

Innovation et concurrence : une analyse dans le cadre d'une industrie en duopole

Babacar NDIAYE

Droit / Economie
Chronologie de l'innovation



*Mes remerciements à tous ceux qui ont contribué à faire du
Groupe de Recherche en Droit, Economie et Gestion
(GREDEG) un environnement privilégié pour la recherche en
général et pour ce travail en particulier.*

Babacar NDIAYE

Innovation et concurrence : une
analyse dans le cadre d'une
industrie en duopole

Editions EDILIVRE APARIS
Collection Universitaire
93200 Saint-Denis – 2011

www.edilivre.com

Edilivre – Éditions APARIS

175, Boulevard Anatole France - Bat A, 1er étage - 93200 Saint-Denis

Tél. : 01 41 62 14 40 – Fax : 01 41 62 14 50 – mail : actualite@edilivre.com

Tous droits de reproduction, d'adaptation et de traduction,
intégrale ou partielle réservés pour tous pays.

ISBN : 978-2-8121-5227-6

Dépôt légal : Avril 2011

INTRODUCTION GÉNÉRALE

La chronologie de l'innovation, appelée « *timing of innovation* » dans la littérature anglo-saxonne, est depuis longtemps au cœur de nombreuses recherches en théorie de l'économie industrielle et en théorie de la décision. Il y a vingt ans, l'article de Reinganum (1989), publié dans le *Handbook of Industrial Organization*, présentait une analyse critique de l'approche de la chronologie de l'innovation en ce qui concerne les modèles d'analyses et les domaines d'application. Selon l'auteur, la chronologie de l'innovation, définie comme l'ensemble des facteurs qui déterminent le processus au cours duquel une innovation est mise en place, joue un rôle central dans la mise en évidence des interactions stratégiques entre les firmes. Dans son interprétation, "*The analysis of the timing of innovation posits a particular innovation (or sequence of innovations) and examines how the expected benefits, the cost of R&D and interactions among competing firms combine to determine the pattern of expenditure across firm and over time, the date of introduction, and the identity of the innovating firm*", (ibid., p. 850). A notre connaissance, cette définition sur la chronologie de l'innovation est la plus complète du point de vue des critères pris en compte depuis que le sujet a été introduit dans la littérature économique. En conséquence, nous retiendrons cette définition dans cet ouvrage.

Toutefois, au cours des dernières décennies, la chronologie de l'innovation en tant que domaine de recherche a connu des bouleversements profonds tant du point de vue des interprétations analytiques que du point de vue des domaines d'application. Ces bouleversements sur le thème de la chronologie de l'innovation sont le point de départ de notre réflexion.

La chronologie de l'innovation concerne une innovation particulière ou plusieurs innovations successives subordonnées les unes les autres. La caractéristique première des travaux dans ce domaine de recherche est la prise en compte du processus temporel au cours duquel l'innovation est examinée depuis la recherche développement (R&D) jusqu'à sa date d'introduction sur le marché. Au cours de ce processus, l'environnement concurrentiel et les interactions stratégiques sont les éléments

fondamentaux à prendre en compte. En effet, deux questions se posent dès lors que l'identité et la position de la firme innovatrice sont des facteurs clés du succès de l'innovation. D'une part, le choix d'entrée sur le marché ("*how to enter*") permet à la firme d'avoir une meilleure connaissance de la dimension et de l'étendue du marché (Schoenecker et Cooper, 1998). Le comportement d'entrée de la firme (*a contrario* le choix de ne pas entrer) influe directement sur le développement (*a contrario* sur le non engagement) de l'innovation. D'autre part, la date d'entrée ("*when to enter*") est aussi une question importante dans la mesure où la position de la firme (first mover ou innovateur versus follower ou imitateur) sur le marché est en général le principal déterminant du niveau de profitabilité de l'innovation (Gaba *et al.*, 2002). Dans ce cas, la stratégie de la firme permet de jouer sur la position concurrentielle.

Notre travail de recherche portera uniquement sur les avantages et les inconvénients d'adopter une position de first mover ou de follower. Les avantages de la firme first mover sont principalement la réalisation de profit identique à celui d'une firme en position de monopole et l'obtention d'un avantage concurrentiel pouvant être durable dans le temps (Lieberman et Montgomery, 1988). Les inconvénients se résument essentiellement à la contrainte de l'irréversibilité des investissements et de l'existence d'incertitude (Bouis *et al.*, 2006). A contrario, la position de firme imitatrice est avantageuse dans la mesure où il y a un bénéfice tiré de l'innovation de la firme first mover en termes de coût de la R&D ou des facteurs de production (Smith *et al.*, 1992). La firme imitatrice profite aussi de l'arrivée d'une information supplémentaire permettant de réduire l'incertitude technologique. L'inconvénient majeur de la position de firme imitatrice est la perte de profit correspondant à la durée pendant laquelle la firme n'est pas encore entrée sur le marché (Odening *et al.*, 2007). De plus, s'il existe des barrières à l'entrée liées à l'efficacité du système de protection des innovations, la firme imitatrice devra faire face à des contraintes réglementaires (Fethke et Birch, 1982). Les attitudes comportementales et stratégiques de la firme sont donc les deux facteurs permettant d'offrir une vue d'ensemble des caractéristiques de l'analyse de la chronologie de l'innovation. Dans cette perspective, la prise en compte d'une industrie en situation de duopole semble la plus adaptée pour représenter les stratégies concurrentielles.

Le travail présenté dans cet ouvrage propose d'analyser de manière systématique les avantages et les inconvénients d'être first mover (resp. follower) en utilisant les modélisations récentes dans le domaine de la chronologie de l'innovation, et en comparant les résultats obtenus à partir de ces différentes modélisations. Cette analyse systématique et

comparative permet, selon nous, d'offrir une approche renouvelée de la chronologie de l'innovation dont la spécificité repose sur les stratégies de la firme en matière d'investissement dans un environnement concurrentiel. Pour atteindre cet objectif, ce travail suit une démarche pédagogique en exposant étape par étape les détails de l'objet de recherche. Dans ce qui suit, trois sections expliquent la structure de notre démarche.

Dans la première section, nous proposons de revenir sur les définitions initiales de la notion de chronologie de l'innovation et sur leurs évolutions. Nous montrons que l'émergence de nouveaux concepts a fortement contribué à l'évolution du domaine de recherche. Dans son acception initiale par exemple, la chronologie de l'innovation était définie essentiellement dans le cadre de la théorie de la décision, et dès lors que l'investissement en R&D était entrepris à une date optimale permettant de récupérer au moins les coûts de la R&D, l'innovation en découlait. Ce mode d'évaluation était basé sur les critères traditionnels d'évaluation comme la Valeur Actualisée Nette ou VAN. Avec la prise en compte dans les développements plus récents : (1) du fait que la décision d'investissement peut être rationnellement avancée ou reportée afin de bénéficier de gains d'informations, par exemple sur les possibilités techniques, et implique que l'investissement en R&D a des chances accrues de se matérialiser en une innovation, (2) de la possibilité de choix stratégiques émanant de concurrents (first mover ou follower) pouvant perturber la décision d'investissement de se transformer en innovation, une approche renouvelée a vu le jour. Cette nouvelle approche sera analysée dans la deuxième section. Nous étudierons alors les différentes étapes qui ont conduit à l'adaptation du contexte initial de la chronologie de l'innovation aux modèles dans lesquels les outils de la théorie des options réelles et de la théorie des jeux sont utilisés. Cette section présente le cœur de notre objet de recherche. Au terme de cette revue, nous aurons une vision d'ensemble qui nous permettra d'être plus opérationnels dans notre démarche de recherche. C'est ainsi que nous pourrons alors proposer dans une troisième section la problématique générale sous forme d'une série de questions de recherche, exposer la méthodologie et indiquer le déroulement exact de notre plan d'ouvrage.

Section 1

La chronologie de l'innovation : un thème développé dans la littérature, mais de manière encore incomplète

Les analyses économiques de la théorie de la décision proposent des définitions multiples sur la chronologie de l'innovation, mais s'accordent

sur le caractère commun de son fondement, à savoir la R&D. Dans ce cadre, une partie de la littérature ne différencie pas les terminologies “*timing of innovation*” et “*timing of investment*” (Barzel, 1968 ; Cripps, 1997). Cette analogie découle du fait que le succès de l’innovation est directement lié à la stratégie d’investissement, laquelle est estimée à partir de critères traditionnels d’évaluation et que l’environnement concurrentiel est souvent négligé (paragraphe 1). Pourtant, selon le cadre de prise de décision et la nature du projet, le décideur ou l’investisseur est prêt à payer afin de pouvoir revenir sur ses choix et maintenir ainsi une décision toujours optimale. Cela justifie la remise en cause des critères traditionnels d’évaluation, puisqu’il y a la prise en compte de la « *valeur d’option* » appelée aussi « *valeur de la flexibilité* » de l’investissement. La prise en compte de la valeur d’option permet de considérer les hypothèses de présence d’incertitude et d’une information croissante, i.e. l’incertitude sur la valeur des variables de décision est décroissante de manière exogène avec le passage du temps, et le décideur est conscient de ce phénomène (paragraphe 2). L’environnement concurrentiel amène alors le décideur à prendre des décisions rationnelles. Chaque firme fait un arbitrage entre le résultat de l’investissement présent et celui de l’investissement futur. Ainsi, l’intégration des stratégies concurrentielles dans le processus de prise de décision devient essentielle pour comprendre les prises de position des firmes (paragraphe 3).

1. Définitions économiques et analogies entre « *timing of innovation* » et « *timing of investment* »

A l’origine, la définition économique de la chronologie de l’innovation était directement liée aux caractéristiques de l’investissement en R&D. L’une des principales règles de prise de décision apprise en finance est celle de la VAN : investir lorsqu’elle est positive et rejeter le projet lorsqu’elle est négative. Dans son article intitulé “*Optimal timing of innovations*”, Barzel (1968, p. 349) considère que la chronologie de l’innovation a pour objectif de déterminer la date à partir de laquelle l’innovation permet de maximiser le profit de la firme et le bien être de la société. Ici, les gains que l’innovation apporte à la collectivité reposent sur les biens de statut collectif pur de l’innovation, i.e. la diffusion de l’innovation doit être rapide et sans coût. Ainsi, le lien triangulaire entre l’utilité tirée de l’innovation par la collectivité, le profit de la firme innovatrice et celui de la firme concurrente (firme imitatrice) soulève de véritables questions sur les sources d’incitation à l’innovation. En définitive, le critère d’investissement retenu dans cette définition est celui de la VAN définie comme la valeur présente des flux de

revenus attendus pendant une certaine période moins le montant de l'investissement initial. De ce point de vue, cette règle de prise de décision est uniquement fondée sur le présent. Tout investissement, dès lors que la VAN est positive, donne lieu automatiquement et immédiatement à une innovation commercialisée sur le marché. Cependant, l'irréversibilité affecte la prise de décision dans un monde incertain parce qu'il est coûteux, voire impossible pour le décideur de revenir à la décision initiale (Bourdieu *et al.*, 1997). Par conséquent, il y a concurrence sur un projet d'investissement unique lorsque celui-ci est entrepris aujourd'hui ou demain (Burger-Helmchen, 2007, 2008). Dans cette perspective, la VAN du projet devient supérieure à son coût du fait de l'augmentation de la valeur d'option liée à l'arrivée d'une information supplémentaire.

En somme les limites de la VAN ont vu le jour avec la prise en compte des différents types d'incertitudes dans la théorie de la décision. Il s'agit de l'incertitude technologique et de l'incertitude de marché. L'incertitude technologique est liée au fait que la firme n'est pas en mesure de déterminer à l'avance les résultats de l'investissement en R&D. Dans ce cas, la probabilité d'atteindre l'objectif que s'est fixé le décideur n'est pas connu. Cet objectif porte en général sur la date d'obtention de l'innovation, et sur le succès de l'investissement en R&D. L'incertitude stratégique, appelée aussi incertitude de marché, est définie à partir de la stratégie concurrentielle adoptée par les firmes rivales, et du succès commercial de l'innovation.

2. Valeur d'option et flexibilité de l'investissement

Dans la mesure où l'investissement est irréversible, il existe une autre limite de la VAN dès lors que la firme dispose de la possibilité d'acquérir une nouvelle information au cours du temps (Bernanke, 1983). Cette opportunité offre au décideur la possibilité de reporter l'investissement du fait de l'existence d'une information croissante (Brennan et Schwartz, 1985). En conséquence, l'incertitude et la flexibilité de l'investissement, i.e. la possibilité pour le décideur de revenir sur sa décision doivent être prises en compte dans les décisions d'investissement (McDonald et Siegel, 1986 ; Ingersol et Ross, 1992).

La valeur d'option est déterminée à partir de la flexibilité du projet d'investissement. Dans sa définition plus générale, la valeur d'option est définie comme le prix que le décideur est prêt à payer afin de pouvoir revenir sur ses choix et maintenir une décision toujours optimale. Avec la prise en compte de l'environnement et de la structure informationnelle, les hypothèses clefs ayant révolutionné les modèles de prise de décision sont la flexibilité totale ou partielle des projets d'investissement et l'existence

d'une information croissante. Ces deux hypothèses sont le point de départ de plusieurs analyses relatives aux stratégies d'investissement notamment dans le cadre de la théorie des options réelles. En effet, lorsque l'incertitude, la flexibilité de l'investissement et le libre choix d'investir sont pris en compte, la date optimale d'introduction de l'innovation n'obéit plus à la règle de la VAN. La firme dispose d'une opportunité d'attente basée sur une information croissante (Dixit et Pindyck, 1994, p. 10). En introduisant le coût d'opportunité de l'investissement, la méthode d'évaluation de la VAN devient erronée pour cause d'ignorance de la possibilité d'investir ultérieurement. L'hypothèse d'information croissante signifie que l'incertitude sur la valeur des variables de décision est décroissante de manière exogène avec le passage du temps, et le décideur est conscient de ce phénomène. Ainsi, la structure informationnelle n'est pas disponible dès la première période de prise de décision, mais se précise étape après étape. Par conséquent, le critère traditionnel d'évaluation de la VAN privilégie ou surestime la valeur des projets à caractère irréversible par rapport aux projets plus réversibles (Henry, 1974). La justification de cette différence d'évaluation est que les investissements irréversibles sont en général supposés avoir une durée plus longue. Ainsi, l'investisseur prend une décision en fonction de l'information présente, i.e. comme si aucune information supplémentaire ne devait lui parvenir par la suite (Favereau, 1989).

Avec les erreurs d'évaluation qui découlent de l'utilisation du critère traditionnel, il s'est révélé nécessaire de construire un nouveau critère considéré comme l'exacte réplique de l'ancien, auquel on ajoute la valeur d'option (Goffin, 2008). Celui-ci permet au décideur de profiter des circonstances favorables et d'éviter les circonstances défavorables dans le but de maintenir ses choix optimaux. A partir de la mise en évidence de ce critère, l'analyse de la chronologie de l'innovation a connu une évolution profonde sur laquelle nous nous appuyons dans ce travail de recherche.

3. Intégration des stratégies concurrentielles

La prise en compte des stratégies concurrentielles s'explique, d'une part, du point de vue des outils utilisés pour analyser la chronologie de l'innovation par l'intermédiaire des critères traditionnels d'évaluation d'un projet d'investissement. Ainsi, l'analyse de la R&D, la prise en compte de l'environnement et les avantages générés par ce dernier dans le futur grâce à l'innovation constituent le point de départ de cette évolution. D'autre part, l'évolution de la chronologie de l'innovation concerne l'intégration des stratégies concurrentielles et la prise en considération d'une nouvelle

approche basée sur des domaines d'application différents. En effet, l'intégration des stratégies concurrentielles permet, en plus de la prise en compte des critères de rentabilité du projet d'investissement, d'introduire des processus de prise de décision et des anticipations croisées dans le but de rationaliser les actions de chaque firme. De plus, la diversité des domaines d'application de la chronologie de l'innovation constitue un fondement permettant de l'analyser sous des aspects différents nous conduisant à des résultats qui ne sont pas forcément identiques (first mover et follower).

Un tel état des lieux fournit les éléments à étudier afin de construire une théorie pertinente de l'analyse de la chronologie de l'innovation. Pour autant, les travaux menés dans ce domaine jusqu'ici ne nécessitent pas de subir des remaniements profonds, mais plutôt une nouvelle approche qui s'inscrit dans la continuité des modèles pionniers. Notre travail de recherche s'inscrit dans une approche remaniée du cadre général de l'analyse de la chronologie de l'innovation. Cette approche vise à inclure et à proposer une nouvelle interprétation des dimensions jusque-là négligées par la littérature qui pourtant caractérisaient déjà les éléments fondamentaux des stratégies d'investissement en situation d'interaction stratégique.

Les types d'interactions stratégiques concevables sont par exemple le cas où les actions des firmes sont séquentielles et le cas où les actions des firmes sont simultanées. Dans le premier cas, le follower choisit une action tout en connaissant le résultat de l'investissement du first mover. En revanche, dans le second cas, le follower agit en même temps que le first mover même si une distinction est faite entre les différentes positions. Dans les deux cas, le type d'information pris en compte reste fondamental. Dans le cadre de la théorie des jeux classiques que nous proposons, l'hypothèse principale retenue est la rationalité forte des agents, et le but est de déterminer le profil de stratégies correspondant à l'équilibre de Nash du jeu. Dans le cadre de l'approche dynamique, les hypothèses retenues (rationalité limitée, apprentissage par imitation et répétition du jeu) permettront de déterminer la stratégie dite *évolutionnairement stable*. En ce sens, nous mettrons en évidence les interactions stratégiques afin de proposer une nouvelle interprétation des stratégies d'investissement.

Section 2

Une approche renouvelée de la chronologie de l'innovation

L'évolution de l'économie industrielle et des outils d'analyse des stratégies d'investissement permet d'expliquer les transformations de l'étude de la chronologie de l'innovation. Dans ce travail, les

développements récents en théorie de la décision feront l'objet d'une attention toute particulière afin d'opérer un retour sur le lien entre la stratégie d'investissement et l'environnement dans lequel sont prises les décisions. Il s'agira de proposer une analyse de la prise de décision dans un environnement certain, puis dans un environnement avec incertitude (Reinganum, 1989). A ce niveau, le renouvellement que nous proposons concerne les modèles stochastiques, et porte plus particulièrement sur la probabilité instantanée de découverte de l'innovation appelée *taux de hasard* ("*hazard rate*") (paragraphe 1). Puis, nous conserverons les hypothèses retenues en situation d'information croissante afin de proposer les développements de la chronologie de l'innovation issus de la théorie des options réelles. A ce niveau, le renouvellement que nous proposons s'inscrit dans le prolongement du modèle de Dixit et Pindyck (1994). Il s'agira de mettre en relation les options *call* et les options *put* (paragraphe 2). Enfin, nous étudierons successivement en quoi il est possible de proposer une approche renouvelée de la chronologie de l'innovation dans le cadre de la théorie des jeux traditionnels ou classiques (paragraphe 3), puis dans le cadre de la théorie des jeux évolutionnistes ou dynamiques (paragraphe 4). Dans le premier cas, l'approche renouvelée porte sur la mise en relation entre le profil de stratégies obtenu à l'équilibre de Nash et la forme prise par la fonction de réaction de chaque firme. La relation qui sera mise en évidence s'appuiera uniquement sur les résultats de l'analyse des deux approches. Dans le second cas, l'approche renouvelée se présente comme un prolongement du précédent. Il s'agira d'utiliser les outils d'analyse de la théorie évolutionniste (dynamique du réplicateur) pour vérifier que l'équilibre de Nash Pareto optimal correspond à l'équilibre dite *évolutionnairement stable*.

Tous ces renouvellements donnent lieu à des avantages et des inconvénients d'être en position de first mover et en position de follower. Dans le cadre de la prise en compte du taux de hasard, il existe un avantage de la firme qui dispose de l'innovation la plus récente. Si celle-ci est la firme first mover, son investissement aboutit à une innovation, et il y a le mécanisme de *winner takes all*. Seulement, l'inconvénient de la position de first mover est l'existence de l'incertitude stratégique et de l'incertitude technologique. A contrario, la position de la firme follower reste moins avantageuse du fait de la prise en compte du processus sans mémoire. Dans le cadre de la théorie des options réelles, les avantages et les inconvénients des différentes positions sont liés à l'existence ou à l'absence de marché potentiel. Lorsque le marché potentiel se trouve avéré, la position du first mover est plus avantageuse que celle du follower. En revanche, lorsque le marché potentiel n'est pas avéré, la position du follower est plus

avantageuse que celle du first mover du fait de l'existence d'une information croissante. Pour chacune de ces positions l'inconvénient majeur est la présence d'incertitude. Dans le cadre de la théorie des jeux traditionnels, comme dans le cadre de la théorie des jeux évolutionnistes, nous mettons en évidence la faiblesse d'être en position de first mover. Cette faiblesse est due au fait que le follower peut profiter de l'innovation obtenue par le first mover en réduisant de manière considérable le montant de l'investissement en R&D nécessaire permettant d'obtenir une innovation. Dans cette perspective, la structure informationnelle considérée permet à la firme follower d'imiter la technologie du first mover.

1. Approche basée sur les modèles de prise de décision

Le choix des modèles de prise de décision est motivé par l'existence des stratégies concurrentielles entre les concurrents. Les modèles que nous considérerons ici sont stochastiques. Le critère essentiel retenu est la prise en compte de l'incertitude technologique matérialisée, d'une part, par une incertitude sur la date d'obtention de l'innovation et, d'autre part, par une incertitude sur le succès de l'innovation. Avec le premier type d'incertitude, la chronologie de l'innovation est décrite par un processus sans mémoire, i.e. il n'existe pas de processus d'apprentissage et l'innovation est découverte sachant qu'elle n'existait pas à une date antérieure (Beath *et al.*, 1989). Avec le second type d'incertitude, l'effort consenti dans la R&D est pris en considération. Ainsi, plus cet effort est important, plus la probabilité de succès de l'innovation est élevée, mais moins que proportionnelle. Dans les deux cas, les facteurs déterminant la chronologie de l'innovation sont identifiés à partir de la date à laquelle l'innovation devient disponible. Cette date dépend de deux facteurs : le montant des dépenses d'investissement et la fonction technique utilisée pour représenter l'investissement (Fetke et Birch, 1982).

Lorsque la date d'innovation est aléatoire, le concept utilisé pour offrir une représentation de la chronologie de l'innovation est d'introduire la fonction de densité conditionnelle appelée le *taux de hasard*. Ce concept est défini comme la probabilité instantanée de découverte d'une innovation. Cette probabilité est conditionnelle car on considère que la firme n'a pas encore obtenu l'innovation. Les deux critères les plus importants permettant de mettre en évidence le taux de hasard sont l'incertitude technologique et les caractéristiques de l'investissement en R&D. L'incertitude technologique reste toujours maintenue. Quant aux caractéristiques de l'investissement, il existe deux effets possibles : l'*effet d'échelle* et l'*effet d'intensité*. La prise en compte du type d'effet est

essentielle dans le cadre des décisions d'investissement en situation d'interdépendance stratégique. L'effet d'échelle signifie que l'investissement en R&D est fixe et réalisé au démarrage du projet (Loury, 1979). Ce terme permet de justifier l'hypothèse selon laquelle la R&D ne peut subir aucune modification au cours du temps, puisque son montant est déterminé en fonction de la taille du projet d'investissement. De ce fait, même si la R&D n'aboutit pas à une innovation, son montant est connu à l'avance. En revanche, l'effet d'intensité de la R&D est défini en terme de flux de dépenses engagé jusqu'à la date de découverte de l'innovation (Lee et Wilde, 1980). Dans ce cas, le montant engagé dans la R&D n'est pas connu à l'avance, puisque l'effet d'intensité de la R&D est continu jusqu'à ce que l'innovation devienne disponible. En procédant à une analyse comparative de ces deux stratégies d'investissement, tout en prenant en compte le nombre de firmes présentes dans l'industrie et les stratégies individuelles d'investissement, nous montrons que les équilibres obtenus sont différents. L'équilibre du niveau d'investissement avec l'effet d'échelle diminue en fonction du nombre de firmes, alors que celui de l'effet d'intensité augmente en fonction du nombre de firmes.

Telle que défini traditionnellement, le taux de hasard est considéré comme étant constant parce qu'il dépend directement du montant investi dans la R&D. Schwartz (1980, p. 246) et Fethke et Birch (1982, p. 273) mettent toutefois en évidence les limites de l'approche du taux de hasard constant. En effet dans l'approche probabiliste, le taux de hasard suit une distribution exponentielle représentant une probabilité d'innovation qui est initialement importante et diminue avec le temps. Même avec la prise en compte d'absence d'hypothèse d'apprentissage, la limite évoquée précédemment laisse une porte ouverte aux modèles probabilistes parce qu'il peut se révéler utile de considérer un processus au cours duquel le taux de hasard augmente avec le temps du fait de l'effort en R&D et de l'environnement concurrentiel. A notre sens, même si la date d'obtention de l'innovation est directement liée à l'effort de R&D, le nombre de firmes présentes dans l'industrie joue un rôle capital. Cette relation se justifie dans la mesure où les interactions stratégiques déterminent la position de chaque firme. Dans ces circonstances, la firme doit savoir *ex ante* la structure du marché dans lequel les interactions sont prises en compte afin de déterminer *ex post* une stratégie de positionnement issue de la chronologie de l'innovation. C'est dans cette perspective que se justifie notre proposition d'analyse comparative des modèles avec un taux de hasard constant et les modèles avec un taux de hasard croissant.

Selon l'existence d'une ou de plusieurs innovations et la présence ou l'absence d'incertitude technologique dans le processus d'innovation, deux

paradigmes sont mis en évidence dans les modèles de la chronologie de l'innovation. Avec le premier paradigme, le modèle est stochastique parce qu'il existe de l'incertitude technologique. Avec ce type de modèle, appelé aussi modèle de course au brevet, l'incertitude technologique prévaut et la qualité de leader n'est que temporaire. Dans ce cas, la firme installée investit un montant faible dans la R&D contrairement à l'entrant potentiel qui voit ses efforts en R&D augmentés au cours du temps. Cette faible incitation à l'investissement de la firme installée est mise en évidence par Arrow (1962) par l'intermédiaire de la notion d'*effet de remplacement*, i.e. la firme installée ou la firme first mover se succède à elle-même et bénéficie uniquement du différentiel de profit issu de l'innovation, alors que la firme follower bénéficie de la totalité des profits générés par l'innovation. Ce postulat permet de justifier l'idée selon laquelle les firmes en situation de concurrence sont plus incitées à innover que les firmes en situation de monopole. En revanche, avec les modèles déterministes, appelés aussi modèles d'enchères, il n'existe pas d'incertitude technologique et la firme installée a tendance à garder sa position de leader au cours du temps (Schumpeter, 1934).

En comparant ces deux paradigmes, Reinganum (1989, p. 904) insiste notamment sur le fait que les modèles de course au brevet semblent plus adaptés pour capturer l'activité de recherche, i.e. une activité pouvant aboutir ou non, et pouvant nécessiter plus de temps et de dépenses que prévus, alors que les modèles d'enchères semblent plus adaptés pour capturer l'activité de développement, i.e. une activité dans laquelle il n'existe plus de l'incertitude technologique. A partir de cette classification, la seconde analyse comparative que nous tentons d'apporter dans le cadre des modèles de prise de décision concerne la relation entre le paradigme pris en compte et le taux de hasard correspondant. Le résultat auquel on aboutit est que le concept de taux de hasard permet, d'une part, de relier les modèles stochastiques et les modèles de course aux brevets et, d'autre part, de décrire la position de la firme innovatrice comme une position de monopole temporaire.

Enfin, trente quatre ans après la définition proposée par Barzel (1968), Baumol (2002) définit, dans le chapitre 12 de son célèbre ouvrage intitulé *The Free-Market Innovation Machine*, la chronologie de l'innovation comme "*the optimal length of time a product or process undergoing continued improvement via research and development (R&D) should continue to be kept in the firm's R&D facilities before it is release to the market*" (Ibid. p. 199). Dans cette approche, l'innovation est analysée comme un processus routinier d'optimisation représenté par un mécanisme de prise de décision continu ("timing decision"). Dans cette perspective,

le choix entre la course pour être le premier à obtenir l'innovation (first mover) et l'attente (follower) est considéré comme au centre de la recherche. En effet, la date d'introduction d'une innovation sur le marché est un choix délibéré (anticipation ou report). Lorsque la date d'introduction de l'innovation est anticipée ("*preemption*"), la firme peut obtenir des profits importants. De même, lorsque la date d'introduction de l'innovation est tardive ("*waiting game*"), la firme peut obtenir des profits significatifs en améliorant le design et/ou la fiabilité de l'innovation. Par rapport à ces deux situations, la stratégie de la firme consiste à trouver, par un processus routinier, la date optimale intermédiaire d'introduction de l'innovation (Riodan, 1992 ; Mariotti, 1992 ; Kapur ; 1995). C'est dans cet esprit que la théorie des options réelles trouve son importance dans la prise en compte des stratégies concurrentielles, dès lors que la présence d'incertitude et l'irréversibilité des investissements restent les deux hypothèses fondamentales.

2. Approche basée sur le développement de la théorie des options réelles

Cette approche est aussi utilisée du fait de la présence des stratégies concurrentielles. La théorie des options réelles est liée à la théorie des options financières mais sa spécificité est de se focaliser sur les actifs réels et non sur les actifs financiers. On emploie le terme « réel » pour la distinguer de la théorie des options financières dont l'actif sous-jacent est un actif financier (par exemple : action, obligation). La théorie renouvelée des options réelles traite des questions relatives aux investissements des actifs réels avec la prise en compte, d'une part, de différents degrés de flexibilité des actifs et, d'autre part, de l'incertitude. Pour engendrer une option réelle, le projet d'investissement doit surtout être risqué, irréversible et flexible (Goffin, 2008, p. 514). Dans ce cadre, une option est définie comme un contrat entre deux parties : l'acheteur, appelé aussi détenteur de l'option, et le vendeur. Le contrat liant les deux parties porte sur un bien (actif sous-jacent) vendu à un prix fixe appelé prix d'exercice. Lorsqu'il s'agit d'une option d'achat, on parle d'une option *call*. Ce concept sera utilisé pour mettre en évidence les stratégies concurrentielles. En revanche, lorsqu'il s'agit d'une option de vente, on parle d'une option *put*.

L'enjeu de notre travail est d'étudier la chronologie de l'innovation issue de la relation entre l'incertitude et l'investissement, puis d'intégrer les éléments stratégiques en termes d'options réelles. D'une part, une partie de la littérature soutient que la présence d'incertitude réduit l'incitation à l'investissement (Cukierman, 1980 ; McDonald et Siegel, 1986 ; Pindyck,

1988 ; Dixit, 1989). La baisse du potentiel d'investissement de la firme est due en grande partie à l'irréversibilité des investissements (Driver, *et al.*, 2006). D'autre part, une autre partie de la littérature considère que l'impact de l'incertitude sur les décisions d'investissement est positif (Sarkar, 2000 ; Wong, 2007). Dans ce cas, l'incertitude issue des stratégies concurrentielles ou de la technologie incite surtout les firmes à investir. De ce point de vue, l'enseignement qui se dégage de cette controverse est que si l'incertitude est un facteur incitatif à la baisse de l'investissement et le déplacement des firmes séquentiel, alors la firme en position de follower garde un avantage sur la firme first mover (Ingersoll et Ross, 1992). Dans cette perspective, la mise en évidence du principe de "*wait and see*" (Smit et Trigeorgis, 2006) permet de remettre en cause la notion standard de calcul de la VAN car celle-ci ne prend en considération que la règle générale d'investissement "*now or never*". De plus, chaque firme choisit une action en fonction de la stratégie de sa rivale. Ainsi, la permutabilité des prises de position des firmes met en lumière le caractère séquentiel des actions des firmes. Dans cette perspective, la firme sera le détenteur de l'option d'investissement (*call*) ou de l'option de désinvestissement (*put*), et sa décision portera sur le choix entre investir ou désinvestir. Dans ce sens, c'est uniquement l'option d'investissement qui détermine la position de la firme, et c'est en tout cas ainsi que nous raisonnerons dans ce travail.

Pour lever le voile sur la problématique des stratégies de positionnement, le concept que nous proposons est fondée sur l'utilisation formelle de la combinaison d'options d'achat et de vente, conformément à Pindyck (1988), Dixit (1989), Abel *et al.*, (1996) et Chaton (2001) afin de compléter la littérature qui généralement ne prend en compte que l'analyse de l'une des deux options (Abel, 1983 ; Caballero, 1991). A cet égard, nous proposons un prolongement du modèle de Nielsen (2002), obtenu à partir de la continuité des modèles de Smets (1991) et de Dixit et Pindyck (1994), pour déterminer les formes prises par les fonctions exercées sur l'option *call* du first mover et du follower. Le cheminement emprunté pour atteindre cet objectif porte sur la densité de fonction exercée sur l'option *call* appelée *intensité de l'option call*. Par conséquent, plus la fonction permettant de matérialiser l'intensité de l'option *call* est croissante, plus la firme est incitée à investir et inversement.

Deux cas distincts seront modélisés dans notre travail en fonction du caractère irréversible de l'investissement : l'irréversibilité partielle et l'irréversibilité totale. Dans le premier cas, l'intensité de l'option *call* est définie sur la base de sa combinaison avec une option de vente, puisque la sortie du marché de la firme peut se faire en récupérant une partie de ses coûts. En revanche, dans le second cas, l'intensité de l'option *call* est

définie sur la base des coûts d'ajustement, puisque la sortie du marché de la firme ne peut se faire que sans récupérer ses coûts. La proposition de modélisation de ces deux cas nous permet d'aboutir au résultat selon lequel, la firme a une incitation plus forte à être en position de first mover quel que soit le niveau d'incertitude. Cela signifie que la forme prise par la fonction de l'intensité de l'option call reste toujours croissante.

3. Approche basée sur la théorie des jeux traditionnels

Cette approche est basée sur la prise en compte des interactions stratégiques entre les firmes. La théorie des jeux est devenue le carrefour des recherches actuelles menées dans le cadre de l'analyse des disciplines d'où se dégagent les interactions stratégiques entre les agents. C'est la situation d'interdépendance entre les agents qui offre un intérêt supplémentaire dans les choix de prise de position. Pour atteindre notre objectif, le travail mené dans cet ouvrage est fondé sur des hypothèses qui nous conduisent à différencier la théorie des jeux traditionnels (classiques) de la théorie des jeux évolutionnistes (dynamiques). L'objectif essentiel de la théorie des jeux traditionnels est de résoudre des problèmes de décision et proposer des solutions optimales qui satisfont les agents en situation d'interdépendance : ces solutions sont les équilibres ou les profils de stratégies à l'équilibre du jeu. Dans cette perspective, nous proposons une approche renouvelée de la chronologie de l'innovation en examinant le raisonnement par lequel les stratégies adoptées déterminent directement les bénéfices ou les profits espérés de chaque agent. Le concept utilisé est ici l'ordre de déplacement séquentiel entre les firmes. Dans le jeu que nous modélisons, les joueurs ou agents sont représentés par des firmes dont les mouvements sont séquentiels, les actions forment une stratégie et les rendements ou paiements sont représentés par les profits issus des différentes stratégies. Le caractère séquentiel des mouvements s'inscrit dans une dynamique de distribution des rôles entre le first mover et le follower (Boyer et Moreaux, 1987). De plus, le jeu modélisé est plus spécifiquement consacré à la théorie des jeux dits *non coopératifs*. Comme le définit Friedman (1991), un jeu est non coopératif si les joueurs ne peuvent pas conclure d'accords entre eux avant de s'engager dans les différentes actions que compose le jeu.

Avec la théorie des jeux traditionnels, chaque joueur est doté d'une rationalité forte appelée aussi hyper rationalité. En d'autres termes, chaque joueur, pris individuellement, agit au mieux de ses intérêts compte tenu des contraintes qu'il subit vis-à-vis des actions du joueur adverse. Le jeu est d'abord analysé sans répétition, et avec cinq étapes successives au cours

desquelles les firmes investissent en R&D puis, selon la structure informationnelle considérée, le follower dispose ou non de la possibilité d'observer le résultat de l'action du first mover avant de choisir une action. Avec la prise en compte de l'incertitude sur le résultat des investissements, les actions des firmes sont identiques à celles d'une population bimorphique, i.e. la présence de deux actions possibles au cours des différentes étapes du jeu (Fernando, 1995). Pour utiliser le concept de déplacement séquentiel entre les firmes, la règle du jeu est la suivante. A la première étape du jeu, la firme first mover a deux actions possibles : « investir » ou « ne pas investir » en R&D. Dans le premier cas, l'investissement aboutit, soit à un succès (innovation), soit à un échec (pas d'innovation) au cours de la deuxième étape du jeu. Lorsqu'il n'y a pas d'investissement en revanche, le first mover garde son ancienne technologie. A la troisième étape du jeu, la firme follower a deux actions : « entrer » ou « ne pas entrer » sur le marché. Si le follower entre sur le marché, deux actions possibles se présentent à nouveau à la quatrième étape du jeu : « investir en R&D » ou « imiter » la technologie de la firme first mover. Lorsque le follower investit, la Nature détermine le résultat de l'investissement à la cinquième étape du jeu, alors qu'avec l'action « imiter », le follower dispose d'un gain certain. Si le follower n'entre pas sur le marché, l'issue du jeu est telle que le paiement du first mover correspond à une situation de monopole et celui du follower est nul. Dans ce cadre, l'objet d'étude principal est la détermination des profils de stratégies correspondant à l'équilibre de Nash qui constitue le point de départ de toutes les analyses qui vont suivre. Dans cette perspective, nous tenterons de mettre en évidence le parallélisme entre les profils de stratégies fondés sur la théorie des jeux traditionnels et la position de la firme fondée sur l'approche des fonctions de réaction. La fonction de réaction de la firme, positive ou négative, se résume à l'analyse des comportements de la firme dans une industrie soumise à la concurrence. Ici le point essentiel consiste à comprendre la réaction d'une firme, i.e. la réponse optimale par rapport à l'action de sa rivale. Le résultat essentiel auquel on aboutit est que la firme follower est incitée à la stratégie d'imitation si le first mover obtient une innovation, et à la stratégie d'investissement en R&D si le first mover n'obtient pas d'innovation.

Par la suite, la répétition sera introduite dans le jeu dans le but de comprendre l'impact de l'horizon considéré sur le comportement stratégique des firmes. Avec la répétition du jeu en horizon temporel fini, notre cheminement suit une récurrence à rebours. En référence à l'exemple classique du dilemme du prisonnier, les analyses théoriques montrent que, aussi longtemps que le jeu est répété en nombre fini de fois (trois fois ou

mille fois), il n'existe qu'un seul équilibre de Nash avec lequel les deux joueurs adoptent la stratégie de la défiance (Guerrien, 1999). De ce fait, nous utilisons le modèle de Wen (2002) pour illustrer les stratégies de positionnement dans le cadre d'une industrie en situation de duopole. Lorsque le jeu est à horizon fini, l'équilibre Pareto optimal auquel on aboutit est la stratégie d'investissement en R&D et la stratégie d'imitation pour la firme first mover et la firme follower respectivement. Cependant, lorsque le jeu est répété en horizon infini, le raisonnement de la récurrence à rebours n'est plus validé. Dans ce cas, les outils d'analyse ne permettent pas de déterminer les profils de stratégies à l'équilibre.

4. Approche basée sur la théorie des jeux évolutionnistes

Dans cette approche, la place qu'occupent les interactions stratégiques est aussi importante. Avec la théorie des jeux évolutionnistes, les principales hypothèses retenues sont : la rationalité limitée des agents, la répétition du jeu et l'apprentissage par imitation. Ici, l'objectif est de résoudre les défaillances (rationalité forte, équilibres multiples) de la théorie des jeux traditionnels. La question qui se pose est de déterminer le processus par lequel les joueurs rationnels doivent suivre avec un raisonnement logique afin de se retrouver dans une position optimale sans pour autant coordonner leurs actions. Avec le processus d'anticipation croisée, il se produit un ajustement des équilibres, puisque chaque joueur agit comme si chaque étape du jeu est la dernière (Kandori *et al.*, 1993).

Parallèlement aux objectifs de la théorie des jeux traditionnels dont la plus importante est la recherche de l'équilibre de Nash, la théorie des jeux évolutionnistes s'oriente vers la recherche d'une stratégie dite *évolutionnairement stable*. Pour déterminer cet équilibre, les concepts utilisés sont la *dynamique du réplicateur* et l'effet stratégique. Considérée comme un « raffinement » de l'équilibre de Nash, la stratégie évolutionnairement stable (SES) est, selon Maynard Smith (1982) la référence centrale de la théorie des jeux évolutionnistes. Bien avant son application dans les sciences sociales, Dawkins (1976) a défini une SES comme une stratégie qui est stable envers elle-même, et supérieure à toutes les autres stratégies du jeu. De ce point de vue, le nombre de joueurs ou d'individus qui adoptent la même stratégie est important : les sujets d'analyse sont deux populations de firmes adoptant séparément des actions identiques au détriment de la situation où les stratégies uniques de deux firmes distinctes sont prises en compte. Cependant, le raisonnement suivi reste le même puisque les choix stratégiques se font par paire d'individus appartenant à chaque population (Maynard Smith et Price,

1973 ; Maynard Smith, 1982). Dans cette perspective, nous tentons de confronter deux logiques : une logique mécanique et une logique stratégique. Traditionnellement, la logique mécanique de la SES est identifiée par rapport à la stratégie adoptée par la proportion la plus importante des individus. A notre sens, cette logique justifie la théorie de l'évolution de Darwin. Au-delà de cette logique, il existe une logique stratégique avec laquelle la SES est directement liée à l'action adoptée entre deux dates successives. De ce fait, on s'affranchît de l'analyse standard de la SES fondée sur le nombre d'individus appartenant à une population. Par conséquent, l'objectif est de recourir à la dynamique déterministe du processus de sélection naturelle appelée la *dynamique du réplicateur*. Introduite par Dawkins (1976) et reprise pour la première fois par Taylor et Junker (1978), la dynamique du réplicateur possède deux caractéristiques : elle s'auto-répique, i.e. se transmet par des gènes de génération en génération et elle détermine le comportement stratégique dans le jeu, i.e. le degré d'adaptabilité de chaque individu. L'intérêt d'utiliser le réplicateur est qu'il agit comme un joueur parce que dans l'approche évolutionniste de la biologie, le degré d'adaptabilité est le paiement espéré tout comme l'utilité espérée de Von Neumann et Morgenstern en théorie des jeux traditionnels. En se basant sur sa représentation formulée sous forme d'équation différentielle (Binmore, 1992), la dynamique du réplicateur permettra d'établir une relation entre les profils de stratégies obtenus à l'équilibre de Nash et la SES. Dans ce cadre, l'équilibre de Nash Pareto optimal correspond à l'équilibre évolutionnaire, c'est-à-dire la position de la firme follower est plus avantageuse que celle du first mover.

Ce tour d'horizon des différentes modélisations et de leurs implications sur facteurs déterminant de la chronologie de l'innovation nous amène à préciser notre question de recherche ainsi que la méthode et les moyens employés afin d'en proposer une analyse en profondeur.

Section 3

Contenu et cheminement du travail

Dans le cadre de cette section, nous reformulons la problématique générale qui sera suivie et déclinons les questions de recherches sous-jacentes (paragraphe 1), détaillons la méthodologie choisie (paragraphe 2), développons les apports principaux (paragraphe 3) et précisons le plan de notre travail (paragraphe 4).

1. Problématique générale et questions de recherches associées

L'objectif de cet ouvrage est d'analyser la chronologie de l'innovation dans une industrie en duopole caractérisée par la prise en compte des mouvements séquentiels des firmes en situation d'interdépendance stratégique. Cet objectif nous conduit à formuler, puis à répondre à une série de questions de recherches sous-jacentes : comment se définissent les motivations de la firme vis-à-vis des positions de first mover et de follower ? Comment la firme détermine sa position sous la contrainte des stratégies de la firme rivale ? Quel rôle joue l'efficacité de la production d'innovation dans la position de la firme ? Selon la modélisation considérée, laquelle des stratégies d'innovation ou d'imitation offre un avantage concurrentiel ? Quel est l'impact de l'incertitude sur les stratégies d'investissement des firmes ? Quel rôle la concurrence joue sur la détermination de la position des firmes ? Quelles sont les incitations de la firme first mover à développer des innovations qui peuvent être instantanément copiées par le follower sans en supporter le coût ? Quelles sont les conséquences du processus d'anticipation croisée sur la rationalité des firmes ? Cette rationalité est-elle toujours vérifiée ? Quels sont les outils permettant de déterminer les profils de stratégies à l'équilibre de Nash, puis à l'équilibre évolutionnaire ? Quelle cheminement faut-il suivre afin d'établir une relation entre les deux équilibres précédents ? Les profils de stratégies obtenus à l'équilibre de Nash correspondent-ils à ceux obtenus avec l'équilibre évolutionnaire ?

A partir des questionnements précédents, le point d'ancrage de notre travail réside dans la détermination de la position de la firme la plus efficace et la justification de la stratégie d'innovation ou d'imitation qui en découle dans le cadre de la théorie de la décision, de la théorie des options réelles, de la théorie des jeux traditionnels et de la théorie des jeux évolutionnistes. Dans la mesure où les résultats sur la position la plus efficace (innovation versus imitation) divergeront selon les cadres d'analyse, nous serons amenés à interpréter ces divergences.

2. Méthodologie de recherche

La méthodologie de recherche est fondée sur la combinaison de différentes modélisations et sur la comparaison des résultats obtenus par celles-ci. Il s'agit de proposer une confrontation des modèles les plus pertinents du point de vue de notre problématique, de faire apparaître les liens qui émergent et de proposer dans la mesure du possible une extension s'inscrivant dans une dynamique de continuité. Notre contribution vise à

mettre en perspective les apports et les limites de la plupart des modèles cités ci-dessus dans le traitement des stratégies d'investissement et de positionnement en situation de concurrence duopolistique. Notre démarche méthodologique se concentre sur la littérature moderne intégrant les aspects temporels du processus d'innovation et les interactions stratégiques (Dasgupta, 1988 ; Reinganum, 1989 ; Dawkins, 1989 ; Beath *et al.*, 1989 ; Binmore, 1992 ; Dixit et Pindyck, 1994 ; Weibull, 1995 ; Nielsen, 2002). Notre travail s'inscrit dans une dynamique éclectique : nous ne cherchons pas effectuer un saut théorique majeur, mais à apporter plutôt notre pierre à l'édifice des différentes modélisations majeures mettant en évidence les stratégies d'investissement. Dans cette démarche, ce travail se situe à l'intersection de plusieurs champs de recherche modélisés : la théorie de la décision, la théorie des options réelles et la théorie des jeux. Les liens qui peuvent être établis entre ces modélisations seront analysés successivement. Ainsi, le dénominateur commun restera la prise en compte des comportements stratégiques des firmes et leurs interactions dans une perspective temporelle.

3. Apports principaux de l'ouvrage

Nous montrons qu'il existe un avantage à être en position de first mover ou en position de follower. Cela nous permettra de discuter et de reconsidérer un certain nombre de résultats considérés comme classiques dans la littérature. A titre d'exemple, on considère bien souvent que la firme qui introduit en premier une nouvelle innovation va réaliser des profits de monopole. Toutefois, nous insistons sur le fait que la durée de la situation de monopole de la firme first mover peut être temporaire ou permanente, ceci selon la réaction de la firme follower. En imitant rapidement l'innovation, le follower affecte la durée au cours de laquelle le first mover reste en position de monopole : l'entrée du follower implique le partage des parts de marché (Lee *et al.*, 2000). Mieux encore, lorsque l'imitation est simultanée, le follower peut obtenir un résultat supérieur à celui du first mover (Baldwin et Chids, 1969 ; Kamien et Schwartz, 1978 ; Gal-or, 1985 ; Teece, 1996 ; Katz et Shapiro, 1987 ; Cannor, 1988). C'est d'ailleurs dans cette perspective que Baumol (2002) explique que le fait d'asseoir sa position concurrentielle sur l'innovation, principal élément de la concurrence, est parfois sans garantie durable. Pour cet auteur, l'hypothèse traditionnelle de la concurrence parfaite concernant la présence de plusieurs firmes sur le marché est simplifiée et réduite à une concurrence duopolistique au cours de laquelle l'objet essentiel est de s'intéresser aux stratégies de positionnement des firmes en termes

d'innovation et en termes d'imitation avec des mouvements séquentiels. Sur ce point, nous franchissons un pas supplémentaire en mettant en évidence les stratégies de la firme la plus efficace selon le type de modélisation considérée.

Enfin, dans l'optique de mettre en évidence les résultats obtenus dans chaque chapitre, nous avons souhaité discuter nos modèles d'analyse et de leurs donner une interprétation économique à partir d'un exemple issu du monde des affaires : le cas de l'industrie aéronautique. Selon le segment de marché considéré, nous tenterons de mettre en évidence l'exemple de la rivalité concurrentielle entre Airbus et Boeing ou entre Embraer et Bombardier. Les discussions et les interprétations qui ressortiront du cas de l'industrie aéronautique, considérée comme un exemple type, permettront de répondre à certaines questions de recherches associées à notre problématique générale dans le but de mieux comprendre les facteurs qui déterminent la chronologie de l'innovation.

4. Plan de l'ouvrage

Cet ouvrage est structuré en quatre chapitres. Le *premier chapitre* analyse la chronologie de l'innovation dans le cadre de la théorie de la décision. Les différents aspects de la description et de la résolution des problèmes de décision que nous introduisons dans ce chapitre ont pris forme depuis la mise en évidence des interactions stratégiques en matière d'innovation dans le domaine de l'économie industrielle. L'apport théorique que nous présentons ici diffère essentiellement de ceux menés habituellement dans ce champ de recherche. Nous avons choisis non pas de traiter la théorie de la décision sous l'angle des méthodes d'analyse et de résolution des problèmes de décision, ce qui nous semble assez classique au vue des outils mathématiques utilisés, mais plutôt de classifier et développer des modèles d'aide à la décision d'investissement en R&D qui paraissent les plus proches des aspects de la chronologie de l'innovation. A l'issue de cette analyse, deux principaux résultats seront illustrés. Avec les modèles déterministes, il existe un avantage de la firme first mover sur la firme follower parce que la position de monopole du first mover est permanente. Avec les modèles stochastiques, la position la plus avantageuse dépend de l'identité de la firme dont l'innovation est la plus récente. Ainsi, la position de monopole du first mover ou du follower est temporaire.

Le *deuxième chapitre* offre une analyse de la chronologie de l'innovation dans le cadre de la théorie des options réelles. Issue de la théorie des options financières, la théorie des options réelles est une

approche permettant d'analyser les options d'investissement (*call*) et de désinvestissement (*put*) sous les hypothèses d'une présence d'incertitude et d'irréversibilité des actifs réels. Dans ce chapitre, l'objectif principal est de proposer une identification formelle de la combinaison des deux options afin de déterminer la position de la firme dont l'incitation à l'exercice de l'option d'investissement est la plus élevée. Cette incitation à l'investissement est appelée *intensité de l'option call*. Celle-ci est définie, d'une part, sur la base de sa combinaison avec une option de vente et, d'autre part, sur la base des coûts d'ajustement. Dans les deux cas, la forme prise par l'intensité de l'option call de la firme first mover et celle de la firme follower sont croissantes. Dans ce chapitre, les avantages des positions sont déterminés en fonction de l'existence ou de l'absence d'un marché potentiel. Si le marché potentiel existe, la position de la firme first mover est la plus avantageuse. En revanche, si le marché potentiel n'existe pas, la position du follower est la plus avantageuse du fait de la prise en compte d'une information croissante.

Le *troisième chapitre* introduit la théorie des jeux traditionnels ou classiques. Ici, les modèles de prise de décision se déroulent en avenir incertain non probabilisable. A ce titre, la théorie des jeux classiques a pour but d'analyser les prises de décision des agents, représentés par des firmes, placés en situation d'interdépendance. Sa principale originalité consiste à prendre en compte l'hypothèse de la forte rationalité de chaque firme, celle-ci étant consciente non seulement de ses propres objectifs, mais aussi de ceux de sa rivale. Le point essentiel consiste à postuler l'interdépendance stratégique entre les firmes. De plus, nous nous intéresserons essentiellement à la théorie des jeux non coopératifs. L'absence d'accords de concertation entre les firmes se justifie d'une façon d'ordre l'égal par une prohibition. Avec ces différentes hypothèses, l'objet d'étude principal de la théorie des jeux classiques est la recherche du profil de stratégies correspondant à l'équilibre de Nash, c'est-à-dire l'équilibre avec lequel aucune firme n'a intérêt à changer unilatéralement de stratégie et ceci afin de s'assurer du maximum des gains possibles. Dans ce chapitre, nous montrons d'abord qu'il existe plusieurs équilibres de Nash avec l'absence de répétition du jeu. Toutefois, lorsque la répétition est introduite, l'équilibre de Nash Pareto optimal correspond à la situation avec laquelle la position de la firme follower devient plus avantageuse que celle du first mover.

Le *quatrième chapitre* introduit la théorie des jeux évolutionnistes. Contrairement à la théorie des jeux classiques qui prend en compte les stratégies individuelles des firmes, la théorie des jeux évolutionnistes prend en considération une grande population de firmes jouant au hasard par

paire comme dans le cadre d'un duopole symétrique. Dans ce chapitre, l'outil d'analyse est de comprendre le déroulement d'un processus au cours duquel les équilibres ne sont plus des données statiques issues d'une méthode de calcul, mais le résultat de l'évolution des stratégies dans le temps. Dans ce cadre, l'objectif sera de mettre en évidence la stratégie évolutionnairement stable, i.e. la stratégie optimale présentement et au cours du temps. Le principal résultat auquel on aboutit relève du fait que l'équilibre de Nash Pareto optimal obtenu avec la répétition du jeu est identique à l'équilibre évolutionnairement stable.

Même si les différentes techniques de modélisation présentées dans cet ouvrage ne couvrent pas, loin de là, l'ensemble de la théorie des jeux non coopératifs, l'esprit dans lequel elles sont étudiées nous paraît bien adapté aux phénomènes d'interactions stratégiques.

La *conclusion générale* sera l'occasion de rappeler les différents résultats auxquels nous sommes parvenus dans cet ouvrage. Nous tenterons alors d'en cerner les points essentiels, puis nous nous interrogerons sur les défauts résiduels de ce travail. Enfin, nous mettrons à jour des questions passées sous silence afin d'identifier des pistes de recherche futures qui pourront faire l'objet de travaux plus approfondis, tant d'un point de vue théorique qu'empirique.

Chapitre 1

Chronologie de l'innovation

dans une industrie en duopole :

une analyse dans le cadre de la théorie de la décision

L'étude de la chronologie de l'innovation à travers la théorie de la décision a donné lieu à différentes modélisations. Au sein de ces modèles, la théorie de la décision tente d'expliquer comment les agents prennent des décisions en présence comme en l'absence d'incertitude. La théorie de la décision permet de déterminer les meilleurs choix, puis de les justifier par un raisonnement logique¹. Ce chapitre propose une analyse critique et une classification des modèles d'aide à la décision liés à la chronologie de l'innovation. De manière plus précise, nous tentons de comprendre, dans le cadre de la concurrence pour l'innovation, les stratégies d'investissement mises en évidence dans les modèles d'aide à la décision.

Dans le cas de la concurrence pour l'innovation, la décision d'investir porte sur la date optimale à partir de laquelle l'innovation devient disponible. Cette date est certaine ou aléatoire. Dans le premier cas le modèle en question est défini comme un modèle *déterministe* alors que dans le second cas on parle de modèle *stochastique*. Dans les modèles déterministes, la relation entre le montant investi dans la recherche et développement (R&D) et la date de réussite de l'innovation est fixe, c'est-à-dire certaine. En d'autres termes, plus l'effort consenti dans la R&D est important, plus la date de réussite de l'innovation est proche. Cette relation fonctionnelle se traduit donc par une fonction décroissante entre les deux variables. En revanche, dans les modèles stochastiques, la relation entre le montant investi dans la R&D et la date de réussite de l'innovation est aléatoire ou probabiliste. Dans ce cas, la date de réussite est une variable aléatoire dont la

¹ Traditionnellement, la théorie de la décision se fonde sur une hypothèse de rationalité forte des agents.

distribution dépend de l'effort consenti. Ainsi, la firme dont l'investissement est plus important a une espérance mathématique de date de découverte plus élevée. Dans les deux types de modèles, on considère que la première firme qui obtient l'innovation est déclarée vainqueur.

La première question à laquelle nous tentons de répondre dans ce chapitre est de savoir si la firme qui acquiert l'innovation (first mover) renforce sa position dans le temps par rapport à l'entrant potentiel (follower) ou si celui-ci peut instaurer une stratégie de concurrence frontale avec la firme innovatrice. Nous proposons d'examiner la concurrence pour l'innovation appliquée, tout d'abord, aux modèles déterministes et, ensuite, aux modèles stochastiques. Pour ce qui est des modèles déterministes, nous nous concentrerons sur la date optimale d'introduction de l'innovation dans le cas de firmes symétriques, puis dans le cas de firmes asymétriques. Si les firmes sont symétriques, chacune est récompensée en fonction de l'effort en R&D. A l'opposé, si les firmes sont asymétriques, alors leurs efficacités à produire et l'incitation à l'innovation sont soumises à *l'effet de remplacement* (Arrow, 1962) : la firme installée (first mover) est moins incitée à innover qu'un entrant potentiel (follower). Cette faible incitation à l'innovation est due au fait que la position de firme leader pour la firme first mover est permanente. Cette firme bénéficie uniquement du différentiel de profit issu de l'innovation, alors que la firme follower bénéficie de la totalité des profits générés par l'innovation. Dans ce contexte, une seconde question se pose. Quelle est la firme la plus incitée à innover : celle dont les coûts sont plus faibles ou celle dont les coûts sont plus élevés ? Nous tenterons de répondre à cette question en nous basant sur « *le prix de l'innovation* », c'est à dire la valeur correspondant à ce que la firme est prête à payer pour obtenir l'innovation (Reinganum, 1989). Nous constaterons alors que la firme follower peut, dans certaines situations, profiter de l'innovation de la firme first mover (Gal-or, 1985 ; Teece, 1996 ; Smith *et al.*, 1992). Cela se justifie par le fait que la firme en position de follower peut imiter la technologie du first mover.

Dans ce chapitre, notre effort portera sur une extension des modèles stochastiques disponibles dans la littérature. En plus des critères traditionnels retenus (présence d'incertitude technologique, relation entre l'effort en R&D et la date d'innovation), nous introduisons une incertitude de marché et un processus de rattrapage au bénéfice de la firme qui dispose de l'innovation la plus récente. Avec ce nouveau cadre d'analyse, nous développerons une analyse comparative des modèles théoriques.

Nous proposerons une typologie classifiée des modèles stochastiques en fonction d'un critère : le *taux de hasard* (« *hazard rate* »). Le taux de hasard est défini comme la probabilité instantanée de découverte d'une

innovation. Les deux critères les plus importants permettant de mettre en évidence le taux de hasard sont l'incertitude technologique et les caractéristiques de l'investissement en R&D. L'incertitude technologique se compose de deux éléments : l'incertitude sur la date d'obtention de l'innovation et l'incertitude sur le succès de l'innovation. Quel que soit le type d'incertitude pris en compte, la date et le succès de l'innovation dépendent de l'effort que la firme consacre à la R&D. Par conséquent, plus cet effort est élevé, plus l'espérance mathématique de la date de découverte est élevée. Ainsi, la probabilité de succès de l'innovation augmente mais moins que proportionnellement par rapport à l'effort en R&D. D'autre part, lorsque l'on s'intéresse aux caractéristiques de l'investissement, il existe deux effets possibles : l'*effet d'échelle* et l'*effet d'intensité*. L'effet d'échelle signifie que l'investissement en R&D est fixe et réalisé au démarrage du projet (Loury, 1979). On constate que, dans ce cas, si la structure industrielle est suffisamment compétitive, il existe une corrélation négative (appelée effet Loury) entre l'effort individuel en R&D de chaque firme et le nombre total de firmes présentes dans l'industrie. En revanche, l'effet d'intensité de la R&D est défini en termes de flux de dépenses engagé jusqu'à la date de découverte de l'innovation (Lee et Wilde, 1980). Dans ce cas, il existe une corrélation positive (effet Lee et Wilde) entre l'effort individuel en R&D de chaque firme et le nombre de firmes présentes dans l'industrie. Ces deux effets sont traditionnellement analysés dans des modèles distincts. Toutefois, Pérez-Castrillo et Verdier (1991) ont proposé un modèle intégrant explicitement les deux types d'effets dans la technologie de la R&D. Avec ce nouveau cadre d'analyse, l'effet de type Loury l'emporte sur l'effet de type Lee et Wilde. A partir de ce résultat, nous tentons d'établir une relation entre la forme prise par le taux de hasard et le type d'effet correspondant.

A partir de ces notions, il sera possible de faire la distinction entre les modèles dont le taux de hasard est constant et les modèles dont le taux de hasard est croissant. Si le taux de hasard est constant, la probabilité de découverte d'une innovation dépend directement du montant investi dans la R&D, et la date optimale d'introduction de l'innovation sur le marché est incertaine. Ici, le taux de hasard est représenté par une distribution exponentielle identique à l'expression prise par le paramètre d'une loi de poisson. Dans ce cas, le modèle de Beath *et al.* (1989) est la référence. De plus, la prise en compte des éléments stratégiques permettra de considérer que le processus de rattrapage est au bénéfice de la firme qui dispose de l'innovation la plus récente. Cela se justifie par le fait que cette firme a un effort en R&D plus important puisque son espérance de date de découverte correspond à l'inverse du taux de hasard.

Si le taux de hasard est croissant, la firme innovatrice se procure la totalité des revenus de l'innovation et la date optimale d'introduction de l'innovation reste toujours incertaine. Dans ce cas, c'est le flux de dépenses qui détermine la forme prise par le taux de hasard. Ici, le modèle de Fethke et Birch (1982) est le plus avancé car il permet de prendre en compte la dynamique stratégique des firmes et le jeu de la course aux brevets. En effet, la mise en évidence de l'incertitude technologique qui caractérise les modèles théoriques avec un taux de hasard croissant permet non seulement de restreindre notre champ d'analyse à la situation où les revenus escomptés de l'innovation sont limités, mais aussi de prendre en considération le seul cas où le rapport entre le profit total escompté de l'industrie et le profit de la firme innovatrice est égal à 1². La principale motivation vis-à-vis de cette restriction est que les modèles théoriques avec un taux de hasard croissant prennent toujours en compte l'hypothèse de « *winner takes all profit* » (Dasgupta et Stiglitz, 1980) au détriment de quelques exceptions avec lesquelles cette hypothèse n'est pas respectée (Stewart, 1983). Dans les modèles sans l'hypothèse de « *winner takes all profit* », la question de l'impact de la position de la firme sur la rentabilité reste ouverte parce que le rapport entre le profit total escompté et le profit de la firme innovatrice peut être inférieur à 1. A partir de cette question ouverte, ce chapitre propose une nouvelle interprétation par l'intermédiaire de l'identité de la firme (innovatrice ou imitatrice).

Enfin, nous analyserons le modèle de Baumol (2002) pour identifier le paradoxe. Ce paradoxe nous enseigne que la fonction de profit de la firme imitatrice dépend de celle de la firme innovatrice. Toutefois, aucune conclusion concrète ne peut être apportée sur l'identité de la firme dont la stratégie est la plus avantageuse. Ce modèle montre que pour une date donnée d'introduction d'une innovation de son concurrent, une firme peut ne pas avoir intérêt à introduire une innovation à la même date, mais à une date antérieure ou postérieure³. Dans ce cas, il est impossible de déterminer la firme leader sur le marché.

Le chapitre est organisé de la façon suivante. Nous présenterons la revue de littérature sur les modèles théoriques de la chronologie de l'innovation (section 1). Ici, nous proposerons de classer les modèles

² Cette considération se justifie par le fait que la firme qui acquiert en premier l'innovation dispose en conséquence de la totalité des profits issus de l'innovation.

³ Ce résultat contraste avec le comportement des firmes dans les modèles de localisation d'Hotelling. Dans ces modèles, le principe de rationalité conduit les firmes à adopter les mêmes stratégies. Qu'il s'agit d'une stratégie de localisation géographique ou de l'introduction d'une innovation, les firmes doivent réduire tous les critères de différenciation.

déterministes en fonction de deux critères : les caractéristiques de la R&D (paragraphe 1) et la relation de symétrie ou d'asymétrie entre les firmes (paragraphe 2). Les modèles stochastiques seront aussi analysés à partir de deux autres critères qui s'ajoutent aux précédents : l'incertitude (paragraphe 3) et le taux de hasard (paragraphe 4). Une relation entre la typologie des innovations et la forme prise par la concurrence sera déduite de l'analyse des modèles (paragraphe 5). Ensuite, la relation fonctionnelle entre l'effort d'investissement et la date d'innovation sera étudiée (section 2). Nous proposerons alors une présentation simple d'un modèle déterministe (paragraphe 1) et d'un modèle stochastique (paragraphe 2). Les résultats de cette présentation seront comparés afin d'analyser le paradoxe du modèle de Baumol (2002) (paragraphe 3). Enfin, la conclusion présentera les principaux résultats obtenus et de nouvelles perspectives de recherche.

Section 1

Analyse des modèles théoriques : revue de littérature

L'évolution des modèles théoriques sur la chronologie de l'innovation est traditionnellement analysée par l'intermédiaire des modèles déterministes et des modèles stochastiques. Nous proposons dans cette section d'analyser ces deux types de modèle en fonction de plusieurs critères.

Les modèles déterministes sont définis à partir de deux critères : les caractéristiques de la R&D et la relation de symétrie ou d'asymétrie entre les firmes. Le premier critère concerne la forme prise par la technologie de l'investissement. Cette technologie prend deux formes possibles : fixe et variable. Dans le premier cas, il existe un effet d'échelle permettant de maintenir le niveau d'investissement réalisé au démarrage du projet. Dans le second cas, il existe un effet d'intensité permettant de mesurer le flux de dépenses engagé jusqu'à la date de découverte de l'innovation.

Le second critère concerne les caractéristiques du modèle étudié. Les firmes sont symétriques si chacune d'entre elles dispose, d'une part, de mêmes capacités d'investissement et, d'autre part, de la même probabilité d'obtenir une innovation. Ici, chaque firme est potentiellement innovatrice et, en général, il existe une seule innovation produite avec des procédés différents (Kamien et Schwartz, 1976 ; Dasgupta et Stiglitz, 1980). Les firmes sont asymétriques si elles sont différenciées par leur position sur le marché (first mover ou follower) ou par leurs capacités d'investissement (monopole ou entrant potentiel). Ici, la firme qui dispose de capacité

d'investissement plus importante a un avantage sur les firmes rivales et, en général les modèles tendent à considérer l'existence de plusieurs innovations (Beath *et al.* 1989 ; Reinganum, 1989).

Les modèles déterministes sont analysés sans présence d'incertitude. Avec ce type de modèle, il existe une relation décroissante entre le montant investi dans la R&D et la date à laquelle l'innovation devient disponible : plus l'effort en R&D est important, plus la date à laquelle l'innovation devient disponible est proche. En d'autres termes, une augmentation de l'effort en R&D permet d'obtenir l'innovation plus tôt que prévue. Il ressort donc que les modèles déterministes ont en définitive les caractéristiques des modèles d'enchères : la firme qui offre le meilleur prix gagne l'enchère.

Les modèles stochastiques nécessitent la définition de deux autres critères qui s'ajoutent aux précédents : le taux de hasard et l'incertitude. Le taux de hasard est la probabilité instantanée de découverte de l'innovation. Cette probabilité est conditionnelle et le taux de hasard dépend directement de l'effort en R&D que la firme consent. Le taux de hasard prend deux formes possibles : constant et croissant. Le taux de hasard est constant lorsque la R&D est mesurée à partir de l'effet d'échelle et un système de protection des innovations inefficace. Dans ce cas, les entrants potentiels peuvent imiter la technologie de la firme innovatrice. Le taux de hasard est croissant lorsque la R&D est mesurée à partir de l'effet d'intensité et le système de protection des innovations efficace. La forme prise par le taux de hasard dépend aussi du type d'incertitude pris en compte. En effet, l'incertitude technologique est distinguée de l'incertitude de marché parce qu'elle est matérialisée, d'une part, par l'incertitude sur la date à laquelle l'innovation devient disponible et, d'autre part, par l'incertitude sur le succès de l'innovation. En général, les espérances mathématiques de date de découverte sont utilisées pour déterminer la date à partir de laquelle l'innovation devient disponible. Par conséquent, plus l'effort en R&D est important, plus l'espérance de date de découverte est proche. L'incertitude de marché est liée à la stratégie concurrentielle adoptée par les firmes. Ce type d'incertitude est introduit pour mettre en évidence les éléments stratégiques entre les firmes. Il ressort de l'analyse que les modèles stochastiques ont les caractéristiques des modèles d'enchères.

Les différents critères d'identification des modèles déterministes et des modèles stochastiques permettent de présenter une évolution de la littérature à travers la Tableau 1. Ce tableau donne une vue d'ensemble des caractéristiques de l'évolution des modèles théoriques les plus pertinents du point de vue de notre approche.

Tableau 1. Evolution des modèles théoriques sur la chronologie de l'innovation

Typologie des modèles	Critères de sélection	Evolution des modèles
Modèles déterministes	- Caractéristiques de la R&D - Fixe - Variable - Relation entre les firmes - Symétrique - Asymétrique	Scherer (1967), Barzel (1968), Dasgupta et Stiglitz, (1980), Reinganum (1989), Pérez-Castrillo et Verdier (1991), Boone (2001)
Modèles stochastiques	- Caractéristiques de la R&D - Fixe - Variable - Relation entre les firmes - Symétrique - Asymétrique - Taux de hasard - Constant - Croissant - Incertitude - Marché - Technologique	Modèles avec taux de hasard constant : Loury (1979), Lee et Wilde (1980), Katz et Shapiro (1985), Beath <i>et al.</i>, (1989), Reinganum (1989), Baumol (2002) Modèles avec taux de hasard croissant : Lucas (1971), Kamien et Schwartz (1976), Fethke et Birch (1982), Reinganum (1985)

Les modèles déterministes ont évolué avec les travaux de Scherer (1967), Barzel (1968), Dasgupta et Stiglitz (1980) et Boone (2001). Avec ces modèles, il n'existe pas d'incertitude. La distinction entre les modèles avec un taux de hasard constant et les modèles avec un taux de hasard croissant est résumée dans la grille des modèles stochastiques. Ici, l'incertitude occupe une place importante. Lorsque l'incertitude technologique est supérieure à l'incertitude de marché, le cadre d'analyse s'oriente vers les modèles de Loury (1979), Lee et Wilde (1980), Beath *et al.*, (1989) et Katz et Shapiro (1985). Dans ce cas, les modèles sont

considérés comme étant à taux de hasard constant. En revanche, lorsque l'incertitude de marché est supérieure à l'incertitude technologique, le cadre d'analyse s'oriente vers les modèles avec un taux de hasard croissant (Lucas, 1971 ; Kamien et Schwartz, 1976 ; Fethke et Birch, 1982).

La Tableau 1 montre que la plupart des modèles théoriques à l'exception de ceux de Reinganum (1985) et de Baumol (2002) prennent en compte, d'une part, l'existence d'une seule innovation et, d'autre part, la présence d'un grand nombre de firmes présentes sur le marché. Dans ces conditions, l'analyse des modèles déterministes et des modèles stochastiques considère la forme prise par la technologie de la R&D comme un critère fondamental permettant de définir les sources d'incitations à l'innovation. A partir de la forme prise par la technologie de la R&D, il existe en général une seule innovation et un grand nombre de firmes. Dans ce cas, on s'intéressera successivement à la firme qui acquiert l'innovation et donc devient leader, puis à son concurrent potentiel (Kamien et Schwartz, 1976). En revanche, lorsqu'il existe plusieurs innovations quel que soit le nombre de firmes, on s'intéresse à la première firme innovatrice en considérant les autres firmes comme follower (Reinganum, 1985). Dans ce cas, on parle d'innovations cumulatives. En recoupant ces deux situations, l'élément que nous retenons dans cette analyse concerne l'identité des deux premières firmes actives qui entrent sur le marché : la première étant le first mover et la seconde le follower.

Dans cette section, l'objectif sera d'analyser d'abord les modèles déterministes par l'intermédiaire de deux critères : les caractéristiques de la R&D (paragraphe 1) et la relation entre les firmes (paragraphe 2). Ensuite, les modèles stochastiques seront analysés à partir de deux autres critères qui s'ajoutent aux précédents : l'incertitude (paragraphe 3) et le taux de hasard (paragraphe 4).

1. Caractéristiques de la R&D

Dans l'analyse des stratégies d'investissement, la R&D est considérée comme le coût de l'investissement. Ce coût dépend de la technologie de la R&D. A partir de cette relation, il existe la distinction entre les coûts fixes et les coûts variables associés à la technologie de la R&D appelés respectivement effet d'échelle et effet d'intensité. L'effet d'échelle signifie que le montant de l'investissement est déterminé avant le lancement du projet, alors que l'effet d'intensité de l'investissement est mesuré en termes de flux de dépenses engagé jusqu'à la date à laquelle l'innovation devient disponible.

A partir de cette distinction, il existe deux thèses en confrontation : celle de Loury et celle de Lee et Wilde. Avec la première thèse, le nombre de firmes présentes dans l'industrie est une fonction décroissante de l'effort individuel en R&D de chaque firme. En revanche, avec la seconde thèse, le nombre de firmes présentes dans l'industrie est une fonction croissante de l'effort individuel en R&D de chaque firme. En d'autres termes, une augmentation du nombre de firmes contribue à la hausse de l'équilibre initial de l'industrie. La spécificité de ces deux thèses, nous conduit donc à des résultats totalement opposés. Les deux effets sont traditionnellement analysés de manière distincte dans les modèles de prise de décision. Cependant, lorsqu'ils sont pris en compte dans un modèle unique, le résultat obtenu avec la première thèse se confirme. Cela signifie que le modèle de Lee et Wilde porte les mêmes caractéristiques que les modèles intégrant les coûts fixes et les coûts variables.

Nous analyserons de manière distincte les spécificités des modèles en termes de coût fixe (paragraphe 1.1) et en termes de coûts variables (paragraphe 1.2), puis ces spécificités seront intégrées dans un modèle unique qui fera aussi l'objet d'une analyse approfondie (paragraphe 1.3).

1.1. Spécificité des modèles en termes de coût fixe

Ici, le coût de la R&D est considéré comme fixe et on essaie de comprendre le processus par lequel une firme peut prétendre acquérir une invention ou une innovation en limitant ses coûts. Dans ce cas, le modèle étudié peut être déterministe ou stochastique. Loury (1979) considère que si le coût de la R&D est « contractuel », alors il est une donnée fixe. En d'autres termes, le coût de la R&D est déterminé avant la mise sur le marché d'un produit ou d'un procédé innovant. Par conséquent, lorsque la structure industrielle est suffisamment compétitive, il existe une corrélation négative entre l'effort individuel en R&D et le nombre de firmes présentes dans l'industrie. Cette relation négative est plus connue sous le nom « effet de type Loury ». De la sorte, une augmentation du nombre de firmes a pour conséquence une baisse de la moyenne des dépenses en R&D de chaque firme. L'augmentation du nombre de firmes présentes dans l'industrie occasionne le déplacement à la baisse du niveau d'équilibre de l'investissement de l'industrie. Comme le modèle de Loury prend comme support l'importance des coûts fixes dans la technologie de la R&D, l'équilibre de l'industrie est donc affecté. « L'effet de type Loury » s'oppose à « l'effet de type Lee et Wilde (1980) ».

1.2. Spécificité des modèles en termes de coût variable

Dans cette situation, le coût de la R&D est variable et mesuré en termes de flux de dépenses. L'un des modèles protagonistes est celui de Lee et Wilde (1980). Dans ce modèle, le coût de la R&D est mesuré par le flux de dépenses engagé par la firme jusqu'à la mise sur le marché de l'innovation (avec succès). La notion de flux de dépenses est importante dans l'identification de la spécificité des modèles en termes de coût variable. Comme la formalisation de ce modèle a des caractéristiques proches de celles des modèles d'enchères, alors il existe une corrélation positive entre l'effort individuel en R&D et le nombre de firmes présentes dans l'industrie. Cette relation est appelée « l'effet de type Lee et Wilde ».

En comparant les deux effets annoncés ci dessus, les résultats sur le niveau d'investissement s'opposent. Avec « l'effet de type Lee et Wilde », l'effort individuel en R&D de chaque firme détermine directement le résultat de l'investissement. Par conséquent, une augmentation du nombre de firmes présentes dans l'industrie a pour effet un déplacement vers le haut de l'équilibre initial. Une autre conséquence de l'augmentation du nombre de firmes sur le marché est qu'elle implique une introduction de l'innovation à une date antérieure à celle prévue initialement. De même, « l'effet de type Loury » conduit à un résultat qui dépend de l'effort en R&D. Dans ce cas, il existe de l'incertitude de marché, mais la firme peut établir une relation entre le résultat attendu de son investissement et l'effort consenti en R&D.

Ces deux effets sont en général analysés de manière distincte dans les modèles de prise de décision, toutes choses égales par ailleurs. Supposons maintenant qu'il soit possible de prendre en considération l'effet d'échelle et l'effet d'intensité dans un modèle unique. De ce fait, quel sera l'impact des deux effets sur l'équilibre de l'industrie ? S'agira-t-il d'un effet global exprimé comme la somme des deux effets ou l'un des effets peut-il l'emporter sur l'autre ? La réponse à ces questions est proposée dans le paragraphe suivant.

1.3. Spécificité des modèles intégrant les coûts fixes et les coûts variables

Contrairement aux modèles de Loury (1979) et de Lee et Wilde (1980), Pérez-Castrillo et Verdier (1991) ont proposé un modèle intégrant explicitement les coûts fixes et les coûts variables dans la technologie de la R&D. En effet, ce modèle s'articule autour de la course aux brevets dont la caractéristique essentielle est de déterminer l'aptitude de la firme à capturer les activités de recherche. Dans cette analyse, la prise en compte des coûts fixes et des coûts variables dans la technologie de la R&D a pour

objectif de résoudre les ambiguïtés sur la relation entre le type de concurrence et la forme prise par la technologie de la R&D. La résolution de cette relation permet aux auteurs de formuler la proposition selon laquelle : si la structure industrielle reste suffisamment compétitive, il existe une corrélation négative entre l'effort individuel en R&D de chaque firme et le nombre de firmes présentes dans l'industrie.

Toutefois, l'existence d'incertitude dans le modèle proposé peut aboutir à des résultats différents ne garantissant pas le succès de l'investissement. En d'autres termes, l'activité de recherche peut, non seulement ne pas aboutir aux résultats attendus, mais elle peut aussi nécessiter plus de dépenses que celles prévues initialement. Le modèle de Pérez-Castrillo et Verdier (1991) est considéré comme une généralisation des conclusions obtenues dans les modèles de Loury (1979) et de Lee et Wilde (1980). Par conséquent, une analyse comparative entre ces trois modèles montre que les conclusions obtenues par Pérez-Castrillo et Verdier (1991) sont plus proches des résultats de Loury (1979) que de ceux obtenus par Lee et Wilde (1980). De ce fait, le résultat obtenu n'est pas la conjonction des deux effets, mais plutôt une domination de l'effet de type Loury sur l'effet de type Lee et Wilde.

Dans le cadre des modèles déterministes et des modèles stochastiques, la prise en compte des coûts fixes et des coûts variables permet d'avoir une vision sur la relation existant entre le montant de la R&D investie et la date d'introduction de l'innovation sur le marché. Même si le type de modèle considéré donne une orientation sur les résultats attendus (date d'innovation), la relation qui existe entre les firmes mérite d'être mise en évidence.

2. Relation entre les firmes

Lorsque les stratégies concurrentielles sont prises en compte, il est important de connaître la relation qui lie les firmes. Cette relation est symétrique (paragraphe 1.2.1) ou asymétrique (paragraphe 1.2.2). Dans le premier cas, les firmes sont homogènes, c'est-à-dire elles possèdent les mêmes caractéristiques et disposent de capacités d'investissement identiques, alors que dans le second cas les firmes sont hétérogènes et investissent en fonction de leur capacité.

2.1. Relation symétrique entre les firmes

Des firmes sont symétriques dans une course à l'innovation si la probabilité d'obtention de l'innovation est la même pour toutes les firmes.

Par exemple en situation de duopole, cette probabilité est égale à $\frac{1}{2}$. Dans les modèles de référence, la relation de symétrie est prise en compte avec le nombre d'innovations présentes dans l'industrie. A ce propos, les analyses de Kamien et Schwartz (1976), et de Fethke et Birch (1982) prennent en compte la chronologie de l'innovation au cours de laquelle des firmes symétriques produisent une seule innovation. Dans ce cas, chacune des firmes est potentiellement innovatrice. Dans ces conditions, la menace concurrentielle détermine la position de chaque firme sur le marché.

Dans les modèles mettant en évidence la relation de symétrie entre les firmes, il existe une particularité. Fethke et Birch (1982) mettent en évidence cette particularité en prenant en compte les facteurs qui déterminent la date d'introduction de l'innovation sur le marché. D'une part, il y a la pression concurrentielle dont l'impact est mesuré grâce au processus concurrentiel à travers lequel, le critère essentiel retenu est le nombre de firmes présentes dans l'industrie. D'autre part, il existe une incertitude sur la stratégie adoptée par les firmes rivales et sur le coût de développement de l'innovation. De ce fait, les firmes sont dans l'incapacité de déterminer avec certitude la date exacte du succès de l'innovation.

En reprenant les sources d'incitation à l'innovation, Boone (2000, 2001) montre que dans le cas de firmes symétriques, l'incitation à l'innovation par la menace concurrentielle est toujours supérieure à l'incitation à l'innovation par le profit. Par conséquent, lorsque la menace concurrentielle est mesurée en termes d'intensité, il existe une forte incitation à l'innovation.

Cette approche s'oppose à celle dans laquelle les firmes sont asymétriques.

2.2. Relation asymétrique entre les firmes

Lorsque la relation entre les firmes est asymétrique, les firmes sont différenciées par leurs positions (first mover ou follower) ou par leurs capacités d'investissement (monopole ou entrant potentiel). En plus de ces critères traditionnels, le nombre d'innovations est aussi pris en compte. Si le modèle prend en compte des innovations cumulatives, la firme en position avantageuse (first mover) a tendance à introduire régulièrement de nouvelles innovations afin de maintenir sa position. De ce point de vue, l'environnement concurrentiel est caractérisé par une économie dans laquelle les firmes mettent en place des innovations successives qui leurs permettent, en cas de succès, d'être en position de monopole temporaire jusqu'à la prochaine innovation (Reinganum, 1985). Dans ces conditions, il

existe une asymétrie entre les firmes puisque le monopole actuel de l'innovation est différencié des entrants potentiels. Par conséquent, quelle que soit la position de la firme – installée ou entrant potentiel – le montant des investissements consacré à la R&D de la présente innovation est moins élevé si la firme anticipe de futures innovations plus importantes.

3. Distinction entre l'incertitude de marché et l'incertitude technologique

Au-delà de l'incertitude sur la date d'introduction de l'innovation, il existe de l'incertitude de marché. Cette incertitude, considérée parfois comme de l'incertitude stratégique, est relative au fait que si plusieurs agents (firmes) s'engagent dans un projet unique, chacun est incertain sur la stratégie de ses concurrents et l'identité du vainqueur à la fin du processus concurrentiel n'est pas connue d'avance. Par conséquent, aucun agent ne connaît avec certitude ce que fait l'autre agent.

La question à laquelle nous tentons de répondre est quel est l'impact des sources d'incertitude sur la stratégie des firmes ? Cette question soulève différentes discussions sur les problèmes relatifs à l'incitation à l'investissement. Elaborée initialement par Loury (1979), c'est dans le modèle de Beath *et al.*, (1989) qu'apparaît réellement l'analyse des différentes sources d'incitations à l'innovation. Ces auteurs distinguent, en effet, l'incitation par le profit et l'incitation par la menace de la concurrence. C'est par l'intermédiaire de la magnitude entre ces deux sources d'incitation que sont analysées traditionnellement les stratégies à la R&D. L'incitation par le profit est toujours présente dans la firme même si cette dernière n'a aucune rivale sur le marché : c'est le fondement principal de tout investissement. Le profit engendré est précisément déterminé en calculant d'abord le montant investi comme si la firme est la seule à être présente sur le marché. La deuxième force qui détermine le montant de la R&D est le fait que la firme est en course avec ses rivales pour être la première à innover. Cette force reflète la différence entre le profit de la firme si elle innove avant ses rivales et le profit qu'elle peut obtenir si sa rivale innove en première position. A partir de ces définitions, nous constatons que le taux de hasard est lié à l'incitation par le profit et à l'incitation par la menace de la concurrence.

L'étude des différentes contributions que l'on retrouve dans la littérature montre en général une prédominance de l'incertitude de marché sur l'incertitude technologique (paragraphe 3.1) ou une prédominance de l'incertitude technologique sur l'incertitude de marché (paragraphe 3.2).

3.1. Prédominance de l'incertitude de marché sur l'incertitude technologique

Lorsque l'incertitude de marché est supérieure à l'incertitude technologique, il devient plus complexe d'établir une relation sur les stratégies des firmes. Dans ce cas, les conjectures sur l'issue de l'investissement sont plus qu'incertaines. A ce propos, Fethke et Birch (1982) considèrent qu'il suffit simplement que le coût de l'innovation soit incertain pour que le comportement stratégique des firmes devienne imprévisible⁴. Dans cette analyse, le coût de l'innovation est déterminé en termes de flux de dépenses et les firmes sont identifiées en fonction de leurs positions sur le marché. Hormis le coût de l'innovation, ce modèle prend aussi en considération le coût de développement de l'innovation. Ce coût est matérialisé par l'ensemble des dépenses engagées après l'introduction de l'innovation sur le marché. De ce point de vue, le coût de l'innovation peut avoir un impact sur l'incertitude de marché. Mieux encore, on peut imaginer la situation où il existe une absence totale d'incertitude sur les coûts de l'innovation. Kamien et Schwartz (1976) ont mis en évidence ce cas de figure par une représentation de l'incertitude dans leur modèle, mais en négligeant son impact sur les autres variables. Ici, l'incertitude technologique est aussi représentée comme une forme de relation stochastique entre le montant des dépenses en R&D et la date éventuelle du succès de l'innovation.

Par rapport à une analyse générale sur la vision du marché, c'est parce qu'il existe des firmes concurrentes dont chacune prise individuellement est incertaine sur la stratégie des autres firmes qu'il existe une prédominance de l'incertitude de marché sur l'incertitude technologique. En revanche, lorsqu'on s'intéresse au résultat ou à la probabilité d'obtention d'une innovation l'incertitude technologique prédomine l'incertitude de marché.

3.2 Prédominance de l'incertitude technologique sur l'incertitude de marché

Cette prédominance se justifie par la place occupée par la variable R&D dans les modèles de chronologie de l'innovation. Dans les modèles stochastiques, la variable R&D a un impact direct sur la date à laquelle l'innovation devient disponible. De plus, la probabilité de découverte d'une innovation augmente avec le montant investi, mais moins que

⁴ L'incertitude sur le coût de l'innovation concerne uniquement le cas où l'effet d'intensité fait l'objet de l'analyse. Cependant, si le coût de l'innovation est une donnée fixe mesurée avant le lancement de l'investissement, l'incertitude sur l'estimation du coût de l'innovation n'est plus prise en compte.

proportionnellement. Dans ce cas, la relation entre le montant engagé dans la R&D et probabilité de découverte est aussi croissante. Ainsi, le choix du modèle permet de justifier la source d'incertitude la plus importante.

En se focalisant sur les facteurs déterminant de la R&D, les modèles de Loury (1979), Lee et Wilde (1980), Beath *et al.*, (1989) et enfin Baumol (2002) présentent un intérêt supplémentaire par rapport aux précédents car ceux-ci prennent en compte l'importance considérable de l'incertitude technologique vis-à-vis de l'innovation stratégique. De ce point de vue, la maîtrise de la technologie est l'élément fondamental permettant la mise en place d'une innovation technologique.

4. Distinction entre taux de hasard constant et taux de hasard croissant

La forme prise par le taux de hasard dépend de la représentation de la distribution du processus étudié. Lorsque l'on s'intéresse à l'identité de la firme qui acquiert l'innovation dans le cadre des modèles stochastiques, le taux de hasard constant dépend uniquement de l'effort en R&D (paragraphe 4.1). Dans ce cas, nous montrerons que l'expression du taux de hasard est identique à la valeur prise par le paramètre utilisé pour représenter un processus qui suit une loi de poisson. En revanche, le taux de hasard croissant sera analysé comme un processus global permettant de prendre en compte simultanément l'effort en R&D et le nombre de firmes présentes dans l'industrie (paragraphe 4.2). Pour chaque type de taux de hasard, nous présenterons l'espérance de la date de découverte de l'innovation (paragraphe 4.3).

4.1. Caractéristiques du taux de hasard constant

La définition du taux de hasard dépend de la prise en compte simultanée de plusieurs critères. En l'absence de processus d'apprentissage par la pratique, un modèle stochastique est caractérisé par un taux de hasard constant si les deux critères suivants sont réunis : (i) la probabilité instantanée de découverte de l'innovation est connue et dépend de l'effort d'investissement et, (ii) il existe une indifférence des firmes vis-à-vis de l'efficacité du système de protection des innovations.

Le premier critère signifie que dans les modèles avec un taux de hasard constant, la relation fondamentale est définie à partir du montant investi dans la R&D et de la date à laquelle l'innovation devient disponible. Plus le montant investi est élevé, plus l'espérance de réussite de l'innovation est importante. De ce fait, la date à laquelle l'innovation devient disponible est obtenue plus tôt que prévue : la relation entre le montant investi dans la

R&D et la date de découverte est décroissante. Dans ces conditions, les firmes sont récompensées en fonction de leurs efforts d'investissement en R&D. Cette relation permet aussi de caractériser les modèles déterministes sauf que dans ce cas il n'y a pas d'incertitude. Avec le second critère, l'identité du vainqueur est celle de la firme qui obtient l'innovation, puis le brevet. De ce fait, lorsque l'effort d'investissement correspond à la stratégie adoptée au cours du tournoi par exemple, alors la récompense est l'obtention du premier brevet de l'invention. Cependant, la détention d'un brevet ne signifie pas forcément l'absence totale d'un mécanisme d'imitation. Dans ce cas, l'entrant potentiel peut imiter la technologie de la firme innovatrice.

La définition précédente montre que le critère essentiel retenu est le montant investi dans la R&D. Comme le confirme l'analyse de Beath *et al.*, (1989), la probabilité de découverte d'un nouveau produit ou d'un nouveau processus dans un intervalle de temps $[t, t + dt]$ est conditionnée non pas par ce qui a été découvert en t , mais dépend uniquement du flux de dépenses engagé par la firme à la date t . Donc, le résultat du processus dépend des capacités d'investissement de la firme. Dans ces conditions, la probabilité conditionnelle de l'introduction de l'innovation dépend seulement de l'intervalle de temps dans lequel la date optimale d'introduction de l'innovation est incluse. En d'autres termes, la date d'introduction de l'innovation est conditionnelle par rapport à la date d'obtention de l'innovation.

Lorsque l'on s'intéresse à l'identité du vainqueur dans la course à l'innovation, le taux de hasard constant est représenté dans Loury (1979) comme une distribution exponentielle. En considérant par exemple qu'une firme dépense un montant noté x en R&D, l'innovation sera obtenue à une date aléatoire $\tau(x)$. Cette date est une variable aléatoire dont la distribution est représentée par une fonction de répartition $F(t, x) = \Pr[\tau(x) \leq t]$. Lorsque cette fonction est continue et dérivable dans l'intervalle t ($t \in [0, +\infty[$), il existe alors une fonction de densité de répartition définie

par $f(t, x) = \frac{\partial F(t, x)}{\partial t}$. En reprenant l'intervalle de temps défini par Beath *et al.*, (1989), le taux de hasard constant est défini par la relation suivante :

$$T(x) = \Pr \left[\tau(x) \in \left[\frac{t, t + dt}{\tau(x)} \right] > t \right] = \frac{F(t + dt, x) - F(t, x)}{1 - F(t, x)} \equiv \frac{f(t, x) dt}{1 - F(t, x)}.$$

Dans cette relation, $T(x) = \frac{f(t, x) dt}{1 - F(t, x)}$ est le ratio qui détermine le taux de hasard constant. Comme $f(t, x) = -\frac{\partial[1 - F(t, x)]}{\partial t}$, alors le taux de hasard peut s'écrire $T(x) = -\frac{\frac{\partial[1 - F(t, x)]}{\partial t}}{1 - F(t, x)}$.

En calculant cette relation par intégration, on a $\log[1 - F(t, x)] = -T(x)t + k$, avec k défini comme la constante d'intégration. Par conséquent, nous avons $F(t, x) = 1 - e^{[-T(x)t + k]}$. Etant donné que la fonction de répartition ne dépend pas de la valeur prise par k donc, $F(t, x) = 1 - e^{[-T(x)t]}$. La fonction de répartition est une distribution exponentielle identique celle d'une loi de poisson dont le paramètre est noté $T(x)$. Analytiquement, le taux de hasard $T(x)$ dépend uniquement de l'effort en R&D.

4.2. Caractéristiques du taux de hasard croissant

Le taux de hasard croissant est aussi défini par plusieurs caractéristiques. Seulement, la spécificité permettant de différencier cette approche de la précédente consiste à préciser qu'en plus de l'effort d'investissement, les firmes sont différenciées par leurs positions concurrentielles. Le taux de hasard est croissant si la firme innovatrice (first mover) dispose de la possibilité de se procurer la totalité des revenus de l'innovation durant un intervalle intégrant la durée légale de protection de l'innovation contre les firmes concurrentes (follower). Ainsi, les critères qui doivent être réunis pour déterminer la forme croissante du taux de hasard sont : (i) l'hypothèse de "*winner takes all profit*" (Dasgupta et Stiglitz, 1980), (ii) une incertitude sur l'intervalle de temps optimal durant lequel l'innovation est introduite sur le marché et, (iii) un système de protection efficace des innovations. Ces critères sont définis sur un marché concurrentiel composé de plusieurs firmes. Avec cette définition, la course à l'innovation est en général matérialisée par l'effet d'intensité parce que la concurrence prend fin lorsqu'une seule firme obtient l'innovation. En effet, le premier critère signifie que la firme first mover ou le vainqueur s'approprie la totalité des revenus de l'innovation. Ce critère rejoint le troisième parce que le vainqueur se trouve dans une situation où les mécanismes de protection des

innovations sont efficaces. De la même manière, la date à laquelle l'innovation devient disponible ne dépend plus uniquement du montant investi dans la R&D, mais elle dépend aussi du nombre de firmes présentes sur le marché. Avec le second critère, l'incertitude sur la date d'obtention de l'innovation implique aussi une incertitude sur le succès de l'innovation.

On considère analytiquement que le nombre de firmes en concurrence pour la course à l'innovation est égale à n . Dans ce cas, la concurrence est de telle sorte que toutes les firmes investissent en R&D de manière indépendante, mais une seule firme acquiert l'innovation. Pour déterminer la forme prise par le taux de hasard de la firme i par exemple, la même méthode que celle utilisée dans le cas du taux de hasard constant est reprise en passant de $i=1$ à $i=1,2,...,n$. En considérant que chaque firme i investit un montant x_i , la date à laquelle la première firme obtient l'innovation est aussi une variable aléatoire notée $\tau(x_1, x_2, \dots, x_n)$ correspondant à la plus petite date parmi celles distribuées à l'ensemble des firmes qui investissent sur le marché. Cette distribution est telle qu'une seule firme obtient l'innovation. Dans ce cas, on a :

$$\tau(x_1, x_2, \dots, x_n) = \min[\tau(x_1), \tau(x_2), \dots, \tau(x_n)].$$

A partir de la variable aléatoire, la fonction de répartition correspondant à la date de découverte de l'innovation par l'ensemble des firmes est donnée par : $F(t, x_1, x_2, \dots, x_n) = \Pr[\tau(x_1, x_2, \dots, x_n) \leq t]$. Avec les probabilités complémentaires, on a $F(t, x_1, x_2, \dots, x_n) = 1 - \Pr[\tau(x_1, x_2, \dots, x_n) > t]$. Comme les variables aléatoires $\tau(x_i)$ sont indépendamment distribuées, on peut écrire par définition :

$$\Pr[\tau(x_1, x_2, \dots, x_n) > t] = \Pr[\min(\tau(x_1), \tau(x_2), \dots, \tau(x_n)) > t] = \prod_{i=1}^n \Pr[\tau(x_i) > t] = e^{\left[-\sum_{i=1}^n T(x_i)t\right]}.$$

En reprenant la fonction de répartition initiale, on a :

$$F(t, x_1, x_2, \dots, x_n) = \Pr[\tau(x_1, x_2, \dots, x_n) \leq t] = 1 - e^{\left[-\sum_{i=1}^n T(x_i)t\right]}.$$

Etant donné que cette fonction de répartition peut prendre la forme d'un logarithme, donc la probabilité instantanée de découverte ou taux de hasard dépend de l'effort en R&D et du nombre de firmes présentes dans l'industrie.

4.3. Espérances de la date de succès et formes prises par le taux de hasard

A partir de la forme prise par le taux de hasard, il est possible de déterminer l'espérance de la date de succès de l'innovation. En considérant par exemple que la R&D est lancée à la date $t = 0$, l'espérance de la date de succès de l'innovation, notée $E[\tau(x)]$, est déterminée à partir de la date de lancement jusqu'à ce que l'innovation soit disponible. Ici, la R&D est mesurée en termes de flux de dépenses prenant fin lorsque l'innovation est obtenue.

Par définition, on a $E[\tau(x)] = \int_0^\infty t \partial F(t, x) = \int_0^\infty t f(t, x) dt = T(x) \int_0^\infty t e^{-T(x)t} dt$.

En appliquant l'intégration par partie $(t, e^{-T(x)t})$, on obtient :

$E[\tau(x)] = \frac{1}{T(x)}$. Lorsque le taux de hasard est constant, l'espérance de la

date de découverte est définie comme l'inverse du taux de hasard. Plus l'investissement est élevé, plus l'espérance de la date de découverte est proche. En revanche, dans le cadre du taux de hasard croissant, l'espérance de succès de l'innovation est aussi une variable exogène déterminée de la

manière suivante : $E[\tau(x_1, x_2, \dots, x_n)] = \int_0^\infty t \frac{\partial F(t, x_1, x_2, \dots, x_n)}{\partial t} dt = \int_0^\infty t e^{\left[-\sum_{i=1}^n T(x_i)t\right]} dt$.

On obtient dans ce cas $E[\tau(x_1, x_2, \dots, x_n)] = \frac{1}{\sum_{i=1}^n T(x_i)}$. En présence de n

firmer, l'espérance de succès de l'innovation est définie comme l'inverse du taux de hasard global. Dans ce cas, le taux de hasard global est la somme des taux de hasard individuel de l'ensemble des firmes que compose l'industrie.

A partir de la comparaison entre les espérances de succès, le nombre de firmes pris en compte joue un rôle important. De ce fait, si l'on passe d'une situation de concurrence avec $i = 1, 2, \dots, n$ à une situation de duopole ($i = 1, 2$), l'espérance de succès de l'innovation devient simple à calculer.

Dans ce cas, le taux de hasard global est la somme des taux de hasard individuel des deux firmes. Par conséquent, le taux de hasard global définissant à l'origine le taux de hasard croissant est approximativement identique au taux de hasard individuel, c'est-à-dire au taux de hasard constant. De manière plus précise, on obtient :

$E[\tau(x_1, x_2)] = \frac{1}{T(x_1) + T(x_2)}$. L'espérance de la date de découverte par au

moins l'une des firmes est égale à l'inverse du taux de hasard des deux firmes réunies. Comme il n'existe qu'une seule firme qui obtient l'innovation dans le cadre de la course à l'innovation, donc si l'effort en R&D est mesuré en termes de flux de dépenses dans le cas où $n = 2$, on obtient $T(x_1) + T(x_2) \equiv T(x_1)$. Il existe un système de *winner take all* et la firme follower pourra être considérée comme une firme inactive (Loury, 1979).

Section 2

Relation fonctionnelle entre l'effort en R&D et la date d'innovation

Dans cette section, la relation fonctionnelle entre l'effort en R&D et la date à laquelle l'innovation devient disponible sera étudiée. Cette relation fonctionnelle est analysée à partir des modèles déterministes (paragraphe 1) et des modèles stochastiques (paragraphe 2). La distinction entre ces deux types de modèle permettra de mesurer l'intérêt de prendre en compte les interactions stratégiques. Dans le cas contraire, nous montrerons qu'il existe un raisonnement que nous qualifions de paradoxe de Baumol en contradiction avec le cadre d'analyse de la course à l'innovation (paragraphe 3).

1. Analyse des modèles déterministes

Dans le cas des modèles déterministes, la méthode traditionnelle permettant de déterminer la date optimale d'innovation est d'introduire la valeur actuelle nette (paragraphe 1.1). Par la suite, nous proposerons une analyse des différentes incitations à l'innovation dans le cas de firmes symétriques (paragraphe 1.2), puis dans le cas de firmes asymétriques (paragraphe 1.3).

1.1. Détermination de la date optimale d'innovation

Ici, nous tentons de déterminer la date optimale d'innovation en situation de concurrence⁵. Etant donné qu'une seule firme acquiert l'innovation selon l'effort d'investissement, la question à laquelle nous tentons de répondre est de savoir quel est l'impact de la concurrence sur la

⁵ Cette analyse est inspirée à partir de Barzel (1968).

valeur actuelle nette (VAN) de chaque firme ? Pour répondre à cette question, on s'appuie sur le calcul de la VAN pour déterminer la date optimale d'introduction de l'innovation. Pour cela, les projets d'investissement de deux firmes seront analysés de manière distincte.

Dans le cadre des modèles déterministes, l'innovation est une donnée exogène du système, c'est-à-dire la firme est certaine d'obtenir une innovation lorsqu'elle investit. Dans ces conditions, nous cherchons à déterminer la date optimale d'introduction de l'innovation à partir de laquelle l'innovation est profitable du fait qu'elle contribue à l'efficacité productive par la maximisation des profits (firme) et à l'efficacité allocative par la création de bien être pour la société (consommateur).

Etant donné que la date d'introduction de l'innovation dépend des conditions du marché, la méthode traditionnelle utilisée pour calculer la VAN nécessite de choisir les variables du modèle et de poser quelques hypothèses de base.

Choix des variables et les hypothèses du modèle :

- i) Le coût de l'investissement ou l'effort en R&D est égal à I .
- ii) L'adoption d'une innovation à une date optimale réduit de k le coût de chaque unité d'output produite.
- iii) La demande de biens X_0 à la date $t = 0$ augmente avec un taux p .
Ainsi, la demande à la date t sera $X_t = X_0 e^{pt}$.
- iv) Le rendement de l'innovation est $h = k$ par unité d'output. Si l'innovation est introduite à la période initiale, le rendement total obtenu pour la période est $hX_0 = S_0$.
- v) Chaque firme qui investit obtient une innovation.
- vi) Le taux d'intérêt constant est donné par r .
- vii) Les paramètres du système sont connus et exprimés en unité monétaire.

Avec ces paramètres, il est possible de calculer la date optimale à laquelle l'innovation sera introduite sur le marché. Etant donné que le bénéfice généré par l'innovation est distribué dans le temps et que l'investissement en R&D peut être reporté, alors la rentabilité de l'innovation est mesurée par la VAN. Notée par R_0 , la VAN de l'innovation est la différence entre le montant espéré des flux de trésorerie escompté de l'investissement (qui augmente à un taux p par période) et le montant investi au début du projet. Pour simplifier le raisonnement, nous considérons que la date à partir de laquelle la firme acquiert l'innovation est identique à la date d'introduction de l'innovation.

En considérant que le flux de bénéfice est escompté à partir de la période t alors par définition la VAN du projet d'investissement est donnée par :

$$R_0 = \int_t^{\infty} S_0 e^{-(r-p)t} dt - Ie^{-rt} \quad (1.1)$$

A partir de l'expression de la VAN, nous prenons uniquement en considération le cas où $p < r$. En effet si $p \geq r$, la VAN est indéterminée. Dans cette analyse, il est essentiel que $p > 0$ pour que l'innovation soit profitable. De ce fait si le pourcentage de hausse de la trésorerie escomptée de l'investissement est inférieur au taux d'intérêt, la VAN est donnée par :

$$R_0 = \frac{S_0 e^{-(r-p)t}}{r-p} - Ie^{-rt} \quad (1.2)$$

Etant donné que le temps t est pris en compte, la date permettant de maximiser R_0 est donnée par :

$$t_m = \frac{\ln r + \ln I - \ln S_0}{p} \quad (1.3)$$

L'expression de l'équation (1.3) est indéterminée lorsque $p = 0$. Si $0 < p$, t_m est négatif et la production diminue avec le temps. En revanche si $p > 0$, la production augmente avec le temps et la date optimale d'introduction de l'innovation est positive. Donc, l'intervalle $0 < p < r$ est pris en compte pour la suite de l'analyse.

Il est possible de constater que plus le montant des flux de trésorerie escomptée est important par rapport au coût de l'investissement, plus la date d'introduction de l'innovation maximisant R_0 est faible. A l'inverse, la date d'introduction de l'innovation maximisant R_0 est élevée, c'est-à-dire la firme a tendance à retarder la date d'introduction de son innovation.

En examinant la date t_m qui maximise la VAN, il est possible d'écrire :

$$S_0 e^{-(r-p)t(m)} = rIe^{-rt(m)} \quad (1.4)$$

En divisant l'expression (1.4) par $e^{-rt(m)}$, on obtient la valeur actuelle de l'investissement notée par $S_0 e^{pt(m)} = rI$.

En utilisant l'expression de (1.4), il est possible de remplacer la valeur de S_0 dans l'équation (1.2). On obtient :

$$R_0(m) = \frac{rIe^{-rt(m)}}{r-p} - Ie^{-rt(m)} = \left(\frac{r}{r-p} - 1 \right) Ie^{-rt(m)}$$

$$R_0(m) = \frac{p}{r-p} Ie^{-rt(m)} \quad (1.5)$$

$R_0(m)$ est la valeur présente maximum à la date 0. Donc, à la date $t(m)$ la valeur présente de l'innovation est donnée par :

$$R_m(m) = \frac{R_0(m)}{e^{-rt(m)}} = I \frac{p}{r-p} \quad (1.6)$$

L'expression (1.6) indique la valeur de l'innovation à la date de son introduction $t(m)$.

A partir de l'équation (1.2), il est possible d'obtenir la date $t(z)$ pour laquelle la VAN de l'innovation sera égale à zéro. En posant $R_0(m) = 0$, on obtient :

$$t(z) = \frac{\ln(r-p) + \ln I - \ln S_0}{p} \quad (1.7)$$

En comparant les expressions (1.3) et (1.7) on constate que $t(z) < t(m)$. La date à laquelle la VAN de l'innovation est égale à zéro est obtenue avant la date à laquelle la VAN de l'innovation est à son maximum. La Figure 1 permet de représenter graphiquement les dates $t(z)$ et $t(m)$ que peut prendre la VAN.

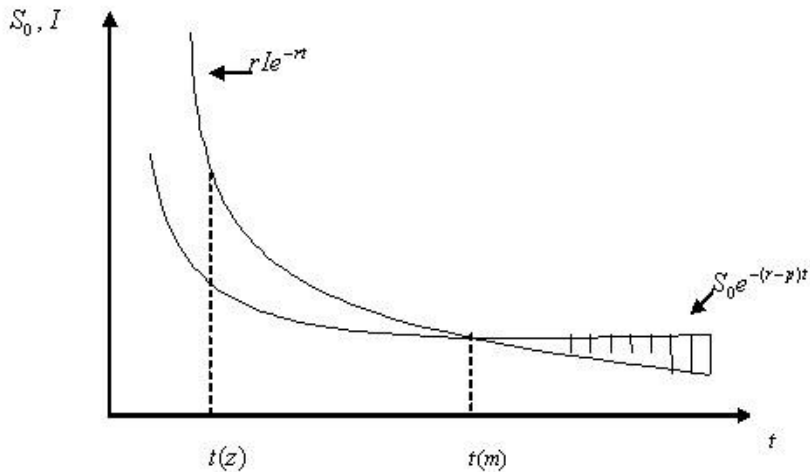


Figure 1. VAN du projet d'investissement

La Figure 1 montre que les deux courbes sont décroissantes. La courbe représentative du bénéfice de l'investissement $(S_0 e^{-(r-p)t})$ décroît plus lentement que la courbe de la valeur présente du coût de l'investissement (rIe^{-rt}) . Cela est dû au fait que le taux de croissance de la demande est positif. En effet, à la date $t(m)$ les deux courbes se rencontrent, c'est-à-dire le bénéfice est égal au coût de l'investissement. A cette date, le profit est au maximum. Pour $t < t(m)$, le coût de l'innovation par unité de temps est supérieur au bénéfice obtenu. Dans ce cas, l'investissement n'est pas rentable. En revanche pour $t > t(m)$ le bénéfice par unité de temps est supérieur au coût de l'investissement. Dans ce cas, l'investissement est rentable. La surface hachurée à droite de $t(m)$ représente la valeur présente nette générée par l'innovation.

Si l'innovation est entreprise à la date $t(z)$ la VAN est nulle, alors qu'à la date $t(m)$ la VAN est au maximum. Par conséquent, même si l'innovation est introduite entre $t(z)$ et $t(m)$ la VAN n'est pas maximisée mais elle est positive.

Dans cette analyse, la date d'introduction de l'innovation est optimale lorsque la valeur présente est maximisée. Dans la situation où il existe une seule firme sur le marché, celle-ci maximise son profit en introduisant l'innovation à la date $t(m)$. En revanche, si la structure du marché est concurrentielle, les firmes sont incitées à introduire leur innovation avant la date $t(m)$. Lorsqu'il y a de la concurrence, l'entrant potentiel peut introduire une innovation entre les dates $t(z)$ et $t(m)$. Cela signifie que le profit maximum correspondant à la date optimale d'introduction de l'innovation n'est pas obtenu parce que la date qui maximise la VAN n'est pas atteinte.

Le principal enseignement de cette analyse est qu'en situation de concurrence les firmes sont incitées à introduire leurs innovations à la date à laquelle la VAN est nulle. Dans ces conditions, le profit maximum n'est pas réalisé. En considérant une situation de duopole, la date à laquelle la firme first mover introduit son innovation, c'est-à-dire la date $t(m)$ peut ne pas être obtenue du fait que l'entrant potentiel est incitée à introduire son innovation entre les dates $t(z)$ et $t(m)$. Même si la date permettant de maximiser la VAN est connue, on peut tenter de comprendre si la relation qui existe entre les firmes peut avoir un impact sur la date à laquelle l'une