévoyons, en recourant à l'induction, est susceptible de se produire mais ne se produira pas nécessairement. La généralisation à partir de l'expérience est valable lorsqu'elle s'appuie sur un raisonnement logique du type : si h alors il est possible que a . Keynes nomme probabilité cette relation entre des propositions, il la définit comme étant logique et n'étant pas subjective, au sens où elle est conditionnée par l'état des connaissances et n'est pas sujette à des opinions particulières.

En fait, l'opposition classique entre probabilité objective et subjective n'éclaire pas sa conception des probabilités qui est très originale, notamment parce qu'elles ne s'expriment pas sous forme numérique. Celle-ci se différencie, de celle de son contemporain Ramsey, pour qui la probabilité attribuée à la proposition a étant donné h est de l'ordre de la croyance personnelle. Ramsey, comme Keynes, soumet la croyance à un impératif de cohérence du raisonnement mais montre qu'elle peut être mesurée, ce qui ouvre sur la théorie calculatoire de la décision en incertitude [Ramsey, 1931 et Hacking, 1975].

* Hacking distingue les probabilités fréquentistes et les probabilités épistémiques. Les premières mesurent la fréquence d'occurrence d'un phénomène aléatoire, les secondes mesurent le degré de connaissance ou de croyance dans une hypothèse ou un jugement [Hacking, 1975].

néoclassique qui, dans la première moitié du xx^e siècle, repose sur l'hypothèse d'information parfaite. Plus généralement, il s'insurge contre l'influence de la pensée cartésienne qui associe raison et certitude. « Depuis Descartes, incertitude et raison ne font plus bon ménage. L'action rationnelle est devenue "l'appellation réservée à

l'action déterminée entièrement par une vérité connue et démontrable" » [Hayek, 1973, p. 11]. Pour Hayek, les individus ne peuvent avoir qu'une connaissance fragmentée des faits en raison de la complexité du monde. Ils sont guidés dans leurs actions par les coutumes, les règles héritées de l'histoire. « L'homme est tout autant un animal obéissant à des règles qu'un animal recherchant des objectifs » [*ibid.*, p. 13]. Autrement dit, les règles sociales guident l'action en économie tout autant que les conséquences anticipées de nos actes. Le duo « connaissance partielle/règle » (Hayek) fait écho au couple « incertitude/convention » (Keynes).

Comme Keynes, le refus de l'omniscience conduit Hayek à attribuer à l'incertitude une influence décisive dans sa conception de l'économie. Cependant, leurs recommandations en politique économique sont largement antagonistes. Pour Hayek, dans une société où l'information est incomplète, fragmentée et où nous dépendons de la division des connaissances, la coordination des actions ou des projets a besoin des prix. Ceux-ci constituent le mécanisme impersonnel idéal de transmission des informations pertinentes. Il est essentiel que les prix donnent des signaux fiables, qu'ils expriment au mieux les pénuries, les opportunités de profit. Pour être fiables, ils ne doivent pas être manipulés ou faussés par l'inflation. Hayek s'oppose donc à l'intervention des pouvoirs publics qui introduirait des modifications artificielles des prix relatifs. Plus largement, Hayek est un partisan du libéralisme : la fragmentation du savoir rend impossible l'action intelligente de l'homme d'État, le retour à l'ordre doit se faire « naturellement », sans entraver les forces spontanées du marché [Dostaler, 2001].

2. Les héritiers : Autrichiens, post-keynésiens, évolutionnistes

Trois courants ont délibérément cherché à exploiter les intuitions de ces pères fondateurs, construisant leur programme de recherche comme une alternative à la théorie néoclassique : les Autrichiens [Lachmann, 1976 ; O'Driscoll et Rizzo, 1985 ; Kirzner, 1997], les post-keynésiens [Minsky, 1975 ; Davidson, 1991 ; Lavoie, 1992] et les évolutionnistes [Nelson et Winter, 1982 ; Dosi, 1991]. Comment contribuent-ils à une économie de l'incertitude ?

Pour ces trois courants, le refus de l'omniscience va de soi. Le savoir des agents est limité et celui du praticien n'est pas le même que celui de l'économiste modélisateur.

Il est donc exclu que les anticipations puissent être rationnelles, même dans la version faible de cette hypothèse. Selon le principe des anticipations rationnelles au sens faible, les agents recherchent l'information jusqu'à un certain point, en font le meilleur usage et arrivent à une prévision parfaite. Rejeter ce principe revient à admettre que l'acquisition de connaissances est plus qu'une affaire d'accumulation d'information, cette dernière s'identifiant à un stock de données préexistant auquel les agents ont plus ou moins accès. La connaissance implique de prendre en compte les capacités d'interprétation et d'organisation de ces connaissances. De plus, certaines connaissances s'acquièrent grâce à l'expérience et à l'apprentissage et non par simple accumulation d'information [Foray, 2000]. Lorsque les personnes sont confrontées à des phénomènes ou à des situations complexes, l'incertitude peut persister en dépit de l'acquisition d'une masse d'informations. Même les situations présentes peuvent être mal connues.

Autrichiens, évolutionnistes et post-keynésiens s'accordent pour considérer que le temps n'est pas un voile, c'est-à-dire qu'il joue un rôle actif et ne peut se réduire à une succession de périodes homogènes comme dans une suite mathématique. Dès lors, la vision mécanique et optimiste de l'équilibre spontané des processus économiques est rejetée.

Les Autrichiens ont développé une réflexion critique visant à alerter les économistes sur les dangers d'une approche où le temps est réduit à une succession de points et où la nouveauté n'a pas sa place [O'Driscoll et Rizzo, 1985]. Pour rendre compte de l'émergence du nouveau et du caractère créateur de l'activité économique, ils mettent en avant des traits du caractère humain, tels que l'imagination, qui sont incarnés par la figure de l'entrepreneur. Chacun se comporte, en partie, comme un entrepreneur, c'est-à-dire qu'il est capable de repérer des opportunités — *alertness* ou vigilance — et de faire des découvertes [Kirzner, 1997]. Les décideurs commettent des erreurs, parce qu'ils ne peuvent pas prévoir parfaitement et parce qu'ils sont victimes de l'illusion monétaire. Les Autrichiens s'appuient sur l'incertitude radicale et la non-neutralité de la monnaie, pour considérer comme négative toute politique économique interventionniste. C'est leur principale différence avec les post-keynésiens.

Les post-keynésiens refusent d'interpréter le message de Keynes à partir d'une logique de l'équilibre. Ils s'élèvent contre la lecture dite hydraulique de Keynes qui conduit au modèle IS-LM. Ils proposent une lecture qualifiée de fondamentaliste où l'instabilité des anticipations joue un rôle central. C'est particulièrement dans le domaine de la monnaie et de la finance que les erreurs d'anticipation entraînent des phénomènes de surréaction et déstabilisent le système. Quand les prévisions trop confiantes des banquiers confortent l'optimisme des entrepreneurs, alors l'économie peut s'emballer et à la suraccumulation succédera une phase de purge [Minsky, 1975]. L'étude des déséquilibres se substitue à celle des équilibres et nous introduit dans un monde de frictions où les ajustements dans le temps peuvent être longs et douloureux.

Les économistes évolutionnistes, comme leur nom l'indique, s'intéressent aux évolutions de long terme. Partant de l'association incertitude-temps, ils explorent de nouveaux champs en économie, en particulier l'innovation et la dynamique des systèmes. Le résultat des processus dynamiques (diffusion, évolution, sélection) ne peut être prédit. Des causes accidentelles survenues au cours du temps peuvent avoir de grandes conséquences. Des effets cumulatifs dans le temps conduisent à des phénomènes d'irréversibilité, entendue au sens d'enfermement dans des situations ou dans des processus que l'on ne peut pas modifier. Le rôle actif du temps se traduit, par exemple, par l'enfermement dans des technologies qui ne sont pas nécessairement les plus efficaces comme l'illustre l'exemple célèbre du clavier Qwerty [David, 1985].

Des acteurs hétérogènes et les difficultés de la coordination

Reconnaître aux agents un savoir limité et admettre que leurs connaissances dépendent de l'expérience et de l'apprentissage, les savoirs se construisant au fil du temps, conduit logiquement à des agents économiques hétérogènes. L'horizon temporel des agents constitue une autre source de différenciation. On peut ainsi distinguer : le très court terme qui caractérise le spéculateur, l'horizon à court ou moyen terme du producteur, et l'horizon à plus long terme du producteur voire même du consommateur quand ils s'engagent dans la décision d'investissement. Cette hétérogénéité des agents est prise en compte différemment selon les courants. Pour les évolutionnistes, les différences d'information sont essentielles à la dynamique des systèmes. Pour les post-keynésiens, la gamme des horizons temporels montre que les interactions macroéconomiques sont riches et complexes. « Il est très difficile d'appréhender

Hasard, sélection et évolutionnisme

Le programme de recherche évolutionniste s'est construit comme une alternative à la théorie néoclassique, postulant l'analogie entre les lois d'évolution de variables économiques et les lois biologiques d'évolution des espèces [Nelson et Winter, 1982]. Les évolutionnistes se focalisent sur la dynamique de processus complexes et non sur l'analyse des équilibres.

Les processus évolutionnistes suivent une trajectoire qui peut aboutir à de multiples équilibres, la sélection de l'un d'entre eux étant dépendante du sentier, c'est-à-dire du chemin parcouru au cours du temps (propriété de « *path dependency* »). Ce chemin, qui n'a rien de linéaire, est soumis à de nombreux aléas, appelés « petits événements historiques », ainsi qu'à des mécanismes de renforcement. Outre le hasard, un principe de sélection entre en jeu pour rendre compte du fait que les objets d'étude que les évolutionnistes privilégient (firmes, techniques et même les institutions) n'évoluent pas selon de simples processus aléatoires. Dans certains modèles, tous les aléas n'ont pas le même impact et certains d'entre eux sont sans effet. Les aléas ne modifient la variable à expliquer que sous certaines conditions. Une sorte de filtre est introduit pour montrer comment, dans les sociétés humaines, les règles sociales réduisent et organisent l'espace des possibilités. Autrement dit, le modélisateur introduit la Nature qui engendre des aléas et la Société qui produit des règles de comportement.

Jean-François Laslier et André Orléan [2002], comme Gilbert Laffond et Jacques Lesourne [1981], montrent comment ce type de modélisation permet d'expliquer le fonctionnement du marché

de l'emploi. Sur ce marché, les salariés sont en situation d'information imparfaite. L'information qu'ils détiennent sur les postes à pourvoir et sur les salaires correspondants est issue de leur lecture de journaux (on suppose dans le modèle que cette information provient de tirages effectués au sein d'échantillons aléatoires). Salariés et employeurs sont supposés se rencontrer de façon aléatoire sur de multiples périodes. À chacune d'entre elles, ils ont la possibilité de contracter ou non, selon leurs attentes personnelles. Leurs exigences sont révisées à chaque période en fonction des situations rencontrées au cours des périodes précédentes. Par exemple, le modèle pose qu'une entreprise dont le poste est occupé à chaque période révisera à la baisse le salaire d'embauche. De même, si un salarié fait face n fois à la même situation, alors il aura une réaction d'un certain type, donné par le modélisateur. L'évolution différenciée des acteurs, selon leur histoire, est saisie dans les modèles évolutionnistes par la séquentialité du modèle.

Laslier et Orléan montrent que cette dynamique conduit à une « auto-organisation » du marché, celui-ci étant finalement organisé autour de deux populations numériquement stables, l'une de salariés, l'autre de non-salariés. Mathématiquement, ils montrent que le système évolue vers une situation avec des états absorbants. Un état est dit absorbant lorsqu'une fois qu'il a été atteint on ne peut plus en sortir. La situation au sein de laquelle tous les agents non employés ont réduit leur exigence salariale au minimum constitue par exemple un état absorbant.

l'attitude face au temps avec un ensemble d'équations simultanées » [Chick, 1997, p. 432].

Pour les Autrichiens, avec des agents hétérogènes, le fonctionnement des marchés devient beaucoup plus compliqué que celui du

marché walrasien. Ils ne reprennent pas à leur compte l'idée d'Hayek selon laquelle les prix sont un guide d'autant plus précieux que la connaissance est limitée. Au contraire, ils doutent que le prix puisse donner aux deux parties l'information suffisante pour se coordonner. En effet, le prix d'un bien s'analysant chez les Autrichiens comme sa valeur actualisée doit intégrer tous les éléments futurs susceptibles d'influencer cette valeur. Par définition, il ne peut pas intégrer les événements imprévisibles et ne résume pas parfaitement toute l'information, une partie étant inaccessible.

Supposons pour illustrer cette question qu'un particulier souhaite acheter un véhicule avec l'intention de le revendre dans quatre ans. Le prix qui correspond à la valeur actualisée de cette automobile doit intégrer les événements futurs de l'industrie automobile et des activités qui lui sont liées. Le prix que l'acheteur est prêt à payer intègre un certain nombre d'anticipations, par exemple la sortie d'un nouveau modèle dans la même gamme et la diffusion d'une nouvelle technique de freinage. Néanmoins, toutes les évolutions ne sont pas prévisibles, il n'imaginera pas l'apparition d'un carburant non polluant qui rendrait obsolètes tous les véhicules ne l'utilisant pas. Les anticipations de l'acheteur ayant toutes les chances de diverger de celles du vendeur, la résolution de ce problème de coordination sera plus complexe que sur un marché walrasien.

Cet exemple aurait pu illustrer le rôle des anticipations dans une perspective keynésienne. Supposons que notre acheteur prévoie que l'industrie automobile fasse prochainement l'objet de bouleversements de taille tels que l'arrivée massive d'un nouveau type de carburant, alors il préférera reporter sa décision d'investissement et conserver sa richesse sous forme de liquidité.

Quand les héritiers divergent

Prendre l'incertitude comme fil conducteur nous a conduits à rassembler ces courants. Néanmoins ce n'est pas une pratique habituelle. En effet, leurs méthodes et leurs thèmes sont divergents et leurs recommandations de politique économique peuvent même être antagonistes.

Les Autrichiens s'opposent fermement à toute formalisation, les évolutionnistes, au contraire, la revendiquent ; la plupart de leurs modèles prédictifs se fondent sur une simulation informatique et s'inspirent souvent des sciences dures.

Logiquement, les Autrichiens se prononcent pour le « *free market* », c'est-à-dire pour la liberté d'entrée sur tout marché où existent des opportunités de profit. Leur confiance dans le marché comme ordonnateur de la vie économique est confortée par leur

méfiance vis-à-vis du pouvoir central dont les capacités de prévision objectives sont, par définition, limitées. Pour les post-keynésiens, en revanche, les marchés ont besoin de l'assistance de l'État pour fonctionner plus efficacement. Pour autant, ils n'ont pas une vision simpliste de l'efficacité des politiques de relance et les post-keynésiens les plus tournés vers l'incertitude radicale, tels Schakle, sont même très sceptiques vis-à-vis de l'interventionnisme.

Conclusion

Le message de Knight et Keynes est clair. L'incertitude est souveraine, son existence impose au théoricien une règle : toute simplification est dangereuse. Il n'existe pas d'outil magique permettant de saisir toute la complexité d'un mode incertain, le raisonnement probabiliste n'a qu'une capacité limitée à rendre compte de phénomènes incertains et de la façon dont les êtres humains font des prévisions. Ne faites pas une confiance aveugle aux probabilités. Ne prêtez pas aux agents économiques qui les utiliseraient une clairvoyance divine. Le message de Hayek est encore plus radical.

À sa façon, chaque courant d'héritiers prend en charge le testament. Les post-keynésiens nous mettent en garde contre une vision épurée des enchaînements macroéconomiques et, en particulier, contre une lecture simpliste qui réduit la *Théorie générale* de Keynes à quelques équations. Dans la même perspective critique, les Autrichiens développent une réflexion de nature philosophique sur le temps. Les évolutionnistes rejettent une vision mécanique des processus de long terme au profit d'une approche où la complexité peut se penser sur le registre de la biologie.

En bref, leur message devient : en économie, il est interdit de penser la coordination, les enchaînements macro, la dynamique d'un système à partir de représentations mécaniques. Ce faisant, ces héritiers, minoritaires, radicalisent le message initial de Knight et Keynes. Rappelons que, selon ces derniers, le théoricien, pour être au plus proche de la réalité, doit poser que les agents font appel à la fois au calcul et à d'autres arguments pour agir en incertitude (jugement intuitif, pour Knight, conventionnel, pour Keynes).

La distinction entre le risque qui est mesurable par les probabilités et assurable, d'une part, l'incertitude, d'autre part, devient un enjeu théorique majeur. Pour les héritiers, ces deux caractéristiques du risque ne s'appliquent pas à l'incertitude. De même, la théorie du choix rationnel, qui prescrit à chacun d'agir de façon à maximiser son bien-être, ne s'applique pas dans les situations dites d'incertitude. Sa validité descriptive et normative est contestée. Cette orientation tranche avec celle des économistes néoclassiques

qui, eux, étendent le champ des probabilités aux situations d'incertitude, préservant ainsi la théorie du choix rationnel. Cependant, dès que le monde des certitudes est abandonné, l'équilibre pose problème. L'édifice néoclassique menace de s'écrouler et les théoriciens de ce courant s'emploient à le fortifier (chapitre II).

II / L'incertitude, un trouble-fête pour les néoclassiques

Dans un environnement de concurrence parfaite, les prix et la qualité des biens sont connus. La théorie néoclassique walrasienne montre alors l'existence d'un prix unique sur le marché. Les agents économiques sont *price takers*, le prix auquel s'établit l'échange résulte de l'action efficace d'un agent fictif, le commissaire-priseur, qui a pour rôle de synthétiser toute l'information disponible et de proposer un prix tel que le volume échangé soit optimal. La levée de l'hypothèse de perfection de l'information, qu'elle concerne les prix, la qualité ou les comportements, modifie profondément la conception suivant laquelle le marché s'équilibre naturellement.

En introduisant l'imperfection de l'information, les théoriciens ont ouvert la boîte de Pandore de l'analyse néoclassique du marché et ont mis en question les principaux résultats de l'économie concurrentielle [Stiglitz, 2002]. Le prix unique d'équilibre disparaît, voire le marché lui-même. L'incertitude est devenue un trouble-fête, ébranlant la construction néoclassique et impulsant la théorie de l'information qui propose des contreforts pour maintenir debout l'édifice. La théorie des jeux prolonge ces résultats en se focalisant sur le rôle des anticipations.

Un point commun rassemble ces analyses : l'incertitude est assimilée à un risque. L'hypothèse, parfois peu réaliste, selon laquelle l'incertitude est probabilisable permet de conserver la modélisation.

1. Quand le manque d'information sème le désordre

L'incertitude sur les prix n'a pas besoin d'être forte pour déstabiliser les résultats de la théorie néoclassique. Ne sera pas évoqué ici le problème posé par l'ignorance du prix du cuivre dans trente

ans (Keynes) mais, plus simplement, le manque d'information sur les prix pratiqués effectivement par un ensemble restreint de vendeurs. Que se passe-t-il quand, comme dans la vie courante, les prix ne sont pas spontanément connus par les acheteurs ? Le prix d'équilibre est-il toujours unique ?

De même, l'incertitude sur la qualité des biens bouleverse la conception de l'échange. Le processus qui permet d'atteindre l'équilibre est grippé, modifier le prix n'induit plus un ajustement automatique des quantités permettant d'optimiser les échanges. Que se passe-t-il quand le commissaire-priseur ne peut plus jouer son rôle ?

Enfin, quand l'incertitude concerne les comportements, apparaissent de nouveaux problèmes de coordination. Ce dernier champ a donné lieu à un foisonnement de recherches.

Stigler et la quête coûteuse des prix

Que se passe-t-il quand les consommateurs ne sont pas parfaitement informés des différents prix pratiqués ? Supposons qu'un agent cherche à acheter une automobile et que les prix affichés diffèrent d'un point de vente à l'autre pour un même modèle de véhicule. Sans doute, une part de l'étalement des prix pratiqués résulte-t-elle des différences de services proposés par les offreurs mais il demeure une part non expliquée qui est imputable, selon George Stigler, à la diversité de coûts de recherche d'information supportés par les consommateurs (*search*).

Stigler [1961] propose de modéliser le comportement d'individus faisant face à une diversité des prix sur un marché décentralisé, c'est-à-dire un marché où il n'existe pas d'agent ou de procédure permettant de centraliser les offres et les demandes. Pour cela, il introduit un coût de recherche d'information dans le programme d'optimisation. Si la dispersion des prix est suffisamment importante au regard des coûts de recherche, il sera avantageux de visiter plusieurs magasins. Dans ce modèle, l'acheteur qui ne connaît pas parfaitement les prix bénéficie néanmoins d'une information importante puisqu'il sait comment sont distribués les prix. Autrement dit, l'acheteur d'une voiture a en tête le tableau des prix mais ne sait pas identifier les vendeurs, de sorte qu'il ignore où sont situés les magasins des vendeurs pratiquant les prix les plus bas. Il connaît également le coût de la recherche qui correspond au temps et à l'effort liés à la visite des différents magasins pour s'informer des prix pratiqués. Cette dernière hypothèse réduit la portée du modèle mais lui est indispensable. En effet, supposer que l'agent ne connaisse pas le coût de recherche implique qu'il doive l'estimer

avant de rechercher l'information elle-même. Mais cette estimation déclenche une recherche d'information spécifique qui doit aussi être évaluée, etc. On entre ainsi dans un processus de régression infini qui empêche toute action. La seule façon de rendre le modèle viable est de poser que le coût de recherche de l'information est connu. [Mongin et Walliser, 1988].

Doté de toutes ces informations, l'acheteur va devoir décider du nombre (n) de recherches à effectuer, il choisit ensuite parmi les n magasins visités celui qui lui propose le prix le plus bas. Étant donné que l'agent économique modélisé a des capacités de calcul illimitées, il détermine le nombre optimal de recherches selon le principe utilitariste suivant : l'agent a intérêt à enquêter tant que le coût marginal de la recherche est inférieur au gain marginal attendu. Dans les modèles plus récents du *search*, l'agent ne fixe pas au préalable le nombre de démarches à effectuer mais compare les prix proposés par les magasins qu'il visite à un seuil de prix qu'il s'est fixé et stoppe sa recherche dès que le prix constaté est inférieur à ce seuil.

• *Quand les vendeurs s'en mêlent.* — Le modèle du *search* de Stigler explique la dispersion des prix par la diversité des coûts de recherche des consommateurs. De nombreux travaux ont, à partir des années 1970, tenté d'intégrer le comportement des vendeurs pour étudier la possibilité d'atteindre ou non un prix d'équilibre. L'incertitude donne un pouvoir — dit de marché — aux agents. Les prix ne s'imposent plus et certains agents ont la capacité de les manipuler. Les comportements deviennent stratégiques [Gabszewicz, 2003]. Peter Diamond est le premier à avoir montré, en 1971, comment l'existence de coûts de recherche positifs, si faibles soient-ils, conduisait, dans certains cas, à l'émergence d'un prix de monopole. La logique générale des travaux portant sur cette question est la suivante.

Supposons que tous les producteurs aient le même coût de production et vendent un produit homogène et que tous les consommateurs aient la même fonction de demande. Ils connaissent, en outre, la fonction de répartition des prix mais ignorent les prix pratiqués par chaque magasin. Le coût de recherche est noté c . Le nombre de magasins est donné et noté n . Si tous les magasins pratiquent le prix de concurrence p_c , alors un magasin a intérêt à pratiquer individuellement le prix $p^* = p_c + \epsilon$ avec ϵ un nombre positif faible. Le consommateur qui s'est rendu par hasard dans le magasin pratiquant ce prix (et connaissant par ailleurs la fonction de distribution des prix) ne sera pas incité à aller ailleurs si p^* est inférieur à $p_c + c$, c'est-à-dire le prix affiché dans tous les autres magasins augmenté

du coût de recherche. On rompt ainsi avec la loi du prix unique. Le prix p_c qui prévaut en concurrence parfaite ne peut plus constituer un équilibre. Le prix p^* , lui non plus, ne constitue pas un prix d'équilibre potentiel. En effet, si tous les magasins s'ajustaient sur p^* , alors, à nouveau, un magasin aurait intérêt à fixer un prix égal à $p^* + \epsilon$, etc. Le seul prix d'équilibre susceptible de s'imposer est le prix de monopole, alors même que coexistent plusieurs magasins sur le marché. L'existence d'un prix d'équilibre unique dépend de plusieurs paramètres, parmi lesquels figure le nombre de magasins présents sur le marché. En effet, reprenons la situation où tous les magasins affichent le prix de monopole. Un magasin peut avoir intérêt à diminuer unilatéralement son prix si cette baisse attire un nombre conséquent de consommateurs. Pour ce faire, la réduction de prix doit être nécessairement supérieure au coût de recherche supporté par les consommateurs, or celui-ci dépend positivement du nombre de magasins. Si ce nombre est élevé, le coût l'est également et les consommateurs sont, en conséquence, peu enclins à rechercher le magasin le « meilleur marché », de sorte que seul prévaut le prix de monopole.

Les résultats de l'analyse néoclassique en information parfaite se trouvent bouleversés, puisque, avec l'introduction d'une incertitude sur les prix, le bien-être sera d'autant plus élevé qu'il y aura moins d'entreprises sur le marché. Les perturbations induites par l'introduction d'une incertitude sur les prix pratiqués viennent ainsi remettre en cause l'efficacité du marché comme mode d'allocation efficace des ressources (pour un *survey*, voir [Stiglitz, 1989 ; Rocheteau, 2001]).

Akerlof et l'asymétrie d'information sur la qualité

Le désordre provoqué, au sein de la théorie néoclassique, par l'introduction d'une incertitude sur les prix s'amplifie lorsque l'on admet que les parties prenantes d'une transaction ne disposent pas de la même information sur la qualité d'un bien. Supposons, pour illustrer ce cas, qu'un acheteur cherche à se procurer un véhicule sur le marché de l'occasion. Il ignore la qualité du véhicule. En revanche, le vendeur la connaît. Tel est le problème étudié par George Akerlof [1970] dans un article fondateur sur les désordres induits par l'existence d'asymétrie d'information.

- *Le « marché des rossignols »*. — Sur un marché, des produits de différentes qualités sont échangés. Supposons que seules deux qualités existent, les bonnes (B) et les mauvaises (M) voitures que l'on appelle communément des « rossignols ». Seuls les vendeurs

connaissent la qualité des véhicules avec certitude. Les acheteurs l'ignorent mais l'incertitude est probabilisable. Cette conception de l'incertitude permet de poser une hypothèse forte : les agents sont supposés connaître la fonction de répartition des qualités (par exemple, ils savent que 40 % des voitures en vente sont de bonne qualité). Notons prob la probabilité d'acheter une bonne voiture et $1-\text{prob}$ la probabilité d'en acheter une mauvaise. Le prix maximal que les acheteurs acceptent de payer s'élève à A_b pour une bonne voiture et A_m pour une mauvaise. Les vendeurs, de leur côté, refusent de baisser leur prix en dessous de V_b pour une voiture de bonne qualité et de V_m pour une mauvaise. Des échanges se réalisent si la valeur attachée aux véhicules par les acheteurs est supérieure ou égale à celle des vendeurs $A_b \geq V_b$ et $A_m \geq V_m$.

Si l'on était en information parfaite, il y aurait deux marchés avec deux prix distincts, l'un pour les voitures de bonne qualité, l'autre pour les voitures de mauvaise qualité. Le prix en vigueur pour les premières serait compris entre V_b et A_b (avec $A_b > V_b$), tandis que celui des rossignols serait compris entre V_m et A_m , la valeur obtenue dépendant du pouvoir de négociation des parties en présence.

Du fait de l'asymétrie d'information, sous un même prix, peuvent se cacher des qualités différentes — bonne *versus* mauvaise, pour simplifier. En conséquence, le prix n'est plus un signal de qualité, à la différence du prix en information parfaite. Comment se comportent les acheteurs sur ce type de marché ? S'ils acceptent un prix élevé, correspondant au prix qu'ils auraient payé en information parfaite pour un véhicule de bonne qualité, ils risquent d'effectuer une mauvaise affaire, puisque, par hypothèse, ils ignorent sa qualité. Rappelons que la possibilité d'acquérir de l'information grâce au contrôle technique, par exemple, est écartée du modèle. Pour éviter que le prix payé pour un véhicule de mauvaise qualité soit trop élevé, les acheteurs proposent au vendeur un prix moyen. Celui-ci est obtenu en pondérant le prix attaché à chaque catégorie de voiture par la part de chaque qualité sur le marché. Ainsi, le prix d'une bonne (mauvaise) voiture est multiplié par la probabilité d'avoir une bonne (mauvaise) voiture. Ce prix se définit comme suit : $\text{prob} \cdot A_b + (1-\text{prob}) \cdot A_m$. Il n'est pas sûr qu'il existe pour ce « prix moyen » des vendeurs, surtout si la proportion de rossignols est élevée. Les vendeurs de véhicules de bonne qualité risquent fort de sortir du marché, refusant de vendre à un prix moyen qu'ils jugent insuffisant.

Akerlof montre comment, dans certaines conditions, le marché peut être conduit à disparaître totalement. En raison de la nature de l'effet de sélection qu'elle engendre, l'asymétrie d'information *ex ante* est qualifiée selon les auteurs de sélection adverse ou d'anti-sélection, qui sont des traductions plus ou moins heureuses de

l'expression anglaise « *adverse selection* ». Dans la suite de l'ouvrage, nous emploierons le qualificatif de « sélection perverse » pour qualifier ce type de sélection par référence aux effets pervers : tout comme un effet pervers correspond à des effets induits négatifs, la sélection perverse conduit, inintentionnellement, à ne conserver que les « mauvais ».

- *Le rationnement du crédit par le banquier mal informé.* — L'effet de sélection perverse apparaît aussi lorsque le déficit d'information touche les offreurs. Joseph Stiglitz et Andrew Weiss [1981] ont élaboré un modèle sur le rationnement du crédit où les banques ne connaissent pas la variable clef, c'est-à-dire la qualité des emprunteurs, celle-ci se définissant par leur probabilité de défaillance. Ils montrent comment, quand la demande excède l'offre, il n'est pas avantageux pour la banque d'augmenter le taux d'intérêt afin d'équilibrer l'offre et la demande.

Supposons que la population d'emprunteurs comprenne des individus ayant des projets très risqués et d'autres des projets peu risqués. L'espérance de gain est la même dans les deux populations et les montants d'emprunt demandés sont identiques. Les emprunteurs ayant des projets peu risqués ont une probabilité de défaillance faible. En cas de réussite, le rendement est modéré, ce qui ne leur permet pas de supporter un taux d'intérêt élevé. Les emprunteurs investissant dans des projets à haut risque, autrement dit ayant une probabilité élevée de défaillance, obtiendront un rendement élevé si leur projet réussit, ce qui leur permet de supporter un taux d'intérêt supérieur. Dans une telle configuration, à la suite d'une hausse du taux d'intérêt en réaction à un excès de demande relativement à l'offre, les « bons emprunteurs » s'en vont, seuls demeurent les « mauvais emprunteurs » dont la probabilité de ne pas rembourser est grande. Au total, le profit de la banque qui augmente son prix d'offre (le taux d'intérêt) sera moindre que si elle avait rationné l'ensemble des emprunteurs en ne leur allouant pas la totalité des capitaux demandés. Sur le marché du crédit, le rationnement s'opère par les quantités offertes et non par le prix, pour que le taux d'intérêt soit maintenu à un niveau acceptable par les emprunteurs à faible risque. La banque aura ainsi évité d'être victime de sélection perverse mais elle le fait en rationnant des emprunteurs prêts à supporter un taux d'intérêt plus élevé et donc en freinant la croissance de l'activité économique, ce qui est sous-optimal.

- *Quand l'ajustement ne peut plus s'effectuer par les prix.* — La logique qui conduit au rationnement du crédit s'observe également sur le marché du travail. L'argument standard voudrait que, lorsque

l'offre de main-d'œuvre excède la demande, les salaires diminuent. Supposons qu'un demandeur d'emploi propose à un employeur d'être embauché à un salaire inférieur à celui en vigueur dans l'entreprise pour sa qualification, l'employeur rejettera son offre, pensant qu'elle correspond à une productivité de l'employé si faible que le coût moyen de production risquerait d'augmenter. Ainsi même s'il existe un équilibre walrasien (c'est-à-dire un vecteur de salaires qui ajuste l'offre et la demande), il ne s'agit pas d'un équilibre concurrentiel. Ici, un déséquilibre sur le marché du travail ne se résout pas par la baisse des salaires.

L'introduction d'une asymétrie d'information avant la réalisation de la transaction remet en cause deux piliers de la théorie économique : les prix indiquent la qualité ; l'équilibre de l'offre et la demande s'obtient par le mécanisme des prix. Lorsque l'incertitude est synonyme d'asymétrie d'information, les prix ne reflètent plus la qualité du produit et, sous un même prix, peuvent se cacher deux biens de qualité différente. En outre, le prix ne constitue plus un mécanisme d'ajustement simple du marché puisque sa variation a un impact sur les comportements. Ce prix peut être classiquement le prix d'un bien mais aussi le taux d'intérêt ou encore le salaire.

Arrow et l'aléa moral

Nous avons vu que l'incertitude pouvait affecter les prix et la qualité. Qu'en est-il des comportements des individus, dans une relation contractuelle ? Dans les modèles ci-dessus, il était admis qu'une des deux parties manquait d'information sur la qualité de l'autre, cette hypothèse n'impliquant nullement que l'agent informé envoie à l'autre, par opportunisme, de fausses informations. Il était également admis que les catégories auxquelles appartenaient les individus étaient prédéterminées. Ainsi, pour Stiglitz et Weiss, la qualité de l'emprunteur était définie une fois pour toutes et ne se modifiait pas en cours de route. Avec la notion d'aléa moral, une nouvelle étape est franchie. Cette notion initialement limitée au domaine de l'assurance désigne, classiquement, le relâchement de la responsabilité d'un individu lorsqu'il a contracté une assurance. Elle traduit l'existence d'une incertitude sur le comportement qui est d'autant plus compliquée à lever que l'assuré est, par définition, dans un monde aléatoire. Deux sources d'incertitude se combinent, l'une vient des aléas engendrés par la Nature et peut être traitée comme un risque probabilisable, l'autre vient du changement de comportement de l'agent. Peut-on toujours se coordonner de façon optimale lorsque l'aléa moral est introduit ?

Kenneth Arrow a été le premier économiste à poser cette question dans un article sur le marché des soins médicaux [1963]. Il s'interroge sur l'existence d'un ensemble de marchés de l'assurance qui donnerait aux agents économiques demandeurs la possibilité de se prémunir contre tous les types d'aléas. Un tel ensemble, que l'on appelle un marché complet, n'existe pas (encadré). Le transfert de certains risques est difficile, voire impossible, dès lors qu'offrir une assurance influence la demande de services fournis¹. Selon Arrow, un système d'assurance efficace requiert que l'élément assuré soit hors du contrôle de l'individu. « Malheureusement dans la vie réelle, une telle séparation est rarement complète. L'incendie de sa maison ou de son entreprise peut être un événement largement incontrôlable par une personne, mais la probabilité qu'il survienne est quelque peu influencée par son manque de prudence, l'action délibérée ne pouvant d'ailleurs être exclue, même si c'est un cas extrême » [Arrow, 1963, p. 130].

L'expression d'aléa de moralité s'est ensuite étendue hors du champ de l'assurance. Elle tend à se confondre avec l'ensemble des comportements opportunistes et englobe notamment la divulgation de fausses informations. Pour l'analyse économique, l'aléa de moralité désigne les situations où l'une des deux parties peut influencer les bénéfices attendus sans que l'autre partie puisse contrôler cette influence. Ce sont des situations où l'hypothèse d'opportunisme d'une des parties ne peut pas être infirmée ou confirmée par l'autre et pour lesquelles on ne peut pas élaborer des contrats complets, c'est-à-dire des contrats qui précisent tout ce qui se passera pour tous les types d'aléas.

• *Action cachée et information cachée.* — Certains auteurs ont considéré par la suite qu'il était plus pertinent de délaissier les termes d'aléa de moralité et de sélection perverse au profit de ceux d'action cachée et d'information cachée [voir Arrow, 1968, 1985]. On parle d'action cachée quand, dans une transaction, une des parties impliquées peut ou non entreprendre des actions que l'autre partie ne peut ni contrôler précisément ni imposer, sachant que ces actions affectent le bien-être de cette dernière. Par exemple, il est difficile pour un concessionnaire automobile de savoir si les mauvaises ventes d'un agent commercial sont dues à sa paresse ou dépendent d'aléas extérieurs tels qu'un taux de chômage élevé dans la région. Dans les modèles d'action cachée, l'effort s'analyse comme une désutilité pour le commercial et comme un *input* susceptible d'augmenter les

1. La démonstration plus générale et formalisée de cette argumentation est ensuite effectuée par Pauly [1968].

Biens et marchés contingents

Arrow a proposé d'étendre le modèle d'équilibre général au cas d'incertitude à partir de la notion de biens contingents. Un bien contingent est un bien disponible à une certaine date, en un certain lieu, si et seulement si une situation spécifiée à l'avance se réalise. Ceci nécessite d'envisager chaque situation possible — chaque « état de la nature » — et de prévoir le prix de chaque bien dans chaque situation, l'utilité d'un bien étant susceptible de varier fortement selon l'aléa qui se réalise. Sur un marché dit contingent, les agents signent aujourd'hui des contrats contingents où ils s'engagent à livrer ou à recevoir un

bien ou un service seulement si un état spécifié à l'avance se réalise. Un tel marché fonctionne « *ex ante* », avant la connaissance de l'état qui se réalise, et permet aux agents grâce à leurs échanges, de s'assurer contre un risque. Par exemple, une personne peut acheter un contrat qui lui assure le remplacement de son véhicule, à une période spécifiée, si son véhicule personnel tombe en panne.

Si ce procédé est possible pour tous les biens et services, dans toutes les éventualités et à tous les moments du futur, alors il existe un système complet de marchés. Cette fiction évacue l'incertitude.

ventes pour son employeur, l'action cachée étant ici induite par la difficulté de mesurer l'effort.

La terminologie d'information cachée est utilisée pour rendre compte de situations simples de sélection perverse où, *ex ante* (avant de conclure la transaction), une information clef est détenue par une seule des parties. Cette terminologie peut aussi désigner des situations où le problème se pose *ex post*, l'exemple canonique étant le service fourni par les experts. Ceux-ci ont une information privée car ils sont seuls capables d'établir un diagnostic en raison de leur compétence. Dans ce cas, la partie non informée peut observer le résultat de l'action mais ne peut pas vérifier si celle-ci était appropriée. L'expert peut avoir sciemment formulé un diagnostic erroné.

Il convient de noter que règne un certain flou dans la littérature autour de la définition de l'information cachée, certains auteurs n'incluant pas la sélection perverse dans cette catégorie.

2. Le coût du retour à l'ordre

Au milieu des années 1970, il est acquis que les résultats fondamentaux de l'économie obtenus dans un contexte d'information parfaite ne résistent plus dès lors qu'une dose d'incertitude, même très faible, est introduite dans les modèles. Dans le cadre très général de la théorie de l'information, de nombreux travaux ont été développés pour démontrer l'intérêt que les agents ont à réduire l'asymétrie

d'information, que celle-ci engendre des problèmes de sélection perverse ou d'aléa moral [Cahuc, 1998]. Que faire pour éviter la sélection perverse ? Comment résoudre les problèmes posés par l'aléa moral ? Tels sont les problèmes auxquels s'efforcent de répondre les théories du signal, du filtre et de l'agence.

Signal et filtre

Comment éviter que l'incertitude ayant pour origine une information cachée sur une caractéristique cruciale d'un bien ou d'un individu ne conduise à un problème de sélection perverse ? Les modèles de référence attachés à cette question sont ceux de Michael Spence [1973] sur le signal et de Michael Rothschild et de Joseph Stiglitz [1976] sur le filtre. Dans ces modèles, il est montré comment la partie qui détient une information décisive a intérêt à ce qu'elle soit divulguée afin d'accroître son gain. Deux types de procédure sont distingués selon que l'initiative provient de l'individu informé ou de la partie non informée. Dans le premier cas, l'individu révèle son information en émettant un signal à destination de l'autre partie (Spence). Dans le second, la partie en quête d'information met en place des clauses contractuelles ou « filtres » qui conduisent l'autre partie à révéler sa caractéristique (Rothschild et Stiglitz).

- *Peut-on se fier à un signal de qualité ?* — Selon Spence, pour résoudre les problèmes que l'information cachée pose au marché du travail, il faut que les offreurs de bonne qualité envoient un signal. Pour que ce dernier soit efficace, il est nécessaire qu'il soit coûteux, de telle façon que seuls les individus compétents aient intérêt à l'émettre. Ainsi, le diplôme peut être discriminant et constituer un indicateur efficace, s'il permet aux employés potentiels de signaler aux employeurs leur « qualité », c'est-à-dire leur productivité. Dans le modèle de Spence, le coût d'obtention d'un diplôme n'est pas le même pour tous les individus, les plus compétents ayant moins de difficultés à l'obtenir supportent un coût inférieur. Spence montre qu'il existe un niveau de diplôme constituant un signal efficace quand tous les individus compétents ont intérêt à obtenir ce diplôme et aucun individu peu compétent.

La vie quotidienne fournit d'autres illustrations de ce dispositif. Par exemple, la garantie attachée à un produit indique sa qualité. En effet, seuls les offreurs de produits de bonne qualité ont intérêt à se signaler, le coût potentiel de la garantie étant dissuasif pour les offreurs de produits de médiocre qualité. Avec le mécanisme du signal, l'équilibre peut être retrouvé. Celui-ci est alors qualifié d'optimum de second rang, puisque le procédé de révélation de

l'information a un coût qui est supporté par l'offreur de biens ou services de bonne qualité.

• *Le contrat à la carte comme filtre.* — Rothschild et Stiglitz [1976] montrent que, sur le marché de l'assurance, c'est la configuration même des contrats proposés qui conduit les assurés à révéler l'information pertinente pour l'assureur. Ils choisissent le contrat en tenant compte de deux paramètres, leur préférence individuelle et la connaissance de leur caractéristique personnelle. Les individus qui se considèrent comme risqués optent pour un contrat caractérisé par une prime élevée et une franchise faible tandis que les individus à faible risque choisissent les contrats avec une franchise élevée et une prime faible. Rappelons que la prime correspond au prix payé par l'assuré pour avoir droit à une protection et la franchise au montant forfaitaire qui restera à la charge de l'assuré en cas de dommage. Cette méthode permet de filtrer la demande et d'éviter l'effet d'éviction des bons risques qui refusent de payer une prime calculée à partir de la moyenne pondérée par la fréquence relative des bons et des mauvais risques.

Plus généralement, par ce mécanisme, l'offreur définit des paires — contrat et consommateurs d'un certain type — qui lui permettent de maximiser ses objectifs. Les contrats sont fixés de telle façon que chaque groupe d'individus demande le bien ou le service qui lui est destiné. Johanne et Steven Salop [1976] qualifient un tel mécanisme d'autosélection. Rothschild et Stiglitz montrent comment la méthode du filtre permet, éventuellement, de retrouver un équilibre unique. Ce retour à l'équilibre, que l'on qualifie de séparateur en raison de la séparation des demandeurs en, au moins, deux catégories, ne s'effectue pas sans perte de bien-être pour la collectivité. En effet, les individus qui présentent peu de risque ont une situation plus défavorable que celle qui prévaudrait en information parfaite car ils acceptent une franchise élevée. Sinon, le contrat ne serait plus discriminant et attirerait également les agents à haut risque. L'autosélection, grâce à la méthode du filtre, permet de surmonter les effets de l'asymétrie d'information mais a un coût éventuellement prohibitif.

De nombreux travaux ont approfondi et prolongé les résultats de ces deux modèles fondateurs non seulement dans le domaine micro-mais également macroéconomique. Le champ d'application de ces théories est très vaste, couvrant des domaines aussi variés que la finance, l'économie industrielle ou encore le marché du travail (pour un *survey*, voir Riley [2001]).

La théorie de l'agence

Dans le cas d'aléa moral au sens large, c'est-à-dire dans le cas où une des parties peut se comporter d'une façon préjudiciable aux intérêts de l'autre, l'incertitude sur le comportement peut faire échouer la coordination. Comment éviter cet écueil ? C'est le problème qui est au cœur de la théorie de l'agence. Celle-ci a pris pour angle d'attaque la délégation de pouvoir appelée relation d'agence. On dit qu'une relation d'agence apparaît entre deux parties quand l'une d'entre elles, appelée l'agent (le mandataire), agit pour l'autre partie, appelée principal (le mandant), à sa demande ou comme son représentant, dans un domaine décisionnel particulier [Ross, 1973]. La théorie de l'agence privilégie les cas où une telle relation induit l'agent à développer un comportement opportuniste, ce qui est source d'incertitude pour le principal [Jensen et Meckling, 1976]. Les conflits d'intérêt se développent avec le manque d'information du principal. Ainsi, lorsque ce dernier ne peut pas observer l'effort fourni par l'agent ou connaître avec précision dans quel contexte une décision a été prise, l'agent est tenté de faire passer son intérêt personnel avant celui du principal. Si l'effort à fournir pour satisfaire le principal réduit le bien-être de l'agent, il va le minimiser. Le problème consiste pour le principal à mettre en place un système qui incite l'agent à entreprendre la bonne action, c'est-à-dire à agir dans l'intérêt du principal. Deux cas sont à considérer : 1) l'agent est indifférent au risque de variation de son revenu, dans ce cas, le principal lui fait supporter l'intégralité de l'incertitude attachée à la mesure de son effort, en lui proposant, par exemple, un contrat où sa rémunération dépend uniquement du résultat ; 2) si l'agent n'est pas enclin à la prise de risque, le contrat est établi de façon à maximiser le revenu du principal. Il comporte alors deux contraintes : selon la première, la rémunération versée devra être telle que l'agent sera incité à fournir un effort élevé, c'est ce que l'on appelle la contrainte d'incitation. La seconde dite contrainte de participation se traduit par le versement d'une rémunération minimale à l'agent, faute de quoi celui-ci aurait intérêt à s'adresser à d'autres mandataires ou à se retirer du marché du travail. La solution qui consiste à verser un revenu fixe (minimum de participation), le reste étant fonction du résultat (incitation) est largement pratiquée, notamment dans la vente. Dans ce deuxième cas, le revenu du principal est toujours inférieur à celui qu'il aurait obtenu dans un monde sans aléa, du fait de la clause de participation. Éviter le comportement opportuniste de l'agent a un coût pour le principal, c'est le coût du retour à l'ordre. (Pour une présentation formelle, voir Kreps [1990].)

• *Le modélisateur et l'incertitude.* — En dépit de résultats proches de la réalité, la modélisation pose des hypothèses fortes sur le traitement de l'incertitude. Le modélisateur distingue deux sources d'incertitude qui sont analysées de façon à permettre le calcul et l'optimisation : le comportement de l'agent, d'une part, les événements indépendants de son action, d'autre part. Lorsqu'elle a pour origine le comportement de l'agent (variable endogène), l'incertitude est décomposée en un petit nombre d'états possibles, au minimum deux. Par exemple, si l'agent est un commercial, le niveau d'effort fourni est élevé *versus* faible. Les événements indépendants de l'agent, qui sont également susceptibles de faire varier le résultat, sont tenus pour aléatoires et décomposables eux aussi en quelques états du monde possibles ; le principal connaît la distribution de probabilités de la réalisation de ces états. Par exemple, le modélisateur pose que la demande des ménages dans la zone de vente du commercial peut être en croissance forte avec une probabilité de 0,4 *versus* en croissance faible avec une probabilité de 0,6.

Le raisonnement a été conduit jusqu'ici à partir de la problématique de l'action cachée, la logique d'ensemble est la même lorsque l'aléa moral a pour origine l'information cachée. L'opération de décomposition des états du monde, essentielle au modèle, est souvent passée sous silence, ce qui est dommage car elle conditionne largement les résultats. Ces modèles supposent que l'agent est capable d'avoir une vision claire et exhaustive de l'ensemble des aléas susceptibles d'affecter la décision. Nous reviendrons sur les enjeux de cette décomposition des états du monde dans le chapitre III, avec le modèle de Savage.

3. L'incertitude stratégique

La théorie des jeux est un outil d'analyse qui reprend à sa façon les questions et les solutions classiques de la théorie de l'information exposées ci-dessus, où dans une visée positive, de nombreux modèles montrent que les pratiques usuelles peuvent être efficaces. Elle va plus loin en posant le problème de la convergence des anticipations. Nous verrons qu'elle aboutit parfois à des résultats originaux où l'incertitude perd son statut de trouble-fête.

La théorie des jeux traite des problèmes de coordination des agents et donne à l'« incertitude stratégique » une place centrale. L'incertitude stratégique apparaît dans les situations où notre bien-être dépend de la stratégie choisie par les autres et où nous ne sommes jamais sûrs de leur choix. Les ingrédients les plus généraux de la modélisation sont les suivants : 1) le nombre d'agents

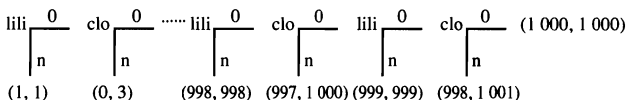
économiques (joueurs) qui interviennent est fixé *a priori* ; 2) il est admis que chaque joueur cherche le gain le plus élevé en anticipant les actions de l'autre joueur ; 3) toutes les actions potentielles des individus sont envisagées ainsi que les gains obtenus dans chaque cas de figure. L'ensemble de ces données constitue ce que l'on appelle les règles du jeu [Cahuc, 1998 ; Guerrien, 2002 ; Rasmusen, 1989].

Dans les jeux les plus simples, il n'y a aucun aléa, l'information est complète, la seule inconnue réside alors dans la stratégie choisie par les autres participants. Dans les jeux à information incomplète, impulsés par Harsanyi dans les années 1960, figure un autre type d'incertitude, avec la présence d'aléas. Ceux-ci peuvent concerner les caractéristiques des joueurs — ces derniers pouvant être des firmes aux coûts de production élevés ou faibles, par exemple — et/ou les règles du jeu, les gains par exemple. L'introduction de ces variables aléatoires s'effectue en ajoutant aux joueurs initiaux un agent fictif appelé « Nature » dont le rôle est de doter les joueurs de caractéristiques. Le plus souvent, dans les jeux à information incomplète, l'information est également asymétrique. Les joueurs connaissent leurs caractéristiques, les autres estiment ces caractéristiques sur la base de leurs connaissances préalables et de leurs croyances. Que les jeux soient complets ou non, la structure de l'information est toujours définie précisément, même pour les éléments aléatoires, ce qui conditionne le calcul qui est au fondement de la théorie des jeux. Le joueur peut ainsi, dès le départ, savoir quelle décision il prendra dans tous les cas de figure. Dit autrement, l'incertitude est ancrée dans le présent.

- *Quand l'incertitude devient souhaitable.* — Une spécificité de la théorie des jeux est de mettre en évidence à quels paradoxes la coordination peut aboutir quand chaque joueur recherche exclusivement son intérêt individuel, les équilibres atteints alors étant très éloignés de l'optimum comme dans l'exemple du mille-pattes de Rosenthal [1981] (voir encadré). Cet exemple revêt un intérêt particulier, ici. Il montre que, selon les principes de l'utilitarisme, les options successives prises par les joueurs devraient les conduire à la pire formule. Douter du comportement d'autrui et supposer que celui-ci puisse être coopératif ou digne de confiance deviennent la solution. L'incertitude stratégique a perdu son statut de trouble-fête et conditionne, dans ce cas particulier, le succès.

Le jeu du mille pattes

Supposons que Lili et Clo jouent à un jeu où elles gagnent une somme dont le montant progresse au cours du jeu.



Le jeu fonctionne de la façon suivante. À chaque étape, l'une des joueuses doit décider si le jeu continue (oui, « o ») ou s'il s'arrête (non, « n »). Elle sait que, si elle décide de continuer, le gain qu'elle obtiendra à l'étape suivante sera inférieur d'une unité à son gain actuel. En revanche, le gain de l'autre joueuse augmentera de 2 unités. Prenons le cas de Lili. Au départ, si elle décide de continuer le jeu, elle offre une opportunité de gain de 2 à Clo qui gagne 3 au lieu de 1, et risque une perte de 1 (gain nul contre un gain de 1). Cette perte n'est que provisoire pour Lili si Clo décide également de continuer. Dans ce cas, Lili gagnera 2 au prochain coup.

Si le jeu ne se termine jamais, les possibilités de gains sont infinies. Mais si le nombre de coups est fini, alors le résultat devient désastreux. Plaçons-nous, en effet, à la dernière étape et effectuons un raisonnement à rebours. Clo a intérêt à dire non car cela lui permet d'obtenir 1001 au lieu de 1000. À l'étape précédente, les gains sont identiques mais Lili a intérêt à dire non, ce qui lui rapporte 999, car elle sait qu'à l'étape suivante son gain serait seulement de 998, Clo ayant intérêt à dire non, etc. Ce raisonnement se répète de proche en proche jusqu'à la première étape où Lili dit « non » de

manière à obtenir un gain de 1 contre un gain nul à la seconde étape. Bien entendu, on sait que, au quotidien, personne ne raisonne à rebours sur un si grand nombre d'itérations. Cet exemple entend simplement illustrer l'inefficacité des décisions non coopératives. En effet, Lili et Clo auraient intérêt à coopérer et à jouer « oui » durant tout le jeu, ce qui les conduirait à un gain très élevé. Ce jeu montre aussi qu'introduire un peu d'incertitude peut être la solution du paradoxe. Il est possible d'introduire une incertitude sur le « moment » où s'arrête le jeu. L'adjonction au modèle d'une *probabilité de fin de jeu* ayant pour effet de modifier le calcul des gains, le jeu ne s'arrête pas nécessairement à la première étape. L'incertitude peut également porter sur le comportement des joueuses. Si, par exemple, Lili pense que Clo est susceptible de répondre « oui » plutôt que « non » à la dernière étape parce qu'elle est irrationnelle et ne maîtrise pas complètement la récurrence à rebours ou parce qu'elle est altruiste, alors le piège du mille-pattes est déjoué et chacune gagne 1 000 [Dupuy, 1989]. Une autre solution pour résoudre ce paradoxe consiste à faire appel à la réputation. Si Clo a acquis la réputation d'être coopérative, alors Lili n'hésitera pas à jouer.

Conclusion

L'incertitude joue un grand rôle dans le renouvellement de la microéconomie. Elle a été introduite comme un grain de sable qui fait perdre à la mécanique des prix et de la concurrence leur capacité à atteindre le meilleur des équilibres. Elle se confond alors avec la

méconnaissance d'une variable clef (le meilleur prix, la qualité particulière) par une des deux parties prenantes de l'échange, dans un monde déterministe où les caractéristiques des biens et des personnes sont prédéfinies. La carence d'information engendre des désordres, elle peut conduire à la perte du prix unique et, même, rendre l'échange impossible. Cependant, ces désordres peuvent être aplanis. Acquérir de l'information ou contraindre les comportements par un contrat sont les principales solutions proposées. Cet exercice est coûteux, voire impossible. Le retour à l'ordre peut s'avérer inaccessible parce que la mise en place de procédures contractuelles, incitatives ou contraignantes est trop coûteuse. La réputation offre, alors, une voie de recherche féconde qui n'a pas été explorée ici [Klein et Leffler, 1981 ; Shapiro, 1982]. Une autre voie de recherche met en exergue des cas où le raisonnement utilitariste des joueurs conduit à une impasse. C'est alors paradoxalement l'introduction du doute sur le comportement d'autrui et le fait de supposer qu'il ne poursuive pas son intérêt mais soit coopératif et digne de confiance qui deviennent la solution. Ces éclairages différents évoquent l'ambivalence de l'incertitude.

Dans tous ces modèles, la carence d'information n'est pas synonyme d'ignorance, bien au contraire. La partie mal informée connaît souvent la distribution de probabilité de la variable méconnue. Autrement dit, les individus sont dans une situation de risque. Cette hypothèse, plus ou moins réaliste selon les situations, permet de maintenir le calcul et la maximisation de l'utilité. Dans certains modèles, des capacités cognitives exceptionnelles sont attribuées au décideur qui lui permettent de suppléer parfaitement à son manque d'information. Elles concernent tant le calcul que la spécification des valeurs des paramètres. *In fine*, la certitude chassée par la grande porte revient par la fenêtre.

Le temps subit également un traitement particulier. Dans le modèle initial du *search*, la décision relative au nombre de commerces à visiter est prise avant toute recherche effective. De même dans la théorie des jeux, les décisions sont prises initialement, sans révision ultérieure. Arrow et Debreu avaient donné le ton, en proposant d'étendre le modèle d'équilibre général au cas d'incertitude et en supposant que les agents prennent toutes les décisions économiques concernant leur comportement futur, une fois pour toutes, au début de leur vie en quelque sorte. Cette vision est cohérente avec le rôle central donné au contrat pour expliquer la coordination sur les marchés des biens (Arrow-Debreu), du travail (théorie de l'agence). Les contrats sont établis en t_0 et tout se passe comme si les agents passaient le reste de leur vie à les honorer. Cette conception du

temps a été critiquée par les économistes néoclassiques eux-mêmes, c'est ainsi que le modèle du *search* a été rapidement dépassé par une approche séquentielle de la recherche d'information.

3. General Information (provide a brief description of the project, including the location, the purpose of the project, and the expected outcomes. Also, provide a brief description of the project's history and the organization's role in the project.)

DEUXIÈME PARTIE

LA QUESTION DE LA RATIONALITÉ EN INCERTITUDE : MAXIMISER OU AGIR RAISONNABLEMENT ?

Comment choisit-on en incertitude ? La règle de maximisation de l'utilité peut-elle être préservée ?

Les premiers modèles qui conservent cette règle et introduisent les probabilités constituent, encore aujourd'hui, la référence incontournable. Une règle simple de décision en incertitude est obtenue au prix d'une modélisation exigeant de la part des individus une cohérence logique contraignante (chapitre III).

Aussi n'est-il pas surprenant qu'à la même époque des chercheurs contestent que la décision puisse être parfaite. Les êtres humains ne sont pas de froids logiciens. Ne seraient-ils pas sensibles au contexte, certaines données du problème ne seraient-elles pas privilégiées ? Autrement dit, les règles de cohérence logique des modèles de décision en incertitude ne sont-elles pas transgressées ? Les tests expérimentaux apportent des réponses qui suscitent une seconde génération de modèles où le théoricien cherche à intégrer cette part de subjectivité qui nous conduit à interpréter des données objectives. N'est-ce pas la rationalité, elle-même, qu'il faudrait redéfinir pour rendre intelligibles les comportements en incertitude, en particulier, lorsque la coordination avec autrui est en jeu (chapitre IV) ?

L'idée que la rationalité peut être limitée incite d'autres courants à développer des conceptions alternatives du comportement en incertitude où les institutions et les habitudes jouent un rôle central (chapitre V).

III / Quand la décision est parfaite

En univers certain, il est admis que le choix des individus résulte de la maximisation de leur satisfaction *via* leur fonction d'utilité. Est-il possible de transposer le modèle de décision en information parfaite à un contexte d'incertitude ? La théorie de l'utilité espérée offre une réponse positive à cette question, en s'appuyant sur les probabilités, cet outil commode pour modéliser en incertitude. Les probabilités associées aux variables aléatoires peuvent être ou non connues du décideur. Si elles sont données, la décision est dite risquée et lui correspond le modèle de l'utilité espérée (UE) de John von Neumann et Oskar Morgenstern [1944]. Que faire quand les probabilités ne sont pas données au décideur, ce qui dans le langage de la théorie de la décision, correspond à l'incertain ? Alors la décision s'appuie sur des probabilités dites subjectives et le modèle correspondant est celui de l'espérance d'utilité subjective (SEU) de Léonard Savage [1954]. Ces constructions théoriques sophistiquées permettent de parvenir à la règle de décision simple : maximiser son espérance d'utilité. Munis de cette règle, les agents peuvent et doivent décider de façon parfaite. Suivant les principes de l'utilitarisme, la règle de décision respecte la diversité des préférences. Avec l'incertitude, le champ des préférences s'élargit et une notion supplémentaire apparaît, celle de l'attitude individuelle vis-à-vis du risque.

1. Le modèle de l'utilité espérée

En univers certain, le choix du consommateur rationnel résulte de la maximisation de l'utilité calculée à partir de paniers de biens. En univers risqué et incertain, la disponibilité des biens n'étant pas garantie, ceux-ci sont dits contingents (encadré, chapitre II, pour une

définition) ; les conséquences des décisions des individus ne sont plus uniques, elles dépendent d'événements aléatoires ou « états de la nature », l'état de la nature qui se réalisera étant inconnu au moment de la prise de décision, elles sont donc hypothétiques. Ceci complique considérablement la règle de décision. Pourtant, le modélisateur a réussi à conserver la maximisation, chère aux utilitaristes, en raisonnant sur des loteries plutôt que sur des biens et en recourant à l'utilité espérée. Comprendre ce modèle suppose donc de se familiariser avec ces notions spécifiques. Nous verrons ensuite quelles contraintes lui sont associées.

Les ingrédients et la logique de la décision

Commençons par présenter l'utilité espérée (UE). Cette notion est apparue en réponse au paradoxe de Saint-Pétersbourg soulevé par un jeu de hasard. Elle a donné son nom au modèle de décision en univers incertain, que nous désignerons par son sigle, le modèle de l'UE.

- *L'utilité espérée et le paradoxe de Saint-Pétersbourg.* — Pour comprendre comment s'effectue la décision en incertitude, partons de Daniel Bernoulli (1700-1782). Celui-ci a proposé de calculer l'espérance de l'utilité des gains et non pas l'espérance des gains, c'est-à-dire le résultat que l'on obtient directement par le calcul mathématique. Il s'est appuyé sur un jeu de hasard présenté ici sous une forme légèrement remaniée.

Un joueur lance une pièce de monnaie aussi longtemps que face apparaît et s'arrête dès qu'il obtient pile. Si pile sort dès le premier coup, le joueur reçoit deux ducats. Si pile ne sort qu'au deuxième coup, il reçoit 2^2 ducats et, si pile ne sort qu'au n ème coup, il reçoit 2^n ducats.

Le problème consiste à évaluer le montant du droit d'entrée maximal qu'un individu rationnel est prêt à payer pour participer à ce jeu. Quel est ce montant ? Un mathématicien proposera de prendre pour référence l'espérance de gain, c'est-à-dire ce qu'il peut gagner en moyenne. Dans cette hypothèse, le joueur serait disposé à payer un droit d'entrée d'un montant égal à :

$$E(G) = \frac{1}{2} 2 + \frac{1}{2^2} 2^2 + \dots + \frac{1}{2^n} 2^n + \dots = 1 + 1 + 1 + \dots + 1 + \dots$$

Le nombre de coups de la partie n'étant pas limité, l'espérance de gains est infinie. Or Bernoulli a constaté que la plupart des joueurs n'acceptaient de verser qu'un montant très faible pour participer à ce jeu bien que son espérance de gain soit infinie. Tel est le

paradoxe de Saint-Pétersbourg. Pour le résoudre, Bernoulli a proposé de prendre en compte l'utilité associée aux différents gains plutôt que de considérer leur espérance mathématique. En effet, le joueur compare le coût de participation au jeu avec ce que le jeu est susceptible de lui rapporter. Or son utilité marginale décline et l'utilité d'un gain hypothétique de 2^n par rapport à 2^{n-1} lorsque n est élevé est dérisoire, en raison certes de l'effet de richesse mais surtout de l'aversion au risque (*infra*). Cette solution qui consiste à passer par la satisfaction procurée par un gain est un des fondements du modèle d'espérance d'utilité. Utiliser la solution préconisée par Bernoulli transforme le jeu de Saint-Pétersbourg comme suit :

$$E(U(G)) = \frac{1}{2} U(2) + \frac{1}{2^2} U(2^2) + \dots + \frac{1}{2^n} U(2^n) + \dots$$

Afin d'exprimer la décroissance de l'utilité marginale, Bernoulli a suggéré de prendre la fonction Log pour représenter l'utilité, ce qui permet d'obtenir une suite dont la somme est bornée et égale à Log 4. Le caractère paradoxal des réponses des joueurs disparaît si l'on admet que ceux-ci procèdent à un traitement non linéaire des résultats, ce que permet la notion d'utilité.

Les fondements axiomatiques du modèle de l'« utilité espérée » ont été établis beaucoup plus tard par Ramsey [1931] et surtout par John von Neumann et Oskar Morgenstern (que l'on notera désormais VNM).

• *Les loteries.* — Qu'appelle-t-on « loterie » dans le modèle de l'UE ? Partons du sens courant du jeu et de la loterie la plus élémentaire, où l'on a une probabilité, notée prob, de gagner le gros lot et une probabilité, notée 1-prob, qui est le complément à l'unité, de ne rien gagner. Dans le modèle de l'UE, la loterie la plus simple comprendra deux lots auxquels sont associées deux probabilités dont la somme doit être égale à un. Sous cette forme très générique, la notion de lot déborde très largement celle du panier de biens utilisée dans la théorie du consommateur. Le lot désigne aussi une conséquence possible d'une décision, autrement dit, un gain contingent. Par exemple, si je décide de vendre des parapluies, le montant des ventes dépendra de la météorologie. Supposons qu'il ne dépende que de cette variable, et que celle-ci se décompose en deux « états » simples, du type, il fait beau *versus* il pleut, alors je peux rendre compte de l'incertitude attachée à la vente de parapluies sous forme d'une loterie et la noter $L[(\text{prob}_1, c_1) ; (\text{prob}_2, c_2)]$, où c_1 est nommé lot ou, indifféremment, conséquence.

Cet exemple peut être prolongé pour rendre compte du choix en incertitude. Un individu a le choix entre vendre des parasols ou des

parapluies, au printemps prochain. Sa décision est fonction des conséquences de son choix, c'est-à-dire de ses recettes futures, or celles-ci dépendent de l'environnement. Ce n'est qu'à la fin de la saison qu'il saura précisément si le temps a été beau ou pluvieux. Néanmoins, il peut attribuer, dès le 21 mars, date de sa décision, des probabilités concernant le climat futur, celles-ci étant connues, et représenter ces éléments sous forme de loterie. Soit L_1 , la loterie correspondant à la vente des parasols. On a $L_1 [(prob_1, C_{11}) ; (1-prob_1, C_{12})]$, avec $prob_1$ la probabilité qu'il fasse beau et $(1-prob_1)$ la probabilité qu'il pleuve. C_{11} désigne le gain obtenu par la vente de parasols s'il fait beau, C_{12} désigne le gain obtenu par la vente de parasols si le printemps est pluvieux. La loterie L_2 correspondant à la vente de parapluies se construit de la même façon.

Ici, le choix de l'individu est modélisé à partir des loteries (choix entre L_1 et L_2) parce qu'il est discret. Dans le cas où le choix est continu, le raisonnement est plus abstrait.

• *La règle de décision.* — L'utilité d'une loterie de type $L_1 [(prob_1, C_{11}) ; (1-prob_1, C_{12})]$ dépend de l'ensemble des conséquences, puisque celle qui se réalisera effectivement est ignorée. L'utilité de chaque loterie s'exprime, donc, comme l'espérance mathématique de l'utilité des conséquences, autrement dit comme la somme des utilités des conséquences pondérées par leur probabilité de réalisation.

L'espérance d'utilité d'une loterie L_j (notée $EU(L_j)$) est la suivante :

$$EU(L_j) = \sum_{i=1}^2 prob_i \cdot U(C_{ji})$$

où $prob_i$ désigne la probabilité associée à la réalisation de l'événement i et $U(C_{ji})$ désigne l'utilité accordée à la conséquence de la loterie L_j lorsque l'événement qui se réalise est i .

C'est la fonction d'utilité espérée ou fonction de VNM sous sa forme la plus sommaire. Dans ce modèle, les probabilités sont de type fréquentiste [von Neumann et Morgenstern, 1947, p. 19], elles sont donc considérées comme objectives. Ce ne sont pas des croyances.

L'agent étant confronté à un choix, il faut étendre ce raisonnement à l'ensemble des loteries. Est sélectionnée celle qui lui procure l'espérance d'utilité la plus grande. Ceci permet de maintenir le principe de la maximisation.

Illustrons ce modèle par l'exemple d'un automobiliste qui a le choix entre deux trajets A et B. Pour chacun de ces trajets, trois états

de la circulation sont envisageables : fluide, chargée ou embouteillée. Les conséquences s'identifient à la durée du trajet. Le tableau suivant donne l'utilité associée à chaque trajet conditionnellement à chaque état de la circulation. Bien entendu, le caractère avantageux d'un trajet est d'autant plus grand que la durée est courte.

	<i>Fluide</i>	<i>Chargée</i>	<i>Embouteillée</i>
Utilité des conséquences de A	100	70	15
Utilité des conséquences de B	90	60	5

Le modèle requiert aussi de disposer des probabilités associées aux états de la circulation. On connaît d'après des études statistiques les probabilités d'encombrement des itinéraires A et B.

	<i>Fluide</i>	<i>Chargée</i>	<i>Embouteillée</i>
Probabilité associée au trajet A	1/3	1/3	1/3
Probabilité associée au trajet B	3/4	0	1/4

L'utilité espérée de chaque trajet est :

$$U(A) = (1/3.100) + (1/3.70) + (1/3.15) = 61,67.$$

$$U(B) = (3/4.90) + (1/4.5) = 68,75.$$

Le conducteur rationnel choisit le trajet B.

Volontairement, c'est un exemple aux enjeux très limités qui a servi, ici, à illustrer le comportement des personnes ayant à choisir en incertitude, selon le modèle de VNM. Le modèle peut être transposé à des choix économiques simples, par exemple, une décision d'investissement entre deux types d'activités. Les tableaux seraient construits selon le même principe : on trouverait en colonne, les différents états de la conjoncture (bon, médiocre, mauvais), en ligne, l'utilité des revenus de chaque investissement et les probabilités de réalisation des différents états.

Les contraintes du modèle

Sous cette règle de décision en apparence simple se cache une construction théorique sophistiquée et contraignante. Sa sophistication tient aux caractéristiques de l'utilité et son caractère contraignant à la lourdeur des axiomes à respecter pour assurer la cohérence logique du choix.

Le modèle de l'UE présente une différence fondamentale par rapport à celui de la décision en univers certain. Pour des raisons

techniques, l'utilité attachée à chaque lot ou conséquence doit être cardinale, l'utilité simplement ordinale ne convenant pas (voir encadré).

Les théoriciens ont imaginé une procédure spécifique fondée sur le questionnement pour mesurer l'utilité attachée à chaque conséquence et, ainsi, construire la fonction d'utilité du décideur. Cette procédure est complexe puisqu'il s'agit de mesurer la satisfaction des individus vis-à-vis de gains hypothétiques. En effet, si l'on n'aime pas le risque, alors on minore l'utilité du gain qui n'est pas certain. La fonction d'utilité doit donc intégrer l'attitude vis-à-vis du risque.

• *L'axiomatique de von Neumann et Morgenstern.* — Les axiomes sont des contraintes logiques nécessaires à la construction du modèle. Différents axiomes (préordre total, continuité, indépendance) ont été introduits par von Neumann et Morgenstern pour passer du choix ponctuel à une fonction de choix continu et permettre que la fonction d'utilité obtenue retranscrive le même ordre que les préférences individuelles, soit :

Si $\forall L_i$ et L_j , on a $L_i \geq L_j$, alors $E[U(L_i)] \geq E[U(L_j)]$.

La notation $\geq$ est utilisée, ici, pour classer des loteries, elle se distingue de la notation usuelle $\geq$ qui s'applique à des nombres. Pour cette dernière, certaines propriétés telles la symétrie et la transitivité sont déjà établies. Ces propriétés n'étant pas acquises pour la relation $\geq$ ont été posées comme axiomes par von Neumann et Morgenstern.

Parmi les axiomes qu'ils introduisent, celui d'indépendance a un statut particulier. Il pose que l'ordre des préférences entre deux loteries L_1 et L_2 restera le même si chacune de ces deux loteries initiales est combinée dans une même proportion avec une troisième L_3 . Contrairement aux attentes de von Neumann et Morgenstern, cette opération n'est pas triviale. Combiner une loterie avec une autre est une opération très différente de l'addition de deux variables. En particulier, si les montants en jeu sont élevés, elle peut conduire les personnes à réviser l'ordre de leur préférence, comme l'ont montré de nombreux travaux expérimentaux (*infra*).

Rappelons que, outre une fonction technique, l'axiomatique a une dimension normative. Elle privilégie une conception du « bon » raisonnement et définit la façon dont les individus doivent formuler leur choix. Lorsque des tests expérimentaux sont menés et que les observations montrent qu'un axiome n'est pas respecté, le modélisateur conclut soit à l'irrationalité des agents, soit à la nécessité de faire évoluer le modèle (chapitre IV).

Cardinalité et ordinalité

La fonction d'utilité en univers certain est ordinale, les valeurs attribuées aux différents biens n'ont pas d'importance, ce qui importe est leur hiérarchie.

Par exemple, pour trois biens : brioche tartine et croissant, on peut avoir indifféremment les utilités suivantes 1 ou 2 :

Biens	Utilité 1	Utilité 2
Brioche	50	2
Tartine	70	4
Croissant	100	5

Le classement entre les biens est respecté dans les deux cas, le croissant étant préféré à la tartine et à la brioche, les valeurs attribuées ne reflètent qu'un simple classement.

En univers risqué, on ne peut plus utiliser une utilité ordinale obtenue par simple classement, le modèle de l'UE impose d'avoir recours à une utilité cardinale, ce qui suppose de mesurer son intensité.

Supposons, en effet, que l'on ait le choix entre : 1) avoir un croissant avec une probabilité 1/2 et une brioche avec une probabilité 1/2 ; 2) avoir une tartine de façon certaine. Il faut effectuer un choix entre les deux loteries suivantes :

$L_1 [(1/2, 100), (1/2, 50)]$ et $L_2 (1, 70)$ si l'on utilise la colonne 1 des utilités ci-dessus. Si l'on avait utilisé la colonne 2 des utilités, on aurait eu le choix entre $L'_1 [(1/2; 5), (1/2; 2)]$ et $L'_2 (1; 4)$. Dans le premier cas, on a :

$EU(L_1) = (1/2.100) + (1/2.50) = 75$ et $EU(L_2) = 70$. On préfère l'option « croissant brioche » plutôt que celle de la tartine. En utilisant les loteries L_1 et L_2 on obtiendrait le choix inverse. Pour conserver l'ordre des préférences entre les loteries, il est nécessaire d'avoir recours à une fonction d'utilité cardinale, la colonne 2 d'utilité étant déduite de l'utilité de la première par une transformation affine (de la forme $ax + b$) et non par une fonction croissante quelconque.

2. L'extension du modèle aux situations d'incertitude

Le modèle d'utilité espérée de VNM propose une théorie de la décision qui est qualifiée de risquée parce les probabilités utilisées dans le modèle sont objectives. Les travaux théoriques pionniers de Léonard Savage ainsi que ceux de Francis Anscombe et Robert Aumann ont étendu le modèle de décision de l'UE à l'ensemble des situations incertaines, en introduisant la notion de probabilité subjective. La règle de décision demeure la maximisation de l'utilité espérée comme dans le modèle de l'UE, cependant la conception du modèle change ainsi que l'axiomatique. L'utilisation des probabilités subjectives pose un nouveau problème au théoricien, celui de leur connaissance. En raison de sa simplicité, la méthode, dite de révélation des probabilités, d'Anscombe et Aumann sera présentée avant le modèle de Savage [Savage, 1954 ; Anscombe et Aumann, 1963].

Dans la seconde moitié du xx^e siècle, l'essor des probabilités subjectives a ouvert de nouveaux horizons à la modélisation des comportements en incertitude. Pour Savage et les néoclassiques, le comportement des individus, dans n'importe quel contexte d'incertitude, est désormais susceptible d'être modélisé. En faisant de toute probabilité une croyance, Savage rend caduque la distinction entre risque et incertitude et le modèle SEU de Savage englobe celui de l'utilité espérée de von Neumann et Morgenstern. L'argumentation de Savage peut s'illustrer ainsi : même dans le cas d'un jeu de dé, où assigner une loi de probabilité apparaît simple, cette apparence est conditionnelle. Le joueur présume que le dé n'est pas pipé, une condition dont il ne peut pas être certain. Autrement dit, tout choix de distribution de probabilité est subjectif, donc toute décision prise sur la base d'une loi de probabilité l'est également.

La généralité du modèle de Savage tient à la double nature des probabilités subjectives qui désignent à la fois 1) les croyances que forment les êtres humains, chacun de leur côté, lorsqu'ils ont une connaissance imparfaite des probabilités objectives ou lorsqu'ils veulent intégrer des informations locales, pour prévoir un phénomène de nature aléatoire ; 2) les probabilités qui sont de purs jugements sur l'occurrence d'un phénomène qui n'est pas récurrent, telle que la probabilité que l'homme aille sur la planète Mars au xxi^e siècle. Savage privilégie la première acception ; la seconde oriente vers Knight et Keynes.

Les probabilités subjectives ou croyances à la Savage sont-elles aussi subjectives et personnelles que les goûts ? Selon Savage, ces probabilités sont personnelles et une réponse affirmative devrait s'imposer mais, aujourd'hui, « ... pour de nombreux microéconomistes, il est donné pour un dogme (philosophie ?) que deux individus ayant accès à la même information formeront nécessairement les mêmes probabilités subjectives. Toute différence dans les évaluations des probabilités subjectives ne peut être que le résultat de différences ayant pour origine l'information » [Kreps, 1990, p. 111]. Cette hypothèse de modélisation, appelée indifféremment hypothèse des probabilités *a priori* communes ou doctrine Harsanyi, permet d'agréger aisément les comportements. Les probabilités subjectives s'analysent comme des copies, en information imparfaite, des probabilités objectives. Pointons le paradoxe suivant : les économistes admettent parfaitement la diversité des goûts et des attitudes vis-à-vis du risque (*infra*) mais, avec la doctrine d'Harsanyi, ils refusent la diversité, lorsqu'il s'agit de prévoir.

Pour étendre le modèle UE, Anscombe et Aumann se focalisent sur les probabilités formulées individuellement qu'ils analysent comme des chances. Les chances peuvent avoir des valeurs connues et objectives comme dans un jeu de hasard ou inconnues et subjectives comme dans une course de chevaux. Ces dernières peuvent être révélées, ce qui peut s'illustrer par cet exemple. S'il m'est indifférent de parier 1 euro sur le cheval 18 ou de parier 1 euro sur le rouge dans une loterie où la probabilité que sorte le rouge est de $2/3$, alors, à mon avis, la chance que le cheval 18 soit le gagnant est de $2/3$. Suivant la même logique, tout individu peut traiter toute probabilité comme une chance et la quantifier en exprimant sa préférence entre des options fictives. Ce procédé simple, sinon simpliste, explique son succès. Pendant plusieurs décennies, ce postulat s'est imposé comme une évidence, contestée aujourd'hui.

La logique du modèle de Savage

En univers risqué, pour simuler la décision, le modélisateur recourt aux loteries où les probabilités sont données, seules les utilités sont à découvrir. En incertitude, pour Savage la décision ne s'analyse plus en termes de loteries. Considérons dans un premier temps pour la commodité de l'exposé qu'elle se structure à partir d'un nombre fini d'actes et d'éventualités possibles. La décision intègre aussi les conséquences qui forment un espace. Tous ces éléments sont représentés dans la matrice ci jointe :

		États de la nature					
		S_1	S_2	...	S_j	...	S_n
Actes	A_1	C_{11}	C_{12}	...	C_{1j}	...	C_{1n}
	A_2	C_{21}	C_{22}	...	C_{2j}	...	C_{2n}
	...	...	...	...	...	...	...
	A_k	C_{k1}	C_{k2}	...	C_{kj}	...	C_{kn}
	...	...	...	...	...	...	...
	A_m	C_{m1}	C_{m2}	...	C_{mj}	...	C_{mn}

— en ligne, figure l'ensemble des actions possibles (A_1 à A_m dans la matrice), par exemple vendre des parasols ou des parapluies ;

— en colonne, figurent les états de la nature (S_1 à S_n), soit la description exhaustive des états possibles de l'environnement qui sont mutuellement exclusifs (il ne peut faire pas beau et pluvieux en même temps). Ces états de la nature sont exogènes, c'est-à-dire hors du contrôle de l'individu qui prend sa décision ;

— dans les cases, figurent les conséquences d'un acte (C_{11} à C_{mn}). Elles sont conditionnelles, puisqu'elles dépendent de la réalisation d'un état particulier de la nature.

Dans le modèle de Savage, les probabilités d'occurrence des états de la nature et l'utilité de chaque conséquence sont à découvrir et à présenter sous une forme quantifiée. Le mode de résolution de ce double problème de quantification comprend deux étapes [Runde, 2000].

La première étape est celle de la quantification de l'utilité espérée. On suppose que la variation d'utilité est fixée à 1 si l'événement S_1 se réalise et à 0 dans les autres cas. Cet artifice permet de résoudre aisément le problème de quantification de l'utilité.

La seconde étape est celle de la révélation des probabilités concernant les états de la nature. Ce second problème peut être traité par une méthode de questionnement fictif. Le procédé proposé, ici, consiste à révéler les probabilités par des préférences entre des options fictives, l'une étant certaine et l'autre incertaine.

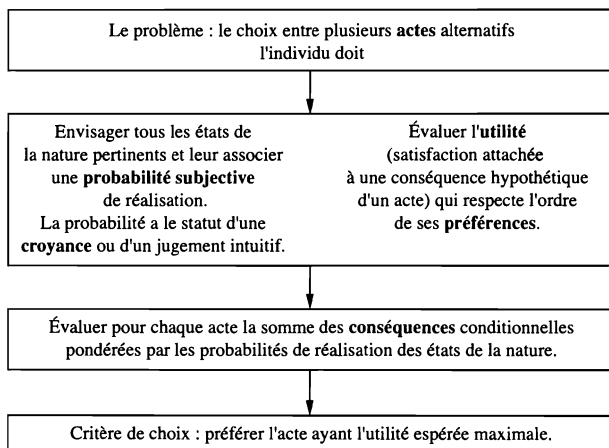
On pose que $\text{prob}_{(S_1)}$ désigne la probabilité attachée à la réalisation de S_1 . Cette probabilité qualifiée par Savage de personnelle est appelée communément subjective et, par extension abusive de langage, croyance. L'ensemble des états du monde est divisé en S_1 et son complément $\bar{S}_1$. Révéler la probabilité subjective suppose que l'on trouve la valeur numérique q (comprise entre 0 et 1) pour laquelle l'individu est indifférent entre obtenir de façon certaine q et obtenir 1 de façon incertaine : $q = \text{prob}_{(S_1)}(1) + [1 - \text{prob}_{(S_1)}](0)$.

D'où $q = \text{prob}_{(S_1)}$. Supposons que le parieur déclare $q = 0,8$ alors il révèle que, pour lui, la probabilité $\text{prob}_{(S_1)}$ que l'état de la nature favorable se réalise est de 0,8. Avec cette présentation, les probabilités sur les événements aléatoires sont représentées numériquement par le truchement de q .

Ces étapes ont été présentées en se plaçant du point de vue du modélisateur qui voudrait reconstituer la procédure de décision, découvrir les paramètres utilisés par le décideur (utilité et probabilités subjectives) et, ainsi, s'assurer que l'acte choisi est bien celui qui maximise l'utilité.

Le décideur, de son côté, n'a pas à respecter ces étapes, puisqu'il a en tête des probabilités subjectives. Le schéma suivant retrace les opérations cognitives du décideur. Sa décision est rationnelle si l'option choisie maximise sa fonction d'utilité espérée. Comme celle-ci est calculée à partir de probabilités subjectives, sa forme abrégée courante est SEU (*subjective expected utility*).

SCHEMA 1. — DÉCOMPOSITION DU COMPORTEMENT
DU DÉCIDEUR RATIONNEL



Le modèle en débat

Dans la modélisation de Savage, avant de décider, il faut structurer le problème, donc le simplifier, afin que l'évaluation des conséquences de chaque acte pour chaque éventualité n'admette qu'une seule valeur. La délimitation de la situation met de côté des effets trop complexes, par exemple, des conséquences peu vraisemblables. Savage n'a pas rendu compte des opérations de construction du problème, celles-ci ne faisant pas seulement appel à la logique. C'est d'ailleurs ce qu'il précise dans son ouvrage où il prend pour exemple une situation banale. Un homme voit, dans la cuisine familiale, un œuf posé à côté d'un bol contenant cinq œufs déjà cassés et préparés pour cuire une omelette. Trois « actes » s'offrent à lui : casser le sixième œuf directement dans le bol, le casser dans une tasse avant de le mélanger aux autres œufs ou ne rien faire. Deux états de la nature sont envisageables : l'œuf est sain ou pourri, ce dernier aléa était fréquent dans les années 1950. L'utilité de chaque acte — l'avantage moins l'effort à fournir — dépend de l'état de la nature susceptible ou non de se réaliser. L'acte qu'il choisira révélera ses probabilités personnelles d'occurrence des différents états de la nature.

Supposons que, par son intervention malheureuse, il ait gâché l'omelette. Il arguerait que la probabilité que l'œuf fût pourri était, selon lui, trop faible pour fonder en raison une mesure de précaution.

Avec cet exemple, les mondes de Savage se présentent comme de tout petits mondes, c'est-à-dire des situations bien délimitées dans le temps et l'espace. Ces mondes sont simples pour au moins deux raisons : 1) seuls les éléments pertinents du problème (états de la nature et conséquences) sont pris en compte, les autres étant négligés, comme le fait que, dans l'exemple choisi, gâcher l'omelette pourrait conduire à ce que toute la famille aille au restaurant, etc. ; 2) les événements aléatoires, qualifiés de façon très suggestive d'états de la nature ou d'états du monde, sont totalement indépendants des actes possibles. L'état de la nature qui se réalise effectivement est, dans cet exemple, indépendant du sujet. Ceci est nécessaire à la viabilité du modèle de décision.

Pour la facilité de l'exposé, le modèle a été présenté comme si l'espace des conséquences était fini. Or c'est inexact. Pour des raisons mathématiques, Savage admet que les actes sont infinis. En revanche, l'espace des états de la nature est considéré par Savage comme étant fini. Ce point a fait l'objet de critiques. Pour Kreps, la viabilité du modèle suppose qu'ils soient infinis [Kreps, 1990, p. 103]. De plus, poser que le décideur dispose d'une liste complète des états exclut que se produisent des événements imprévus [Schackle, 1990]. Une autre question se pose : comment ces états de la nature sont-ils construits ? Savage a noté que les états du monde résultent d'une construction personnelle du décideur. Il est regrettable que les économistes l'oublient souvent. Préférant gommer cette dimension subjective du modèle de décision, ils considèrent que ces états s'imposent d'un point de vue extérieur, comme le font la pluie ou le beau temps.

Le débat porte surtout sur l'axiomatique de Savage qui sous-tend le modèle. Certains axiomes sont analogues à ceux du modèle UE, d'autres sont spécifiques [Munier, 1984 ; Granger, 1997]. Les critiques se sont focalisées sur l'un d'entre eux. Il s'agit de l'axiome dit d'indépendance des préférences, selon lequel les préférences sur les conséquences sont indépendantes des croyances sur les états de la nature. Autrement dit, la conséquence hypothétique d'un acte est évaluée sans prendre en compte la probabilité de réalisation de l'état du monde qui lui correspond. Selon ce cadre d'analyse, la décision rationnelle repose sur deux piliers indépendants : les préférences et les croyances, ou plus précisément la préférence pour la conséquence la plus satisfaisante et la croyance que tel état du monde a telle probabilité de se réaliser. Cette construction qui renvoie au découpage cartésien sujet/environnement permet de maintenir la

vision de la décision, chère aux utilitaristes, où celle-ci se conforme à des buts qui ont été, au préalable, clairement définis et demeurent stables. Ce découpage sujet/environnement est critiqué au XX^e siècle par les courants philosophiques non cartésiens.

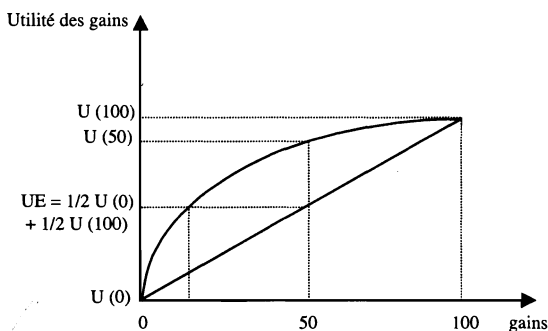
3. La peur et le goût du risque

Tous les individus n'ont pas la même attitude face au risque, certains sont peu enclins à la prise de risque, d'autres au contraire aiment le risque. Comment la fonction d'utilité espérée intègre-t-elle ces différences d'attitude vis-à-vis du risque ? Peut-on les mesurer ?

Attitude vis-à-vis du risque et utilité espérée

Un agent a de l'aversion pour le risque s'il préfère à toute loterie le gain de son espérance mathématique reçu avec certitude.

Par exemple, Jeanne a le choix entre 1) participer à un jeu de pile ou face où le gain est de 100 euros si la pièce tombe sur pile, de 0 sinon et 2) gagner avec certitude 50 euros. Ayant de l'aversion pour le risque, elle préférera gagner 50 euros avec certitude plutôt que de participer au jeu. Cette aversion peut être représentée graphiquement.



Le graphique montre que l'utilité d'un gain que Jeanne est certaine d'obtenir, $U(50)$, est supérieure à l'utilité attendue du jeu, c'est-à-dire à l'utilité espérée, avec $UE = [1/2 U(0) + 1/2 U(100)]$.

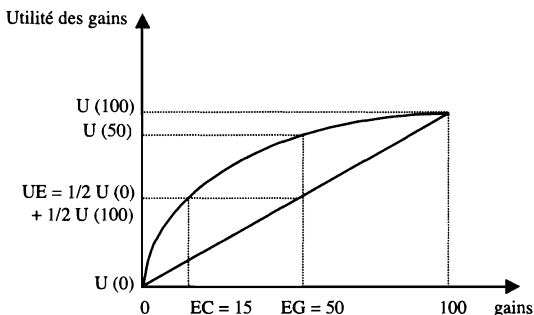
Soit $U(50) > [1/2 U(0) + 1/2 U(100)]$.

Plus généralement, le graphique illustre le fait que l'utilité de gains certains excède l'utilité de gains incertains. La corde à la courbe passant par les points ayant pour abscisse 0 et 100 est tracée pour que soit respecté le traitement linéaire des probabilités sur l'axe des ordonnées. La concavité de la fonction d'utilité traduit le fait que le joueur « minore » la satisfaction d'un gain incertain. Plus la fonction d'utilité est concave, donc plus elle s'éloigne de la corde, plus l'agent éprouve de l'aversion au risque.

Quand la fonction d'utilité et la corde sont confondues, alors l'agent est qualifié d'« indifférent au risque ». Si Jeanne est indifférente au risque, alors il lui est indifférent de participer au jeu de pile ou face ou de recevoir avec certitude 50 euros. Quand la fonction d'utilité est convexe, c'est-à-dire sous la corde, alors l'agent a le goût du risque. Dans ce cas, Jeanne attribue au gain de 50 euros qu'elle est certaine d'obtenir une utilité inférieure à celle qu'elle attend du jeu, soit $U(50) < [1/2 U(0) + 1/2 U(100)]$ (pour une présentation plus complète [Kast, 2002]).

- *Équivalent certain et prime de risque.* — Le modèle de l'UE permet de mesurer l'attitude vis-à-vis du « risque au sens faible », en se référant à la notion d'équivalent certain. L'équivalent certain d'une loterie (EC) correspond au montant qui, s'il était donné de façon certaine à un individu, lui procurerait la même utilité que celle obtenue par sa participation à cette loterie. Par exemple, il est indifférent pour Elsa de jouer à la loterie $L_1[(1/2;100),(1/2;0)]$ et de recevoir de façon sûre 15 euros. Son équivalent certain est le gain qui, obtenu de façon certaine, lui procure la même satisfaction que la loterie, il est donc de 15 euros. Ceci peut se représenter graphiquement.

Pour les agents ayant de l'aversion au risque, l'équivalent certain est toujours inférieur à l'espérance de gain, notée EG. Quand ils ont du goût pour le risque, c'est l'inverse. On peut comparer la plus ou moins grande aversion pour le risque des individus en comparant la valeur de leur équivalent certain pour une loterie donnée. L'équivalent certain permet de calculer la prime de risque. La prime de risque R associée à une loterie L se définit comme l'écart entre l'espérance mathématique de cette loterie et son équivalent certain. Soit $R(L) = EG(L) - EC(L)$ avec $EG(L)$ l'espérance du gain de la loterie L et $EC(L)$ l'équivalent certain de cette loterie. La prime de risque correspond au montant $R(L)$ qu'un agent est prêt à payer pour échanger la loterie L contre son équivalent certain. C'est ce type de calcul qui explique qu'un individu ayant de l'aversion au risque choisit de s'assurer et accepte de payer une prime plus ou moins élevée selon son aversion au risque. Si la prime demandée par



l'assurance est inférieure ou égale au montant que l'agent est prêt à payer, calculé comme ci-dessus par référence à son équivalent certain, alors il s'assure. Dans le cas contraire, il ne s'assure pas.

D'autres mesures

Le modèle d'utilité espérée définit un type d'aversion au risque que l'on nomme l'« aversion faible ». Au quotidien, les individus étant plus sensibles à la variation du risque qu'à son niveau, il est utile d'introduire une autre mesure appelée l'« aversion forte ».

Alors que l'aversion faible a été définie en comparant une perspective risquée à une perspective certaine, l'aversion forte ou aversion pour un accroissement de risque se définit en comparant deux perspectives inégalement risquées. Il existe différents moyens de mesurer l'accroissement de risque. L'utilisation de la variance pour comparer des risques est très utilisée en finance, bien qu'elle ne soit pas sans défaut [Gayant, 2001].

L'attitude des individus est également influencée par la quantité d'information disponible sur les perspectives incertaines.

Prenons l'exemple d'un jeu de hasard où deux urnes, A et B, contiennent des boules blanches et noires. Si la boule tirée est blanche, le joueur a gagné. Le joueur sait que l'urne A contient autant de blanches que de noires mais ignore comment se répartissent les boules dans l'urne B. Si le joueur préfère que le tirage se fasse à partir de l'urne A, alors il préfère une situation moins incertaine à une situation plus incertaine. Cette préférence est qualifiée d'aversion pour l'incertitude ou pour l'ambiguïté.

Conclusion

Le modèle de l'utilité espérée permet de retrouver une règle de décision, simple, où l'incertitude n'apparaît pas comme un trouble-fête, empêchant la maximisation. Ce modèle connaît un succès considérable. En effet, le cadre de la décision en univers risqué et incertain est, à certains égards, peu exigeant : les individus font face à des éventails de choix clairs et n'ont pas besoin de beaucoup d'informations ; ils décident isolément ; ils ont un étalon unique, l'utilité.

De nombreux courants théoriques novateurs liés à la théorie de l'information se fondent sur ce modèle pour expliquer les comportements. C'est particulièrement le cas des théories de l'agence et des jeux. De plus, en raison de sa grande généralité, il s'applique à des champs très variés qui vont de la décision d'investissement à l'économie dite du crime, en passant par le mariage, où les enjeux moraux sont très différents. Sa portée a conduit certains économistes à déclarer dépassée la distinction de Knight entre risque et incertitude, du fait de la possibilité ouverte par Savage d'utiliser des probabilités subjectives, même en contexte d'incertitude [Arrow, 1951 ; Hirshleifer et Riley, 1979 ; Lucas, 1981].

Très loué, le modèle est aussi très critiqué. Les critiques émanent à la fois de logiciens préoccupés par les axiomes, d'économistes ayant une autre conception de l'incertitude, de psychologues conduisant des tests expérimentaux, etc. Le chapitre suivant prolonge la réflexion en présentant les critiques qui émanent à la fois de logiciens et de psychologues.

IV / La mise à l'épreuve du choix rationnel

Les modèles de l'utilité espérée, qu'ils utilisent les probabilités objectives (risque) ou subjectives (incertitude) ont subi l'épreuve de la réalité. Une réalité particulière puisqu'il s'agit surtout d'expérimentations sur des personnes faisant face à des choix fictifs où les gains ne sont pas certains. Ce chapitre montre comment les modèles pionniers du choix rationnel en incertitude ont été mis en cause par des tests et comment, en retour, les modèles ont été modifiés afin d'être plus proches des comportements observés lors des tests. La large gamme des nouveaux modèles traduit l'effort des économistes pour que le modèle mathématique de la décision rationnelle en incertitude intègre des résultats obtenus par les psychologues. Enfin, nous verrons que des théoriciens de la décision ont ouvert une autre voie. Dans ses travaux pionniers, Herbert Simon, à partir notamment de l'étude des comportements dans les organisations, a rejeté l'idée que le processus de décision en incertitude suivrait la règle de maximisation de la satisfaction.

1. Les premiers paradoxes expérimentaux

Selon les modèles initiaux, rappelons-le, le décideur fait face à un choix bien défini entre des loteries (risque) ou des actes (incertitude) ayant des conséquences conditionnées par différents états aléatoires. Le choix en incertitude, à la différence d'un choix risqué, ne peut pas se fonder sur des probabilités objectives. L'approche classique considère que les probabilités sont données en univers risqué (modèle de l'UE de von Neumann et Morgenstern). En incertitude, les probabilités ne sont pas données. Ce sont des croyances appelées probabilités subjectives (modèle SEU de Savage) (chapitre III).

L'approche classique a fait l'objet de nombreuses critiques. L'interaction entre les paradoxes expérimentaux et les progrès mathématiques est manifeste dans les modèles de décision en incertitude.

Les scientifiques parlent de paradoxe expérimental ou empirique lorsqu'un modèle s'avère incompatible avec les faits, c'est-à-dire que des résultats théoriquement impossibles sont observés empiriquement et que cette divergence n'est imputable ni à la structure logique du modèle dont la cohérence interne est assurée, ni à la qualité des tests [Walliser, 1995]. Confrontés à un paradoxe expérimental, les économistes ont deux réactions possibles, selon qu'ils adoptent une approche normative ou descriptive. L'approche normative conduit à considérer que ce sont les faits qui sont en cause, autrement dit les agents ne sont pas rationnels, ils commettent des erreurs de choix et de jugement qui les conduisent à des choix incohérents ; dans l'approche descriptive, c'est le modèle qui est mis en cause, car il perd son pouvoir prédictif et il convient de réviser ses axiomes. C'est cette seconde réaction qui domine aujourd'hui pour le modèle de l'utilité espérée. Les paradoxes expérimentaux, dont les plus célèbres sont ceux d'Allais et d'Ellsberg, ont conduit à un foisonnement de travaux qui proposent de nouveaux modèles, se substituant ou englobant le modèle de base.

Le paradoxe d'Allais

Le paradoxe développé par Maurice Allais [1953], prix Nobel d'économie en 1988, remet en cause l'axiome dit d'indépendance du modèle de l'UE de VNM. Le terme d'indépendance résume l'idée selon laquelle les préférences des individus entre deux loteries ne doivent pas être modifiées lorsqu'elles sont combinées dans les mêmes proportions avec une troisième loterie.

L'ordre des préférences est-il invariant lorsque les variables subissent une modification en apparence neutre ? Sous cette formulation technique, se cache un enjeu majeur du modèle de la décision : les préférences concernant des gains donnés sont-elles stables, quelles que soient les probabilités attachées à ces gains ? Les agents révisent-ils l'utilité qu'ils attachent à un gain quand les probabilités de recevoir ce gain se modifient, ou lorsque le gain cesse d'être certain pour devenir hypothétique ? Allais a proposé deux couples de loteries judicieusement construites pour montrer le non-respect de cet axiome et, ce faisant, réfuter un des piliers du modèle de l'espérance d'utilité.

• *Le test.* — Le test consiste à présenter à différents individus deux options, chacune se décomposant en deux loteries, et à leur demander de révéler leurs préférences :

Option A

Loterie L_1 : recevoir 100 millions de façon certaine.

Loterie L_2 : 10 chances sur 100 de gagner 500 millions, 89 chances sur 100 de gagner 100 millions, 1 chance sur 100 de ne rien gagner.

Option B

Loterie L'_1 : 11 chances sur 100 de gagner 100 millions, 89 chances sur 100 de ne rien gagner.

Loterie L'_2 : 10 chances sur 100 de gagner 500 millions, 90 chances sur 100 de ne rien gagner.

Quatre couples de réponses sont envisageables : (L_1, L'_1) , (L_2, L'_2) , (L_1, L'_2) et (L_2, L'_1) , néanmoins seuls les couples (L_1, L'_1) et (L_2, L'_2) sont compatibles avec l'axiomatique de VNM. Calculons les espérances d'utilité associées aux différents gains :

— loterie L_1 : $U(100)$

— loterie L_2 : $0,1 U(500) + 0,89 U(100) + 0,01 U(0)$

— loterie L'_1 : $0,11 U(100) + 0,89 U(0)$

— loterie L'_2 : $0,1 U(500) + 0,9 U(0)$

Si un individu déclare préférer la loterie 1 à la loterie 2, alors $U(100) > 0,1 U(500) + 0,89 U(100) + 0,01 U(0)$

$\Leftrightarrow 0,11 U(100) > 0,1 U(500) + 0,01 U(0)$

$\Leftrightarrow 0,11 U(100) + 0,89 U(0) > 0,1 U(500) + 0,9 U(0)$.

Selon ces transformations, les individus qui ont préféré L_1 à L_2 doivent préférer L'_1 à L'_2 . Cette logique du choix a été mise en évidence à partir du calcul des espérances des utilités, il est possible de l'exposer directement à partir des loteries. En effet, le passage de L_1 à L'_1 et de L_2 à L'_2 résulte d'une combinaison linéaire des loteries initiales comme suit : $L'_1 = L_1 - 0,89 (1,100) + 0,89 (1,0)$. De même, pour L'_2 . Avec ce second mode d'exposition, on retrouve l'axiome d'indépendance du modèle de l'utilité espérée de VNM.

Les résultats de l'enquête menée par Allais montrent que de nombreux individus choisissent le couple (L_1, L'_2) . Ce choix n'est pas conforme à l'axiome d'indépendance. Les individus ne sont pas pour autant irrationnels. Dans le premier choix, ils préfèrent recevoir de façon sûre un montant très élevé (100 millions) plutôt que de choisir une loterie qui inclut un résultat nul, bien que celui-ci soit associé à une très faible probabilité. Dans le second choix, les probabilités sont sensiblement identiques (0,1 contre 0,11) et les individus sont, alors, frappés par la différence de gains affichée (500 millions contre 100 millions).

Le paradoxe d'Ellsberg

La méthode est la même que précédemment, où des expériences répétées sur des jeux adéquats montrent que les parieurs ne respectent pas un axiome. Tandis que le paradoxe d'Allais remet en cause la fonction d'utilité de VNM, celui de Daniel Ellsberg [1961] met en cause le modèle SEU et la représentation de l'information sous forme de probabilité. En effet, cette dernière fait l'objet de déformation, dans certains contextes. Ellsberg montre comment la nature de l'information disponible pour formuler les probabilités en incertitude joue sur la confiance qu'ont les individus dans leur jugement et donc sur leurs préférences. On retrouve ici le rôle de la confiance dans le jugement souligné dans le premier chapitre (Knight, Keynes).

• *Le test.* — Supposons qu'une urne contienne 90 boules, dont 30 rouges (R) et 60 noires (N) ou jaunes (J) dans une proportion inconnue. Dans un jeu à deux tirages, on demande à différents individus de parier sur la couleur de la boule tirée au cours de chaque tirage. Les propositions de paris et les gains correspondants sont présentés dans les tableaux ci-dessous.

Pour le premier tirage :

	30 boules	60 boules	
	Rouges	Noires	Jaunes
Pari I	100	0	0
Pari II	0	100	0

Les joueurs doivent dire s'ils préfèrent le pari I où ils gagnent si la boule tirée est rouge au pari II où la boule gagnante est la boule noire.

Pour le second tirage :

	30 boules	60 boules	
	Rouges	Noires	Jaunes
Pari III	100	0	100
Pari IV	0	100	100

Dans ce second tirage, on demande aux joueurs s'ils préfèrent le pari III où les boules gagnantes sont « rouges ou jaunes » au pari IV où elles sont « noires ou jaunes ».

Quatre couples de paris sont envisageables (I, III) (II, III) (I, IV) et (II, IV) mais seuls les couples (I, III) et (II, IV) sont compatibles avec l'axiomatique de Savage. Pourtant le couple de paris le plus fréquemment choisi est I, IV. Le premier choix révèle que $\text{prob}(N) < \text{prob}(R)$, où $\text{prob}(N)$ désigne la probabilité — subjective — de tirer une boule noire. Le second choix révèle que $\text{prob}(N \text{ ou } J) > \text{prob}(R \text{ ou } J)$, ce qui implique que $\text{prob}(N) > \text{prob}(R)$, d'où le paradoxe.

Les réponses des joueurs s'expliquent ainsi. Lors du premier choix, les joueurs connaissent la probabilité de tirer une boule rouge, celle de tirer une boule noire peut prendre toutes les valeurs comprises entre 0 et 2/3. Du fait du manque d'information pour formuler des probabilités, les joueurs préfèrent le premier pari. Dans le second cas, c'est la probabilité de tirer une boule noire ou jaune qui est connue, donc c'est la dernière solution qui est préférée.

Avec ce test, Ellsberg conteste l'idée que les agents manquant d'information statistique pourraient toujours formuler correctement des probabilités subjectives. Son paradoxe remet en cause l'optimisme des économistes pour lesquels la distinction de Knight entre risque et incertitude n'était plus utile, le théoricien pouvant toujours poser que le raisonnement à partir de probabilités subjectives est aussi rigoureux que si les probabilités sont objectives.

• *Aversion à l'ambiguïté.* — Ce paradoxe ouvre la voie à une nouvelle notion, introduite par Ellsberg, l'ambiguïté. L'ambiguïté, une notion qui fait débat, est définie, ici, de manière stricte comme l'incertitude sur les probabilités créée par un manque d'informations pertinentes alors que ces informations pourraient être connues [Frisch et Baron, 1988]. Dans une définition plus large, l'aversion à l'ambiguïté désigne la préférence des agents pour la chose probable et précise relativement à la chose vraisemblable (probable mais imprécise). Se pose, en filigrane, la question de la représentation que se font les agents d'un environnement complexe. L'ambiguïté ouvre sur différents travaux qui vont des plus radicaux, avec le refus d'utiliser des probabilités lorsque leur valeur est ignorée, aux plus classiques, avec des modélisations parfois très proches de celle de l'utilité espérée. Ainsi certains modèles intègrent un coefficient d'aversion à l'ambiguïté, de sorte que le modèle de l'utilité espérée devient un cas particulier où ce coefficient est nul ([Viviani, 1994], pour un *survey*).

2. La décision optimale se personnifie

Les paradoxes ont déclenché un foisonnement de travaux expérimentaux de psychologie cognitive destinés à mieux comprendre le comportement du décideur en incertitude. Daniel Kahneman, prix Nobel d'économie 2002, et Amos Tversky qui sont les pionniers dans cette voie [1979], ont établi une théorie du choix risqué dite théorie « des perspectives » (*prospect theory*). Élaborée pour répondre à des paradoxes expérimentaux, elle constitue une alternative au modèle de l'utilité espérée.

Les comportements saisis par l'expérience

Les expériences conduites par Kahneman et Tversky font apparaître des régularités de comportement. Ils en déduisent plusieurs effets qui complexifient le comportement de choix en incertitude.

- *L'effet de certitude.* — L'effet de certitude vient du fait que la différence entre un gain certain et un gain probable nous paraît bien plus importante qu'une différence comparable dans la gamme des probabilités intermédiaires. Cet effet contribue à biaiser le choix. Kahneman et Tversky l'ont mis en évidence, en reprenant une expérience, proposée par Allais, qui montre comment l'axiomatique du modèle de VNM est transgressée.

Question 1 : Préférez-vous gagner A (0,8 ; 4 000) ou B (1 ; 3 000) ? Une grande majorité préfère B.

Question 2 : Préférez-vous gagner C (0,2 ; 4 000) ou D (0,25 ; 3 000) ? La majorité s'inverse et préfère C.

Ce comportement est paradoxal, l'ordre des préférences a été inversé alors que, pour passer de A à C et de B à D, les probabilités ont simplement été multipliées par 0,25. Cette opération conduit la majorité des parieurs à modifier l'ordre de leurs préférences, contrairement à l'axiomatique et à la logique mathématique. Ici, encore, les parieurs ne sont pas irrationnels et leurs choix s'expliquent par la psychologie de la décision en incertitude. Quand la probabilité de gagner est devenue faible, notre sensibilité à la valeur du gain s'est accrue tandis que notre sensibilité aux probabilités a diminué. En résumé, l'effet de certitude rend compte du fait que la sensibilité aux probabilités n'est pas linéaire.

- *L'effet miroir.* — L'effet miroir vient du fait que les individus se comportent différemment face à un gain ou à une perte. Confrontés à des jeux de hasard incluant des pertes, nous préférons

une perte hypothétique à une perte garantie, sachant que l'espérance de perte hypothétique est supérieure à la perte certaine. Autrement dit, nous sommes prêts à prendre le risque de participer au jeu, même si nous avons conscience que celui-ci nous fait encourir une perte en moyenne plus élevée que la perte certaine. Par exemple, si nous transposons l'exemple précédent nous serons nombreux à préférer l'option A' (0,8 ; - 4 000) plutôt que B' (1 ; - 3 000), bien que l'espérance de perte dans le cas A' soit de 3 200. Une façon d'interpréter mathématiquement ce résultat est de dire que les pertes certaines sont surpondérées. L'effet miroir exprime que l'aversion pour l'incertitude, qui caractérise le comportement de la majorité des agents face à des gains, se renverse en « préférence pour l'incertitude » face à des pertes. En fait, la tendance à minimiser la probabilité associée à une perte relève plus de la répugnance à perdre que du goût pour le risque. Cet effet précise la forme de la fonction qui exprime notre attitude en univers incertain. Elle est concave pour les gains potentiels mais convexe pour les pertes potentielles. Elle se présente, donc, sous une forme en S.

- *L'effet de construction.* — Les individus en situation de choix ont tendance à reformuler les données du problème pour le simplifier. L'effet de construction provient de biais qui se produisent lors de la mise en forme du problème et lors du traitement des données. Ces biais sont liés au fait que nous n'avons pas une vue globale des conséquences de nos choix. Par exemple, nous nous focalisons sur une séquence particulière quand nous sommes confrontés à un choix séquentiel. De même, nous raisonnons en variation, d'où l'importance du référent, le choix du référent étant susceptible d'influencer notre vision du problème. Ces opérations conduisent à effectuer des choix qui ne respectent pas nécessairement ceux auxquels aboutirait le statisticien.

Le modèle alternatif de Kahneman et Tversky

Kahneman et Tversky ont élaboré leur théorie à partir de l'observation des comportements lors de jeux de hasard. Néanmoins, ils ne parlent pas de loterie. En effet, leur objectif est de saisir les transformations que les agents opèrent sur les données, c'est-à-dire d'intégrer comment les loteries initiales sont recodées et simplifiées pour donner des « perspectives ». Leur théorie, dite des « perspectives », repose sur une décomposition du processus de décision en deux : 1) la phase du formatage où sont perçues et reformulées les données objectives du problème (*editing*) ; 2) la phase de l'évaluation au cours de laquelle la satisfaction retirée des gains et des pertes

Quand le choix est influencé par le contexte

Kahneman et Tversky ont demandé à un grand nombre de médecins de répondre au problème suivant :

Imaginez que votre pays se prépare à combattre l'arrivée d'une maladie rare qui pourrait entraîner la mort de 600 personnes. Nous vous proposons deux programmes de lutte contre cette maladie ; les conséquences de chacun de ces programmes ont été estimées scientifiquement avec exactitude. Si vous optez pour le programme A, 200 personnes seront sauvées, si vous adoptez le programme B, vous avez la probabilité 1/3 de sauver 600 personnes et une probabilité 2/3 de ne sauver personne. Quel programme préférez-vous ? La majorité des médecins choisit A, refusant le programme « risqué ». Le même problème formulé avec un autre référent a été posé à un autre groupe de médecins. Si vous adoptez le programme C, 400 personnes

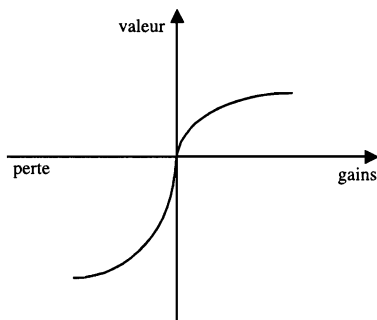
mourront, si vous adoptez le programme D, vous avez une probabilité de 1/3 que personne ne meure et une probabilité de 2/3 de voir 600 personnes mourir. Dans ce cas, la majorité choisit D, donc le risque. Pourtant, il s'agit de deux versions du même problème. La seule différence est la suivante : dans la première version, la mort de 600 personnes est la référence et les résultats du programme choisi sont évalués en termes de gains (nombre de vies sauvées) ; dans la deuxième version, la situation avant la maladie constitue la référence et les programmes sont évalués en termes de perte (nombre de vies perdues). Dans cet exemple, la modification de la présentation renverse les préférences, puisqu'elle se conjugue avec l'effet miroir (forme en S de la fonction) et avec celui de certitude (surpondération des résultats certains).

est quantifiée et les probabilités sont pondérées. Pour évaluer chaque perspective, la valeur de chaque conséquence $v(x_i)$ est multipliée par le poids de la décision $\phi(\pi_i)$, soit $\sum_i \phi(\pi_i) \cdot v(x_i)$. On choisit, ensuite, la perspective qui a la valeur maximale. Les propriétés des fonctions ϕ et V se déduisent des effets étudiés ci-dessus.

• *La fonction de valeur.* — La fonction de valeur $V(x)$ assigne à chaque conséquence un nombre $v(x_i)$ qui reflète sa valeur subjective. Les conséquences sont définies relativement à une origine qui peut être 0 comme dans une échelle classique et $V(x)$ mesure la valeur des écarts à cette référence. Sa représentation graphique montre une forme en S et une pente plus forte pour les pertes que pour les gains, ce qui traduit le fait que le dépit éprouvé par une perte d'un montant x est plus intense que la satisfaction d'un gain d'un même montant.

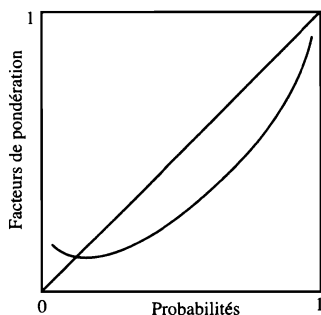
• *La fonction de transformation des probabilités.* — La fonction de transformation (ou pondération) des probabilités notée ϕ est une originalité de ce modèle. « Les poids de décision ne sont pas des probabilités : ils n'obéissent pas aux axiomes des probabilités et ne

FONCTION DE VALEUR



doivent pas s'interpréter comme des mesures de degré de croyance » [Kahneman et Tversky, 1979, p. 280]. Ils rendent compte de l'existence d'une interaction entre les préférences et les probabilités ; ils essaient de capter autant l'impact d'un aléa sur l'attrait d'une conséquence qu'inversement l'influence d'une conséquence sur la probabilité. Les probabilités peuvent être surpondérées ou sous-pondérées, ce qui est différent de la surestimation/sous-estimation de probabilités mal connues ; d'ailleurs, pour montrer la différence entre ces deux concepts, les probabilités sont données dans la plupart des expériences, comme l'illustre l'exemple de la maladie rare présenté en encadré.

FONCTION DE TRANSFORMATION DES PROBABILITÉS



La représentation graphique permet de visualiser les différentes propriétés de la fonction. Certaines d'entre elles concernent toutes les valeurs des probabilités, d'autres sont spécifiques aux valeurs extrêmes. Pour traduire l'effet de certitude inspiré du paradoxe d'Allais, Kahneman et Tversky ont introduit la propriété dite de non-proportionnalité, suivant laquelle plus les probabilités sont faibles, moins leurs variations sont correctement perçues par les individus : ainsi le ratio 0,001/0,002 sera perçu comme étant supérieur à 1/2 et, donc, plus grand que, par exemple, 0,4/0,8.

La seconde propriété, dite de sous-certitude, admet que la somme des poids des décisions (les probabilités transformées) associés à un événement et à son complément est strictement inférieure à l'unité. D'autres propriétés traduisent la tendance à surpondérer les probabilités lorsqu'elles prennent des petites valeurs, comme le saisit le graphique. Enfin, à proximité des valeurs extrêmes, les comportements deviennent imprévisibles et la fonction cesse d'être définie. En effet, les expériences montrent que nous avons tendance soit à assimiler les valeurs très proches de 0 ou de 1 à des certitudes, soit, au contraire nous refusons de les faire tendre vers la certitude, comme l'exprime, par exemple, la formule : « Il n'y a pas de risque 0. »

• *Des précurseurs.* — Différents éléments de la théorie développée par Kahneman et Tversky étaient déjà présents dans des travaux antérieurs. Harry Markowitz a été le premier à définir l'utilité comme une fonction ayant une forme en S, à appréhender les gains hypothétiques de façon relative et à mettre en évidence la différence de comportement des agents devant les gains (peur du risque) et les pertes (prise de risque) [Markowitz, 1952]. De même, l'idée d'utiliser des poids plutôt que des probabilités dans une fonction d'utilité était déjà connue.

Les travaux de Kahneman et Tversky qui associent ces deux éléments ont joué un rôle considérable dans la théorie du choix en incertitude. Ces deux auteurs sont devenus eux-mêmes les précurseurs de deux courants : l'un regroupe de nombreuses modélisations du choix en incertitude mieux aptes à traduire les comportements réels que ne le fait l'utilité espérée ; l'autre remet en cause le rôle central de l'optimisation dans la décision. Ces courants sont présentés successivement dans les sections 3 et 4 de ce chapitre.

3. Une nouvelle lignée de modèles

Bernoulli avait cherché à résoudre le paradoxe de Saint-Pétersbourg par une réflexion sur la fonction d'utilité. Les réflexions sur l'utilité, menées au xx^e siècle, ont conduit, en particulier, à remettre en cause le postulat de transitivité des préférences individuelles, comme dans la théorie du regret initiée par David Bell [1982]. Ces modèles font écho aux critiques formulées par les sociologues et les économistes institutionnalistes concernant l'hypothèse de stabilité des préférences et leur détermination hors contexte.

Cependant, les aspects novateurs des modèles concernent davantage les probabilités que les préférences. La réflexion sur les « probabilités », qui est le pendant de celle de Bernoulli avec sa réflexion sur l'utilité, avait d'ailleurs été envisagée au même siècle par le comte de Buffon, mais cette piste n'avait pas donné de suite. Dès les années 1970, des économistes ont cherché à résoudre les paradoxes expérimentaux en se concentrant sur la fonction de probabilités, comme Kahneman et Tversky. En dépit du caractère séduisant de la modélisation proposée par ces derniers, celle-ci a été rapidement abandonnée. Il est très vite apparu qu'appliquer leur modèle, sans introduire de nouvelles contraintes de modélisation, conduisait à des résultats incohérents [Fishburn, 1978]. Les auteurs l'ont d'ailleurs reconnu eux-mêmes. Il fallait donc opérer autrement pour intégrer un traitement non linéaire des probabilités.

Cette section se concentre sur les travaux où les « probabilités » sont devenues non additives. Nous montrerons, tout d'abord, comment l'axiomatique a évolué afin de résoudre le paradoxe d'Allais. Ensuite, nous présenterons les modèles de décision où les probabilités sont transformées et non additives. Nous reprendrons le clivage usuel entre les modèles du risque et ceux de l'incertitude.

L'économiste, le mathématicien et le psychologue

Rappelons que, dans les jeux expérimentaux, l'axiome le plus souvent transgressé est celui d'indépendance, selon lequel l'ordre de préférence des individus entre deux loteries n'a pas de raison de changer lorsqu'elles sont combinées linéairement avec une troisième.

Pour approfondir les causes de la transgression de cet axiome, les théoriciens de la décision ont fait appel à la notion mathématique de comonotonie, utilisée dans la couverture contre les risques. Deux variables comonotones sont des variables qui, dans chaque état de la nature, offrent des résultats qui sont, soient tous les deux meilleurs, soit tous les deux moins bons [Gayant, 2001]. La combinaison

de deux variables comonotones augmente le risque, en revanche la combinaison de deux variables non comonotones permet de le réduire. C'est ce que font les gestionnaires de portefeuille en proposant à leurs clients qui ont de l'aversion pour le risque, de mélanger des titres financiers procycliques et des titres contracycliques. Comme leur nom l'indique, les premiers varient dans le même sens que l'activité économique, les seconds dans le sens inverse. Par exemple, les résultats des firmes de l'automobile et ceux des firmes pétrolières varient en sens inverse, que l'économie soit en croissance ou en récession.

Revenons aux paradoxes expérimentaux et, en particulier, au paradoxe d'Allais qui illustre la remise en cause de l'axiome d'indépendance (*supra*, section 1). Le renversement de l'ordre des préférences entre deux loteries s'explique par la modification du risque encouru, après que les loteries initiales aient été combinées avec une nouvelle loterie. Choisir le couple L_1 et L'_2 , comme le font la plupart des gens, ne peut plus être vu comme incohérent, puisque la combinaison des loteries initiales avec une troisième loterie qui donne L'_1 et L'_2 ne respecte pas la comonotonie. Si l'on restreint le postulat de l'invariance de l'ordre des préférences entre deux loteries aux seules combinaisons qui respectent la comonotonie, alors le paradoxe disparaît, puisque le décideur ne viole plus l'axiome défini par le logicien.

En bref, introduire la comonotonie, comme contrainte définissant l'axiome de l'indépendance de façon plus stricte, conduit à restreindre le champ d'application de cet axiome et à résoudre bon nombre de paradoxes expérimentaux. Ici, le changement de l'outil mathématique permet à l'économiste de rapprocher le modèle des comportements en incertitude observés par le psychologue.

Quand les « probabilités » sont non additives et la décision risquée

Un nouveau modèle, fondé sur des « probabilités » non additives, a vu le jour pour capter l'idée que les probabilités objectives ne sont pas perçues telles quelles par les individus mais font l'objet de transformations. Ce modèle de décision en univers risqué est celui de l'utilité dépendante du rang des conséquences, désigné par le sigle RDEU (*rank dependent expected utility*). Son nom provient du fait que la transformation des probabilités dépend du rang occupé par les conséquences quand celles-ci sont ordonnées par utilité croissante. Initialement développé par Quiggin [1982], sous le nom d'utilité anticipée, le modèle de l'utilité dépendante du rang connaît plusieurs variantes.

Le modèle générique de décision en univers risqué RDEU se construit à partir de deux principes : 1) le décideur raisonne en variation d'utilité et non à partir de l'utilité absolue, ce qui suppose de classer l'utilité des conséquences. Suivre l'ordre croissant de l'utilité des conséquences est nécessaire pour montrer que les préférences influencent la perception des probabilités ; 2) le décideur utilise des probabilités cumulées et non des probabilités simples. Ces principes permettent de rendre compte de la déformation des probabilités observée lors des tests expérimentaux. Depuis Kahneman et Tversky, on sait que la fonction de transformation vise à sous-pondérer les probabilités, de telle sorte que $V(X) < EU(X)$. La non-additivité des « probabilités » traduit l'idée qu'elles font l'objet, dans la vie courante, d'un traitement non linéaire. Quand les « probabilités » ne sont pas additives, le décideur peut transformer les probabilités objectives dont il dispose, il peut, par exemple, les sous-pondérer systématiquement, puisque, par définition, il n'est pas tenu de respecter la condition du modèle de l'UE selon laquelle la somme des probabilités doit être égale à l'unité.

Le modèle RDEU à l'œuvre

Pour simplifier la présentation de la logique du modèle, supposons, dans un premier temps, que seuls trois états du monde sont envisagés, auxquels sont associées les probabilités π_1, π_2, π_3 .

Soit ϕ , la fonction de transformation définie de $[0,1]$ dans $[0,1]$. Dans ce modèle, la fonction $V(X)$ représentative des préférences se présente ainsi :

$$(1) V(X) = u(x_1) + \phi(\pi_2 + \pi_3) [u(x_2) - u(x_1)] + \phi(\pi_3) [u(x_3) - u(x_2)].$$

Cette formulation s'interprète comme suit : le décideur raisonne en variation d'utilité. Il commence par évaluer l'utilité minimale qu'il est sûr de recevoir, soit $u(x_1)$. Puis, il fait le produit de l'accroissement possible de son utilité $[u(x_2) - u(x_1)]$ par la probabilité transformée $\phi(\pi_2 + \pi_3)$ de recevoir au moins cet accroissement d'utilité. Ensuite, il associe l'ultime variation de son utilité avec la probabilité transformée $\phi(\pi_3)$ de la recevoir. Les poids décisionnels qui déforment les probabilités cumulées changent selon l'ordre des conséquences.

Le modèle RDEU englobe celui de l'utilité espérée de von Neumann et Morgenstern. Pour le montrer, partons de l'équation (1) du modèle qui peut, après quelques transformations, se réécrire comme suit :

$$(2) V(X) = u(x_1) [\phi(\pi_1 + \pi_2 + \pi_3) - \phi(\pi_2 + \pi_3)] + u(x_2) [\phi(\pi_2 + \pi_3) - \phi(\pi_3)] + u(x_3) \cdot \phi(\pi_3).$$

La fonction φ n'a pas les propriétés de l'additivité mais, si l'on suppose que les probabilités sont transformées linéairement, alors les poids décisionnels peuvent s'additionner et l'équation (2) se réécrit :

$$(3) V(X) = u(x_1) \cdot \varphi(\pi_1) + u(x_2) \cdot \varphi(\pi_2) + u(x_3) \cdot \varphi(\pi_3).$$

Dans cette hypothèse, du fait des caractéristiques de la fonction, on a nécessairement $\varphi(\pi_1) = \pi_1$, etc. et la formule n'est autre que la formule classique de calcul de l'espérance de l'utilité. Cette dernière est devenue un cas particulier du modèle RDEU.

Illustrons le modèle de l'utilité dépendante du rang, en prenant l'exemple d'un automobiliste ayant le choix entre deux trajets A et B, sachant que leur durée est conditionnée par l'état aléatoire de la circulation. Il est admis que, pour ces deux trajets, les probabilités d'encombrement sont des données objectives que l'automobiliste déforme. Les données du problème pour le trajet A sont synthétisées dans les deux premières lignes du tableau suivant.

	<i>États de la circulation</i>		
	<i>Embouteillée</i>	<i>Chargée</i>	<i>Fluide</i>
$U(x_i)$	15	70	100
Probabilité (π_i)	0,7	0,2	0,1
$U(x_i) - U(x_{i-1})$	15	55	30
Probabilités cumulées	1	0,3	0,1
Probabilités « cumulées transformées »	1	$\varphi(0,3)$	$\varphi(0,1)$
Utilité (RDEU)	15	$\varphi(0,3) \cdot 55$	$\varphi(0,1) \cdot 30$

L'automobiliste commence par évaluer l'utilité minimale qu'il est sûr d'obtenir et qui correspond au trajet embouteillé. Puis il multiplie l'accroissement de son utilité dans le cas où la route est simplement chargée par la probabilité transformée d'avoir une route au pire chargée (la transformation porte sur la probabilité cumulée d'avoir une route chargée et celle d'avoir une route fluide). Bien entendu, la variation d'utilité obtenue quand la route est fluide est multipliée par la probabilité transformée d'avoir une route fluide. Enfin, il calcule l'utilité espérée dépendante du rang, soit, dans cet exemple :

$$15 + \varphi(0,3) \cdot 55 + \varphi(0,1) \cdot 30.$$

Le conducteur procède de la même façon pour construire la fonction de représentation du trajet B et choisit l'option lui procurant la plus grande utilité.

En généralisant à n états de la nature, la formulation du modèle RDEU devient :

$$(4) V(X) = u(x_1) + \varphi \left(\sum_{i=2}^n \pi_i \right) [u(x_2) - u(x_1)] + \dots \\ + \varphi \left(\sum_{i=j+1}^n \pi_i \right) [u(x_{j+1}) - u(x_j)] + \dots + \varphi(\pi_n) [u(x_n) - u(x_{n-1})]$$

Dans ce chapitre, nous supposons constamment que l'ensemble des événements est fini et discret, alors que dans les modèles non additifs, comme dans celui de Savage, l'ensemble est infini et continu, de sorte que la formule la plus générale s'exprime à l'aide d'intégrale.

Quand la mesure de l'incertitude fait appel aux capacités

Que devient ce type de modèle quand les probabilités ne sont pas « données » au départ, c'est-à-dire dans une situation d'incertitude ? La plupart des approches à la mode, au tournant du XXI^e siècle, peuvent se décliner en référence au modèle connu sous le nom de l'espérance d'utilité à la Choquet ou CEU (*Choquet expected utility*), initié par Schmeidler [1989]. Dans ce modèle, le décideur raisonne non plus à partir de probabilités mais à partir de capacités. Cependant, la logique du modèle présente de nombreuses analogies avec celle du modèle RDEU.

- *Une réponse au paradoxe d'Ellsberg.* — Le modèle d'utilité espérée à la Choquet a pour ambition de répondre au paradoxe expérimental d'Ellsberg. Ce paradoxe, rappelons-le, montre que l'on ne raisonne pas de la même façon, selon que les probabilités sont connues ou non. Dans ce dernier cas, on s'appuie sur des croyances. Comment mesurer les croyances des individus dans l'occurrence des événements ? Pour les mesurer, le modélisateur a fait appel à la notion de capacité définie par Choquet. Une capacité est une fonction, notée ici $w(\cdot)$, dont les propriétés sont moins contraignantes que celles d'une fonction de probabilité. La capacité de Choquet permet, notamment, de modéliser un raisonnement familier, du type : « La possibilité que l'événement A se réalise est au moins de 60 %. » Par exemple, un décideur pessimiste, ne sachant pas quelle est la vraie loi de probabilité, affectera à chaque événement possible la capacité w égale à la valeur minimale qu'il accorde à la réalisation de l'événement.

La formulation du modèle est analogue à celle du modèle RDEU. Ici encore, les capacités ne sont pas additives. Le décideur calcule son utilité en commençant par la valeur minimale qu'il obtiendra dans le pire des états et ajoute les suppléments d'utilité hypothétique en les pondérant par la fonction w de mesure subjective de leur occurrence. En raisonnant par analogie avec le modèle RDEU, on

pourrait montrer que la fonction représentative des préférences selon Choquet englobe celle du modèle de base de l'espérance subjective d'utilité, pour n'en faire plus qu'un cas particulier, celui où les capacités sont additives.

La similitude entre les deux modèles non additifs, le modèle RDEU et celui de l'espérance d'utilité à la Choquet, n'est pas seulement formelle. L'exemple d'Ellsberg pourrait s'interpréter, comme son auteur le suggérait, à partir de la notion de poids décisionnels. Les probabilités qui sont obtenues en appliquant le principe suivant lequel, faute de mieux, on suppose l'équiprobabilité, pèseront moins dans la décision que les probabilités connues avec précision. Inversement, l'exemple de l'automobiliste utilisé pour illustrer le modèle RDEU pourrait être transposé ici. Un conducteur qui a une connaissance imprécise des probabilités attachées aux différents états de la circulation raisonnera en termes de capacité ou de plausibilité minimale, donnée par la fonction $w(\cdot)$, pour choisir entre les trajets A et B.

- *Les limites de l'ambiguïté.* — C'est une conception très particulière de l'incertitude que retiennent les modèles de décision. Elle se restreint à l'ambiguïté, donc au manque d'information. L'approche par le paradoxe d'Ellsberg nous invite à concevoir les difficultés de la décision en incertitude à partir d'un choix entre des loteries, c'est-à-dire dans un cadre bien défini où seule manque une donnée du problème, la répartition des boules dans l'urne.

L'incertitude ne peut pas se résumer à l'ambiguïté et, plus généralement, à l'imperfection de l'information, elle a d'autres origines, notamment la complexité de la situation (chapitre 1). Or, avec la complexité et le besoin de clarifier le problème, il est nécessaire d'analyser la décision comme un processus global et de ne pas se limiter à la phase du choix. Quand l'incertitude se fait plus envahissante, c'est le mode même de raisonnement par l'optimisation qui est remis en cause.

4. Simon et la complexité

Pour H.A. Simon, prix Nobel d'économie en 1978, quand les individus sont mis dans des situations simples, à l'occasion de tests de psychologie cognitive, ils se comportent, dans l'ensemble, comme le prévoit le modèle de l'utilité espérée et l'on peut parler de rationalité parfaite ou substantielle, c'est-à-dire définie par les conséquences. Dès que l'on s'écarte de ces situations, ce modèle n'explique plus

Schackle et le débat sur la décision en incertitude*

Dans les années 1950, George Schackle (1903-1992) reproche aux économistes qui modélisent la décision rationnelle, leur conception du temps et de l'incertitude. Pour Schackle, la décision ne peut pas s'analyser à partir d'une urne déjà remplie, c'est-à-dire d'un modèle que l'on peut se représenter aisément, l'incertitude n'étant pas simplement un problème d'acquisition d'information dans un monde fini. Il associe l'incertitude à la liberté d'imaginer, à la créativité du choix et à l'indétermination du futur. Il estime, comme Knight et Keynes, que la décision en incertitude se heurte aux limites de la connaissance dans un monde indéterminé.

Par exemple, selon Schackle, l'homme d'affaires doit fournir un effort d'imagination pour concevoir l'espace des mondes possibles. Néanmoins, envisager tous les états du monde possibles est un exercice impossible, des hypothèses résiduelles demeureront toujours inexplorées (ce sont les « *unforeseen contingencies* »). Pour décider au mieux, il doit se focaliser sur deux hypothèses, celle qui donne la meilleure et celle qui aboutit à la pire conséquence, et évaluer leur

degré de possibilité. Schackle qualifie ces possibilités de « degré de surprise potentielle » et les substitue aux probabilités, leur somme n'étant pas nécessairement égale à l'unité.

En raison de sa position critique à l'égard de la vision mécanique du choix, Schackle n'a pas développé de modèle concurrent de celui de l'utilité espérée et, dans les années 1960, c'est l'approche de Savage qui s'est imposée dans le débat sur le comportement en incertitude. Au tournant des années 1980, la conception de la décision en incertitude proposée par Schackle suscite, de nouveau, l'intérêt. Ainsi, Shafer [1976] a développé un modèle de décision qui part de l'idée qu'en incertitude radicale le décideur ne peut pas se représenter tous les états du monde. De même, les fonctions de croyance (BEL) et de doute du modèle de Shafer renvoient à l'idée de surprise potentielle de Schackle.

* Une sélection des travaux de Schackle est parue sous le titre de *Time, Expectations and Uncertainty in Economics* [1990].

les comportements et l'on ne peut plus se référer à la rationalité parfaite.

De la simplicité à la complexité

Quand H.A. Simon parle de situation simple ou plutôt « *very simplest* », [1976, p. 143], il a comme référence : 1) le modèle de Savage où le cadre est donné, les alternatives sont fixées à l'avance, le décideur peut saisir la totalité des éléments pertinents pour son choix ; 2) un monde où l'information est précise et parfaitement disponible. Dans toutes les autres situations, qualifiées de complexes, le théoricien doit rendre compte de la capacité des individus à délibérer. Ce terme très général recouvre une série d'éléments inséparables de la décision : l'intuition que Simon associe à l'expérience et

qui permet de reconnaître une configuration déjà rencontrée dans le passé ; la qualification de la situation qui permet de structurer le problème ; la construction d'alternatives qui renvoie aux capacités humaines d'anticiper et d'inventer ; l'élaboration d'une solution au problème. Autrement dit, la théorie de la décision doit incorporer une théorie de la recherche portant à la fois sur la construction du problème et sur sa solution.

L'existence de comportements se différenciant selon le contexte le conduit à distinguer deux modes d'évaluation de la rationalité. Dans une situation simple, la rationalité est jugée au regard du choix effectué, ce choix devant être conforme à la règle de décision donnée par le modèle. Simon la qualifie de rationalité substantielle. Il voit dans le modèle de Savage un cas particulier d'explication du comportement, dans des situations où le processus de délibération est évident. Dans les situations complexes, la rationalité ne s'analyse plus à partir du choix et en référence à la règle mais au regard du processus de délibération. Ce dernier n'a rien d'évident et Simon parle, alors, de rationalité procédurale, l'important se déroulant dans la « boîte noire » du décideur. La priorité donnée au cognitif le conduit à considérer que l'incertitude existe seulement dans l'esprit du décideur [Simon, 1976].

Bien que sa théorie soit surtout positive, elle conserve une inspiration normative. Certes, il ne s'agit plus, pour lui, de proposer un mode de décision parfait, mais d'aider à la décision stratégique. Le processus de délibération étant coûteux, il est recommandé aux décideurs de ne pas perdre leur temps à envisager toutes les éventualités. Mieux vaut s'arrêter de les explorer, lorsqu'une solution satisfaisante a été atteinte (*satisficing*).

Du modèle à la règle

Avec la complexité, les individus sont confrontés à l'incertitude de la situation présente et aux difficultés de sa perception, sachant qu'ils souffrent plus d'un excès d'information que d'un manque, contrairement aux idées reçues. Les difficultés posées par la perception de la situation sont aggravées lorsque les personnes ont à se coordonner. Ceci est manifeste dans un collectif. En effet, la façon dont les individus se saisissent d'un problème et le structurent dépend de la fonction occupée dans une organisation [March et Simon, 1958]. Il sera difficile de faire partager la même vision du problème aux membres de l'organisation. On en déduit l'intérêt des règles qui guident le comportement dans un collectif. Une autre conception de l'action se profile. L'action ne se déduit plus de la

décision rationnelle — c'est-à-dire maximiser son utilité — mais de la règle en usage dans l'organisation.

Une source majeure d'inspiration

La théorie de Simon est prolifique à la fois en termes de publications (son premier ouvrage date de 1947) et de thématiques. Il n'est donc pas étonnant qu'elle ait inspiré de nombreuses recherches de nature très différente.

À un pôle, figurent des modèles se rapportant à la rationalité substantielle. Dans ces modèles inspirés du « *search* » (chapitre II), l'individu introduit dans sa fonction d'utilité les coûts de recherche qu'il doit minimiser. Notons que Simon a, finalement, réfuté la paternité de ces modèles sophistiqués qui nécessitent des capacités de calcul supplémentaires. En effet, cette hypothèse est contradictoire avec les limites des individus dans leur capacité de calcul.

À l'autre pôle, figurent des travaux sur la rationalité procédurale. Pour Simon, la théorie du processus de délibération qui rendrait compte de la rationalité procédurale ne peut pas avoir l'élégance des modèles de décision fondés sur la maximisation. Elle ressemblera plutôt à la biologie moléculaire avec sa riche taxinomie de mécanismes [1976, p. 145-146], comme le montre le très (trop ?) grand nombre de définitions de la rationalité [March, 1978]. Les travaux développés dans cette lignée se focalisent sur les procédures de décision dans les organisations et les règles. Ils retiennent l'idée que la rationalité est limitée par les capacités de collecte, de traitement de l'information et de calcul. La plupart des économistes ont tendance à se focaliser sur cette limite que Simon a mise en relief et qui leur est familière [Laville, 1998]. Cette conception de la rationalité limitée fait, néanmoins, l'impasse sur d'autres raisons pour lesquelles l'être humain ne se comporte pas en maximisateur. À la difficulté d'explorer le champ des possibles s'ajoute celle de concevoir le futur (Schackle et les économistes « autrichiens »). En outre, le conflit de jugement entre l'efficacité et l'équité peut placer le décideur devant un dilemme qui fait obstacle à la décision rationnelle. Certains courants contemporains aspirent à en rendre compte. Leur présentation fait l'objet du chapitre suivant.

Conclusion

Depuis le milieu du xx^e siècle, nombreux sont les économistes qui ambitionnent de saisir le comportement en incertitude, à partir de la décision optimale [voir les *surveys* de Cohen et Tallon, 2000 ; Hirshleifer et Riley, 1979 ; Shoemaker, 1982 ; Starmer, 2000 ;

Willinger, 1990 ; ainsi que les recueils thématiques de Diamond et Rothschild, 1978 ; Hey, 1997]. Le débat sur cette question s'est construit autour du modèle de l'espérance de l'utilité. Ce chapitre a montré comment les formalisations novatrices sont soucieuses de diminuer la distance entre le modèle et les comportements en incertitude, tels qu'ils ressortent de l'expérimentation. Rappelons que le propre d'un modèle étant d'expliquer, il se doit de simplifier. En économie, le modélisateur simplifie d'autant plus volontiers qu'il laisse aux sociologues une série de questions, comme celle de la détermination des préférences. Avec les probabilités non additives, un pas a été franchi pour rapprocher le modèle de l'enquête. Néanmoins, il est clair, pour tous, qu'il est utopique de vouloir saisir toute la complexité du processus de décision, celui-ci dépendant de ce qui se passe « dans la tête du décideur »... et du contexte [Laville, 2000].

D'autres courants de recherche posent que l'incertitude et la décision optimale ne vont pas de pair, comme le suggérait Armen Alchian, dès 1950. Quelles sont les propositions des économistes qui, comme H.A. Simon, étudient le comportement en incertitude, en sortant du cadre d'analyse traditionnel de maximisation des préférences ? Tel est l'objet du chapitre suivant.

V / Les institutions : un guide pour l'action

De nombreux économistes mettent en avant le rôle des institutions (normes, règles, conventions) parce qu'elles guident les comportements et, en les délimitant, aident à prévoir comment se comporteront les autres. Les institutions réduisent l'incertitude, en la domestiquant, et aident à agir raisonnablement, c'est-à-dire d'une façon qui n'est pas irrationnelle sans être pour autant conforme à la cohérence du choix prescrite par Savage.

Les courants présentés ici — néo-institutionnaliste, économie des conventions et évolutionniste — partagent l'idée que les agents s'accommodent de l'incertitude et reprennent l'hypothèse de rationalité limitée. Ceci étant, leurs développements divergent, en raison de leur conception différente de l'incertitude. Pour les néo-institutionnalistes, l'incertitude s'identifie à l'existence de phénomènes aléatoires, ce qui leur permet de ne pas s'éloigner du risque. Pour les autres, l'incertitude trouve son origine dans la complexité de la situation et la pluralité des informations. Paradoxalement, une information pléthorique peut créer de l'incertitude au lieu de la réduire. Comment éviter cet effet pervers ? Pour y répondre la théorie de la décision doit s'ouvrir sur l'action.

1. Williamson et l'incomplétude des contrats

Recourir au marché est coûteux et il peut être préférable, dans certains cas, de recourir à l'organisation [Coase, 1937]. Tel est le point de départ de l'analyse d'Olivier Williamson. L'échange, dans le monde réel, s'inscrit dans la durée et comprend trois moments : 1) la rédaction du contrat qui spécifie les conditions de l'échange ; 2) l'échange proprement dit ; 3) le suivi du contrat. Selon la nature des biens, ces opérations sont plus ou moins complexes et

coûteuses. Le coût d'une transaction sur le marché est comparé au coût de gestion de la même opération au sein d'une organisation et la solution la moins onéreuse est privilégiée. Plus généralement, les individus recherchent, parmi les modes de gouvernance d'une transaction, la formule la mieux adaptée, le critère de choix étant la minimisation des coûts de coordination, que ce soit sur le marché ou au sein de la firme. Les modes de gouvernance vont du marché à la firme en passant par un ensemble de contrats types (franchise, sous-traitance, etc.) voire de contrats *ad hoc*. Tel est, très grossièrement résumé, l'objet de la théorie développée par Williamson.

Le corpus d'hypothèses

Les hypothèses sur le comportement des agents économiques privilégiées par Williamson sont au nombre de deux : 1) la rationalité est limitée. Cette hypothèse, reprise de Simon, signifie pour Williamson que la connaissance de l'environnement et les capacités de prévision des agents sont limitées, ce qui pose problème pour rédiger des contrats complets ; 2) les individus sont opportunistes. Cette hypothèse, caractéristique de Williamson, signifie que les agents sont mus par leur intérêt personnel et peuvent recourir à la ruse ou à la tromperie.

À ces hypothèses clefs, s'ajoutent les déterminants des coûts de transaction dont les plus importants sont les suivants : 1) la spécificité des actifs (un actif spécifique perd de sa valeur lorsqu'il est utilisé pour un autre usage) qui rend notamment importante l'idée d'une continuité de la relation entre les agents. Par exemple, si la relation entre un sous-traitant et un donneur d'ordre prend fin, les actifs spécifiques du sous-traitant perdent de leur valeur ; 2) la fréquence des transactions ; 3) l'incertitude, Williamson donne à ce dernier terme un sens très précis. Un environnement incertain se caractérise par la fréquence des événements fortuits et imprévisibles auxquels les parties prenantes de la transaction doivent s'adapter.

Les comportements des agents induits par les hypothèses (rationalité limitée, opportunisme) vont avoir des effets d'autant plus importants que la fréquence de la transaction, la spécificité des actifs et l'incertitude seront élevées. Ces trois éléments en amplifiant les effets des hypothèses comportementales contribuent à accroître les coûts de transaction. Ainsi le coût de rédaction d'un contrat entre deux agents sera d'autant plus élevé que les actifs mis en jeu seront spécifiques. En effet, la présence d'un actif spécifique dans l'échange offre un support pour un comportement opportuniste du donneur d'ordre. Il importe de prévenir ces comportements

en insérant des clauses adaptées dans le contrat, ce qui renchérit le coût de sa rédaction [Williamson, 1985].

Une conception très personnelle des contrats

Dans les cas complexes, indépendamment de son coût, le contrat n'a qu'une validité limitée. S'inspirant du modèle d'Arrow-Debreu, les économistes parlent de contrats complets pour désigner un contrat qui spécifie tous les actes à accomplir dans toutes les éventualités susceptibles de se réaliser (*supra*, chapitre II). Pour Hart et Holmström [1987], il est possible mais très coûteux de négocier un contrat complet de long terme parce que cela suppose que les contractants envisagent toutes les éventualités possibles et se mettent d'accord sur les actes à accomplir dans chacune de ces éventualités. Par définition, le contrat complet gère l'incertitude en la ramenant à des cas connus aléatoires.

Pour Williamson, cet exercice est illusoire. « En raison de la rationalité limitée, tous les contrats complexes sont inévitablement incomplets » [Williamson, 1990, p. 12], les agents ne pouvant pas envisager toutes les éventualités. En outre, la complexité ouvre sur des interprétations divergentes, ce qui limite le rôle de la justice. Celle-ci n'apparaît plus comme le garant de l'application d'un contrat mais comme un cadre pour exposer les conflits. La complexité et l'incomplétude offrent bien évidemment un espace propice au développement de comportements opportunistes. Se reposer sur les promesses pour résoudre l'incomplétude des contrats n'est pas réaliste pour Williamson, ce qui est cohérent avec l'importance qu'il donne à l'opportunisme des agents. Puisque les contrats complexes sont incomplets, les parties prenantes doivent chercher des garde-fous, telles que des garanties ou des cautions, pour que l'engagement soit crédible [Williamson, 1983].

L'institution : la solution pour domestiquer l'incertitude

Pour Williamson, l'organisation hiérarchique est souvent la réponse la plus satisfaisante que l'on puisse donner à l'incomplétude des contrats. L'illustration la plus évidente est fournie par le contrat de travail où l'accord porte non pas sur des actions détaillées mais sur l'acceptation d'une autorité. L'employé se soumet aux directives de l'employeur, pourvu que ce dernier respecte certaines limites, au demeurant assez vagues. Des coûts de rédaction du contrat sont ainsi économisés. Plus généralement, l'organisation hiérarchique est une alternative avantageuse par rapport au marché lorsque des dispositifs sont mis en place. Par exemple,

l'introduction de mécanismes incitatifs, tels que l'intéressement des salariés, réduit le risque de comportement opportuniste. La spécialisation des individus atténue les problèmes posés par la rationalité limitée en permettant aux agents de traiter une information plus étendue, etc. Au total, l'organisation et les contrats incomplets sont présumés plus aptes que le marché et les contrats contingents complets pour domestiquer l'incertitude.

• *Une incertitude peut en cacher une autre.* — Quand Williamson parle d'incertitude c'est dans un sens très précis, celui déjà évoqué des perturbations de l'environnement auxquelles il faut s'adapter. Mais l'incertitude est bien plus largement présente dans sa construction : 1) l'incertitude stratégique que Williamson qualifie d'incertitude comportementale est renforcée par l'hypothèse d'opportunisme [Williamson, 1985] ; 2) la rationalité limitée renvoie aux limites de la connaissance et au concept d'incertitude épistémique que l'on peut trouver chez Knight, Keynes et Hayek (cf. chapitre II). Le croisement de ces différentes sources d'incertitude débouche sur des coûts élevés de transaction et sur les limites de la justice à faire respecter les contrats, d'où la nécessité de rendre économiquement crédibles les engagements et le recours à l'organisation synonyme de hiérarchie.

Un terme reste à clarifier, celui d'institution. La littérature anglosaxonne lui donne un sens très extensif : l'institution inclut toutes les régularités de comportement qu'elles soient fondées sur des règles contractuelles, des routines, etc. Elle englobe aussi l'organisation, dès lors que celle-ci s'analyse comme un ensemble de règles cohérentes. Dans la littérature d'Europe continentale, l'institution prend souvent un sens plus restrictif et désigne, alors, les règles dont le tiers garant est une émanation de l'autorité publique, généralement la justice. Le scepticisme de Williamson sur la capacité de la justice à faire respecter un contrat complexe le conduit à ne pas isoler les règles de droit des autres.

Le travail de Williamson est ambivalent. D'un côté, l'incertitude ne fait pas obstacle au calcul économique et permet le choix du mode de gouvernance optimal ; en ce sens, il est très néoclassique. De l'autre, quand il affirme que les contrats complexes sont « inévitablement incomplets » [Williamson, 1990], il refuse l'idée que les agents économiques puissent parfaitement délimiter et traiter l'incertitude. Celle-ci n'est pas transformée en un risque mais simplement domestiquée.

2. L'économie des conventions et les repères partagés

Pour l'économie des conventions, agir avec autrui constitue une source majeure d'incertitude. Comprendre comment les agents économiques peuvent se coordonner en incertitude est une des visées de ce courant qui a émergé en France dans les années 1980.

Guide de lecture des conventions

Il est commode de partir de la définition de David Lewis [1969, 1983]. Celle-ci s'est imposée comme une référence incontournable bien qu'elle ne soit pas à l'origine du programme de l'économie des conventions. Reprenons sa définition. Une régularité R dans l'action ou dans l'action et la croyance constitue une convention dans la population P si, et seulement si au sein de P, les six conditions suivantes sont remplies (ou tout au moins sont remplies le plus souvent) :

- 1) chacun se conforme à R,
- 2) chacun croit que les autres se conforment à R,
- 3) cette croyance donne à chacun une bonne et décisive raison de se conformer à R,
- 4) tous préfèrent une conformité générale à R plutôt qu'une conformité légèrement moindre que générale — notamment plutôt qu'une conformité de tous sauf une personne (...),
- 5) R n'est pas la seule régularité possible à remplir ces deux conditions. Il existe au moins une alternative R' (...),
- 6) Pour finir, les différents faits énumérés dans les conditions 1) à 5) sont affaire de connaissance commune ou mutuelle : tout le monde les connaît, tout le monde sait que tout le monde les connaît et ainsi de suite (...).

Lewis [1983, p. 164-165 ; 1993, p. 12-13].

• *Le point focal à l'origine de la convention.* — Des éléments d'explication du fondement d'une convention peuvent être trouvés chez Thomas Schelling [1960]. Il montre que, lorsque les agents ne peuvent pas communiquer, les problèmes de coordination peuvent être résolus grâce à certains mécanismes cognitifs simples et à certains détails saillants qui s'imposent comme une convention. La meilleure façon d'exposer le point de vue des pionniers est de s'appuyer sur leurs exemples.

Deux individus ne se connaissant pas devaient se retrouver dans une ville sans rendez-vous précis. Que font-ils ? Chacun se rend à la gare centrale à midi pour retrouver l'autre. Bien qu'*a priori* il n'y ait aucune raison pour que cette solution se dégage plutôt qu'une autre,

l'expérience montre que les choix des agents convergent fréquemment. Pour ce faire, il faut que les agents, guidés par quelques détails saillants, aient un même cadre de référence sur lequel s'appuyer. Ainsi « les individus ne réfléchissent pas sur un mode purement abstrait ou conjectural : les possibles ne sont pas tous également possibles, le monde n'est pas un ensemble homogène d'hypothèses » [Boyer, Orléan, 1991, p. 236] et la convergence des anticipations des agents résulte de la capacité d'attraction de certaines solutions dites de point focal [Schelling, 1960].

Un autre exemple, le jeu des villes, confirme l'importance du cadre commun pour se coordonner en incertitude. On demande à chacun des participants qui ne doivent pas communiquer entre eux de répartir en deux sous-ensembles une liste initiale de onze villes des États-Unis. Pour que les joueurs gagnent, il faut que tous proposent la même répartition des villes. Ce jeu est un exemple de coordination en incertitude que les théoriciens des jeux qualifient de pure coopération puisque chacun a intérêt individuellement à coopérer. Les joueurs originaires des États-Unis réussissent assez bien lorsque les villes choisies peuvent se répartir selon un critère géographique familier. La référence au Mississippi, par exemple, permet de construire un sous-ensemble composé des villes à l'est et à l'ouest du fleuve. Les joueurs ne connaissant pas la géographie américaine réussissent plus rarement à s'accorder sur une répartition identique des villes, faute de disposer d'un cadre commun de connaissance où puiser des références offrant une solution simple. Cet exemple montre comment la convention est attachée à un groupe particulier et ne constitue pas une solution universelle à un problème. Comme le suggèrent ces illustrations, il ne s'agit pas de copier, d'imiter l'autre mais de converger vers la même solution.

Les débats

L'hypothèse de connaissance commune de Lewis a suscité des débats sur les anticipations sans limite. Prise au pied de la lettre, cette hypothèse suppose des capacités de raisonnement peu réalistes : pour réussir à se coordonner les personnes doivent simuler le raisonnement d'autrui « je sais que x connaît les faits (degré 1)... je sais que x sait que je connais les faits (degré 2) », etc. Cette spécularité sans limite impliquant des opérations cognitives très lourdes a conduit dans un premier temps à souligner la proximité de Lewis avec les néoclassiques. Sa définition de la convention a été utilisée comme repoussoir pour montrer, lorsqu'il s'agit d'expliquer la coordination, à quelle impasse aboutit une conception hyperrationnaliste où les agents sont capables de tout prévoir. Dans des jeux *ad*

hoc, comme celui du mille-pattes (chapitre II), l'hyperrationalité des joueurs empêche toute coopération bien qu'elle soit avantageuse pour les deux. La solution à ce type de jeu consiste à opposer au couple hyperrationalité et égoïsme de la théorie du choix rationnel le couple rationalité limitée et altruisme de l'économie des conventions [Dupuy, 1989]. S'engager dans une relation en sachant qu'autrui est susceptible d'opportunisme, c'est donner un gage révélant son intention de coopérer ; c'est la confiance dans l'autre qui déclenche l'action avec autrui et non le calcul de son bien-être ou le contrat [Favereau, 1996].

D'autres approches ont été développées qui sont compatibles avec le modèle de Lewis, dès lors que l'hypothèse de connaissance commune ou mutuelle est vue différemment. Lewis avait d'ailleurs lui-même suggéré dans ses travaux une interprétation plus réaliste de cette hypothèse, en notant que la connaissance commune peut être seulement potentielle [Lewis, 1969, p. 32, 63-64]. Dans la pratique, l'acteur n'accéderait, s'il en fait l'effort, qu'au premier degré des anticipations mutuelles, cette opération cognitive étant analogue au raisonnement que l'on mobilise lorsque l'on veut retrouver une personne dans un espace public. Chacun recherche dans l'environnement un repère, une chose faisant saillie, en se plaçant de son propre point de vue et en essayant de se placer du point de vue de l'autre pour tester la pertinence du repère.

• *D'où viennent les repères partagés ?* — La convention permet de gérer l'incertitude issue de situations complexes qu'elle simplifie en donnant des repères partagés. D'où viennent-ils ? Pas seulement des habitudes répond l'économie des conventions. Certains travaux valorisent la dimension matérielle des repères [Thévenot, 1989]. D'autres soulignent l'importance des croyances collectives dans la constitution des repères [Orléan, 1994] ou mettent en avant l'idée que le cadre commun aux acteurs intègre les jugements de valeur qui sont propres à un collectif [Boltanski, Thévenot, 1991]. Cette démarche a conduit à analyser la pluralité des façons de juger la qualité du travail et du produit [Eymard-Duvernay, 1989 ; Salais, Storper, 1993]. Pour Salais et Storper, la qualité « standard-générique » définit un ensemble de produits destinés à un vaste marché anonyme et fabriqués selon une convention de standardisation, connue de tous. L'incertitude sur la qualité est traitée comme un risque. Dans le monde de production qui lui est symétrique, la qualité est « spécialisée-dédiée », c'est-à-dire que le produit est fabriqué de façon originale pour un client particulier. L'évaluation de la qualité est problématique, faute de repères extérieurs aux agents. Elle reposera sur un accord mutuel. Dans d'autres

configurations, l'incertitude sur la qualité est surmontée grâce à des informations codifiées, comme la certification de qualité.

Ce sont, comme l'indique le titre ce chapitre, les institutions et, avec elles, les règles qui figurent en tête des repères partagés. Les règles sont analysées sous un nouveau jour. Elles sont traitées non comme des contraintes souvent inefficaces mais comme des ressources précieuses pour coopérer. Suivre une règle cesse d'être un acte passif et devient une activité cognitive où chacun anticipe la façon dont il sera évalué. Ceci rejoint l'idée que les règles dans une organisation s'analysent comme un point focalisant les attentes de conformité [Kreps, 1990]. Voir dans les règles un guide pour l'action plutôt qu'une contrainte conduit à s'intéresser à leur mise en œuvre dans les organisations. Certaines « règles fermées » sont très prescriptives, tandis que d'autres sont ouvertes, laissant une marge d'interprétation et d'adaptation à ceux qui les appliquent [Reynaud, 2002].

• *Ordre privé ou public ?* — À première vue, ces travaux valorisent l'ordre privé qui se construit sans recourir à une autorité extérieure. Pourtant, ce n'est pas du tout leur objectif. L'économie des conventions a un projet institutionnaliste bien différent de celui de Williamson. Le comportement des personnes n'est pas seulement guidé par le souci du gain et de l'efficacité mais aussi par des principes non utilitaristes, tels que la loyauté, l'équité. La convention permet de réduire les conflits, renforce la confiance dans les contrats implicites. Ce faisant, elle facilite la coordination en incertitude.

• *Quel apport pour la compréhension des phénomènes concrets ?* — L'économie des conventions entend montrer que la coordination est possible grâce à l'existence d'un cadre commun. Celui-ci intègre, selon les situations, les saillances de l'environnement (Schelling), les signaux de qualité (néoclassiques), des données contextuelles (philosophes pragmatistes), les pratiques communes fondées sur l'existence de précédents (Lewis). Une dimension essentielle sous-tend ces arguments : les acteurs impliqués délimitent un problème de façon commune ou tout au moins compatible, sans envisager toutes les éventualités, ce qui rejoint la notion de rationalité limitée. Prenons l'exemple de la relation de travail qui est *a priori* très conflictuelle. Comment les employeurs et les salariés s'accordent-ils pour gérer les aléas de la demande ? Dans les sociétés capitalistes modernes, les règles du travail s'organisent à partir de deux conventions : 1) la convention salaire-productivité où l'employeur s'engage à court terme sur un taux de salaire, le niveau d'effort du salarié et donc la productivité

du travail étant tenus pour stables ; 2) la convention de chômage dite keynésienne où il est admis par les deux parties que l'ajustement aux aléas de la demande se fait par le volume d'emploi. Ces deux conventions qui sont indissociables conduisent à ce que l'aléa économique soit essentiellement traité hors de l'entreprise, sous forme de chômage. Les règles de travail n'ont pas toujours eu ces soubassements conventionnels. D'autres conventions ont existé. Ainsi, dans le modèle paternaliste, l'emploi était tenu pour stable, l'ajustement de l'offre à la demande s'effectuait par la variation de la productivité [Salais, 1989].

L'analyse du premier couple illustre l'importance des repères partagés et objectivés. Ainsi, la convention keynésienne de chômage n'existerait pas sans les repères donnés par le salaire horaire et la mesure du temps de travail. La convention s'analyse comme un mode de traitement des aléas et/ou d'évaluation de la qualité qui est soutenu par des points focaux sur lesquels convergent, à un moment donné, les parties prenantes.

3. Les routines vues par les évolutionnistes

Le projet des économistes évolutionnistes est de construire une alternative à la théorie néoclassique. Ceux-ci considèrent que, du fait de la complexité de l'environnement, les agents économiques n'ont pas devant eux l'éventail des possibilités pour effectuer leurs choix et s'appuient sur des routines conçues comme une aptitude à exécuter une action répétée, dans un contexte donné [Cohen *et al.*, 1996, p. 683]. Les façons routinières d'exécuter des tâches ou de prendre des décisions évitent — pour le meilleur ou pour le pire — de s'interroger longuement sur le comportement à privilégier et réduisent l'espace des conduites possibles. Elles conduisent à négliger, dans les situations banales et répétitives, la possibilité que surviennent des aléas perturbateurs. Ainsi, les routines domestiquent l'incertitude.

Quels sont les fondements conceptuels des routines ?

La conception des routines a nettement évolué. Initialement, elles s'analysent comme : 1) une capacité individuelle d'exécuter une tâche répétitive, autrement dit, comme une compétence acquise en mobilisant un nombre limité de connaissances codifiées et surtout tacites ; 2) une règle de décision simple qui conduit à des automatismes, par exemple le budget consacré à la recherche développement doit être de 6 % du total des dépenses. Ces comportements font

peu appel au raisonnement et excluent en particulier les choix délibérés [Nelson et Winter, 1982].

Cette conception des routines repose sur deux référents théoriques : 1) les « habits » (coutumes, normes, habitudes) qui prolongent une tradition sociologique et évolutionniste amorcée par Veblen ; 2) les heuristiques (les processus mentaux de résolution d'un problème qui en captent la structure et conduisent à se focaliser sur une réponse possible, au lieu d'envisager toutes les possibilités) qui renvoient aux travaux de Simon et à l'intelligence artificielle [Mangolte, 1998]. Les évolutionnistes associent aux routines une conception de l'action intégrant l'expérience car celle-ci permet l'apprentissage (*learning by doing*). L'apprentissage est lui même le résultat d'un processus d'essais-erreurs et d'adaptation. Il faut souligner, en outre, l'influence du béhaviorisme dans les travaux fondateurs ; les règles-routines conduisent à déterminer le comportement par le contexte selon la séquence : si l'on est dans le cas S, alors la régularité de comportement R s'impose.

• *L'évolution des routines vers les heuristiques.* — Avec l'influence du courant cognitiviste, les évolutionnistes se sont polarisés sur le deuxième pilier conceptuel évoqué ci-dessus, c'est-à-dire les heuristiques. La conception de la routine est devenue plus sophistiquée et extensive. La routine tend à se définir, alors, comme un mode de résolution de problèmes nouveaux, elle devient dynamique et la résolution de problème devient un processus délibératif. Les agents adoptent un comportement de « *search* » où la recherche de solutions nouvelles procède par tâtonnement, essais-erreurs.

C'est pour montrer que les routines sont efficaces dès lors qu'elles induisent une solution appropriée au problème posé que les évolutionnistes accordent une importance croissante aux heuristiques. Se profilent alors deux figures de l'agent : 1) celle proposée par Brian Arthur à partir de l'exemple d'une personne qui, placée devant un tableau non signé, le perçoit comme un Veermer par assimilation et rapprochement de formes comme dans un jeu de puzzle [Arthur, 1992, p. 9]. Ici, le problème est exploré dans son contexte ; 2) l'informaticien programmant sur son ordinateur une procédure type, c'est-à-dire un algorithme contenant une liste finie d'instructions à suivre dans un ordre donné, ce qui permettra ensuite que d'autres individus, en suivant la liste des instructions, traitent de la même manière un problème similaire [Dosi, Egidi, 1991]. Ici, le problème est exploré indépendamment du contexte. On a donc deux conceptions éloignées, voire antinomiques de la résolution d'un problème, avec d'un côté une routine dite dynamique qui correspond à un mode d'exploration tacite et située et, de l'autre, une

procédure codifiée, hors contexte et qui ouvre vers une possible modélisation.

Pour illustrer cette problématique, prenons l'exemple d'un chauffeur de bus qui reçoit une multiplicité d'informations, venant de sources différentes. Outre sa propre perception de son environnement, un ordinateur de bord lui indique, *via* un satellite, sa position dans le trafic routier. De plus, il doit respecter de très nombreuses règles, puisque à celles du code de la route s'ajoutent des règles très contraignantes spécifiques à sa compagnie, telles que la règle des encadrants qui prescrit de respecter une distance minimale entre le bus qui le précède et celui qui le succède ou, encore, celle qui prescrit la ponctualité... Ce chauffeur de bus n'est pas à la recherche d'informations, au contraire, il est surchargé d'informations, susceptibles d'être contradictoires, et de règles, également susceptibles d'être contradictoires. Comment peut-il agir efficacement ? Pour cela, il doit trier parmi tous ces éléments qui guident sa conduite¹. Comme dans les modèles évolutionnistes de comportement sur les marchés, le chauffeur a, grâce à son expérience du trajet, construit son propre répertoire de routines, ce qui le conduit à opérer une sélection au sein des informations reçues et des règles prescrites. Les routines apparaissent, alors, comme un art de suivre les règles de l'organisation qui sont en arrière-plan [Reynaud, 2001].

Des agents homogènes au sein de la firme et hétérogènes sur les marchés

Les compétences acquises par apprentissage et les routines constituent le patrimoine immatériel des organisations. Les compétences globales des membres de l'organisation excèdent la somme des compétences individuelles et constituent un actif hautement spécifique [Coriat, Weinstein, 1995]. Chaque organisation développe un savoir tacite et une capacité d'adaptation différenciée à la nouveauté. La coordination se fait alors automatiquement [Nelson et Winter, 1982], les comportements des membres d'une organisation étant parfaitement prévisibles. Cependant tout n'est pas déterminé, le processus de sélection des routines au sein d'une firme étant pour partie inintentionnel. Les routines sont soumises à une logique d'évolution qui conjugue les actions délibérées et le hasard.

1. Cet exemple est librement inspiré d'une étude réalisée par Hatchuel A. et al. (1997), « Des autobus bien tempérés », in Moisdon J.-C. (éd.), *Du mode d'existence des outils de gestion*, Seli Arlan.

À l'homogénéisation des comportements qui s'opère au sein de la firme, les évolutionnistes opposent la diversité des comportements sur le marché. Du fait de l'hypothèse de rationalité limitée et du rôle accordé à l'expérience et à l'apprentissage, les agents se différencient, ils n'ont pas tous la même représentation de l'environnement ni les mêmes schèmes de traitement de l'information. Quand les routines sont définies comme des capacités de recherche individuelle de solution à un problème qui s'appuient sur des anticipations différentes, elles apportent de la complexité et des dynamiques sans équilibre.

Cette problématique aboutit à des modélisations originales où l'agent représentatif disparaît au profit d'agents différenciés. Arthur et d'autres représentants du courant cognitiviste l'ont appliquée à un marché financier artificiel, c'est-à-dire faisant appel à la simulation informatique. L'hypothèse d'anticipations communes est rejetée comme point de départ et remplacée par le principe suivant : chaque agent a son propre modèle de prévision, le teste, le rejette s'il ne donne pas de bons résultats. En bref, les agents créent en permanence de la nouveauté, donc de l'incertitude et de l'instabilité [Arthur et *al.*, 1997 ; Tordjman, 1997].

<i>Comment domestiquer l'incertitude ?</i>			
	<i>Économie néo-institutionnaliste</i>	<i>Économie des conventions</i>	<i>Économie évolutionniste</i>
Sources majeures d'incertitude	Les perturbations exogènes de l'environnement et l'opportunisme des agents	Complexité et singularité de la situation et incertitude sur le comportement d'autrui	Complexité et nouveauté de la situation et évolution aléatoire des routines et des techniques
Capacités des acteurs pour domestiquer l'incertitude	Les agents évaluent les coûts de transaction pour choisir le mode de gouvernement des comportements le plus efficace	La convention est un système d'attentes mutuelles permettant de se coordonner ; les conventions sont en arrière-plan des règles	Les acteurs utilisent des routines ou des programmes de routines pour traiter l'information et résoudre le problème ; les règles sont en arrière-plan des routines

Conclusion

Plusieurs courants théoriques, en économie, s'intéressent aux régularités de comportements, celles ci ayant pour avantage insigne de réduire l'incertitude. Ces régularités prennent, selon chaque courant, des appellations particulières qui renvoient à des conceptions différentes. Williamson part des individus et du marché pour aboutir à une construction où l'organisation hiérarchique et les contrats domestiquent l'incertitude, en bridant la liberté des agents. Il rejoint sur ce terrain la théorie de l'agence. Pour l'économie des conventions, les individus sont au contraire insérés initialement dans un contexte social qui les aide à se coordonner. Les règles communes aux participants s'interprètent en prenant appui sur des modèles conventionnels d'évaluation qui délimitent les attentes mutuelles. Pour les évolutionnistes, chaque agent opère une sélection parmi les règles, d'où émergent un nombre limité de routines. Si l'on suppose qu'il agit au sein d'une organisation, alors son comportement est largement prédéterminé. S'il est en position d'offreur ou de demandeur sur un marché, alors il détient un répertoire personnel de routines. *In fine*, économistes des conventions et évolutionnistes, à partir de fondements différents, se rejoignent pour proposer une lecture originale des règles.

TROISIÈME PARTIE
L'EFFICACITÉ
DE CERTAINS COMPORTEMENTS
EN INCERTITUDE

L'incertitude ouvre sur une large gamme de comportements. Dans les chapitres précédents, nous avons vu que chaque courant de pensée s'intéresse à un comportement type en incertitude, par exemple, maximiser son bien-être ou suivre des règles. Cette troisième partie privilégie l'analyse de trois comportements qu'adoptent fréquemment les individus pour se protéger de l'incertitude : l'imitation, l'assurance, l'attentisme. Quelle est leur efficacité ? Est-il efficace d'imiter autrui, en pensant qu'il est mieux informé et que l'on peut ainsi réduire l'incertitude (chapitre VI) ? L'assurance permet-elle de protéger les individus de dommages futurs, dans n'importe quelle situation ? Est-il préférable de décider aujourd'hui ou de reporter à demain sa décision en attendant de nouvelles informations ? (chapitre VII).

VI / Le mimétisme, de l'intérêt individuel aux infortunes du collectif

Imiter le comportement des autres, lorsqu'on a des doutes sur la façon d'agir, semble naturel. Pourtant la fable du mouton de Panurge rappelle que les effets d'un comportement grégaire peuvent être pervers. Keynes s'est plu à le souligner à propos de l'épargne et des marchés financiers [1936]. Il faut, cependant, attendre les années 1990 pour que le mimétisme soit introduit dans l'analyse économique. Ceci traduit un tournant. Admettre que les agents puissent former des anticipations fondées sur le comportement ou l'opinion d'autrui et non sur les déterminants fondamentaux tourne le dos à la théorie des anticipations rationnelles. Rappelons que, dans sa version forte, la théorie des anticipations rationnelles pose que les agents décident isolément les uns des autres, utilisent au mieux toute l'information et forment parfaitement leurs anticipations car ils connaissent le vrai modèle de l'économie.

Le thème du mimétisme en économie est, aujourd'hui, consacré par le terme de « cascade informationnelle », introduit dans un des modèles pionniers. Notons que les travaux sur le mimétisme ne l'associent pas toujours à l'incertitude, certains modèles, à la suite de Thorstein Veblen, le dérivant du conformisme social.

Dans les travaux présentés ici, les agents imitent pour réduire l'incertitude, entendue comme le manque d'information pertinente. Tous ces travaux privilégient une même question : quels sont les effets du comportement mimétique, quand il se généralise ?

1. Keynes à l'avant-garde

Le mimétisme chez Keynes est associé aux difficultés de la prévision, particulièrement manifestes sur le marché financier. Pour en

rendre compte, Keynes a distingué deux modes d'évaluation du cours des actions.

Comportement d'entreprise et de spéculation

L'investisseur qui adopte le comportement dit d'entreprise pour évaluer des titres financiers « prévoit le rendement escompté des capitaux pendant leur existence entière » [1936, p. 174]. Cette opération est nécessaire pour la détermination de la valeur fondamentale d'une action (c'est-à-dire la somme de la valeur actualisée des cash-flows anticipés). Ici, l'évaluation est menée à partir d'éléments objectifs.

Suivant un second mode d'évaluation, l'investisseur, qualifié alors de spéculateur, néglige sa propre opinion, pour s'intéresser uniquement aux anticipations des autres intervenants sur le marché financier. Ceci le conduira à acheter un titre s'il anticipe que le « marché » est durablement orienté à la hausse. Ici, l'évaluation fait appel au mimétisme.

Keynes fournit deux justifications au mimétisme : 1) si l'investisseur est susceptible d'avoir besoin de liquidité, dans le court terme, et donc de revendre son portefeuille, il est de son intérêt de se conformer à l'opinion moyenne du marché. « Il serait absurde, en effet, de payer 25 pour un investissement dont on juge que la valeur correspondant au rendement escompté est 30, si l'on estime aussi que trois mois plus tard le marché l'évaluera à 20 » [1936, p. 170] ; 2) comme la prévision de long terme est un exercice difficile, voire impossible, il est préférable de se comporter en « spéculateur ». « Ceux qui s'y attellent (prévision objective) sont sûrs de mener une existence beaucoup plus laborieuse et de courir des risques plus grands que ceux qui essayent de deviner les réactions du public plus exactement que le public lui-même ; et, à égalité d'intelligence dans les deux activités, ils risquent de commettre, dans la première, des erreurs beaucoup plus désastreuses » (p. 172).

Quand l'imitation tourne mal

La spéculation chez Keynes ouvre sur un comportement mimétique bien particulier. Imiter, ici, ne consiste pas à copier le comportement d'autrui mais à deviner ses anticipations pour prévoir les retournements du marché. L'exemple canonique est celui du jeu du concours de beauté. Gagne celui qui a voté pour le plus beau visage, c'est-à-dire le visage qui a reçu le plus de suffrages. Un joueur n'a pas intérêt à voter selon son propre goût mais pour ce qu'il estime être le goût moyen. Le processus est le même pour tous les joueurs

de sorte que le référent n'est pas le *goût moyen*, c'est-à-dire la moyenne statistique des goûts individuels, mais le *goût commun*, soit le résultat des anticipations mutuelles. Le *goût commun* qui, transposé aux questions financières, devient l'opinion du marché, se construit à partir d'un jeu spéculaire de devinettes où chacun anticipe ce que l'autre anticipe. *Goût moyen* et *goût commun* peuvent être très divergents. Cet exercice d'anticipations mutuelles fait émerger des évaluations artificielles, c'est-à-dire déconnectées des fondements objectifs, dès lors que chacun néglige son opinion personnelle. Transposé à la vie économique et généralisé, le mimétisme peut devenir pervers. Les vagues de spéculations sur les marchés financiers et leurs conséquences sur la vie économique en témoignent. Si Keynes est un des premiers économistes à parler du mimétisme, sa conception sophistiquée est difficile à modéliser. Les modèles les plus standard ne s'y réfèrent pas, leur conception de l'imitation étant triviale : imiter c'est copier.

2. Un modèle de cascade informationnelle

Le terme de cascade informationnelle a été introduit par Sushil Bikhchandani et *al.* [1992], en référence aux situations où l'incertitude vient d'un manque d'information. Ils ont illustré ce concept par l'exemple de la soumission d'un article dans un journal à comité de lecture. Supposons qu'un premier *referee* le rejette et que le second auquel l'article est soumis apprenne qu'il a déjà été refusé, il y a alors de fortes chances pour qu'il rejette également l'article, ne pouvant apprécier parfaitement sa qualité. Un troisième apprenant l'avis donné par les deux précédents aura le plus grand mal à l'accepter et ne le regardera peut-être même pas... Pour un projet d'article de qualité moyenne, on conçoit aisément qu'un mécanisme opposé aurait pu se déclencher si l'ordre des *referees* avait été différent. C'est ce type de mécanisme dont rend compte Abhigjit Banerjee [1992], avec l'exemple de deux restaurants. L'un est plein, l'autre vide, alors qu'ils sont de qualité comparable, ce résultat dépendant du choix des premiers clients.

Les ingrédients du modèle

La formalisation que proposent Bikhchandani *et al.* est présentée, ici, en la transposant à un exemple de la vie ordinaire, où l'enjeu est également peu important. Supposons que la route qui relie la ville de Montpellier à celle de Palavas puisse être impraticable à la suite d'inondations, un orage venant de survenir. Les automobilistes qui

ont l'intention d'aller à Palavas ne disposent pas d'une information parfaitement fiable sur l'état de la route. Vont-ils imiter les autres conducteurs et, s'ils le font, quelle est la probabilité que ce comportement soit optimal ?

Les automobilistes doivent décider s'ils se rendent ou non à Palavas. Chaque individu a accès à deux sources d'information : la première, que l'on qualifie de « signal privé », provient de sources objectives mais elle est plus ou moins fiable ; la seconde résulte de l'observation du comportement d'autrui. Dans ce modèle séquentiel, l'ordre suivant lequel les individus choisissent est exogène et connu de tous. Chaque automobiliste observe le comportement de tous ceux qui le précèdent.

Posons $V = 1$, l'utilité retirée par chaque automobiliste quand la route est praticable, $V = 0$, correspond à la valeur retirée par l'automobiliste quand elle est impraticable. Ces éventualités sont *a priori* équiprobables. Le coût associé à l'option « aller à Palavas » est, pour tous, par hypothèse de $1/2$; le coût de l'option « rester à Montpellier » est égal à 0. L'individu qui n'a aucune information sur l'état de la route est donc indifférent et choisit, au hasard, une de ces deux branches de l'alternative.

Soit X_i , le signal reçu par le i -ème individu. Il peut être : B la route est bonne ou M elle est mauvaise. La probabilité de recevoir un signal fiable est la même dans les deux cas. Ainsi, la probabilité de recevoir le signal B si la route est réellement bonne est identique à la probabilité de recevoir le signal M si la route est réellement mauvaise, elle est notée *prob*, avec *prob* $> 1/2$. Poser que les agents reçoivent un signal B ou M dont la qualité est tirée de façon aléatoire est un artifice de modélisation, qui permet de simuler l'activité d'interprétation d'une information à laquelle nous nous livrons, avec un raisonnement du type : « Cette radio ne donne pas toujours des informations fiables, je ne leur accorde qu'une valeur relative. » Ce doute est saisi par les probabilités. Celles-ci sont récapitulées dans le tableau suivant :

	$Prob(X_i = B/V)$	$Prob(X_i = M/V)$
$V = 1$	<i>prob</i>	<i>1-prob</i>
$V = 0$	<i>1-prob</i>	<i>prob</i>

Quand $V = 0$, alors $Prob(X_i = B/V) = 1\text{-prob}$. Cette probabilité s'interprète ainsi : l'automobiliste i reçoit avec une probabilité de *1-prob* un signal X_i lui indiquant que la route est bonne alors qu'elle est mauvaise.

Le processus séquentiel est le suivant :

— le premier automobiliste prend la route pour Palavas s'il reçoit le signal B, il reste à Montpellier s'il reçoit le signal M ;

— le second observe le comportement du premier afin d'en retirer de l'information. Supposons que le premier soit parti, cela signifie qu'il a reçu le signal B. En conséquence, si le second automobiliste reçoit, lui-même, le signal B, il part à Palavas sans hésiter, son choix étant conforté par celui du conducteur précédent. En revanche, si le second automobiliste reçoit le signal M, les deux informations sont contradictoires. Dans ce cas, il est posé que ces deux sources d'information se neutralisent, ce qui place l'automobiliste dans une situation d'indifférence, la probabilité qu'il parte à Palavas étant de $1/2$;

— pour le troisième individu, trois cas sont envisageables : 1) ses deux prédécesseurs sont partis, il en fait de même, indépendamment de la nature du signal qu'il reçoit en privé. Son choix marque le début d'une « cascade d'adoptions » ; 2) ses deux prédécesseurs sont restés, alors la réception du signal B est sans effet, il imite ses prédécesseurs et son rejet marque le début d'une autre cascade ; 3) l'un des automobilistes est resté, l'autre est parti. Ce troisième automobiliste se retrouve alors dans la situation du premier, il effectue son choix en fonction du signal qu'il a reçu en privé. Si cette séquence caractérisée par l'alternance des comportements se poursuit, alors le quatrième individu sera dans la même situation que le second, le cinquième dans celle du troisième, etc.

Ce type de modèle permet de répondre aux questions : quelle est la probabilité d'occurrence d'une cascade ? et quelle est la probabilité qu'elle soit bonne ?

Les cascades sont-elles inévitables ? souhaitables ?

Pour répondre à ces questions, construisons l'arbre suivant qui résume les différents choix possibles des deux premiers conducteurs et les probabilités correspondantes.

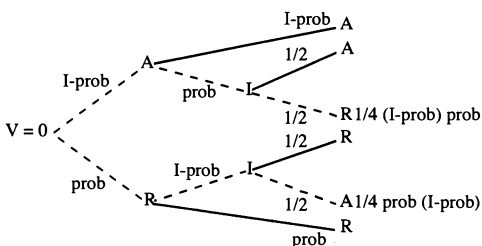
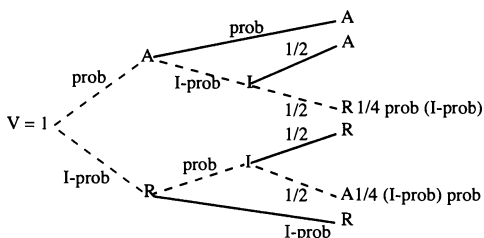
Les branches de l'arbre donnent le choix du premier puis du deuxième individu. L'avant-dernier chemin se lit ainsi : je (le conducteur n° 1) reçois une information me disant que la route est mauvaise et elle est vraiment mauvaise (la fréquence de ce signal est prob). Je reste à Montpellier (R). Tu (conducteur n° 2) reçois un signal indiquant que la route est bonne alors qu'en réalité elle est mauvaise (fréquence 1-prob), le signal et l'observation se neutralisent, chaque option est équiprobable ($1/2$) et tu décides d'aller à Palavas (A).

Il est aisé de déduire du circuit en pointillé, tracé sur le graphique, la « probabilité d'indifférence ». Ce terme désigne la probabilité

Valeur de chaque
éventualité

Choix du premier
individu

Choix du second
individu



A, aller à Palavas
R, rester à Montpellier
I, indifférent

probabilité *a priori* de réalisation = $1/2$

que les deux agents fassent des choix différents, par exemple A pour le premier et R pour le second. Cette probabilité est égale à $\text{prob}(1-\text{prob})$, ce qui s'obtient en faisant la somme des probabilités indiquées à l'extrême droite du graphique. En conséquence, par complémentarité, la probabilité d'apparition d'une cascade après deux séquences est égale à $1-\text{prob}(1-\text{prob})$, soit $1-\text{prob} + \text{prob}^2$.

La généralisation des résultats précédents à n individus conduit à une probabilité d'indétermination s'élevant à : $(\text{prob}-\text{prob}^2)^{n/2}$ et donc à une probabilité d'apparition de cascade qui admet pour valeur le complément à l'unité de la précédente, soit : $1-(\text{prob}-\text{prob}^2)^{n/2}$.

L'intérêt de ces calculs est de montrer que l'obtention d'une cascade dépend du nombre d'individus et de la valeur de prob. Rappelons que prob mesure la fréquence des signaux de bonne qualité. Si prob est proche de 1, très vite, tous les individus font le même choix. Dans le cas où prob tend vers $1/2$, la convergence — ou formation d'une cascade — est plus lente. On déduit de la formule ci-dessus que la probabilité d'éviter une cascade diminue de façon exponentielle avec le nombre d'individus. Ainsi, même pour un signal très peu significatif, tel que $\text{prob} = 1/2 + \epsilon$, la probabilité d'éviter une cascade, après que dix individus ont choisi leur itinéraire, est inférieure à 10 %.

Le mimétisme conduit-il à un résultat satisfaisant ? Le résultat dépend de la nature du signal reçu par les premiers entrants ; si, dans les premières séquences, dominent les signaux erronés, alors se constitue une cascade préjudiciable. Dès qu'une cascade a débuté, les choix des agents ne peuvent plus être guidés par le signal qu'ils ont reçu en propre. Puisqu'ils se sont contentés de copier, leur choix ne véhicule plus aucune information pertinente. Un nouveau venu ne pourra pas réorienter le choix du suivant. Ce modèle montre que, dans un univers d'information imparfaite sur la qualité, une attitude rationnelle des agents peut conduire, après un petit nombre d'entrées sur le marché, l'ensemble de la population à choisir unanimement une option erronée.

En résumé, quelle est la leçon de ce modèle ? Dans une situation incertaine, observer le comportement de ses proches pour en déduire de l'information, est, *a priori*, rationnel. En fait, ceci peut conduire à négliger sa propre information lorsqu'elle oriente dans un autre sens que celui du collectif, bien que cette information personnelle puisse être un meilleur guide.

Le modèle des cascades concerne des situations simples, ce qui justifie le choix que nous avons fait de retenir un exemple familier pour l'illustrer. Appliqué à des situations économiques plus complexes, il ne rend compte, le plus souvent, que de certains aspects des mécanismes en jeu. Par exemple, si l'on s'intéresse aux problèmes de diffusion des technologies, le choix que font les individus pour les ordinateurs PC compatible IBM plutôt qu'Apple s'explique, pour partie, par le mimétisme provoqué par l'incertitude sur la qualité du bien. L'analyse est plus complexe que dans l'exemple de Palavas car d'autres phénomènes interviennent, tels que les effets de réseau qui viennent modifier l'utilité que l'on peut attendre du bien. Ainsi, si les agents adoptent tous une même technologie, alors que celle-ci est initialement de qualité inférieure — la cascade est dite erronée —, ce ne sera pas simplement en raison

d'effets mimétiques mais également en raison des effets positifs de réseau.

Prenons un autre exemple, celui de la vague des fusions et acquisitions de ces dernières années. L'essor des acquisitions peut s'expliquer par l'adoption d'un comportement mimétique des firmes. Les dirigeants, en raison de l'incertitude attachée aux effets positifs réels d'une fusion, ont pu être incités à acquérir une autre entreprise de leur secteur parce qu'ils avaient constaté que d'autres l'avaient déjà fait avant eux. Le comportement des premières firmes ayant fusionné véhicule de l'information pour les autres du type : « Si cette firme a fusionné avant moi, c'est parce que les décideurs doivent avoir des informations sur les bénéfices attendus de la fusion (quantification des économies d'échelle, etc.) que je ne détiens pas, j'ai donc intérêt à les imiter »... Une cascade erronée peut apparaître ici assez facilement, car l'incertitude ne se dénoue pas aisément en raison du temps nécessaire pour évaluer la réussite de l'acquisition. Mais ici encore, le mimétisme n'apporte qu'un éclairage partiel sur les mécanismes à l'œuvre, d'autres facteurs plus complexes rentrent, bien sûr, en jeu, tels que la déréglementation des marchés financiers qui a favorisé les mouvements de capitaux.

3. Confiance dans son jugement ou dans le marché ?

La probabilité d'adopter l'opinion d'autrui est une fonction croissante du mimétisme. Cette évidence est manifeste dans le modèle séquentiel de Bikhchandani et *al.* D'autres modèles ont été développés en parallèle, où la probabilité évolue au gré des rencontres aléatoires des individus, d'emblée présents sur le marché. Certains sont purement évolutionnistes [Kirman, 1993], d'autres introduisent des variables d'inspiration keynésienne pour rendre compte de la complexité du processus de décision, en incertitude. André Orléan propose un modèle, appliqué au marché financier, d'inspiration à la fois évolutionniste et keynésienne [1990, pour la première version].

L'inspiration évolutionniste

Pour présenter la logique du modèle d'Orléan, supposons qu'un investisseur formule une probabilité *a priori* sur la valeur d'un titre financier. Il l'actualise ensuite selon un processus de révision qui tient compte des informations qu'il reçoit personnellement et de l'opinion dominante du marché. Le poids attaché à ces deux sources

d'information dépend du degré de confiance que l'individu a dans son jugement relativement à celui qu'il a à l'égard de l'opinion moyenne. Les probabilités *a posteriori* émergent au terme d'un processus d'interaction.

L'influence exercée par le marché est plus ou moins grande. Si les investisseurs ont un degré de confiance dans leur opinion relativement élevé, alors le marché agrège efficacement les opinions personnelles et les agents se rapprochent progressivement de la valeur fondamentale au terme d'un processus d'apprentissage. Ici, imiter est efficace. À l'autre pôle, si la confiance qu'ont les investisseurs dans leur propre jugement est relativement faible, ils ne se fient qu'à l'opinion du marché, ce qui déclenche un mécanisme d'autorenforcement (qualifié de rétroaction positive par les évolutionnistes). Quand le nombre de participants grégaires est élevé, alors la confiance que chaque investisseur place dans l'opinion du marché est élevée. Lorsqu'un individu modifie son opinion fixée *a priori*, sous l'influence de l'opinion du marché, il vient grossir la proportion d'individus qui forment l'« humeur du marché ». Ce changement est susceptible, à son tour, d'influencer un autre agent, etc. Pour Orléan, il s'agit d'une dynamique spéculative où le cours de l'action se déconnecte de sa valeur fondamentale, ce qui conduit à la création d'une bulle ou, au contraire, à l'effondrement du cours. Le mimétisme, lorsqu'il se généralise, ne peut conduire qu'à des excès, du fait de la polarisation des opinions sur une valeur « artificielle ».

Un modèle post-keynésien

Dans le modèle d'Orléan, les participants au marché financier sont d'emblée présents, à la différence du modèle des cascades. La non-séquentialité permet de rendre compte du jeu de l'influence réciproque des individus, ce qui renvoie à l'autoréférence et à la conception keynésienne du mimétisme. L'introduction d'un paramètre mesurant la confiance relative dans son jugement manifeste, également, le souci de capter la complexité de l'évaluation. Rappelons que la réflexion de Keynes sur les probabilités [1921] conduit à distinguer deux dimensions dans un jugement sur une éventualité : 1) la possibilité d'apparition de cette éventualité ; 2) la confiance que le décideur place dans cette conjecture, ce degré de confiance dépendant de la quantité et de la qualité des informations disponibles.

4. Le risque systémique

Dans une société capitaliste moderne, la spéculation, issue d'un comportement mimétique est devenue une réalité quotidienne. Sur un marché spéculatif, de nombreux agents veulent être les premiers à anticiper les retournements des tendances. Ceci explique que l'on passe brutalement d'un marché où les prix sont orientés à la hausse à un marché baissier. Si les spéculateurs ont pris des risques pour financer leurs engagements, cette tendance peut être dévastatrice. Un célèbre exemple de bulles spéculatives, celui de la « *mania* » des tulipes, résume bien les caractéristiques d'un marché spéculatif et d'une crise systémique.

Du bulbe de tulipe à la bulle

La folie mimétique pour les tulipes ou « *tulipomania* » a sévi en Hollande au XVII^e siècle. La tulipe, bien rare et cher, avait un caractère ostentatoire évident. Charles Mackay ¹ raconte comment, dans les années qui ont suivi l'arrivée en 1559 d'une cargaison de bulbes de tulipes chez un collectionneur hollandais, la tulipe est devenue un signe de distinction et a été recherchée par tous les nobles et les bourgeois hollandais. Cet accroissement de la demande a induit une élévation considérable des prix. Ainsi, un individu dépensa pour un bulbe *Semper Augustus*, qui était une variété rare, quatre mille six cents florins, soit l'équivalent de quatre cents bœufs, plus une diligence, deux juments grises ainsi qu'une bride et un harnais. Au mimétisme social s'ajouta le mimétisme spéculatif. La perspective de plus-value attira des spéculateurs faisant augmenter le prix des bulbes et de nombreuses fortunes se construisirent sur le négoce de tulipes. En 1636, année de l'éclatement de la bulle, les transactions sur les tulipes étaient non seulement le fait de collectionneurs fortunés mais également de professionnels de la finance, développant des techniques favorisant la spéculation. Selon la technique classique du marché à terme, l'échange de contrats permettait d'opérer sur le marché, en ne déposant que 10 à 20 % de la valeur de la transaction, lors de l'achat du contrat. Intervenait des spéculateurs qui achetaient à terme, sans disposer de la totalité des fonds, dans la perspective de financer l'acquisition de bulbes par leur revente le jour même où le contrat viendrait à exécution, l'opération devant dégager une plus-value. Quand les prix commencèrent

1. Mackay C. (1841), *Extraordinary Popular Delusions and the Madness of Crowds*, rééd., New York, Harmony Books, 1980. Cet exemple est cité par Tvede Lars (1994), *La Psychologie des marchés financiers*, Sefi.

à chuter, de nombreuses fortunes s'écroulèrent, les spéculateurs étant dans l'impossibilité d'honorer leurs contrats. Devant l'ampleur des dégâts, le gouvernement fut amené à stipuler que seuls les contrats postérieurs à novembre 1636 étaient valides et seraient remboursés à hauteur de 10 % de leur valeur, soit le minimum déposé pour acheter un contrat à terme. Néanmoins, l'éclatement de la bulle entraîna une forte récession économique en Hollande.

Quand le pessimisme s'installe en cascade

Dans les périodes de flambée de prix, s'instaure tôt ou tard le doute de chacun sur la poursuite de la tendance haussière. Ce doute conduit à un renversement brutal de la « cascade informationnelle », quand les acteurs, qui anticipent désormais des baisses de prix, deviennent vendeurs. Les conséquences du retournement des anticipations sont amplifiées par la prise de risque financier, l'endettement exerçant un effet de levier. Sur un marché à terme, il est possible d'acquérir un bien en ne disposant initialement, c'est-à-dire à la date de la signature du contrat, que d'une partie des fonds ou de vendre à terme un bien sans le posséder initialement. Quand les anticipations se retournent, passant de la hausse à la baisse, les acheteurs à terme manquent de liquidité pour honorer leur contrat. D'où des baisses brutales des prix d'autres actifs, à la suite de vente, en panique, des biens patrimoniaux des spéculateurs.

Les banques jouent souvent un rôle clef dans la propagation des cascades. Dans les phases haussières, les banquiers, victimes du climat d'euphorie, sous-évaluent le risque de défaillance de l'emprunteur en consentant des prêts au-delà des normes. Ce comportement est analysé sous le nom de myopie. Dans les phases baissières, les spéculateurs qui avaient aisément contracté des emprunts bancaires ne peuvent pas les rembourser. Lorsque les banques constatent que la qualité d'une partie de leurs créances s'est fortement détériorée, elles deviennent très méfiantes à l'égard de la masse des emprunteurs, même s'ils présentent peu de risque. Le pessimisme succède à l'optimisme [Kindleberger, 1978].

De nombreux autres mécanismes de propagation interviennent dans nos sociétés, du fait de la globalisation financière. Bref, un retournement des anticipations sur un marché particulier peut se répercuter à l'ensemble d'une nation, voire, aujourd'hui, d'une région du monde. C'est le risque de système [Aglietta, 1995].

L'exemple de la *tulipomania* illustre la nécessité d'une intervention des pouvoirs publics pour limiter la propagation de la crise à l'ensemble du système. Bien souvent, même si l'intervention des

pouvoirs publics, en tant que prêteur en dernier ressort, est limitée quantitativement, elle s'accompagne d'un effet d'annonce qui la rend efficace. En réduisant l'incertitude des dirigeants des institutions financières quant à l'accès à la liquidité, les pouvoirs publics peuvent ralentir voire stopper la cascade de pessimisme.

Conclusion

Dans les modèles mimétiques, copier est rationnel et permet de réduire l'incertitude. La justification de l'imitation repose sur une idée simple. Le comportement humain qui est le résultat des choix réfléchis des individus est un vecteur d'information dont tirent partie les agents peu informés. L'incertitude est donc, dans ces modèles, initialement réduite à un simple problème informationnel. Mais l'agrégation des comportements individuels et la constitution de cascades peuvent conduire à l'émergence d'« états du monde totalement nouveaux », imprévisibles. On débouche alors sur une vision du comportement humain et, avec elle, de l'incertitude beaucoup plus large selon laquelle le comportement des agents peut modifier l'environnement. L'action est susceptible de créer de la nouveauté et, ce faisant, d'engendrer de l'incertitude.

VII / Risque avéré ou incertitude scientifique, des gestions différentes

Le risque, entendu dans l'acception courante de dangerosité, recouvre des aléas de nature distincte selon que le danger est avéré ou non. L'économiste réserve le qualificatif de risque à l'ensemble des aléas avérés et prévisibles. L'autre ensemble regroupe les dangers purement hypothétiques dont l'occurrence et l'amplitude sont indéterminées scientifiquement, les domaines concernés étant surtout la santé et l'environnement.

Au sein des risques avérés, certains peuvent faire l'objet d'un contrat d'assurance, d'autres non. Pourquoi ? Y répondre suppose de comprendre la logique de l'assurance. Le principe de précaution est associé à la gestion des dangers relevant du second ensemble. Comment les économistes contribuent-ils au débat sur ce principe ?

1. Les trois figures de la prudence

Les individus ont toujours cherché à se protéger du risque, les modalités ont néanmoins évolué passant de formes de solidarité spontanée à des dispositifs impersonnels. On peut, suivant la suggestion de François Ewald [1997], distinguer trois figures de la prudence qui sont valorisées en Occident à différentes périodes de l'histoire moderne. Ces trois figures se construisent à partir des thèmes suivants : d'où vient le risque ? Comment se prémunit-on contre le risque ? Où se situe la responsabilité humaine ?

Au XIX^e siècle, la figure de la prudence est pensée sur le registre personnel à partir de la responsabilité pour faute et de la prévoyance individuelle. Avec la responsabilité pour faute, seul le dommage provoqué par l'erreur humaine donne droit à réparation. Quand surviennent des aléas où la faute n'est pas patente, c'est la fatalité qui

est invoquée et c'est l'épargne individuelle qui permet de se protéger.

Au xx^e siècle, la figure de la prudence est pensée sur le registre collectif de l'assurance. Avec le développement des statistiques, l'aléa cesse d'apparaître comme un événement fortuit pour devenir un cas identifié et prévisible. La protection individuelle repose sur la technique de l'assurance qui, en mutualisant les demandes de garantie, permet de reporter la charge sur un tiers anonyme. Le champ de la protection s'étend. Même les auteurs d'accident peuvent être indemnisés, comme l'illustre l'assurance tous risques automobile. La prévention collective se développe, en particulier dans la santé publique, avec les programmes de vaccination. En passant de la responsabilité pour faute à la responsabilité pour risque, le droit à l'indemnité s'étend puisqu'il n'est plus nécessaire de démontrer la faute. La thèse d'Ewald sur l'État-providence souligne le progrès que représente, au tournant du xx^e siècle, le passage à une figure de la prudence dominée par des procédures impersonnelles. Il n'est plus nécessaire de prouver la faute de l'employeur pour que la victime d'un accident du travail ait droit à réparation [Ewald, 1986].

Le lien interpersonnel s'efface derrière l'anonymat de la loi des grands nombres et la notion statistique de risque. L'aléa indésirable se détache de nos actes, devient une abstraction. Le chômage, par exemple, tend à être pensé comme un risque qui s'inscrirait dans la nature des choses tout comme la vieillesse. Il cesse d'être pensé comme une responsabilité des employeurs, il est externalisé avec le couple cotisation-indemnisation.

À l'aube du xxi^e siècle, se profile une troisième figure de la prudence, celle de la précaution. Le principe de précaution désigne la conduite à adopter pour gérer un risque non avéré, susceptible d'être un danger majeur pour la collectivité, qui n'est qu'au stade de l'hypothèse. La précaution se distingue de la prévention, ce second terme étant réservé aux risques dont l'existence est établie. Rappelons qu'il existe des dangers qui reposent sur un déterminisme strict. Ainsi la pollution des eaux par les rejets de l'industrie chimique est un risque connu auquel sont associés le principe du pollueur-payeur et une réglementation *ad hoc*. Le réchauffement climatique suscité par le mode de vie contemporain est d'une autre nature, puisque la complexité du phénomène fait obstacle à la validation des hypothèses émises par les experts ainsi qu'au repérage précis des déterminants et des conséquences. Le principe de précaution concerne la gestion de ce type d'incertitude, qualifiée de scientifique.

La nouvelle figure de la prudence se construit également à partir de la notion de développement durable, qui synthétise le souci de

protéger les générations futures de risques majeurs. Ceux-ci n'ont plus rien à voir avec la fatalité mais ont, au contraire, pour origine certaines activités humaines dont la dangerosité est en partie inconnue. Les risques les plus préoccupants ont pour origine l'interaction entre notre modernité et la complexité des phénomènes naturels.

Après la responsabilité pour faute du XIX^e siècle, puis pour risque du XX^e siècle, émerge ce que l'on pourrait appeler « la responsabilité pour défaut de précaution ». La vigilance face à des dangers potentiels majeurs et la rapidité de réaction aux indices de dangerosité deviennent des valeurs dominantes. Ainsi, avec l'affaire du sang contaminé, on a reproché aux pouvoirs publics en France de ne pas avoir introduit des mesures simples, comme l'utilisation de questionnaires, pour mieux identifier les risques associés aux donneurs. Gérer l'attente et l'arrivée de nouvelles informations devient une question cruciale pour les décideurs publics.

2. Le monde de l'assurance

La naissance de l'assurance est concomitante avec les débuts de la statistique. Quelles relations, le monde de l'assurance entretient-il avec celui de l'incertitude probabilisable ?

Le recours à l'assurance suppose que les demandeurs éprouvent de l'aversion au risque. S'assurer permet d'échanger une situation incertaine contre une situation certaine, tout en acceptant une diminution de sa richesse. S'assurer nécessite de trouver une contrepartie. Dans certains cas, le transfert de risque se fait directement et l'échange de risque peut même avoir lieu entre des partenaires également risquophobes, comme l'illustre l'exemple suivant lié aux données climatiques. Le directeur d'une station de ski pourra souhaiter qu'il neige beaucoup pour que la saison dure longtemps tandis que l'entrepreneur de BTP résidant dans la même région préférera qu'il ne neige pas pour faciliter l'entretien des routes. Il peut être dans leur intérêt de négocier un contrat de type *swap* (échange en anglais). L'entrepreneur payera une somme S à la station de ski si l'enneigement est inférieur à la moyenne saisonnière. Dans le cas contraire, ce sera la station de ski qui le dédommagera. Les deux parties peuvent ainsi compenser leurs risques, chacune des parties s'engageant à assurer l'autre, le manque à gagner de l'un étant atténué par le transfert d'une partie des gains d'aubaine de l'autre. Chaque assuré lisse le résultat de son activité, en renonçant à un gain aléatoire et, en parallèle, en évitant une perte hypothétique.

Longtemps circonscrite aux prix (prix des matières premières, taux de change, etc.), la liste des variables faisant l'objet de tels contrats s'est récemment allongée avec l'introduction des données climatiques. Cependant, son extension se heurte à des limites strictes. Elle suppose, notamment, le respect de différentes conditions : l'aléa doit être mesurable, la mesure doit être observable par tous, sa variation doit être indépendante de l'action des personnes.

Le plus souvent, le transfert de risque ne se fait pas selon un principe aussi simple. L'offre d'assurance repose, généralement, sur la mutualisation des risques.

Le champ de l'assurance est-il illimité ?

La mutualisation des aléas constitue un des principes clefs de l'assurance. Un grand nombre de personnes verse une prime pour un risque donné, ce qui permet de dédommager le petit nombre d'entre elles qui subit ce dommage. Ce principe renvoie à la loi des grands nombres. Partons de l'exemple classique du lancer de dé non pipé. Pour un petit nombre de lancers, certains chiffres sortent plus fréquemment que d'autres mais, pour un nombre élevé de lancers, la fréquence d'apparition tend à être la même pour chaque face. Le seuil des grands nombres est atteint, la fréquence observée se confond avec l'espérance mathématique. Cette loi suppose que le résultat de chaque lancer soit indépendant des précédents.

Les conditions d'existence de l'assurance déduites de la loi des grands nombres s'énoncent ainsi : pour un risque donné, si les cas de dommage sont indépendants les uns des autres et si le nombre d'individus soumis à ce risque est élevé, alors le risque peut être mutualisé et donc assuré. Lorsque l'indépendance des sinistres n'est pas respectée, comme dans les cas d'épidémie, d'inondation etc. alors seule la puissance publique peut indemniser les victimes.

Le véritable ennemi de l'assurance est l'aléa moral (cf. chapitre II). Quand la réalisation d'un événement aléatoire dépend largement de l'effort fourni par l'assuré, alors les possibilités de développer un comportement opportuniste sont trop importantes et l'assureur refusera de prendre en charge cet aléa. Une compagnie d'assurance ne garantira jamais à un entrepreneur un montant de bénéfice préalablement fixé. En pratique, la possibilité d'aléa moral est souvent surmontée et n'interdit pas l'assurance. Des clauses ou des mécanismes de tarification, tels que la franchise, incitent l'assuré à prendre soin de l'objet assuré.

• *Quid des événements singuliers ?* — Le monde de l'assurance est celui des aléas identifiés, répétitifs, qui peuvent être aisément

saisis par des lois de probabilités. En bref, c'est le monde de l'incertitude probabilisable au sens de Knight. Le transfert du risque, d'une personne à une autre ou d'une institution à une autre, est conditionné par l'existence de séries statistiques permettant d'estimer la fréquence d'occurrence d'un sinistre.

Toutefois, cette définition souffre d'exception. L'assurance couvre des cas singuliers, comme le montre l'exemple suivant. Lors de la campagne des présidentielles en France, au printemps 2002, une compagnie d'assurance a proposé à Robert Hue, candidat du PCF, de souscrire une assurance au cas où il n'atteindrait pas le seuil de 5 %, seuil qui ouvre droit au remboursement des frais engagés dans la campagne. Hue a refusé. Cet exemple montre que des cas atypiques peuvent s'inscrire dans le champ de l'assurance. En effet, 1) la garantie offerte n'allait pas inciter le candidat à relâcher son effort pour obtenir un pourcentage élevé de voix, autrement dit la question de l'aléa moral est évacuée ; 2) l'événement, quoique singulier, peut faire l'objet de prévisions (sondages, historique des résultats électoraux, etc.). Seule la mutualisation est remise en cause par le très faible nombre de candidats susceptibles de souscrire cette assurance. Pour y pallier, l'assureur fixe une prime élevée. En outre, il accepte de prendre un risque et de puiser dans ses bénéfices pour compenser la perte éventuelle occasionnée par ce type de contrat. Quand la mutualisation est difficile en raison du petit nombre de personnes concernées et que les montants en cause sont importants, alors la répartition des risques ne peut pas s'effectuer au sein des sinistrés potentiels et ce sont les actionnaires et des compagnies d'assurance organisées en syndicat qui prennent en charge le risque, en le fractionnant. Les limites du monde de l'assurance sont atteintes. Au-delà, s'ouvre le monde de la spéculation financière.

3. L'attentisme et le principe de précaution

Le principe de précaution concerne la gestion des risques en incertitude scientifique, c'est-à-dire quand l'existence d'un danger majeur pour la collectivité est seulement hypothétique, en l'absence de consensus scientifique. Il conduit à s'interroger sur l'existence d'une norme internationale de gestion de ce type de risque.

Attendre pour améliorer ses connaissances, dans un contexte où l'information future est susceptible de réduire l'incertitude radicale, est à la base du principe de précaution. L'attentisme n'est pas un comportement familier pour les économistes qui préfèrent étudier l'épargne. Néanmoins, différents travaux ont montré son importance dans les situations d'incertitude.

Keynes l'avait évoqué dans la *Théorie générale* avec le motif de spéculation [1936, chapitre XV]. Les investisseurs qui anticipent une hausse des taux conservent de la liquidité et reportent leur décision d'investissement dans des titres financiers. Comme le mimétisme, l'attentisme, pour Keynes, est susceptible de produire des effets pervers au niveau agrégé, la thésaurisation pouvant faire obstacle à la reprise. À la fin du xx^e siècle, les économistes ont associé l'attentisme à la notion d'irréversibilité et de valeur d'option, attendre devenant un comportement susceptible d'être optimal.

La valeur d'option d'Arrow-Fisher

Les notions d'irréversibilité et de valeur d'option sont apparues en économie de l'environnement afin de mesurer l'impact d'aménagements irréversibles sur un espace naturel. Elles ont été introduites par Kenneth Arrow et Anthony Fisher [1974] et par Claude Henry [1974].

Arrow et Fisher ont proposé d'illustrer leur modèle par un problème de développement d'une région forestière. Ils partent d'un modèle de décision séquentielle très simple, composé de deux projets alternatifs et de deux périodes (T1 et T2) : préserver une région forestière qui est à l'état naturel, par exemple en en faisant une zone de détente, sans investissement particulier (P_1), ou développer économiquement la région, par exemple en exploitant la forêt (P_2). Ils admettent que les décideurs en savent moins au début de la période T1 qu'à la fin de cette première période, date à laquelle ils pourront réviser leur programme de développement en fonction des connaissances acquises sur l'intérêt du projet. Ils pourront décider soit d'exploiter toute la région, soit de stopper l'extension de la zone d'exploitation forestière. Le projet P_2 peut, donc, être interrompu si les informations recueillies pendant la première période conduisent à réviser la comparaison des avantages/coûts et à revoir à la baisse les avantages comparatifs de cette option. Dans ce dernier cas, on regrettera d'autant plus de l'avoir entrepris qu'une partie de la forêt aura été endommagée, de sorte que l'option « préservation et zone de détente » verra sa surface réduite, par rapport à la surface initiale.

Leur modèle montre qu'introduire l'incertitude sur les avantages du projet « développement » exerce un effet sur le critère d'investissement. Ils concluent que « les bénéfices attendus d'une décision qui exerce des effets irréversibles devraient être ajustés (à la baisse) pour refléter la perte d'option qu'elle implique » (p. 319). Arrow et Fisher énoncent un principe économique de précaution. La comparaison coût/avantage de deux projets alternatifs doit refléter la perte de la valeur d'option, si le projet P_2 est sélectionné. Cette perte

d'option est indépendante de l'attitude vis-à-vis du risque. Elle se produit quand des décisions exercent sur l'environnement des effets qui ne pourront pas être supprimés, par impossibilité technique ou parce que le coût serait exorbitant.

Dans ce modèle, c'est parce que l'on commence à mettre en œuvre un projet que l'on en sait plus sur ses avantages/coûts futurs. Les avantages/désavantages se découvrent parce qu'ils dépendent de nos actes. Le décideur ne peut pas disposer, à l'étape initiale, de toutes les informations pertinentes. Il agit puis apprend, il doit donc agir prudemment.

Les usages du modèle

La notion d'irréversibilité s'est étendue, au-delà des questions de l'environnement, aux situations où l'arrivée potentielle de nouvelles informations est susceptible de faire regretter la décision initiale [Pindyck, 1991 ; Dixit et Pindyck, 1994 ; pour un *survey*, Bourdieu *et al.*, 1997]. L'acquisition d'information potentielle conduit à attribuer une prime ou valeur d'option à la décision qui préserve la flexibilité et laisse des « degrés de liberté », la décision la plus flexible étant celle qui conduit aux conséquences les moins irréversibles. La valeur d'option peut se définir comme la valeur associée à la possibilité de reconsidérer, ultérieurement, une décision, lorsque l'on sera mieux informé (encadré). Quand un entrepreneur réalise un investissement irréversible, alors il perd (exerce) l'option de le réaliser ou non dans le futur. Cette option perdue représente un coût d'opportunité qui s'ajoute au coût de l'investissement.

Le modèle de décision est séquentiel. Le choix se fait parmi différents scénarios dynamiques où les résultats, à chaque étape, sont susceptibles de changer, en raison de l'acquisition d'information, de l'apprentissage, du progrès technique, etc. La séquentialité est bien adaptée à des investissements dans des domaines qui, comme l'activité pharmaceutique ou pétrolière, se découpent en phases, de la phase amont de découverte d'une molécule ou d'un gisement pétrolier à la phase aval de commercialisation. Pour ce type d'activité, il est possible d'évaluer, par analogie, les probabilités associées aux différentes branches de l'arbre de décision, pour chaque étape. Celles-ci seront établies à partir des fréquences observées dans le passé. Certes, le médicament recherché est nouveau mais l'organisation de la recherche ne l'est pas, de sorte que le raisonnement probabiliste peut s'appliquer. Dans l'industrie pharmaceutique, le délai moyen nécessaire pour trouver une molécule nouvelle, puis le délai requis pour réaliser les tests avant d'obtenir l'autorisation de mise

La valeur d'option

Un entrepreneur envisage d'effectuer un investissement, d'un montant I . Le cash-flow immédiat associé à cet investissement s'élève à CF . Le montant des cash-flows ultérieurs est incertain et s'élève soit à $1,5 CF$, soit à $0,5 CF$. La probabilité d'occurrence de chacune de ces éventualités est de $1/2$. Le taux d'actualisation sélectionné par l'entrepreneur est k , le facteur d'escompte est β , avec $\beta = 1/(1 + k)$.

Quelle règle doit-il adopter pour décider d'effectuer ou non cet investissement ?

Une première solution consiste à se fonder sur le critère classique de la valeur actuelle nette (VAN)

$$VAN = -I + \sum_{i=0}^n CF_i \beta^i \text{ soit}$$

$$VAN = -I + \frac{CF}{1 - \beta} \text{ quand } n \rightarrow \infty$$

Supposons maintenant qu'à la fin de la période initiale, soit levée l'incertitude

sur les cash-flows. L'entrepreneur peut alors décider d'effectuer l'investissement seulement si le cash-flow, par période, est de $1,5 CF$.

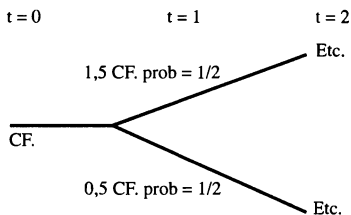
VAN^* , la valeur actualisée et espérée de cette stratégie alternative est :

$$VAN^* = \frac{1}{2} \beta \left(-I + \sum_{i=0}^n 1,5 CF_i \beta^i \right)$$

$$\text{soit } VAN^* = \frac{1}{2} \beta \left(-I + \frac{1,5 CF}{1 - \beta} \right) \text{ quand } n \rightarrow \infty$$

Sa règle de décision résulte de la comparaison entre VAN^* et VAN . Si $VAN^* > VAN$, il a intérêt à attendre la période 1 avant d'effectuer son investissement.

La différence entre VAN^* et VAN s'appelle la valeur d'option. Intégrer la valeur d'option durcit le critère d'acceptabilité du projet et incite à la prudence dans un contexte d'évolution des connaissances et d'irréversibilité.



Inspiré de Treich [2001], d'après Dixit et Pindyck [1994].

sur le marché du médicament sont connus, les probabilités de réussite, à chaque étape, le sont également.

• *Le protocole de Kyoto, une démarche exemplaire.* — Pour les problèmes liés à l'environnement ou à la santé publique, qui constituent le champ privilégié d'application du principe de précaution, calculer la valeur d'option, donc, les avantages de l'attentisme, devient une tâche très difficile. En effet, les décideurs publics font

face à des questions d'une très grande complexité, comme celle du changement climatique. En raison des multiples interactions entre l'atmosphère, l'océan, la couverture nuageuse... aucune base scientifique ne permet de s'accorder sur les probabilités de réalisation des différents scénarios. Aussi, les décideurs publics doivent-ils prendre au sérieux tous les scénarios proposés par les climatologues, des plus optimistes aux plus pessimistes. Ils ne disposent pas d'une règle qui permettrait de prendre la décision optimale. Cependant la référence à la valeur d'option est un guide. Elle enseigne que si l'incertitude scientifique est grande et si la recherche scientifique est susceptible d'apporter des connaissances nouvelles, alors la décision doit être flexible et se fonder sur une approche séquentielle.

L'idée d'un calendrier optimal dans la gestion des risques hypothétiques est associée à celle de flexibilité. Ce calendrier doit éviter deux écueils. Le premier est de prendre des mesures trop tardivement, ce qui peut transformer le dommage initial en une catastrophe. Le second est de prendre des mesures excessives trop tôt, ce qui ralentirait le progrès dans nos sociétés [Treich, 2001]. Le protocole de Kyoto s'en inspire. La stratégie consiste à commencer à réduire les émissions de gaz à effet de serre à l'échelle d'une génération (le très court terme pour les climatologues), cette orientation pouvant être révisée au profit d'un scénario plus fondé scientifiquement. Le protocole de Kyoto permet aux décideurs de garder ouverte la gamme des scénarios qui va du laisser-faire à la réduction massive des émissions de gaz, ce dernier scénario restant économiquement supportable, dès lors que des premiers pas ont été effectués dans ce sens.

Conclusion

Ce chapitre a privilégié, parmi les figures de la prudence, deux comportements, l'assurance et l'attente. Ceux-ci renvoient à deux mondes distincts. Le monde de l'assurance est celui de l'incertitude probabilisable, donc du risque. Dans le monde de l'attente, lorsque l'incertitude est scientifique, font défaut les informations permettant de spécifier les valeurs des paramètres décisifs de la décision. Pour les économistes, attendre n'est pas synonyme d'immobilisme bien au contraire. Attendre permet d'apprendre, d'acquérir des informations, voire de produire des connaissances nouvelles. Dans ces deux mondes, la relation action/connaissance est de nature différente. Dans le premier, la connaissance préalable est suffisante pour agir efficacement, dans le second elle ne l'est pas.

L'assureur auquel il est demandé de prendre en charge le risque d'incendie d'une usine calculera la prime à partir d'un modèle

utilisant des probabilités objectives d'occurrence d'un incendie établies à partir de l'observation de cas similaires. Néanmoins, son jugement personnel complétera le calcul. Dans les situations singulières et complexes, les référents s'inversent. Les éléments subjectifs l'emportent. Attribuer des valeurs aux paramètres qui orientent la décision, en se fondant sur des expériences similaires, est un exercice illusoire. Que deviennent, alors, les outils de l'analyse économique ? En incertitude scientifique, certes, le calcul de la valeur d'option ne permet pas d'aboutir à un résultat significatif, néanmoins, les travaux des économistes sur l'attentisme fournissent un guide au décideur et concourent à l'ensemble des réflexions sur le principe de précaution. L'analyse économique atteint ses limites qu'il convient de reconnaître.

Conclusion générale

L'incertitude conduit, nécessairement, les économistes à traiter de questions fondamentales concernant les rapports entre décision, action et connaissance. Les problématiques adoptées par les économistes contemporains sont source de divergences profondes, dont on peut rendre compte à partir de la distinction de Knight. Pour les néoclassiques, l'incertitude est probabilisable et peut donc être traitée comme un risque. Le plus souvent, elle s'identifie à un manque d'information dans un monde fini, où l'environnement peut être décrit à partir de catégories simples (les « états »). Cette conception, centrale en économie de l'information, rend réaliste le recours aux probabilités. Pour les autres courants économiques, l'incertitude n'est pas probabilisable. Le monde est trop complexe, l'usage des probabilités bute sur la singularité ou l'imprévisibilité des situations.

Ce clivage se renforce lorsque les économistes étudient la décision. D'un côté, selon le modèle de Savage (*subjective expected utility*), chacun, dans son for intérieur, peut traiter toute situation incertaine à partir de probabilités subjectives, de sorte que la distinction de Knight entre risque et incertitude est tenue pour inutile. De l'autre, l'introduction des institutions, comme guide de l'action et de la coordination, conduit à faire reculer, voire disparaître, l'usage des probabilités et, avec elles, le calcul économique.

Nous n'avons pas voulu rester sur ce constat radical. Nous avons pris le parti de montrer comment ces courants peuvent aussi se compléter. Ainsi, les néoclassiques, par leurs travaux sur les vertus de l'attentisme dans des situations incertaines et irréversibles, contribuent au débat théorique sur le principe de précaution. Cependant, entre le modèle de décision et son utilisation par les pouvoirs publics, la distance est immense. La souligner fait écho aux

réticences de nombreux économistes devant l'usage excessif des probabilités.

Remarquons que cet usage masque un problème de fond, celui de l'agent représentatif. Les travaux les plus récents sur le choix en incertitude font appel aux probabilités non additives pour intégrer l'idée que les personnes déforment les probabilités. Les agents économiques acquièrent ainsi plus de personnalité, se différenciant non seulement par leurs préférences mais aussi par leurs croyances sur les mondes futurs. Cette orientation met en cause la notion d'agent représentatif, ce que refusent la plupart des théoriciens néoclassiques. L'agent représentatif est préservé de deux façons. Le théoricien postule que les individus ayant accès aux mêmes informations forment les mêmes probabilités subjectives. Dans la plupart des modèles, les probabilités sont tout simplement données au décideur et donc présumées objectives. En refusant la personnalisation des agents, le modélisateur résout, ainsi, les problèmes d'agrégation des comportements.

La très faible utilisation des probabilités subjectives, au sens de personnelles (Savage), hors des modèles strictement centrés sur la décision, montre que la distinction proposée par Knight entre risque — c'est-à-dire incertitude probabilisable — et incertitude radicale — c'est-à-dire non probabilisable — a encore de beaux jours devant elle.

Repères bibliographiques

- AGLIETTA M. (1995), *Macroéconomie financière*, La Découverte, coll. « Repères », Paris (3^e édition, 2001).
- AKERLOF G. (1970), « The Market for "Lemons" : Quality Uncertainty and the Market Mechanism », *Quarterly Journal of Economics*, 84, 488-500.
- ALCHIAN A. A. (1950), « Uncertainty, Evolution and Economic Theory », *Journal of Political Economy*, 58, 211-221.
- ALLAIS M. (1953), « Le comportement de l'homme rationnel devant le risque », *Econometrica*, 21, 503-546.
- ANSCOME F., AUMANN R. (1963), « A Definition of Subjective probability », *Annals of Mathematical Statistics*, 34, 199-205, reproduit in DEY J., *The Economics of Uncertainty*, 1997, II, Edward Elgar, Cheltenham, Royaume-Uni.
- ARROW K. (1951), « Alternative Approaches to the Theory of Choice », in *Risk-Taking Situations*, *Econometrica*, 19, 404-437.
- ARROW K. J. (1963), « Uncertainty and the Welfare Economics of Medical Care », *American Economic Review*, 53, 941-973 ; trad. fr. in *Théorie de l'information et des organisations*, Dunod, Paris, 2000.
- ARROW K. J. (1968), « The Economics of Moral Hazard Further Comment », *American Economic Review*, 58, 537-539 ; trad. fr. in *Théorie de l'information et des organisations*, op. cit.
- ARROW K. J. (1985), « The Economics of Agency », in PRATT J., ZECKHAUSER R. (dir.), *Principal and Agents, The Structure of Business*, Harvard Business School Press, Boston Mass., 37-51.
- ARROW K., FISHER A. (1974), « Environmental Preservation, Uncertainty and Irreversibility », *Quarterly Journal of Economics*, 88 (2), 312-319.
- ARTHUR W. B. (1992), « On Learning and Adaptation in the Economy », *Working Paper*, Santa Fe Institute, 92-07-038.
- ARTHUR W. B. et al. (1997), « Asset Pricing Under Endogenous Expectations in an Artificial Stock Market », in ARTHUR et al. (dir.), *The Economy as an Evolving Complex System II*, Perseus Books, Reading, Mass., 15-44.
- BANERJEE A. V. (1992), « A Simple Model of Herd Behavior », *The Quarterly Journal of Economics*, CVII, 797-817.
- BARRÈRE A. (1985), *Keynes aujourd'hui : théories et politiques*, Economica, Paris.
- BATEMAN B. W. (1996), *Keynes's Uncertain Revolution*, Ann Arbor, The University of Michigan Press.
- BELL D. (1982), « Regret in Decision Making under Uncertainty », *Operations Research*, 30, 961-981.
- BIKHCHANDANI S., HIRSHLEIFER D., WELCH I. (1992), « A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades », *Journal of Political Economy*, 100, 992-1026.

- BOLTANSKI L., THÉVENOT L. (1991), *De la justification. Les économies de la grandeur*, Gallimard, Paris.
- BOURDIEU J., CŒURÉ B., SÉDILLOT B. (1997), « Investissement, incertitude et irréversibilité », *Revue économique*, 48, 23-53.
- BOYER R., ORLÉAN A. (1991), « Les transformations des conventions salariales entre économie et histoire », *Revue économique*, 42, 233-272.
- CAHUC P. (1998), *La Nouvelle Micro-économie*, La Découverte, coll. « Repères », Paris, nouvelle édition.
- CHIAPPORI P.-A. (1997), *Risque et assurance*, Flammarion, Paris.
- CHICK V. (1997), *Entretien*, in SNOWDON B., VANE H., WYNARCZYK P., *La Pensée économique moderne*, Ediscience International.
- COASE W. (1937), « The Nature of the Firm », *Economica*, 4, 386-405, trad. fr. (1987), *Revue française d'économie*, II.
- COHEN M. *et al.* (1996), « Routines and Other Recurring Action Patterns of Organisations : Contemporary Research Issues », *Industrial and Corporate Change*, 5 (3), 653-698.
- COHEN M., TALLON J.-M. (2000), « Décision dans le risque et l'incertain : L'apport des modèles non additifs », *Revue d'économie politique*, 110 (5), 631-681.
- COMBEMALE P. (2003), *Introduction à Keynes*, La Découverte, coll. « Repères », nouvelle édition.
- CORIAT B., WEINSTEIN O. (1995), *Les Nouvelles Théories de l'entreprise*, Le Livre de poche, Paris.
- DAVID P. (1985), « Clio and the Economics of Qwerty », *American Economic Review*, 75, 332-337.
- DAVIDSON P. (1991), « Is Probability Theory Relevant for Uncertainty ? A Post-Keynesian Perspective », *Journal of Economic Perspectives*, 5, 129-143.
- DIAMOND P. A. (1971), « A Model of Price Adjustment », *Journal of Economic Theory*, 3, 156-168.
- DIAMOND P., ROTHSCILD M. (1978), *Uncertainty in Economics : Readings and Exercices*, Academic Press, New York.
- DIXIT A., PINDYCK R. (1994), *Investment under Uncertainty*, Princeton University Press, Princeton.
- DOSI G. (1991), « Perspectives on Evolutionary Theory », *Science and Public Policy*, 18, 353-361.
- DOSI G., EGIDI M. (1991), « Substantive and Procedural Uncertainty », *Journal of Evolutionary Economics*, 1, 145-168.
- DOSTALER G. (2001), *Le Libéralisme de Hayek*, La Découverte, coll. « Repères », Paris.
- DUPUY J.-P. (1989), « Convention et common knowledge », *Revue économique*, 2, 361-400.
- ELLSBERG D. (1961), « Risk, Ambiguity, and the Savage Axioms », *Quarterly Journal of Economics*, 75, 643-669 reproduit in J. D. DEY, *The Economics of Uncertainty*, 1997, vol. II, Edwar Elgar, Cheltenham, Royaume-Uni.
- EWALD F. (1986), *L'État-providence*, Grasset, Paris.
- EWALD F. (1997), « Le retour du malin génie. Esquisse d'une philosophie de la précaution », in GODARD O. (éd.), *Le Principe de précaution dans la conduite des affaires humaines*, INRA, Paris.
- EYMARD-DUVERNAY F. (1989), « Conventions de qualité et formes de coordination », *Revue économique*, 40, 329-359.
- FAVEREAU O. (1996), « L'incomplétude n'est pas le problème, c'est la solution », in B. REYNAUD (dir.), *Les Limites de la rationalité*, tome 2, La Découverte, Paris, 219-234.
- FISHBURN P. C. (1978), « On Handa's "New Theory of Cardinal Utility" and the Maximization of Expected Return », *Journal of Political Economy*, 86, 321-324.
- FORAY D. (2000), *L'Économie de la connaissance*, La Découverte, coll. « Repères », Paris.
- FRISH R., BARON D. (1988), « Ambiguity and Rationality », *Journal of Behavioural Decision Making*, 1, 149-157.
- GABSZEWICZ J. (1994, 2003 nouv. éd.), *La Concurrence imparfaite*, La Découverte, coll. « Repères », Paris.
- GAYANT J.-P. (2001), *Risque et décision*, Vuibert, Paris.

- GILBOA I., SCHMEIDLER D. (1989), « Maximin Expected Utility with Non-unique Prior », *Journal of Math. Economics*, 18 (2), 141-153.
- GRANGER T. (1997), « Axiomatique des choix dans l'incertain », in SIMON Y. (dir.), *Encyclopédie des marchés financiers*, Economica, Paris, 186-217.
- GUERRIEN B. (2002), *La Théorie des jeux*, Economica, Paris.
- HACKING I. (1975), *The Emergence of Probability*, Cambridge University Press, Cambridge.
- HART O., HOLMSTRÖM B. (1987), « The Theory of Contracts », in T. BEWLEY (dir.), *Advances in Economic Theory*, Fifth World Congress, Cambridge University Press, Cambridge, 70-155.
- HAYEK F. (1973-1979), *Law, Legislation and Liberty*, Routledge, Londres, trad. fr. (1980-1983), *Droit législation et liberté*, PUF, Paris.
- HENRY C. (1974), « Investment Decisions under Uncertainty : The Irreversibility effect », *The American Economic Review*, 64, 1006-1012.
- HEY J. D. (dir.) (1997), *The Economics of Uncertainty*, Edward Elgar Publishing Company, Cheltenham.
- HIRSHLEIFER J., RILEY J. (1979), « The Analytics of Uncertainty and Information – An Expository Survey », *Journal of Economic Literature*, 17, 1375-1421.
- JENSEN M. C., MECKLING W. H. (1976), « Theory of the Firm : Managerial Behavior, Agency Costs and Capital Structure », *Journal of Financial Economics*, 3, 305-360.
- KAHNEMAN D., TVERSKY A. (1979), « Prospect Theory : An Analysis of Decision under Risk », *Econometrica*, 47, 263-291.
- KAST R. (2002), *La Théorie de la décision*, La Découverte, coll. « Repères », Paris, nouvelle édition.
- KEYNES J. M. (1921) [1973], *A Treatise on Probability*, Macmillan, Londres.
- KEYNES J. M. (1936), *The General Theory of Employment, Interest and Money*, Macmillan, Londres, in *Collected Writings*, vii.
- KEYNES J. M. (1937), « The General Theory of Employment », *Quarterly Journal of Economics*, 51, p. 209-23, in *Collected Writings*, vol. XIV ; trad. fr. par N. JABKO, *Revue française d'économie*, vol. V, 4^e trimestre, 1990.
- KINDLEBERGER C. P. (1978), *Manias, Panics and Crashes : A History of Financial Crisis*, Basic Books.
- KIRMAN A. (1993), « Ants, Rationality and Recruitment », *Quarterly Journal of Economics*, 108, 137-156.
- KIRZNER I. M. (1997), « Entrepreneurial Discovery and the Competitive Market Process : an Austrian Approach », *Journal of Economic Literature*, XXXV, 60-85.
- KLEIN B., LEFFLER K. (1981), « The Role of Market Forces in Assuring Contractual Performance », *Journal of Political Economy*, 89, 615-641.
- KNIGHT F. H. (1921) [1971a], *Risk, Uncertainty and Profit*, Houghton Mifflin, Boston, et University of Chicago Press, Chicago.
- KREPS D. M. (1990), « Corporate Culture and Economic Theory », in J. E. ALT, K. A. SHEPSON (eds), *Perspective on Positive Political Economy*, Cambridge University Press, Cambridge.
- KREPS D. M. (1990), *A Course in Microeconomic Theory*, Harvester Wheatsheaf ; trad. fr. *Leçons de théorie microéconomique*, PUF, Paris, 1996.
- LACHMANN L. (1976), « From Mises to Schackle : An Essay on Austrian Economics and the Kaleidic Society », *Journal of Economic Literature*, 14, 54-62.
- LAFFOND G., LESOURNE J. (1981), « Markets Dynamics and Search Processes with Information Cost », communication au Congrès européen de la Société d'économétrie, Amsterdam.
- LASLIER J.-F., ORLÉAN A. (2002), « Le marché élémentaire », in LESOURNE J., ORLÉAN A., WALLISER B., *Leçons de microéconomie évolutionniste*, Odile Jacob, Paris.
- LAVILLE F. (1998), « Modélisations de la rationalité limitée : de quels outils dispose-t-on ? », *Revue économique*, 49, 335-365.
- LAVILLE F. (2000), « La cognition située : une nouvelle approche de la rationalité limitée », *Revue économique*, 51, 1301-1331.

- LAVOIE M. (1985), « La distinction entre l'incertitude keynésienne et le risque néoclassique », *Économie appliquée*, 37, 493-518.
- LAVOIE M. (1992), *Foundations of Post-Keynesian Economic Analysis*, Edwar Elgar, Aldershot.
- LEWIS D. (1969), *Convention, A Philosophical Study*, Cambridge University Press, Cambridge Mass.
- LEWIS D. (1983), « Languages and Language », repris in *Philosophical Papers*, Oxford University Press, p. 163-188, trad. fr. partielle (1993), « Langages et langage », *Réseaux*, 62, 11-18.
- LUCAS R. (1981), *Studies in Business Cycles Theory*, MIT Press, Cambridge.
- MANGOLTE P.-A. (1998), *Le Concept de « routine organisationnelle » entre cognition et institution*, thèse de doctorat en sciences économiques, université de Paris-XIII.
- MARCH J. G. et SIMON H. A. (1958), *Organisations*, Wiley, New York ; trad. fr., *Les Organisations*, Dunod, Paris, 1999, 2^e éd.
- MARCH J. G. (1978), « Bounded Rationality, Ambiguity, and the Engineering of Choice », *The Bell Journal of Economics*, 9, 587-608.
- MARKOWITZ H. (1952), « The Utility of Wealth », *Journal of Political Economy*, 60, 151-158.
- MINSKY H. P. (1975), *John Maynard Keynes*, Columbia University Press, New York.
- MONGIN P., WALLISER B. (1988), « Infinite Regressions in the optimising Theory of Decision » in MUNIER B., *Risk Decision and Rationality*, 435-457.
- MUNIER B. (1984), « Quelques critiques de la rationalité économique dans l'incertain », *Revue économique*, 35, 65-86.
- NELSON R., WINTER S. (1982), *An Evolutionary Theory of Economic Change*, Harvard University Press, Harvard.
- NEUMANN (VON) J., MORGENTHAU O. (1944), *Theory of Games and Economic Behavior*, Princeton University Press, Princeton, 2^e éd. 1947.
- O'DRISCOLL G., RIZZO M. (1985), *The Economics of Time and Ignorance*, Basil Blackwell, Oxford.
- ORLÉAN A. (1990), « Le rôle des influences interpersonnelles dans la détermination des cours boursiers », *Revue économique*, 41, 839-868.
- ORLÉAN A. (1994), « Vers un modèle général de la coordination économique par les conventions », in A. ORLÉAN (éd.), *L'Analyse économique des conventions*, PUF, coll. « Économie », Paris, 9-40.
- PAULY M. (1968), « The Economics of Moral Hazard », *The American Economic Review*, 58, 531-537.
- PINDYCK R. (1991), « Irreversibility, Uncertainty and Investment », *Journal of Economic Literature*, 29, 1110-1148.
- QUIGGIN J. (1982), « A Theory of Anticipated Utility », *Journal of Economic Behavior and Organization*, 3, 323-343.
- RAMSEY F. (1931), « Truth and Probability », in R. B. Braithwaite (éd.), *The Foundations of Mathematics and other Logical Essays*, Routledge and Kegan, Londres.
- RASMUSEN E. (1989), *Games and Information An Introduction to Game Theory*, Basil Blackwell, Cambridge, Mass.
- REYNAUD B. (2001), « "Suivre des règles" dans les organisations », *Revue d'économie industrielle*, 97, 53-68.
- REYNAUD B. (2002), *Operating Rules in Organizations, Macroeconomic and Microeconomic Analyses*, Palgrave Macmillan, Londres.
- RILEY J. G. (2001), « Siver Signals : Twenty-Five Years of Screening and Signaling », *Journal of Economic Literature*, 39, 432-478.
- RIVAUD-DANSET D. (1998), « Le traitement de l'incertitude en situation. Une lecture de Knight », in *Institutions et conventions. La réflexivité de l'action économique*, Éditions de l'EHESS, Paris, 23-49.
- ROCHETEAU G. (2001), « Théorie de la recherche », in C. JESSUA et al. (dir.), *Dictionnaire des sciences économiques*, PUF, Paris, 782.
- ROSENTHAL R. (1981), « Games of Perfect Information, Predatory Pricing, and the Chain Store Paradox », *Journal of Economic Theory*, 25, 92-100.

- ROSS S. A. (1973), « The Economic Theory of Agency : The Principal's Problem », *The American Economic Review*, 63, 134-139.
- ROTHSCHILD M., STIGLITZ J. (1976), « Equilibrium in Competitive Insurance Market : an Essay on the Economics of Imperfect Information », *Quarterly Journal of Economics*, 80, 629-649.
- RUNDE J. (2000), « Shackle on Probability », in *Economy as an Art of Thought in Memory of Schackle*, E. E. PETER, S. F. FROWEN (dir.), Routledge, Londres.
- SALAI S. (1989), « L'analyse économique des conventions du travail », *Revue économique*, 40, 199-240.
- SALAI S., STORPER M. (1993), *Les Mondes de production. Enquête sur l'identité économique de la France*, Éditions de l'EHESS, Paris.
- SALOP J., SALOP S. (1976), « Self Selection and Turn Over in the Labor Market », *Quarterly Journal of Economics*, XC, 619-628.
- SALOP S., STIGLITZ J. (1982), « The Theory of Sales : a Simple Model of Equilibrium Price Dispersion with Identical Agents », *The American Economic Review*, 72, 1211-1230.
- SAVAGE L. J. (1954), *The Foundations of Statistics*, John Wiley & Sons (New York, Dover Publications, 1972).
- SCHACKLE G.L.S. (1990), *Time, Expectations and Uncertainty in Economics : Selected Essays*, J. Ford (éd.), Edward Elgar, Aldershot.
- SCHELLING T. (1960), *The Strategy of Conflict*, Oxford University Press, Oxford.
- SCHMEIDLER D. (1989), « Subjective Probability and Expected Utility without Additivity », *Econometrica*, 57, 571-587.
- SHAFFER G. (1976), *A Mathematical Theory of Evidence*, Princeton University Press, Princeton.
- SHAPIRO C. (1982), « Consumer Information, Product Quality, and Seller Reputation », *The Bell Journal of Economics*, 13, 20-35.
- SHOEMAKER P. J. H. (1982), « The Expected Utility Model : its Variant, Purposes, Evidence and Limitations », *Journal of Economic Literature*, XX, 529-563.
- SIMON H. A. (1947), *Administrative Behavior*, Macmillan, New York.
- SIMON H. A. (1976), « From Substantive to Procedural Rationality », *Method and Appraisal in Economics*, S. J. LATSIS (éd.), Cambridge University Press, Cambridge, 129-148, reproduit dans H. A. SIMON (1982), *Models of Bounded Rationality*, 2, The MIT Press, Cambridge, Mass.
- SPENCE M. (1973), « Job Market Signaling », *Quarterly Journal of Economics*, 87, 355-379.
- STARMER C. (2000), « Developments in non-Expected Utility Theory : the Hunt for a Descriptive Theory of Choice under Risk », *Journal of Economic Literature*, XXXVIII, 332-382.
- STIGLER G. J. (1961), « The Economics of Information », *Journal of Political Economy*, 69, 213-225.
- STIGLITZ J. (1989), « Imperfect Information in the Product Market », in SCHMALENSEE et WILLIG (dir.), *The Handbook of Industrial Organization*, Elsevier Science Publishers, Inc, New York, 771-847.
- STIGLITZ J. (2002), « Information and the Change in the Paradigm in Economics », *The American Economic Review*, 92, 460-501.
- STIGLITZ J., WEISS A. (1981), « Credit Rationing in Markets with Imperfect Information », *The American Economic Review*, 71, 393-410.
- THÉVENOT L. (1989), « Équilibre et rationalité dans un univers complexe », *Revue économique*, 40, 147-198.
- TORDJMAN H. (1997), « Spéculation, hétérogénéité des agents et apprentissage : un modèle de "marché des changes artificiel" », *Revue économique*, 48, 869-897.
- TREICH N. (2001), « Le principe de précaution est-il économiquement acceptable ? », *Inra sciences sociales*, juillet, 6, 1-4.
- VIVIANI J.-L. (1994), « Incertitude et rationalité », *Revue française d'économie*, IX, 105-146.
- VIVIANI J.-L. (1994), *Ambiguïté*, Cahiers de Recherche CRIEGE, université de Paris-XIII, n° 94-103.

- WALLISER B. (1995), « Les paradoxes de la décision rationnelle », *Annales des Ponts et Chaussées*, 76, 42-51.
- WILLIAMSON O. E. (1983), « Credible Commitments : using Hostages to Support Exchange », *The American Economic Review*, 73 (4), 519-540.
- (1985), *The Economic Institutions of Capitalism*, The Free Press, New York.
- WILLIAMSON O. E. (1990), « The Firm as a Nexus of Treaties : an Introduction », in AOKI M. *et al.* (dir.), *The Firm as a Nexus of Treaties*, Sage Publications, Londres.
- WILLINGER M. (1990), « La rénovation des fondements de l'utilité et du risque », *Revue économique*, 41, 5-48.

Table

Introduction	3
---------------------------	----------

PREMIÈRE PARTIE LES MONDES DE L'INCERTITUDE

I / L'incertitude souveraine, une vision partagée par des courants hétérogènes	6
1. <i>L'incertitude « épistémique » : l'apport d'économistes-philosophes</i>	<i>6</i>
Le risque et l'incertitude selon Knight	7
– Un exemple de complémentarité entre probabilité objective et jugement	8
– Le profit, à mi-chemin entre le hasard et la bonne prévision	9
« Keynes sans l'incertitude, ce serait comme Hamlet sans le prince »	9
– Le poids des anticipations	11
– Anticipations et conventions	12
Hayek : quand l'incertitude et l'ordre spontané sont indissociables	12
2. <i>Les héritiers : Autrichiens, post-keynésiens, évolutionnistes</i>	<i>14</i>
L'incertitude et le rôle actif du temps vont de pair	15
Des acteurs hétérogènes et les difficultés de la coordination	16
Quand les héritiers divergent	18
Conclusion	19

II / L'incertitude, un trouble-fête	
pour les néoclassiques	21
1. <i>Quand le manque d'information sème le désordre</i>	21
Stigler et la quête coûteuse des prix	22
– Quand les vendeurs s'en mêlent	23
Akerlof et l'asymétrie d'information sur la qualité	24
– Le « marché des rossignols »	24
– Le rationnement du crédit par le banquier mal informé	26
– Quand l'ajustement ne peut plus s'effectuer par les prix	26
Arrow et l'aléa moral	27
– Action cachée et information cachée	28
2. <i>Le coût du retour à l'ordre</i>	29
Signal et filtre	30
– Peut-on se fier à un signal de qualité ?	30
– Le contrat à la carte comme filtre	31
La théorie de l'agence	32
– Le modélisateur et l'incertitude	33
3. <i>L'incertitude stratégique</i>	33
– Quand l'incertitude devient souhaitable	34
Conclusion	35

DEUXIÈME PARTIE
LA QUESTION DE LA RATIONALITÉ
EN INCERTITUDE :
MAXIMISER OU AGIR RAISONNABLEMENT ?

III / Quand la décision est parfaite	40
1. <i>Le modèle de l'utilité espérée</i>	40
Les ingrédients et la logique de la décision	41
– L'utilité espérée et le paradoxe de Saint-Petersbourg	41
– Les loteries	42
– La règle de décision	43
Les contraintes du modèle	44
– L'axiomatique de von Neumann et Morgenstern .	45
2. <i>L'extension du modèle aux situations d'incertitude</i>	46
Le statut des probabilités subjectives	47
La logique du modèle de Savage	48
Le modèle en débat	50

3. <i>La peur et le goût du risque</i>	52
Attitude vis-à-vis du risque et utilité espérée	52
– Équivalent certain et prime de risque	53
D'autres mesures	54
Conclusion	55
IV / La mise à l'épreuve du choix rationnel	56
1. <i>Les premiers paradoxes expérimentaux</i>	56
Le paradoxe d'Allais	57
– Le test	58
Le paradoxe d'Ellsberg	59
– Le test	59
– Aversion à l'ambiguïté	60
2. <i>La décision optimale se personifie</i>	61
Les comportements saisis par l'expérience	61
– L'effet de certitude	61
– L'effet miroir	61
– L'effet de construction	62
Le modèle alternatif de Kahneman et Tversky	62
– La fonction de valeur	63
– La fonction de transformation des probabilités	63
– Des précurseurs	65
3. <i>Une nouvelle lignée de modèles</i>	66
L'économiste, le mathématicien et le psychologue	66
Quand les « probabilités » sont non additives	
et la décision risquée	67
Le modèle RDEU à l'œuvre	68
Quand la mesure de l'incertitude fait appel	
aux capacités	70
– Une réponse au paradoxe d'Ellsberg	70
– Les limites de l'ambiguïté	71
4. <i>Simon et la complexité</i>	71
De la simplicité à la complexité	72
Du modèle à la règle	73
Une source majeure d'inspiration	74
Conclusion	74
V / Les institutions : un guide pour l'action	76
1. <i>Williamson et l'incomplétude des contrats</i>	76
Le corpus d'hypothèses	77
Une conception très personnelle des contrats	78
L'institution : la solution pour domestiquer	
l'incertitude	78
– Une incertitude peut en cacher une autre	79

2. <i>L'économie des conventions et les repères partagés ...</i>	80
Guide de lecture des conventions	80
– Le point focal à l'origine de la convention	80
Les débats	81
– D'où viennent les repères partagés ?	82
– Ordre privé ou public ?	83
– Quel apport pour la compréhension des phénomènes concrets ?	83
3. <i>Les routines vues par les évolutionnistes</i>	84
Quels sont les fondements conceptuels des routines ? .	84
– L'évolution des routines vers les heuristiques	85
Des agents homogènes au sein de la firme et hétérogènes sur les marchés	86
Conclusion	88

TROISIÈME PARTIE L'EFFICACITÉ DE CERTAINS COMPORTEMENTS EN INCERTITUDE

VI / Le mimétisme, de l'intérêt individuel aux infortunes du collectif	90
1. <i>Keynes à l'avant-garde</i>	90
Comportement d'entreprise et de spéculation	91
Quand l'imitation tourne mal	91
2. <i>Un modèle de cascade informationnelle</i>	92
Les ingrédients du modèle	92
Les cascades sont-elles inévitables ? souhaitables ?	94
3. <i>Confiance dans son jugement ou dans le marché ?</i>	97
L'inspiration évolutionniste	97
Un modèle post-keynésien	98
4. <i>Le risque systémique</i>	99
Du bulbe de tulipe à la bulle	99
Quand le pessimisme s'installe en cascade	100
Conclusion	101

VII / Risque avéré ou incertitude scientifique, des gestions différentes	102
1. <i>Les trois figures de la prudence</i>	102
2. <i>Le monde de l'assurance</i>	104
Le champ de l'assurance est-il illimité ?	105
– <i>Quid</i> des événements singuliers ?	105
3. <i>L'attentisme et le principe de précaution</i>	106
La valeur d'option d'Arrow-Fisher	107

Les usages du modèle	108
– Le protocole de Kyoto, une démarche exemplaire .	109
Conclusion	110
Conclusion générale	112
Repères bibliographiques	114

L'incertitude dans les théories économiques

Reconnaître l'importance de l'incertitude met profondément en question l'analyse économique. Les premiers économistes à s'emparer de ce thème, Knight et Keynes, ont lancé le débat. Dans certaines situations, l'incertitude peut être traitée par les probabilités : on parle alors de risque. Dans d'autres cas, les probabilités ne sont d'aucun recours. Cette distinction permet de caractériser deux démarches. La première, majoritaire, retient les probabilités pour représenter l'incertitude, la seconde, qui regroupe des courants hétérogènes, postule que l'incertitude n'est pas probabilisable et refuse toute vision mécanique de l'économie.

Cet ouvrage présente les modalités et les conséquences de la prise en compte de l'incertitude dans l'analyse économique.

DANS LA MÊME COLLECTION

Économie de la connaissance • Économie de la firme
• Économie des fusions et acquisitions • Économie
de l'innovation • Économie des logiciels • Économie de
la propriété intellectuelle • Économie de la réglementation
• Économie des réseaux • Introduction à Keynes
• La gouvernance de l'entreprise • Le libéralisme de Hayek
• Macroéconomie : l'investissement • La nouvelle
microéconomie • La théorie de la décision • La théorie
économique néoclassique...

Collection

R E P È R E S

Plus de 300 synthèses à jour, rédigées par des
spécialistes reconnus en économie, gestion, histoire,
sociologie, etc. ➔ Liste à la fin de ce livre.

Pour en savoir plus :
www.collectionreperes.com

Nathalie Moureau

est maître de conférences
en économie à l'université
Montpellier-III et chercheur
associée au MATISSE
(université Paris-I).

Dorothée Rivaud-Danset

est professeur d'économie
à l'université de Reims
et chercheur au Centre
d'économie de l'université
de Paris-Nord (CEPN).

Elles sont spécialisées
dans des domaines
où l'incertitude joue
un rôle central : le marché
de l'art pour la première,
le financement
des entreprises pour
la seconde.



ISBN 2-7071-3851-7



9 782707 138514