



Allocation stratégique d'actifs et ALM pour les régimes de retraite

Alaeddine Faleh

► To cite this version:

Alaeddine Faleh. Allocation stratégique d'actifs et ALM pour les régimes de retraite. Gestion de portefeuilles [q-fin.PM]. Université Claude Bernard - Lyon I, 2011. Français. NNT : . tel-00689907v3

HAL Id: tel-00689907

<https://theses.hal.science/tel-00689907v3>

Submitted on 11 Sep 2012 (v3), last revised 11 Sep 2012 (v4)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

THÈSE
présentée
devant l'UNIVERSITÉ CLAUDE BERNARD - LYON 1 – ISFA
pour l'obtention du
DIPLÔME DE DOCTORAT
Spécialité sciences actuarielle et financière

présentée et soutenue publiquement le 13/05/2011
par Mr. Alaeddine FALEH

Allocation stratégique d'actifs et ALM pour les régimes de retraite

Directeur de thèse : Pr. Frédéric PLANCHET
Dr. Didier RULLIERE (co-directeur)

Composition du jury :

Président : Pr. Jean-Claude AUGROS
Professeur à l'Université Claude Bernard Lyon1 (I.S.F.A.)

Rapporteurs : Pr. Jean-Paul LAURENT
Professeur à l'Université Paris 1 Panthéon-Sorbonne

Pr. François DUFRESNE
Professeur à l'Université de Lausanne (Suisse)

Directeur de thèse : Pr. Frédéric PLANCHET
Professeur associé à l'Université Claude Bernard Lyon1 (I.S.F.A.)

Dr. Didier RULLIERE
Maître de conférences à l'Université Claude Bernard Lyon1 (I.S.F.A.)

Suffragant : Pr. Jean-Pierre AUBIN
Professeur à l'Ecole Polytechnique de Paris

Dr. Guillaume LEZAN
Caisse des Dépôts et Consignations (CDC)

-Ecole Doctorale Sciences Economiques et de Gestion-
-Institut de Science Financière et d'Assurances (Univ. Lyon I) -
-Laboratoire de Sciences Actuarielle et Financière (équipe d'accueil n°2429)-

RÉSUMÉ

La présente thèse s'intéresse aux modèles d'allocation stratégiques d'actifs et à leurs applications pour la gestion des réserves financières des régimes de retraite par répartition, en particulier ceux partiellement provisionnés. L'étude de l'utilité des réserves pour un système par répartition et *a fortiori* de leur gestion reste un sujet peu exploré. Les hypothèses classiques sont parfois jugées trop restrictives pour décrire l'évolution complexe des réserves. De nouveaux modèles et de nouveaux résultats sont développés à trois niveaux : la génération de scénarios économiques (GSE), les techniques d'optimisation numérique et le choix de l'allocation stratégique optimale dans un contexte de gestion actif-passif (ALM).

Dans le cadre de la génération de scénarios économiques et financiers, certains indicateurs de mesure de performance du GSE ont été étudiés. Par ailleurs, des améliorations par rapport à ce qui se pratique usuellement lors de la construction du GSE ont été apportées, notamment au niveau du choix de la matrice de corrélation entre les variables modélisées. Concernant le calibrage du GSE, un ensemble d'outils permettant l'estimation de ses différents paramètres a été présenté.

Cette thèse a également accordé une attention particulière aux techniques numériques de recherche de l'optimum, qui demeurent des questions essentielles pour la mise en place d'un modèle d'allocation. Une réflexion sur un algorithme d'optimisation globale d'une fonction non convexe et bruitée a été développée. L'algorithme permet de moduler facilement, au moyen de deux paramètres, la réitération de tirages dans un voisinage des points solutions découverts, ou à l'inverse l'exploration de la fonction dans des zones encore peu explorées.

Nous présentons ensuite des techniques novatrices d'ALM basées sur la programmation stochastique. Leur application a été développée pour le choix de l'allocation stratégique d'actifs des régimes de retraite par répartition partiellement provisionnés. Une nouvelle méthodologie pour la génération de l'arbre des scénarios a été adoptée à ce niveau. Enfin, une étude comparative du modèle d'ALM développé avec celui basé sur la stratégie *Fixed-Mix* a été effectuée. Différents tests de sensibilité ont été par ailleurs mis en place pour mesurer l'impact du changement de certaines variables clés d'entrée sur les résultats produits par notre modèle d'ALM.

MOTS-CLÉS

Gestion actif-passif, Allocation d'actifs, Génération de scénarios, Optimisation stochastique, Programmation stochastique, Régime de retraite

ABSTRACT

This thesis focuses on the strategic asset allocation models and on their application for the financial reserve management of a pay-as-you-go (PAYG) retirement schemes, especially those with partial provision. The study of the reserve utility for a PAYG system and of their management still leaves a lot to be explored. Classical hypothesis are usually considered too restrictive for the description of the complex reserve evolution. New models and new results have been developed over three levels : economic scenario generation (ESG), numerical optimization techniques and the choice of optimal strategic asset allocation in the case of an Asset-Liability Management (ALM).

For the generation of financial and economic scenarios, some ESG performance indicators have been studied. Also, we detailed and proposed to improve ESG construction, notably the choice of the correlation matrix between modelled variables. Then, a set of tools were presented so that we could estimate ESG parameters variety.

This thesis has also paid particular attention to numerical techniques of optimum research, which is an important step for the asset allocation implementation. We developed a reflexion about a global optimisation algorithm of a non convex and a noisy function. The algorithm allows for simple modulating, through two parameters, the reiteration of evaluations at an observed point or the exploration of the noisy function at a new unobserved point.

Then, we presented new ALM techniques based on stochastic programming. An application to the strategic asset allocation of a retirement scheme with partial provision is developed. A specific methodology for the scenario tree generation was proposed at this level. Finally, a comparative study between proposed ALM model and Fixed-Mix strategy based model was achieved. We also made a variety of a sensitivity tests to detect the impact of the input values changes on the output results, provided by our ALM model.

KEY-WORDS

Asset-Liability Management, Asset Allocation, Scenario Generation, Stochastic Optimization, Stochastic Programming, Retirement Scheme

REMERCIEMENTS

Je souhaite tout d'abord remercier vivement Prof. Frédéric PLANCHET pour avoir dirigé ma thèse. Je suis particulièrement reconnaissant de la confiance qu'il a su m'accorder dès les premiers kilomètres de ce marathon, ainsi que de ses précieux conseils.

Je veux également remercier Dr. Didier RULLIERE, qui a co-dirigé cette thèse, et dont les suggestions ont été particulièrement profitables au développement de cette thèse.

J'adresse chaleureusement mes remerciements aux membres du jury auxquels je suis reconnaissant pour l'honneur qu'ils me font de juger ce travail.

Cette thèse a également reçu un soutien considérable, à son commencement, de la part du Prof. Jean Claude AUGROS. Je le remercie d'avoir été toujours disponible pour répondre à mes questions.

Mes vifs remerciements s'adressent à Mme. Edith JOUSSEAUME (directrice du département des Investissements et de Comptabilité à la Caisse des Dépôts et Consignation) pour m'avoir accordée sa confiance et son encouragement : des atouts qui ont permis le déroulement de mon travail, au sein de son département, dans les meilleures conditions.

Ma profonde gratitude à Dr. Guillaume LEZAN, mon responsable scientifique à la Caisse des Dépôts et Consignation, qui a passé bien des heures dans la relecture de cette thèse tout en l'enrichissant avec ses pertinentes suggestions.

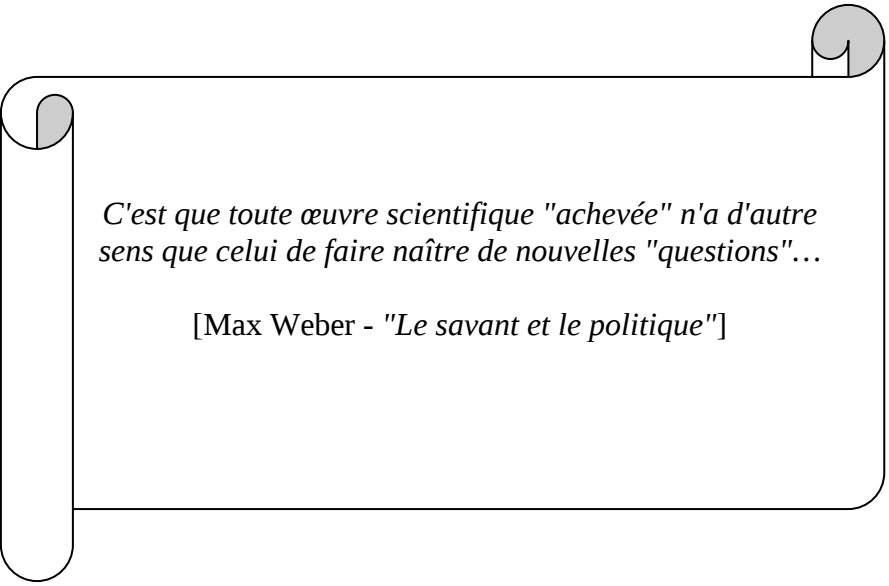
Cette thèse a bénéficié d'un précieux soutien, à son commencement, de la part de Mr. Pierre LEYGUE (Caisse des Dépôts et consignations) envers qui je reste toujours reconnaissant.

J'adresse un remerciement tout particulier à ma chère femme Asma pour la patience dont elle a fait preuve pendant ces dernières années de travail et l'exploit de me supporter durant cette période.

Merci également à ma famille pour avoir consenti l'éloignement et bien d'autres sacrifices en vue de l'aboutissement de cette thèse.

Finalement, je remercie tous ceux qui ont participé, de près ou de loin, à l'accomplissement de ce travail.

À mes deux chers parents



C'est que toute œuvre scientifique "achevée" n'a d'autre sens que celui de faire naître de nouvelles "questions"...

[Max Weber - "Le savant et le politique"]

Sommaire

Sommaire.....	1
Introduction générale.....	4
Partie I : Les générateurs de scénarios économiques (GSE).....	17
Introduction	18
Chapitre 1 : Etude des GSE.....	19
I- Présentation théorique	19
I-1 Problèmes théoriques de modélisation liés aux GSE.....	19
I-2 Littérature sur les structures des GSE.....	24
II- Mise en œuvre d'un GSE.....	30
II-1 Structure schématique de projection de scénarios pour un GSE.....	30
II-2 Méthodologies pour la génération de scénarios économiques.....	34
II-3 Probabilité risque-neutre Vs probabilité réelle.....	36
II-4 L'exclusion de valeurs négatives pour les modèles de taux.....	37
III- Mesure de qualité.....	38
III-1 Mesure qualitative de la qualité d'un GSE.....	38
III-2 Mesure quantitative de la qualité d'un GSE.....	40
III-3 Application numérique.....	45
Conclusion du chapitre 1.....	49
Chapitre 2 : Construction d'un générateur de scénarios économiques.....	50
I- Conception et composantes	50
I-1 Conception théorique	50
I-2 Composantes	53
II- Calibrage.....	73
II-1 Présentation du contexte de l'étude	73
II-2 Calibrage du modèle de l'inflation.....	76

II-3 Calibrage du modèle des taux réels.....	79
II-4 Calibrage du modèle des actions.....	83
II-5 Calibrage du modèle de l'immobilier.....	89
II-6 Calibrage du modèle de crédit.....	91
III- Projection des scénarios	93
III-1 Éléments sur la mise en œuvre.....	93
III-2 Tests du modèle : paramètres et matrice de corrélation.....	96
III-3 Résultats de projection.....	99
Conclusion du chapitre 2.....	110
Conclusion de la partie I.....	110
Partie II : L'allocation stratégique d'actifs dans le cadre de la gestion actif-passif (ALM).....	112
Introduction.....	113
Chapitre 1 : Les modèles d'ALM classiques (déterministes).....	118
I- Immunisation du portefeuille.....	118
I-1 Adossement des flux de trésorerie.....	119
I-2 Adossement par la durée.....	120
II- Modèles basés sur le surplus.....	125
II-1 Modèle de Kim et Santomero [1988].....	125
II-2 Modèle de Sharpe et Tint [1990]	128
II-3 Modèle de Leibowitz [1992].....	130
II-4 Durée du Surplus.....	137
Conclusion du chapitre 1.....	139
Chapitre 2 : Allocation stratégique d'actifs et stratégie <i>Fixed-Mix</i>.....	140
I- Etude de l'allocation d'actifs dans le cadre de la stratégie <i>Fixed-Mix</i>.....	143
I-1 Présentation.....	143
I-2 Application.....	143

II- Proposition de méthodes numériques	168
II-1 Introduction.....	168
II-2 Une grille à pas variable par subdivision systématique.....	173
II-3 Une grille à pas variable avec retraitage possible.....	188
II-4 Applications numériques.....	191
II-5 Conclusion.....	211
Conclusion du chapitre 2.....	212
Chapitre 3 : Allocation stratégique d’actifs et modèles d’ALM dynamique.....	213
I - L’allocation stratégique d’actifs dans le cadre des modèles classiques d’ALM dynamique.....	214
I-1 Présentation générale.....	214
I-2 Techniques d’assurance de portefeuille	215
I-3 Programmation dynamique.....	223
II- L’allocation stratégique d’actifs dans le cadre de la programmation stochastique	227
II-1 Introduction	227
II-2 Présentation de la programmation stochastique.....	228
II-3 Illustration de la programmation stochastique avec recours dans le cas de la planification de la production.....	230
II-4 Formulation mathématique de la programmation stochastique avec recours.....	236
II-5 Résolution des programmes stochastiques.....	240
III- Nouvelle approche d’ALM par la discrétisation des scénarios économiques.....	242
III-1 Méthode des quantiles de référence pour le GSE.....	242
III-2 Modèle d’optimisation basé sur la programmation stochastique.....	247
III-3 Discussion des résultats.....	253
Conclusion du chapitre 3.....	260
Conclusion de la partie II	260
Conclusion générale.....	261
Index des graphiques.....	267
Bibliographie.....	271
Table des matières.....	283

INTRODUCTION GÉNÉRALE

Objectifs et enjeux de la thèse

Le champ d'étude offert par les régimes de retraite est très étendu et peut être segmenté selon différentes clés telles que notamment : régimes publics *versus* régimes d'entreprises, gestion en répartition¹ *versus* gestion en capitalisation², régimes à cotisations définies *versus* régimes à prestations définies, etc... Trois situations peuvent être recensées pour un régime donné : régime provisionné (cas où le régime est tenu à la couverture totale des engagements souscrits par ses cotisants et retraités actuels), régime partiellement provisionné (couverture d'une partie seulement des engagements souscrits par ses cotisants et retraités actuels) et enfin régime non provisionné.

Pour apprécier la solidité prudentielle d'un régime de retraite, les experts se sont longtemps limités à une approche binaire, selon laquelle il convenait de tout provisionner dans le cas d'un système par capitalisation et rien dans le cas d'un système par répartition. Des travaux récents mettent en évidence l'utilité de réserves pour un système par répartition (cf. Delarue [2001]).

Le fonctionnement financier des régimes de retraite par répartition, en particulier en France, peut être schématisé comme suit : les flux de cotisations permettent de régler les flux de prestations, ensuite le surplus permet, le cas échéant, d'alimenter une réserve destinée à régler une partie des prestations futures. Cette même réserve peut se voir prélever le solde technique débiteur s'il s'avère que les cotisations sont insuffisantes pour régler les prestations. Dans le même temps, la réserve est placée sur les marchés financiers et allouée selon différentes classes d'actifs.

La gestion actif-passif d'un régime par répartition partiellement provisionné peut être basée sur l'optimisation de la valeur de la réserve, compte tenu des contraintes liées au passif qu'il doit respecter. C'est dans ce cadre que le choix d'une « bonne » allocation stratégique joue un rôle essentiel dans le pilotage actif-passif d'un régime de retraite, étant donné que les réserves contribuent à part entière à la solidité du régime.

Cependant une difficulté consiste à définir la « meilleure » stratégie de placement de ces réserves sur les marchés financiers, notamment dans un contexte de fortes incertitudes économiques, comme c'est actuellement le cas du fait de la récente crise financière.

La présente thèse est consacrée aux modèles d'allocation stratégiques d'actifs et à leurs applications pour la gestion des réserves financières des régimes de retraites sur le long-terme.

¹ Mode d'organisation des systèmes de retraite fondé sur la solidarité entre générations. Les cotisations versées par les actifs au titre de l'assurance vieillesse servent immédiatement à payer les retraites. L'équilibre financier des systèmes de retraite par répartition est fonction du rapport entre le nombre de cotisant (population active, taux de croissance des revenus) et celui des retraités. Le système français de retraite est fondé sur le principe de la répartition.

² Mode d'organisation des systèmes de retraite dans lequel les cotisations d'un assuré sont placées à son nom durant sa vie active (placements financiers et immobiliers, dont le rendement varie en fonction des taux d'intérêt) avant de lui être restituées sous forme de rente après l'arrêt de son activité professionnelle. La constitution du capital peut s'effectuer à titre individuel ou dans un cadre collectif (accord d'entreprise). En France, seuls les systèmes de retraite dits sur-complémentaires (ex. PREFON, COREM, le PERP ou plan d'épargne populaire) fonctionnent selon le principe de la capitalisation à l'exception du RAFP.

Deux points particuliers font l'objet d'un examen approfondi : d'une part les générateurs de scénarios économiques (GSE) utilisés pour la modélisation des fluctuations des actifs financiers sur le long-terme puis, dans un deuxième temps, les modèles proposés pour l'élaboration de l'allocation d'actifs elle-même. A chaque fois les aspects théoriques précéderont la mise en œuvre opérationnelle sur des exemples.

Concernant les GSE, nous mettons en évidence leurs principales composantes, que ce soit au niveau de la conception théorique ou à celui de la mise en œuvre. Le choix de ces composantes est lié à la vocation finale du générateur de scénarios économiques, que ce soit en tant qu'outil d'évaluation des produits financiers (*pricing*) ou en tant qu'outil de projection et de gestion des risques. Par ailleurs, nous développons une étude sur certains indicateurs de mesure de la performance du GSE comme un outil en amont du processus de prise de décision : la stabilité et l'absence de biais. Une application numérique permettant d'illustrer ces différents points est présentée.

Notre objectif final est de proposer un GSE adapté aux spécificités de la gestion actif-passif d'un régime de retraite. Différents modèles de diffusion et d'estimation de leurs paramètres ont ainsi été envisagés pour le panel des variables financières et macro-économiques qui interviennent dans ce contexte. Cette partie est illustrée par une comparaison des résultats du générateur de scénario proposé avec ceux obtenus avec le modèle d'Ahlgrim et al. [2005].

A ce niveau, notre étude se distingue par l'adoption d'une démarche novatrice dont l'objectif est de tenir compte de l'instabilité temporelle de la corrélation entre les rentabilités des actifs. Autrement dit, nous tenons compte dans la modélisation des possibilités d'évolution de cette corrélation selon différents régimes. L'idée principale est de partir de l'approche classique de changement de régime (cf. Hamilton [1989, 1994] et Hardy [2001]), qui concerne souvent les rendements moyens ainsi que les volatilités, pour supposer que les corrélations basculent elles aussi d'un régime à un autre.

Nous introduisons un nouveau concept qui est celui de la « corrélation à risque » correspondant à la matrice de corrélation dans le régime de crise (régime à forte volatilité des marchés). Cette matrice reflètera *a priori* la perception subjective de l'investisseur en fonction des risques auxquels il est exposé. La matrice de corrélation dans le régime à faible volatilité sera considérée comme celle reflétant la stabilité des marchés (estimée par exemple à partir de données historiques hors périodes de crises). Une illustration numérique de cette démarche est présentée.

Concernant les modèles d'allocation d'actifs, notre étude est axée sur la comparaison des modèles disponibles selon les hypothèses sous-jacentes de « rebalancement » des portefeuilles : nous faisons la distinction entre une gestion « statique » (pour laquelle les poids reviennent périodiquement à ceux de l'allocation stratégique de long-terme - cas par exemple des modèles à poids constants *Fixed-Mix*) et une gestion dite « dynamique », pour laquelle les poids peuvent s'écarter définitivement de l'allocation stratégique initiale selon des règles de gestion prédéfinies.

Nous verrons que l'approche dynamique présente l'avantage théorique de la robustesse face aux changements de régime des marchés. L'autorisation du changement des poids des différentes classes d'actifs, sur la base d'une règle de gestion bien définie, constitue *a priori* un élément intéressant. Cela en effet permet l'ajustement des expositions aux différentes classes d'actifs suite à l'évolution des conditions de marché.

La gestion dynamique de portefeuille sur le long terme reste un domaine de recherche relativement peu exploré, par comparaison avec l'importance des travaux déjà réalisés sur les aspects à court terme. Nous pouvons néanmoins citer, par exemple, les études de Hainaut et al. [2005], Hainaut et al. [2007], Rudolf et al. [2004] et Yen et al. [2003], sachant que ces différents auteurs font part des difficultés pratiques d'implémentation, en raison de la complexité des modèles proposés.

Cette réflexion sur la construction de modèles d'allocation d'actifs applicables dans une optique prévisionnelle à long-terme nous conduira à étudier une approche innovante fondée sur les techniques de « programmation stochastique » (cf. Birge et Louveaux [1997]). Il s'agit d'une version adaptée d'une technique déjà utilisée dans le domaine de l'ingénierie pour la planification de la production (cf. Dantzig et al. [1990], Escudero et al. [1993]). L'objectif sera la mise en évidence et l'étude des caractéristiques de cette approche.

Comme déjà mentionné, cette thèse accorde également une attention toute particulière aux techniques numériques de recherche de l'optimum, qui demeurent des questions essentielles pour la mise en place d'un modèle d'allocation. Le point de départ sera notre constat d'un temps de calcul significatif dû simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Dans ce cadre, nous présentons un algorithme d'optimisation globale d'une fonction non convexe et bruitée. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement.

L'ensemble de ce travail se place dans le contexte des régimes de retraite par répartition en France : par suite les différentes applications proposées tiendront compte des spécificités de ces régimes. De même, la faisabilité opérationnelle est un des objectifs « cibles » que nous gardons à l'esprit tout au long de cette étude. Enfin, les trois années de préparation de cette thèse ont donné lieu à la rédaction de trois articles, portant respectivement sur les générateurs de scénarios économiques (cf. Faleh, Planchet et Rullière [2010] « Les générateurs de scénarios économiques : de la conception à la mesure de la qualité », Assurances et gestion des risques, n° double avril/juillet, Vol. 78 (1/2)), les techniques numériques d'optimisation (cf. Rullière, Faleh et Planchet [2010] « Un algorithme d'optimisation par exploration sélective », soumis) et l'allocation stratégique d'actifs avec les techniques de programmation stochastique (cf. Faleh [2011] « Un modèle de programmation stochastique pour l'allocation stratégique d'actifs d'un régime de retraite partiellement provisionné », soumis).

Importance de l'allocation stratégique d'actifs

L'allocation stratégique d'actifs d'un régime de retraite (ou d'une compagnie d'assurance en général) est souvent définie comme une étape d'un processus plus général de la gestion actif-passif, en particulier comme l'étape en aval de l'appréhension des risques et en amont de l'allocation tactique d'actifs. Compte tenu de ce positionnement, l'allocation stratégique d'actifs vise, soit à confirmer l'optimalité de la structure de l'actif existant de la réserve, soit à proposer une structure optimale d'actifs de cette réserve qui permette au régime d'atteindre un

certain objectif de performance financière tout en respectant ses engagements avec un niveau de confiance donné. La caractérisation de ‘stratégique’ vient d’une part de l’horizon temporel auquel s’appliquent les études d’allocation stratégique, d’autre part du nombre limité de classes d’actifs considérées dans ces études, généralement limité entre trois et dix au maximum.

Après avoir déterminé l’ensemble des allocations d’actifs possibles (allocations à tester), la résolution du problème de détermination de l’allocation stratégique d’actifs pour le régime de retraite étudié passe en pratique par trois étapes principales. La première étape consiste à **générer des trajectoires** pour chaque classe d’actifs (actions, obligations, immobilier, etc.). La deuxième consiste à projeter chaque allocation possible en fonction de **la règle de gestion** du dispositif. La troisième étape vise à déterminer la valeur de **la fonction objectif** pour chacune de ces allocations. Seront susceptibles d’être retenues alors celles qui à la fois maximisent la fonction objectif et respectent les contraintes fixées.

L’importance de l’allocation stratégique pour les investisseurs à long terme est mise en évidence dans l’étude de Brinson et al. [1991] comme le montre le graphique suivant :

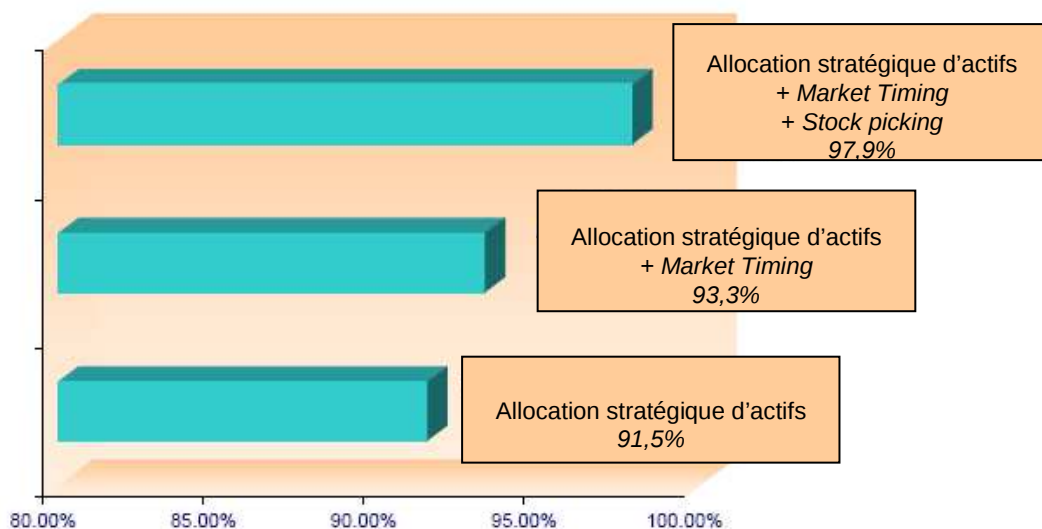


Fig. 1 : Pourcentages de la performance globale expliqués par certaines composantes de l’allocation d’actifs selon l’étude de Brinson et al. [1991]

Selon cette étude, la gestion passive d’un indice de référence répliquant le jeu de poids stratégiques du portefeuille suffit pour obtenir l’essentiel de la performance espérée de ce portefeuille (autrement dit l’évolution de la rentabilité d’un portefeuille est à 91,5 % le résultat de l’évolution des classes d’actifs sur lesquelles il est investi). L’allocation tactique qui permet de bénéficier des inefficiences temporaires du marché des titres et qui peut avoir pour source soit une sélection optimale de titres différents de ceux de l’indice (*stock picking*) soit le choix optimal des dates d’achat et de vente de ces titres (*market timing*), n’explique que près de 6,5 % de la performance globale obtenue *in fine* sur le portefeuille.

L’allocation stratégique doit tenir compte de l’évolution dans le temps des rendements des actifs et de leur structure de dépendance. Cela passe par la sélection de la dynamique de chaque variable pour ensuite effectuer l’estimation des paramètres du modèle choisi (le

calibrage du modèle). Nous aboutissons ainsi à la mise en place d'un modèle de génération de scénarios économiques de long terme permettant la projection aussi bien de la valeur de l'actif que celle du passif (le cas échéant).

Allocation stratégique et générateurs de scénarios économiques

La projection sur le long terme des valeurs de marché des actifs financiers et des variables macro-économiques, souvent appelée « génération de scénarios économiques », constitue une phase cruciale dans le processus d'allocation stratégique des fonds d'une compagnie d'assurance ou d'un fonds de retraite. Elle est un élément central de l'évaluation des provisions pour les garanties financières sur des contrats d'épargne dans le cadre de la directive Solvabilité II (cf. Planchet [2009]).

Un générateur de scénarios économiques (GSE) s'avère ainsi un outil important d'aide à la décision dans le domaine de la gestion des risques, en permettant d'obtenir des projections dans le futur des valeurs des éléments qui sont présents dans les deux compartiments du bilan de la société : les actifs (actions, obligations, immobilier,..) et le passif (provisions techniques, dettes financières,..). L'obtention de ces valeurs passe par la projection d'autres variables macro-économiques et financières telles que les taux d'intérêt, l'inflation des prix, l'inflation des salaires et le taux de chômage. A titre d'exemple, les prix futurs des obligations sont déduits à partir de la projection des taux d'intérêt (cf. Ahlgrim et al. [2008]), de même les prestations futures du régime de retraite peuvent être indexées sur l'inflation des prix et/ou l'inflation des salaires (cf. Kouwenberg [2001]).

Du point de vue opérationnel, la mise en œuvre de ces projections s'est basée initialement sur des méthodes déterministes, dans des logiques de scénarios, pour tester le comportement des objectifs techniques dans différentes situations considérées comme caractéristiques. Grâce au développement concomitant de l'outil informatique et des techniques de simulation, il est devenu possible de générer à moindre coût de très nombreux scénarios en tenant compte des interactions entre les multiples sources de risque et de leurs distributions, dans une perspective probabiliste.

La construction et la mise en œuvre d'un GSE passent par les quatre étapes suivantes (voir par exemple Hibbert et al. [2001] ou Ahlgrim et al. [2005]) : la première étape consiste en l'identification des sources de risque prises en compte et des variables financières à modéliser (taux d'intérêt, inflation, rendement des actions, etc.) qu'on appellera dans la suite les variables du GSE. Ensuite est effectué le choix du modèle pour la dynamique de chacune de ces variables. La troisième étape consiste à sélectionner une structure de dépendance entre les sources de risque de façon à obtenir des projections cohérentes. Ensuite l'estimation et le calibrage des paramètres des modèles retenus doivent être effectués. Enfin l'analyse des résultats obtenus de chaque GSE se fait en termes probabilistes, en analysant la distribution d'indicateurs clés tels que le surplus, ou la valeur nette de l'actif.

Construire un GSE pertinent pour l'ensemble des problématiques techniques n'est pas chose facile, et en pratique les objectifs des décisionnaires en termes de choix de gestion ont une incidence sur la manière de structurer le générateur (cf. Ahlgrim et al. [2008]). Deux niveaux de difficultés peuvent être distingués : celui relatif à la conception théorique d'un GSE et celui relatif à sa mise en œuvre pratique. Il est donc particulièrement important de définir avec précision les principaux éléments qui caractérisent un GSE ainsi que les indicateurs de mesure de sa qualité et de sa performance.

Allocation stratégique dans le cadre de l'ALM

De son côté, la gestion actif-passif, ou *Asset Liability Management* (ALM), consiste dans une méthode globale et coordonnée permettant à une société et notamment à un régime de retraite partiellement provisionné, de gérer la composition et l'adéquation de l'ensemble de ses actifs et passifs. Les techniques utilisées pour sa mise en place diffèrent particulièrement en fonction de la nature des engagements du régime.

Récemment, la gestion actif-passif s'est imposée pour les fonds de pension comme une approche de gestion des risques qui tient compte des actifs, des engagements et aussi des différentes interactions existantes entre ces deux parties (cf. Adam[2007]). Les gérants des fonds doivent déterminer les stratégies admissibles qui garantissent avec une probabilité suffisante que la solvabilité du fonds est assurée (compte tenu des prestations attendues). La solvabilité est définie comme la capacité du fonds à payer les prestations sur le long terme.

La solvabilité du fonds à une date fixée est souvent mesurée par le ratio de financement (le *funding ratio*). Il s'agit du ratio de la valeur des actifs financiers du fonds par rapport à la valeur des engagements actualisés à un taux de rendement choisi. Le sous financement (*under-funding*) aura lieu lorsque ce ratio est inférieur à 1. Une autre façon d'exprimer le sous financement est de dire que le surplus est négatif. Ce dernier représente en fait l'écart entre la valeur des actifs et la valeur des engagements actualisés.

Notons aussi que la définition de la notion d'actif ou d'engagement diffère selon l'approche utilisée (comptable ou économique) et selon le type de régime (capitalisation, répartition provisionné ou pas) : cela influe sur le calcul du ratio de financement.

Les fonds des régimes de retraite sont plus ou moins exposés à plusieurs facteurs de risque, dont le plus important pour le gérant de fonds est le risque de sous financement (*risk of under-funding*) : c'est le risque que la valeur des engagements soit supérieure à la valeur des actifs. Les sources de ce risque peuvent être classées en deux catégories : les risques financiers (risque des actions, risque de taux, risque de crédit,...) et les risques actuariels (risque de longévité, risque de taux technique supérieur aux taux observés sur le marché,...).

En effet, le niveau du ratio de financement change au cours du temps, essentiellement à cause des fluctuations de la valeur de marché des différentes composantes du bilan. Comme conséquence, les régimes de retraite provisionnés sont, par exemple, amenés à rebalancer leur allocation d'actifs et/ou à ajuster le taux de cotisation de leurs affiliés afin de mieux contrôler le changement du niveau de ce ratio.

Dans la littérature, les modèles d'ALM sont généralement classés en trois groupes présentés chronologiquement comme suit.

Le premier groupe contient les modèles d'adossement (ou *matching*) et d'immunisation par la duration (cf. Macaulay [1938], Redington [1952]). Ces modèles se basent sur le fait que les investissements sont essentiellement effectués dans des obligations. Ceci nous permet d'obtenir, soit un adossement des flux de trésorerie des actifs financiers à ceux du passif (*matching*), soit un adossement de la duration de l'actif à celle du passif (immunisation par la duration).

Ces techniques étaient utilisées jusqu'au milieu des années 80 et avaient comme inconvénients principaux la considération du risque de taux comme seule source de risque pour le fonds, ainsi que la nécessité d'un rebalancement périodique du portefeuille en ré-estimant à chaque fois la duration du passif, qui change continûment du fait du changement des taux d'intérêts et de l'écoulement du temps. Pierre [2009] présente une approche de la couverture de passif qui vise à résoudre les problèmes mentionnés ci-dessus. L'actif sera constitué dans ce cas d'un portefeuille de couverture de type taux (obligations ou dérivés) couplé à un portefeuille de rendement. Selon Pierre [2009], cette architecture permet une flexibilité suffisante pour permettre à l'actif de s'adapter aux mise à jour de la valeur des engagements lors de la revue des hypothèses actuarielles ayant permis leur évaluation.

Le deuxième groupe contient les modèles basés sur la simulation de scénarios déterministes et sur la notion de surplus (Kim et Santomero [1988], Sharpe et Tint [1990], Leibowitz et al. [1992]). Les modèles de surplus ont pour objet la minimisation du risque de perte du surplus (mesuré par la variance de la rentabilité du surplus) sous contraintes de rentabilité et de poids des actifs. Ils sont des modèles mono-périodiques, ce qui limite leur utilité en pratique pour des problèmes d'allocation sur le long terme.

Le troisième groupe de modèles utilise les techniques de simulation stochastique (*Monte Carlo*) pour modéliser l'évolution des différents éléments, que ce soit au niveau des actifs financiers et des engagements, ou au niveau des variables de marché et des variables démographiques (cf. Frauendorfer [2007], Munk et al. [2004], Waring [2004], Martellini [2006]). Ainsi les lois de probabilité associées aux résultats du fonds de retraite sur le long terme peuvent être estimées. A ce niveau, nous nous proposons de distinguer deux sous-groupes de modèles d'ALM basés sur les techniques stochastiques. L'élément clé de distinction sera si oui ou non les poids des différents actifs reviennent périodiquement à ceux de l'allocation stratégique définie initialement (si oui, les modèles seront appelés modèles à poids constants ou stratégie *Fixed-Mix*).

- Pour le premier sous-groupe de modèles à poids constants et malgré les avancées réalisées avec ces techniques (surtout au niveau de l'implémentation informatique), l'aspect dynamique de l'allocation stratégique reste encore marginalisé. En fait, ces modèles permettent de comparer des allocations constantes dans le temps (statiques) indépendamment des opportunités liées aux évolutions inter-temporelles des marchés (cf. Merton [1990], Dempster et al. [2003], Infanger [2002], Brinson et al. [1991] et Kouwenberg [2001]).
- Le deuxième sous-groupe de modèles, et le plus récent, est principalement inspiré de la théorie du choix de la consommation et de portefeuille développée par Merton [1971]. Il s'agit des modèles d'allocation dynamique ou inter temporels. Par exemple, à partir de la définition de la fonction objectif pour l'investisseur, ces modèles permettent la détermination d'une trajectoire des poids des différents actifs jusqu'à la date d'échéance (l'ajustement des poids est fonction des évolutions projetées du marché et de la règle de gestion prédéfinie). L'allocation stratégique retenue sera l'allocation optimale d'actifs à la date initiale t_0 . Le cadre d'utilisation de ces modèles récents se heurte au problème d'implémentation vu la complexité des outils mathématiques employés (cf. Hainaut et al. [2005], Hainaut et al. [2007], Rudolf et al. [2004] et Yen et al. [2003]).

Le nécessaire positionnement de l'analyse des risques financiers par rapport au cadre Solvabilité II

Solvabilité II est une directive-cadre qui a pour but de mettre à jour le système européen de solvabilité des entreprises d'assurance. Cette directive vise en particulier une meilleure adaptation des capitaux propres exigés des compagnies d'assurances et de réassurance avec les risques que celles-ci encourent dans leur activité. Sa mise en application est prévue au 31 octobre 2012 lorsque les discussions entre la commission européenne et le parlement européen auront abouti et que cette directive aura été retranscrite dans les législations nationales par chaque parlement.

Dans cette directive, une formule standard permet le calcul du capital nécessaire pour couvrir les risques supportés par les assureurs et les réassureurs notamment suite à un choc provoqué par un événement exceptionnel (ce niveau de capital est également appelé capital de solvabilité requis ou SCR en anglais pour *Solvency Capital Requirement*). Cela est dans le but de contrôler la probabilité de ruine à un an et de la limiter à moins de 0,5 %. Une autre alternative proposée aux assureurs sera l'utilisation d'un modèle interne complet (basé sur leur structure de risque spécifique). Dans ce cas, une validation de l'autorité de contrôle sera requise préalablement à la détermination effective du SCR à partir du modèle interne.

Le calcul du SCR tient compte d'une panoplie de facteurs de risque (taux, action, *spread*, concentration, liquidité, etc.) agrégés sous forme de modules et ce en fonction du domaine d'activité de la compagnie. Le projet de spécifications techniques pour la cinquième étude³ quantitative d'impact (QIS5) précise que l'analyse des risques se décline en quatre catégories principales que l'on retrouve au sein des différents modules :

- le risque de marché, provenant de l'incertitude associée à la valeur et aux rendements des actifs financiers.
- le risque de souscription, provenant de l'incertitude liée à la mesure des engagements pris par l'assureur en vie, en santé et en non vie.
 - le risque de contrepartie, lié à au défaut potentiel des contreparties.
 - le risque opérationnel comprenant l'ensemble des risques associés aux procédures de gestion interne et aux conséquences d'un dysfonctionnement à ce niveau.

Chacune de ces grandes catégories de risques se retrouve éclatée dans les différents sous modules qu'il faut ensuite agréger, de manière à tenir compte de la dépendance entre les risques (cf. graphique 2).

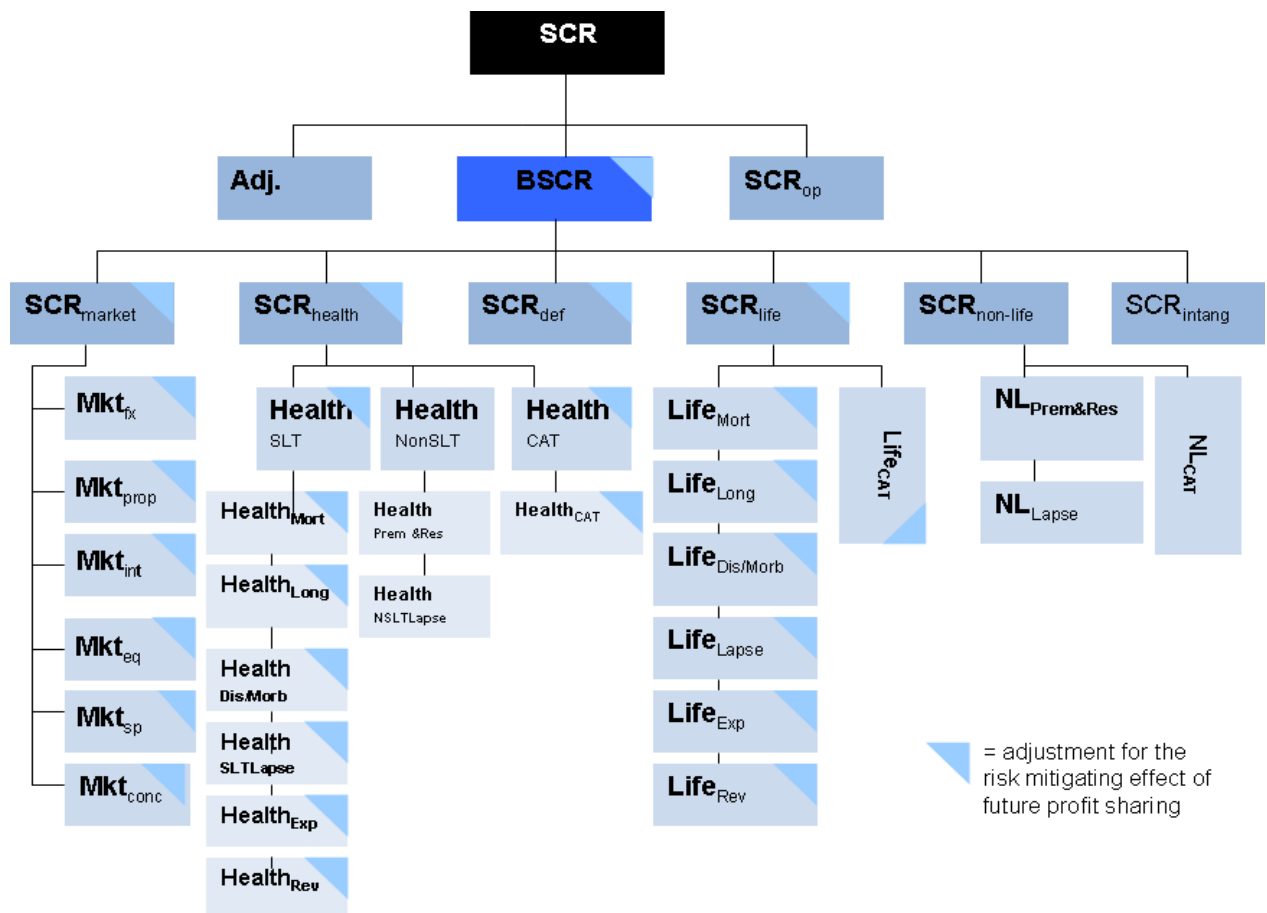
L'application de Solvabilité II aux opérations de retraite semble problématique, étant donné la particularité de ces dernières :

- Horizon de gestion : La probabilité de ruine à un an n'est pas forcément une mesure de risque pertinente pour juger de la solvabilité de l'activité de retraite, pour laquelle les engagements se font sur le long terme. Ceci laisse donc une grande marge de pilotage pour les régimes de retraite susceptibles d'amortir sur le long terme les risques à court terme.

³ Afin de mesurer leurs implications concrètes, les directives européennes doivent faire l'objet d'une étude d'impact. C'est dans ce cadre que s'inscrit la cinquième étude d'impact (QIS 5 acronyme anglais pour *Quantitative Impact Studies* 5) pour Solvabilité II.

- Politique d'investissement : Une gestion adaptée aux contrats retraite tournée vers des actifs risqués induit un besoin en capital de solvabilité élevé. La détention d'actifs réels ou risqués, tels que les actions ou l'immobilier, est fortement pénalisée dans la formule standard, alors qu'il s'agit de supports qui présentent des avantages notables sur le long terme pour optimiser le couple rendement/risque.

Toutefois, le positionnement par rapport à la grille d'analyse des risques proposée par Solvabilité II est incontournable compte tenu de la référence que constitue ce référentiel (même si le présent travail n'est pas dans un contexte de Solvabilité II). De ce fait, l'analyse des risques supportés par un régime de retraite sera effectuée dans la suite conformément à cette grille de lecture.



Source : Projet de spécifications techniques pour la cinquième étude quantitative d'impact (QIS5)

Fig. 2 : Les modules de la formule standard du SCR

Présentation détaillée de la thèse

Cette thèse est composée de deux parties. La première partie est constituée de deux chapitres. La deuxième partie est constituée de trois chapitres.

La partie I traite des générateurs de scénarios économiques (GSE).

Le chapitre 1 de la partie I. Il reprend largement les points discutés dans l'article de Faleh et al. [2010].

Au niveau de ce chapitre, notre objectif est d'exposer l'état de l'art en matière d'identification des risques à intégrer à un GSE. Nous proposons ainsi différents critères de classement ainsi qu'une étude comparative des caractéristiques théoriques et des composantes de différents modèles de GSE.

L'intérêt est porté, dans un second temps, sur les aspects relatifs à la mise en œuvre d'un GSE. Nous présentons donc les caractéristiques des différentes structures schématiques de projection de scénarios utilisées en pratique. Toujours dans le cadre de la mise en œuvre, nous nous inspirons des travaux de Mitra [2006] et de Zenios [2007] pour recenser différentes méthodologies de génération des scénarios (les modèles ayant trait à la dynamique des variables du GSE).

Ensuite, nous étudions une série d'indicateurs, qualitatifs et quantitatifs, pour la mesure de la qualité d'un GSE tout en étudiant leurs limites. Nous nous référons pour cela au travail de Kaut et al. [2003]. Enfin, les principaux éléments exposés tout au long de ce chapitre sont illustrés à travers une application numérique.

Le chapitre 2 de la partie I. Un générateur de scénario adapté aux spécificités de la gestion ALM d'un régime de retraite est construit. Les éléments suivants relatifs à ce GSE sont détaillés : sa conception théorique, les modèles de diffusion retenus, les approches de calibrage de ses paramètres ainsi que sa mise en œuvre.

Notre objectif à ce niveau est d'utiliser le GSE construit en tant que moyen pour illustrer un certain nombre de résultats obtenus, tout en gardant à l'esprit les spécificités de la gestion actif-passif d'un régime de retraite. Quelques améliorations par rapport à ce qui se pratique usuellement sont apportées.

En particulier, notre étude se distingue par l'adoption d'une démarche novatrice dont l'objectif est de tenir compte de l'instabilité temporelle de la corrélation entre les rentabilités des actifs. Nous nous inspirons de l'approche classique de changement de régime (cf. Hamilton [1989], [1994] et Hardy [2001]), qui concerne souvent les rendements moyens ainsi que les volatilités, pour étendre son principe et supposer un changement de régime au niveau des matrices de corrélations aussi.

D'un autre côté, différents modèles de diffusion et de calibrage sont envisagés pour le panel des variables financières et macro-économiques qui interviennent dans ce contexte. Nous comparons les résultats d'estimation des paramètres du GSE avec ceux mentionnés dans l'étude d'Ahlgrim et al. [2005]. Ensuite, nous étudions les résultats de projection obtenus tout en proposant un certain nombre de tests pour juger de la cohérence de ces résultats.

La seconde partie est consacrée à l'élaboration de l'allocation d'actif elle-même.

Le chapitre 1 de la partie II. Il est question dans ce chapitre d'analyser les approches classiques (ou déterministes) de gestion actif-passif à savoir les modèles basés sur la notion de « duration » et de « surplus ». Nous détaillons différentes approches tout en mettant en évidence la différence entre elles. L'objectif de ce chapitre est de revoir l'état de l'art en matière de modèles d'ALM déterministe. Après avoir mis en évidence les limites de ces modèles, nous passons à l'étude de modèles plus élaborés et plus sophistiqués (objet des chapitres ultérieurs).

Le chapitre 2 de la partie II. Nous considérons ici une approche récente d'allocation stratégique d'actifs basée sur la stratégie dite « à poids constants » ou *Fixed-Mix* (cf. Kouwenberg [2001]).

Nous proposons une modélisation du régime-type de retraite et étudions certains critères d'allocation stratégique d'actifs en fonction du type du régime (provisionné, partiellement provisionné, etc.) sont étudiés. Nous passons ensuite à l'illustration de la stratégie *Fixed-Mix* avec une application sur les réserves d'un régime de retraite partiellement provisionné. Les résultats obtenus sont discutés et différents tests de sensibilité sont mises en place : ces tests sont liés principalement à l'impact des hypothèses de rendement ou de corrélation retenues.

Les difficultés rencontrées lors de la mise en place de la stratégie *Fixed-Mix*, dues essentiellement à la multiplicité du nombre de classes d'actifs considérées et au nombre de scénarios économiques simulés, nous ont mené à nous pencher sur les aspects d'optimisation numérique. Le point de départ est notre constat d'un temps de calcul significatif dû simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Dans ce cadre, une présentation détaillée des travaux menés lors de la rédaction de l'article de Rullière et al. [2010] est effectuée. Un algorithme d'optimisation globale d'une fonction non convexe et bruitée est présenté. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement et compare les résultats obtenus avec l'algorithme de Kiefer-Wolfowitz- Blum (cf. Blum [1954], Kiefer et Wolfowitz [1952]).

Le chapitre 3 de la partie II.

Les modèles classiques d'ALM dynamiques sont explorés, notamment les techniques d'assurance de portefeuille (cf. Perold et Sharpe [1988]) et les techniques de programmation dynamique (cf. Cox et Huang [1989], Merton [1971]). A ce niveau, les techniques d'assurance de portefeuille basées sur la notion de CPPI ou *Constant Proportion Portfolio Insurance* (cf. Perold et Sharpe [1988]) sont mises en place et certains résultats relatifs à ce modèle sont étudiés. De même, les principes des techniques de programmation dynamique et leurs limites sont également mises en évidence.

Nous nous penchons par la suite sur une approche innovante fondée sur les techniques de « programmation stochastique » (cf. Birge et Louveaux [1997]). Il s'agit d'une version adaptée d'une technique déjà utilisée dans le domaine de l'ingénierie pour la planification de la production (cf. Dantzig et al. [1990], Escudero et al. [1993]). Dans ce cadre, nous mettons en place un modèle d'ALM dynamique basé sur les techniques de programmation stochastique.

Nous proposons, au cours d'une illustration numérique, une nouvelle méthodologie de génération de scénarios économiques que nous appelons méthodologie « des quantiles de référence ». Cette dernière permet de partir d'une structure linéaire de génération de scénarios (telle que décrite dans le chapitre 2 de la partie I) pour réduire la dimension du problème rencontré avec la stratégie *Fixed-Mix* tout en tenant compte de la corrélation entre les distributions des différentes variables projetées. Cela s'insère dans le cadre de la recherche d'une vision à la fois simplifiée, réelle et dynamique des stratégies possibles pour l'allocation d'actifs sur le long terme.

A travers la même application numérique, nous comparons certains résultats relatifs aux deux approches d'allocation stratégique d'actifs : celle basée sur la stratégie *Fixed-Mix* et celle basée sur les techniques de programmation stochastique. Nous testons également la sensibilité de cette deuxième approche par rapport au changement de certains de ses paramètres, toutes choses étant égales par ailleurs. Cette étude sur les techniques de programmation stochastique, dans le contexte de la gestion des réserves des régimes de retraite (en particulier ceux partiellement provisionnés), reprend largement les points évoqués dans l'article de Faleh [2011].

Tout au long de ce travail, nous mettons en évidence le lien effectif entre les recherches académiques et les besoins des systèmes de retraite en matière de gestion des risques et de respect des engagements vis-à-vis des futurs retraités.

Partie I :

**LES GÉNÉRATEURS DE SCÉNARIOS
ÉCONOMIQUES (GSE)**

Introduction

Au départ, l'analyse des risques futurs pour les compagnies d'assurance et les fonds de retraite (ou de pension, selon la terminologie anglo-saxonne) se faisait « à la main » en essayant de répondre à la question « *et si jamais..?* » (étude de scénarios déterministes). Ceci a été suivi par des études basées sur l'adossement du passif par l'actif que ce soit au niveau des flux de trésorerie futurs ou au niveau des durations de ces deux compartiments du bilan. Grâce au développement de l'outil informatique et aux techniques de simulation de *Monte Carlo*, il est devenu possible de générer des milliers de scénarios économiques tenant compte des corrélations entre les différentes sources de risque pour ensuite analyser ces résultats en termes probabilistes.

Les travaux académiques antérieurs à 1984 traitent souvent une partie du problème rencontré par les actuaires. Les travaux sont concentrés uniquement soit sur les actions, soit sur les taux d'intérêt, soit sur l'inflation. Il y a peu de travaux qui mettent toutes ces variables dans un seul modèle en tenant compte des différentes interactions entre elles. Le travail de Wilkie [1984] constitue sans doute une référence à ce niveau. Ce professeur a présenté pour la première fois un modèle général qui inclut toutes les variables macro économiques et financières. De même, son modèle a été simple de point de vue de son implémentation, ce qui l'a rendu populaire et pendant deux décennies il a été la référence de tous les modèles postérieurs proposés.

Ce type de modèle, appelé par la suite « générateurs de scénarios économiques GSE » (en anglais *Economic Scenario Generator* ou ESG) permet par exemple de prendre en compte l'horizon long d'investissement d'un fonds de retraite provisionné et de contrôler l'influence des évolutions futures sur le choix de ses paramètres techniques (taux de cotisation, taux de prestation, etc.). Ce contrôle est souvent effectué dans le but de garantir un écart positif permanent entre la valeur des actifs financiers à une date donnée et la valeur des engagements actualisés à la même date via un taux d'intérêt de référence (l'écart est appelé le surplus). Il est noté à ce titre que les GSE constituent le cœur des modèles de gestion actif-passif appartenant à la génération des modèles d'ALM stochastiques. Le choix de l'allocation stratégique vient dans un deuxième temps refléter la fonction d'utilité de l'investisseur de long terme.

Le premier chapitre de cette partie reprend largement les points discutés dans l'article de Faleh et al. [2010]. Notre objectif est d'exposer l'état de l'art en matière d'identification des risques à intégrer à un GSE. Nous proposons ainsi différents critères de classement ainsi qu'une étude comparative des caractéristiques théoriques et des composantes de différents modèles de GSE.

L'intérêt est porté, dans un second temps, sur les aspects relatifs à la mise en œuvre d'un GSE. Nous présentons donc les caractéristiques des différentes structures schématiques de projection de scénarios utilisées en pratique. Toujours dans le cadre de la mise en œuvre, nous nous inspirons des travaux de Mitra [2006] et de Zenios [2007] pour recenser différentes méthodologies de génération des scénarios (les modèles ayant trait à la dynamique des variables du GSE).

Ensuite, nous étudions une série d'indicateurs, qualitatifs et quantitatifs, pour la mesure de la qualité d'un GSE tout en étudiant leurs limites. Nous nous référons pour cela au travail de Kaut et al. [2003]. Enfin, les principaux éléments exposés tout au long de ce chapitre sont illustrés à travers une application numérique.

Chapitre 1 : Etude des GSE

I- Présentation théorique

Cette section s'intéresse aux problèmes théoriques de modélisation liés aux GSE et présente une revue de la littérature concernant leurs différentes structures. L'objectif étant de mettre en évidence les composantes essentielles d'un générateur de scénarios économiques au niveau de sa conception théorique.

I-1 Problèmes théoriques de modélisation liés aux GSE

Comme mentionné ci-dessus, certaines problématiques liées à la conception théorique des GSE sont exposées. Ces dernières traitent de questions qui permettent l'amélioration de la performance des résultats et le développement du modèle pour l'adapter au contexte d'étude. Les problèmes de modélisation dans les GSE sont en premier lieu liés aux choix du modèle de la structure par terme des taux d'intérêt (STTI). La modélisation de la dynamique des rendements des actions vient dans un deuxième temps susciter elle aussi pas mal de questions.

I-1-1 Modèles des taux d'intérêt

En se basant sur les travaux de Roncalli [1998] et Décamps [1993], deux grandes classes de modèles de taux pour l'évaluation d'actifs financiers peuvent être présentées : les modèles d'absence d'opportunité d'arbitrage (AOA) et les modèles d'équilibre général.

I-1-1-1 *Les modèles d'absence d'opportunité d'arbitrage (AOA)*

L'évaluation dans le cadre des modèles d'AOA est de nature purement financière. Elle repose entièrement sur l'hypothèse d'absence d'opportunité d'arbitrage. L'utilisation de cette hypothèse fondamentale a permis de mettre en évidence deux approches d'évaluation par arbitrage :

- La première approche considère le prix des instruments financiers de taux comme fonction de variables d'état. Ces variables d'état –généralement le taux court pour les modèles univariés (cf. Vasicek [1977]), ou le couple (taux court, taux long) pour les modèles bivariés (cf. Brennan et Schwartz [1982]) sont supposées de dynamique exogène. Dans de tels modèles, l'hypothèse d'absence d'opportunité d'arbitrage exprime que la prime de risque du marché⁴ est indépendante de la maturité du titre considéré. Le prix d'un produit obligataire est ensuite obtenu comme solution d'une équation aux dérivées partielles. La résolution probabiliste de ces équations permet une meilleure interprétation financière des formules d'évaluation.
- La deuxième approche d'évaluation par arbitrage est celle proposée par Ho et Lee [1986] et Heath, Jarrow et Merton [1992]. Le point clef de cette approche est la prise en compte de toute l'information contenue dans la structure de taux initiale en considérant comme donnée exogène la dynamique simultanée de taux ayant différentes maturités appelée aussi dynamique des taux terme contre terme. Cette

⁴ Supplément de rendement exigé par les investisseurs pour avoir assumé le risque de détenir des actifs risqués plutôt que des actifs sans risque.

dynamique est choisie de façon à ce que l'hypothèse d'absence d'opportunité d'arbitrage soit respectée. Contrairement à l'arbitrage traditionnel la prime de risque du marché n'est pas spécifiée de façon exogène mais elle est définie implicitement par la dynamique des taux terme contre terme. Par ailleurs, aucune hypothèse sur la forme spécifique des prix des produits obligataires comme fonction d'une variable d'état n'est faite. Enfin, la dynamique de prix des zéro-coupons caractérise entièrement le modèle de courbe de taux.

Les modèles d'absence d'opportunité d'arbitrage sont particulièrement appropriés pour évaluer les prix des produits dérivés. Comme les dérivés sont évalués à partir des actifs sous jacents, un modèle qui capte explicitement les prix de marché de ces actifs est *a priori* plus performant qu'un modèle qui ne les prends pas en compte. Ce point a été remarqué notamment par Hull [2000] et Tuckman [2002] qui constatent que le recours à l'hypothèse AOA permet une évaluation des produits dérivés plus plausible que les approches d'équilibre.

Ho et Lee [1986] présentent un modèle en temps discret dans le cadre des hypothèses d'absence d'opportunité d'arbitrage. Ils supposent un *drift* (ou une tendance) de la variable dépendant du temps de façon à pouvoir répliquer les prix de toutes les obligations observées sur le marché.

L'équivalent du modèle de Ho-Lee en temps continu, pour le taux d'intérêt r_t , est :

$$dr_t = \theta(t)dt + \sigma dB_t$$

Avec B_t un mouvement brownien, (ce qui implique que pour deux dates t_1 et t_2 : $B(t_2) - B(t_1) \sim N(0, |t_2 - t_1|)$). Le *drift* dépendant du temps dans le modèle de Ho et Lee, $\theta(t)$, est sélectionné de façon à ce que les taux d'intérêt espérés convergent vers les anticipations données par le marché et soient reflétées dans la structure de taux observée initialement. Ce *drift* est lié principalement aux taux *forwards* implicites.

Heath, Jarrow et Morton [1992] (HJM) généralisent l'approche d'absence d'opportunité d'arbitrage : ainsi est prise en compte l'intégralité de la structure par terme et pas seulement le processus suivi par le taux court. Ils présentent alors la famille des processus de taux *forwards* $f(t, T)$ de la manière suivante :

$$df(t, T) = \mu(t, T, f(t, T))dt + \sigma(t, T, f(t, T))dB_t$$

$$\text{Avec : } f(t, T) = - \frac{\partial \ln P(t, T)}{\partial T}$$

$P(t, T)$: prix à la date t d'un zéro coupon de maturité T .

μ et σ : respectivement la tendance et la volatilité des taux *forwards*.

Le taux *forward* $f(t, T)$ est le taux, déterminé aujourd'hui, à une date future t et sur une durée future (période) $T - t$. HJM remarquent que le *drift* des taux *forwards* peut être exprimé en fonction des volatilités, ce qui conduit à faire de la volatilité le facteur prépondérant dans l'évaluation des prix des produits dérivés.

I-1-1-2 Les modèles d'équilibre général

L'évaluation dans le cadre d'un modèle d'équilibre général ne nécessite pas d'hypothèse sur les dynamiques de prix ou de taux. Ces dynamiques sont obtenues de manière endogène. Cette approche est celle de Cox, Ingersoll et Ross [1985] et de Campbell et al. [2005]. En effet, les modèles d'équilibre typiques se basent sur les anticipations des mouvements futurs des taux d'intérêt de court terme et non pas sur la courbe de taux observée à la date initiale. Ces mouvements peuvent être donc dérivés à partir d'hypothèses plus générales sur des variables d'état qui décrivent l'ensemble de l'économie. En utilisant un processus des taux courts, il est possible de déduire le rendement d'une obligation de long terme en déterminant la trajectoire espérée des taux courts jusqu'à la maturité de cette obligation. L'intérêt d'une telle approche est triple :

- s'assurer que les processus étudiés sont cohérents avec un équilibre général,
- analyser la déformation de la courbe des taux en fonction des chocs sur les variables économiques sous jacentes,
- fournir une spécification fondée pour l'expression des prix ou des risques de marchés utilisés dans la valorisation par arbitrage.

Un des principaux avantages du modèle d'équilibre est que les prix de différents actifs traditionnels ont des formules analytiques explicites (*closed-form analytic solutions*). Un autre avantage est que les modèles d'équilibre sont relativement faciles à utiliser. Mais, les modèles d'équilibre de structure par terme génèrent des prix de produits de taux qui sont potentiellement incohérents avec ceux observés sur le marché à la date initiale. Même si les paramètres de ces modèles peuvent être calibrés avec une précision significative, la structure par terme résultante peut générer des prix éloignés de ceux observés initialement sur le marché.

La formulation mathématique générale des modèles d'équilibre, dans le cas où il est considéré qu'un seul facteur explique ces mouvements, est la suivante :

$$dr_t = \kappa(\theta - r_t)dt + \sigma r_t^\gamma dB_t$$

Ce type de modèle continu est basé sur un seul facteur stochastique : le mouvement du taux d'intérêt instantané (court terme) r_t . La formule générale incorpore le phénomène dit « de retour à la moyenne ». Pour le comprendre, nous considérons le cas où le niveau actuel des taux courts r_t est supérieur au niveau d'équilibre anticipé θ . Dans ce cas, le changement du taux est prévu être négatif (baisse de r_t) de façon à converger vers θ . Dans l'autre cas, celui où le niveau de r_t est inférieur à θ , seule une augmentation des taux courts permet la convergence vers θ . Ainsi, quel que soit le niveau observé des taux courts, nous supposons une évolution de ce niveau vers θ .

La vitesse de retour à la moyenne est donnée par le paramètre κ . La deuxième partie de cette formule incorpore la volatilité inconnue des taux d'intérêt au cours du temps. Le dernier terme, dB_t , constitue la variation du mouvement brownien, de moyenne nulle et de variance dt .

L'incertitude est intégrée à travers le paramètre de volatilité σ :

- Si $\gamma > 0$: la volatilité des taux d'intérêt est liée à leur niveau.
- Si $\gamma = 0$: le modèle est équivalent au modèle de Vasicek [1977].
- Si $\gamma = 0,5$: le modèle est le processus proposé par Cox, Ingersoll, Ross [1985] appelé le modèle CIR.

Chan et al. [1992] estiment les paramètres de cette classe de taux d'intérêt et déterminent, à partir de données mensuelles allant de 1964 jusqu'à 1989, que la valeur de γ est approximativement de 1,5. Ces modèles sont appelés les modèles d'équilibre général car les investisseurs évaluent le prix de l'obligation en se référant principalement aux anticipations futures des taux d'intérêt et non à la courbe des taux observée initialement. En utilisant la trajectoire simulée du taux courts jusqu'à la maturité de ces obligations, il est possible de déterminer le rendement des obligations de long terme. Pour déterminer toute la structure par terme, l'investisseur pourra évaluer les prix des obligations de différentes maturités à partir de l'évolution anticipée des taux courts sur la durée de vie restante de l'obligation.

$$P(t, T) = E \left[\exp \left(- \int_t^T r_u du \right) \right]$$

Avec $P(t, T)$ le prix d'une obligation zéro coupon à l'instant t (qui paie 1 euro à la date T).

Un des avantages les plus remarquables des modèles d'équilibre, est que les prix des obligations et les prix de certains autres dérivés de taux ont des formules analytiques explicites. Vasicek [1977] et CIR [1985] partent de la formule ci-dessus pour retrouver les prix des obligations :

$$P(t, T) = A(t, T) e^{-r_t B(t, T)}$$

Avec $A(t, T)$ et $B(t, T)$ sont des fonctions des paramètres connus κ , σ et θ . Ainsi, étant donné une réalisation de r_t , les taux exigés sur différentes maturités peuvent être obtenus.

$$R_r(t, T) = -\log \{ P(t, T) \} / (T - t)$$

L'inconvénient majeur des modèles d'équilibre est que la STTI (cf. page 21) résultante peut être incohérente avec les prix de marché observés, même si les paramètres sont parfaitement calibrés (cf. Ahlgrim et al. [2005]).

Hull et White [1990] utilisent le principe de *drift* dépendant du temps pour Ho et Lee [1986] pour présenter une extension du modèle d'équilibre de Vasicek [1977] et de CIR [1985]. Le modèle à un seul facteur de Hull et White est :

$$dr_t = \kappa(\theta(t) - r_t)dt + \sigma dB_t$$

Le tableau suivant illustre les points clefs de ces modèles :

	Modèles d’Absence d’Opportunité d’Arbitrage (AOA)		Modèles d’Equilibre Général
	<i>Approches à prime de risque exogène</i>	<i>Approches à prime de risque endogène</i>	<i>Approches entièrement endogènes</i>
Courbe des taux initiale	- Utilisée.	- Utilisée.	- Non utilisée.
Dynamique des taux	Exogène pour des taux de maturité spécifié (court, long,...).	Exogène dans le cas des taux terme contre terme.	Endogène : Anticipation des taux court terme futurs à partir d’hypothèses générales sur les variables qui décrivent l’ensemble de l’économie.
Prime de risque	Exogène et indépendante de la maturité.	Endogène : déduite à partir de la dynamique des taux terme contre terme.	Endogène.
Prix des produits de taux (obligations, etc.)	- Prix fonction de variables d’état (taux courts, taux longs, etc.). - Prix solution d’équations aux dérivés partielles.	- Aucune hypothèse sur la forme spécifique des prix comme fonction de variables d’état.	- Prix déterminés à partir des projections de taux courts.

Tab.1 : Comparaison des caractéristiques des modèles de taux d’intérêts

Une autre problématique théorique relative aux différentes variables modélisées, en particulier les taux d’intérêt, concerne le choix du nombre et de la nature des facteurs à utiliser pour la modélisation. Dans le cas particulier des taux d’intérêt, le choix définitif de ces deux éléments dépend du contexte d’application du GSE (projection de grandeurs réelles, évaluation de produits dérivés,...). La nature des facteurs peut être de sources différentes : elle peut être par exemple liée à l’horizon du taux (taux court, taux long) ou à la structure de la courbe des taux (facteur de courbure, de translation,...). A ce titre, Date et al. [2009] montrent la significativité statistique du choix de deux facteurs (par exemple les taux courts et les taux longs) pour expliquer l’évolution de la courbe de taux.

I-1-2 Modèles de rendement des actions

Concernant les actions, différentes problématiques liées à l’appréhension de la dynamique de leurs rendements montrent l’insuffisance de certains modèles adoptés, jusqu’à une période récente, par les sociétés d’assurance et les fonds de retraite. Souvent les rendements des

actions sont assumés suivre un mouvement brownien. En particulier, Black et Scholes [1973] supposent que les rendements suivent un mouvement brownien géométrique. Ceci implique que sur chaque sous période de temps, les rendements des actions sont distribués selon une loi normale et ils sont indépendants, les prix des actions suivent un processus lognormal. Dans ce cadre, si nous supposons que S_t est le prix des actions à l'instant t et que S_{t_0} est le prix des actions à une date antérieure t_0 , alors :

$$\log(S_t / S_{t_0}) \sim N(\mu(t - t_0), \sigma^2(t - t_0))$$

pour une moyenne μ et une volatilité σ .

Ce modèle lognormal, simple et pratique, fournit des approximations raisonnables sur des périodes courtes mais il est moins adapté aux problématiques de long terme. L'examen des données historiques relatives aux rendements des actions permet de constater, par exemple, que l'hypothèse d'une distribution normale ne permet pas de prévoir des valeurs extrêmes de rendement tel que réalisé dans le passé (cf. Ahlgrim et al. [2005], Mandelbrot [2005] et Hardy [2001]).

Un nombre important de modèles alternatifs a été proposé. Alexander [2001] recense une variété de ces modèles, incluant les processus autorégressifs généraux à volatilité conditionnellement hétéroscédastique (GARCH, cf. Bollerslev [1986]) et les analyses par composante principale (cf. Roncalli [1998]). De même, certains chercheurs proposent l'adaptation d'un modèle de changement de régime dans lequel les rendements des actions peuvent être simulés sous l'un des deux régimes suivants : un premier régime avec une hypothèse de volatilité relativement faible et une moyenne relativement élevée des rendements des actions et un deuxième régime avec une hypothèse de volatilité relativement élevée et une moyenne relativement faible de ces rendements. Une probabilité de transition entre ces deux régimes peut être déterminée *a priori* (cf. Hardy [2001]). L'hypothèse d'avoir plusieurs régimes est également envisageable. D'autres alternatives de modélisation sont proposées notamment dans le cadre des modèles discontinus (cf. Merton [1976]).

I-1-3 Autres classes d'actifs

Concernant les autres classes d'actifs, tels que l'immobilier et les produits dérivés, malgré les différents cadres d'hypothèses proposés pour leur projection, elles se heurtent souvent aux problèmes d'un historique peu profond, d'une liquidité insuffisante et de données confidentielles (cas des fonds de couverture). En matière de projection de grandeurs réelles sur le long terme, ces actifs sont souvent traités avec prudence et ne suscitent pas la grande part de l'intérêt des décideurs, en particulier dans le cas de l'allocation stratégique d'actifs d'un fonds de retraite où la priorité est souvent donnée aux actions, aux obligations et au monétaire (cf. Campbell et al. [2001]). Ceci ne remet pas en cause le potentiel que présentent ces actifs en tant que source de performance et/ou de couverture supplémentaire pour le portefeuille financier de la société ou du fonds de retraite (cf. Ahlgrim et al. [2005]).

I-2 Littérature sur les structures des GSE

La littérature sur les GSE est abondante. Nous proposons ici de classer les différents modèles en fonction de la structure de dépendance entre les variables, en distinguant deux catégories : structure par cascade et structure basée sur les corrélations. De même, pour chacun des

modèles cités, nous préciserons l'objectif qui lui était associé lors de sa conception. A ce titre, il est noté que l'utilisation d'un GSE a souvent pour finalité soit la projection sur le long terme et la prise de décision dans le cadre de la gestion des risques (dans ce cas, l'intérêt porte sur des grandeurs et des valeurs réelles) soit l'évaluation des prix d'équilibre des produits financiers sur le court terme, dit *pricing*, afin de déterminer la stratégie de marché convenable (achat, vente, etc.).

I-2-1 Modèles à structure par cascade

Une structure par cascade est définie comme une structure dans laquelle nous partons de la détermination de la valeur d'une variable (par exemple l'inflation) pour ensuite déduire les valeurs des autres variables (taux réels, rendements des actions, etc.). Jusqu'au début des années 1980, les travaux académiques traitent souvent une partie seulement du problème rencontré par les actuaires : les travaux sont concentrés sur chacune des classes d'actifs financiers (les actions, les taux d'intérêt, l'inflation,..) indépendamment des éventuelles interactions entre elles. L'intégration de ces interactions et l'étude du choix du modèle de dynamique des actifs financiers devient nécessaire pour garantir la cohérence des projections par rapport à un contexte donné.

Le travail de Wilkie [1986] marque de ce point de vue un changement majeur. Il présente pour la première fois un modèle général qui inclut toutes les variables macro économiques et financières. Ce modèle a l'avantage d'être simple à implémenter, raison pour laquelle il est rapidement devenu populaire et considéré comme la référence de tous les modèles proposés durant les deux décennies postérieures, malgré ses nombreuses limites.

La première version du modèle de Wilkie a été appliquée dans le cadre de la mesure de la solvabilité d'une société d'assurance par la *Faculty of Actuaries* [1986]. De même, un des premiers domaines d'application du modèle de Wilkie en actuariat était l'évaluation des engagements indexés sur les actions : dans ce cas nous supposons que les prestations dépendent des prix futurs des actions et que la réserve est principalement investie en obligations. De façon générale, ce modèle est plutôt cohérent avec des logiques de besoin de capital et de projection de valeur (cas de la gestion actif-passif par exemple) qu'avec des logiques de *pricing*.

Wilkie se base sur une structure par cascade, telle que décrite ci-dessus : il postule que l'inflation est la variable indépendante – « la force motrice » – du modèle dont la détermination se fait en premier lieu pour ensuite en dériver les valeurs des autres variables, principalement les dividendes, les revenus de dividende, les taux d'intérêt et la croissance des salaires. Le graphique 3 illustre le principe d'une structure par cascade. Dans ce cadre, Wilkie utilise un modèle autorégressif de premier ordre pour l'inflation. En 1995, il met à jour ce premier modèle en gardant les principes de sa structure par cascade mais en optant cette fois-ci pour une modélisation de l'inflation par un processus ARCH (*Autoregressive Conditional Heteroscedasticity*, cf. Engle [1982]). Ceci a été justifié, selon Wilkie, par la capacité de ce type de processus à tenir compte des caractéristiques des distributions historiques des données observées sur le marché de la Royaume-Uni depuis 1919.

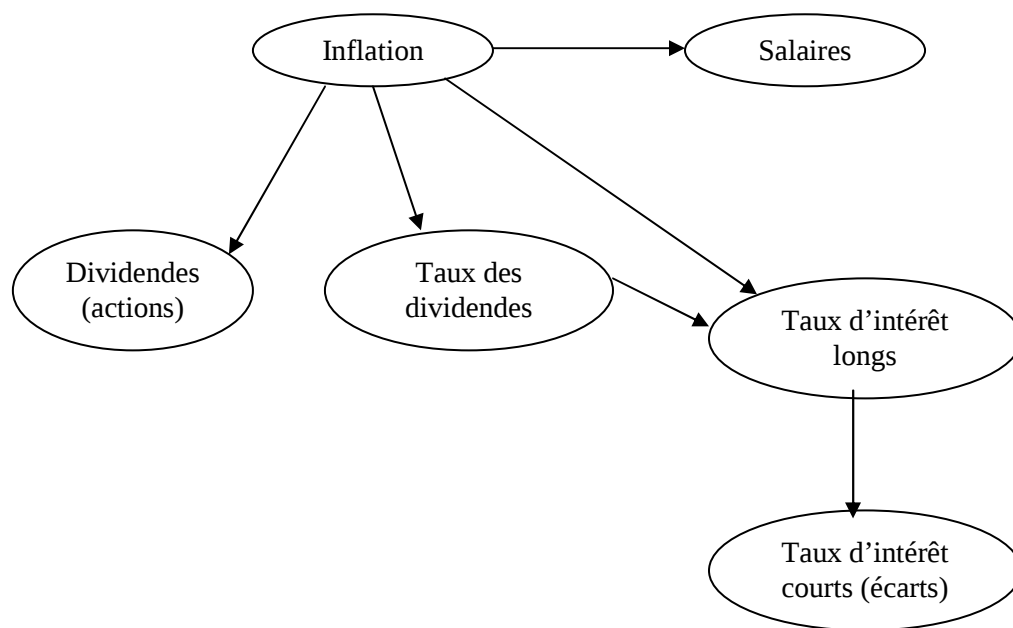


Fig.3 : Structure par cascade dans le modèle de Wilkie [1986]

Toutefois, l'approche de Wilkie a depuis été remise en cause, notamment du fait de sa faible capacité prédictive : il s'agit d'un modèle utilisant un grand nombre de paramètres, dont l'estimation est délicate et qui empêche de fournir des projections pertinentes. Au surplus, le modèle de Wilkie se prête mal à l'évaluation des prix des actifs dérivés, ce qui constitue un handicap important. Il est possible de se référer sur ces points à Rambaruth [2003].

Les problèmes liés au modèle de Wilkie ont également été discuté par Daykin et Hey [1990] et Huber [1995] : certains des paramètres sont instables dans le temps et une corrélation croisée significative entre les résidus des variables projetées est constatée. Le modèle des indices de prix pour l'inflation n'a pas de résidus normaux et ne permet pas la projection de périodes à chocs irréguliers avec des valeurs élevées de l'inflation. De même, la probabilité d'avoir des valeurs négatives de l'inflation avec ce modèle est élevée.

Par ailleurs, Mulvey et Thorlacius [1998] décrivent un modèle de génération de scénarios économiques appelé CAP:Link (développé commercialement par la société Towers Perrin). Ce modèle est basé sur une structure par cascade de ces variables dont la force motrice est supposée être le taux d'intérêt nominal. Il est appliqué principalement dans la gestion actif-passif de long terme. Parmi les variables clés modélisées, nous trouvons l'inflation des prix et des salaires, les taux d'intérêts de différentes maturités (réels et nominaux), le taux de rendement et les taux de dividendes des actions et les taux de change.

Les variables financières sont déterminées simultanément pour différentes économies dans un cadre d'hypothèses générales. Ce modèle s'applique ainsi aux portefeuilles de pension et d'assurance. Les dynamiques des variables sont identiques pour tous les pays alors que les paramètres sont adaptés aux spécificités de chacun d'entre eux. Les auteurs soulignent que, de façon générale, les GSE remplissent au moins l'une de ces trois fonctions suivantes : la prévision, l'évaluation (*pricing*) et l'analyse du risque. Ils considèrent que l'élément clé d'un GSE est le modèle de taux d'intérêt et supposent donc que les taux longs et les taux courts

sont corrélés à travers leurs termes de bruit blanc et que l'écart entre eux est contrôlé par un terme de stabilisation.

D'autres modèles de structure par cascade ont été développés pour l'Australie (cf. Carter [1991]), l'Afrique du Sud (cf. Thomson [1994]), le Japon (cf. Tanaka et al. [1995]) et la Finlande (cf. Ranne [1998]). L'élément commun de ces modèles réside donc dans le fait que le concepteur part de la spécification d'une structure en cascade du modèle à travers les hypothèses sur les liens de causalité entre les variables. En effet, cette structure en cascade permet un seul sens de causalité et exige du modélisateur le choix des liens les plus pertinents de point de vue économique. Par exemple, dans le modèle de Wilkie, la valeur de l'indice des prix permet de déduire la valeur de l'indice des salaires et non l'inverse. Le deuxième sens de causalité est supposé être faible (ou secondaire) sur le long terme.

I-2-2 Modèles basés sur les corrélations

La structure basée sur les corrélations repose quant à elle sur l'idée de permettre aux données disponibles (historiques) de déterminer une structure de corrélation simultanée entre les variables pour ensuite les modéliser et les calibrer en fonction de cette structure. Autrement dit, cette dernière est déterminée essentiellement à travers l'estimation des relations de dépendance observées simultanément dans le passé entre les variables modélisées (par exemple la corrélation linéaire, observée dans le passé entre les rendements des actions et l'inflation, est à retenir et à respecter lors de la projection dans le futur de ces variables). Ainsi, dans ce type de modèle, les données historiques disponibles sur les variables permettent de déduire la structure de dépendance entre elles. Les principaux modèles de GSE en littérature se sont basés sur cette structure.

En adoptant cette structure par corrélation, Campbell et al. [2001] présentent une approche dont l'application a été effectuée dans le cadre de la détermination de l'allocation stratégique d'actifs pour un investisseur de long terme, en particulier les fonds de pension. De son côté, Kouwenberg [2001] se base sur cette structure pour développer un modèle de génération de scénarios qui s'appuie sur un schéma d'arborescence pour la projection des scénarios. L'auteur compare l'effet du choix du schéma de projection sur l'allocation d'actifs optimale dans le cadre d'une gestion actif-passif d'un fonds de pension allemand. La structure d'arborescence retenue par Kouwenberg [2001] est plus adaptée à une série de modèles dynamiques de gestion actif-passif basées sur les techniques de programmation stochastique. Les caractéristiques de ce schéma de projection de scénarios sont présentées de façon détaillée à la section II de ce chapitre.

Hibbert et al. [2001] présente un autre modèle qui génère des valeurs cohérentes, selon les auteurs, de la structure par terme des taux (taux nominaux, réels et d'inflation), des rendements des actions et des revenus de dividendes. Le modèle peut être utilisé pour générer des trajectoires potentielles de chacune de ces variables dans un cadre de modélisation financière et en considérant les différentes corrélations. Hibbert et al. [2001] fournit notamment une revue intéressante des taux d'intérêt, des taux d'inflation et des rendements des actions sur les cent dernières années. Leur modèle est présenté comme un outil de planification et de prise de décisions pour les investisseurs sur le long terme et non comme un outil d'évaluation des produits dérivés (ou *de pricing*).

Ahlgrim et al. [2005] proposent enfin un modèle de GSE qui a le mérite d'être soutenu par la *Casualty Actuarial Society* (CAS) et de la *Society Of Actuaries* (SOA), deux associations professionnelles reconnues aux Etats-Unis. Ahlgrim et al. [2005] partent essentiellement de la

critique de deux points du modèle de Wilkie [1995] : la relation entre l'inflation et les taux d'intérêt est jugée incohérente et le traitement des rendements des actions par une approche autorégressive semble trop simplificateur au regard de l'historique observé. Ils proposent des processus alternatifs en justifiant leurs choix par des *backtesting*⁵ sur des données historiques profondes. Le modèle d'Ahlgrim et al. [2005] rejoint le modèle de Hibbert et al. [2001] en se présentant comme un modèle de projection de valeurs sur le long terme et de gestion des risques.

Le modèle d'Ahlgrim et al. [2005] peut être représenté par les relations suivantes :

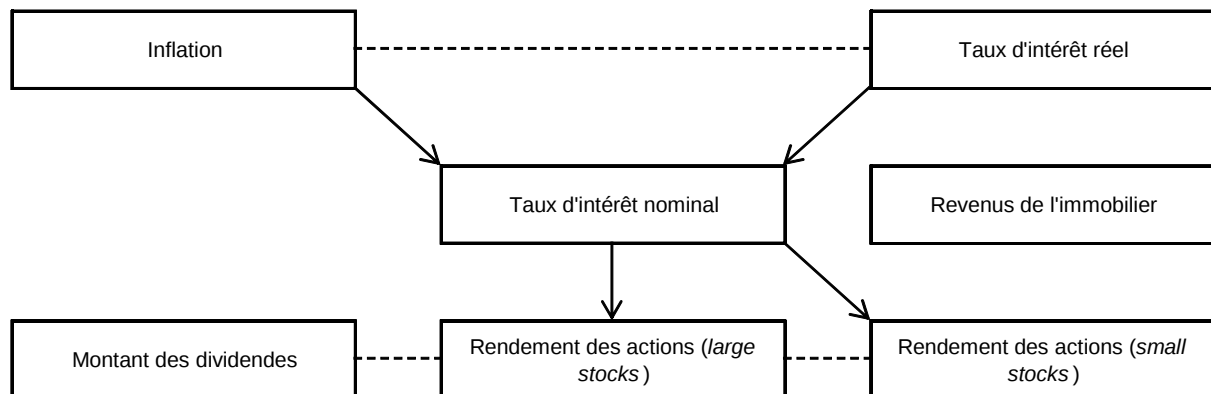


Fig. 4 : Structure du modèle d'Ahlgrim et al. [2005]

Ce graphique illustre le rôle prépondérant de l'inflation et du taux d'intérêt réel dans le modèle.

D'autres modèles se sont basés sur l'hypothèse que les marchés peuvent être soit faiblement efficaces (c'est-à-dire les prix sur le marché reflètent toutes les informations relatives aux prix antérieurs de l'actif) soit fortement efficaces (c'est-à-dire les prix sur le marché reflètent toutes les informations disponibles sur l'actif). De telles approches sont souvent appropriées pour la modélisation de court terme, en particulier pour des fins d'évaluation de produits dérivés. Pour le long terme, l'approche a moins de valeur, puisqu'elle ne tient pas compte des fondamentaux macro-économiques. Smith [1996] et Dyson et Exley [1995] présentent des modèles basés sur les principes du marché efficace pour le cas de la Grande Bretagne. Souvent, ce type de modèle cherche à exclure les opportunités d'arbitrage et ne suppose pas un retour à la moyenne pour les rendements de ces variables.

Une critique commune à tous les modèles ci-dessus, à structure par corrélation, est qu'ils sont fortement dépendants des données sur lesquelles ils sont basés. Autrement dit, si les rendements futurs relatifs à chaque variable du modèle possèdent des caractéristiques significativement différentes de celles observées sur la période historique d'estimation, le GSE pourra conduire à des projections non pertinentes. La prise en compte des avis subjectifs des experts sur le marché (sociétés de gestion, banques d'investissement, etc.) pour fixer ces niveaux futurs de dépendance constitue une source alternative d'alimentation de ces modèles.

⁵ Le *Backtesting* est le test d'une stratégie sur le passé et sur un panel d'actifs financiers.

Nous le voyons, le panorama des modèles proposés dans la littérature est large. Toutefois, quelques-uns de ces travaux peuvent être synthétisés comme suit :

Structure Objectif	Cascade	Corrélation
Projection et gestion des risques	Wilkie [1986, 1995] Mulvey et Thorlacius [1998]	Campbell et al. [2001] Kouwenberg [2001] Hibbert et al. [2001] Ahlgrim et al. [2005]
Evaluation (<i>pricing</i>)	-	Smith [1996] Dyson et Exley [1995]

Tab.2 : Classement des principaux GSE cités dans la littérature

Après avoir choisi la structure théorique du GSE (structure par cascade ou structure basée sur les corrélations), nous arrivons à l'étape de sa mise en œuvre pratique. Cette étape nécessite elle aussi des choix à effectuer que ce soit au niveau du calibrage des différents paramètres du GSE ou au niveau de la génération des trajectoires possibles de ses variables. Concernant le calibrage d'un GSE, différentes techniques peuvent être citées se basant sur les données historiques, les données de marché ou les avis des experts. Certaines de ces techniques seront exposées dans le chapitre 2 de cette partie (relative au modèle de GSE proposé). La section suivante s'intéresse à des problématiques liées à la génération des scénarios.

II- Mise en œuvre d'un GSE

Deux questions principales peuvent être posées lors de la mise en œuvre d'un GSE : d'une part, sous quelle forme schématique devons-nous représenter l'évolution dans le temps des scénarios futurs des variables financières et macro-économiques (inflation, rendement des actions,...) ? Et d'autre part, quelle méthodologie devons-nous adopter pour la génération de ces scénarios ? La première question concerne la structure de projection des scénarios futurs des différentes variables, tandis que la deuxième a trait au choix du modèle d'évolution des valeurs des variables du GSE (processus stochastique, *Bootstrapping*, etc.) : nous notons que ces deux éléments sont cependant liés. Cette section vise à faire l'inventaire (non exhaustif) des différentes possibilités offertes face à ces deux problématiques.

En effet, comme mentionné au début de cette partie, notre objectif sera de mettre en évidence les principaux éléments qui participent à l'amélioration de la qualité d'un GSE. Les états de sortie de ce dernier influencent directement les décisions prises en matière de gestion des risques ou d'évaluation des produits financiers. La détermination de la structure de projection des scénarios et le choix de la méthodologie de leur génération présentent deux étapes inévitables lors de la construction d'un GSE. Ils interviennent, en particulier, au niveau de la mise en œuvre opérationnelle du GSE, d'où l'intérêt de les étudier de façon détaillée et séparée.

II-1 Structure schématique de projection de scénarios pour un GSE

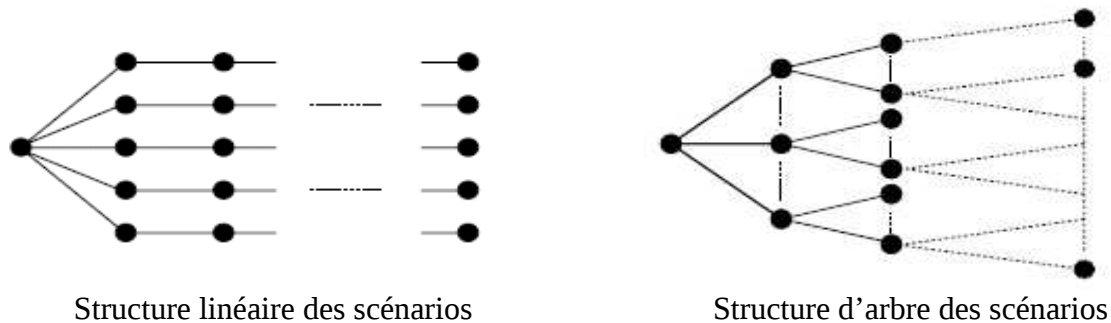
Comme mentionné ci-dessus, le choix de la structure de projection, appelée aussi structure schématique de projection, pour chacune des variables du GSE peut être considéré comme une problématique à part. Elle concerne la définition du schéma graphique de transition entre deux valeurs successives, observées à la date t et $t+1$, de la même variable. Ainsi à chaque variable du GSE peut correspondre une structure de projection particulière. Afin de simplifier l'analyse, nous supposons dans la suite que la structure de projection choisie est la même pour toutes les variables. De même, nous définissons un nœud comme la réalisation possible de la variable modélisée à une date donnée. Une trajectoire correspond ainsi à l'ensemble des nœuds successifs qui forment un scénario futur possible d'évolution de la variable financière ou macro-économique.

Dans cette section, les caractéristiques des structures de projection les plus utilisées en pratique sont présentées avec détails. En particulier, deux principales structures peuvent être avancées à ce stade : la structure de projection linéaire d'une part (cf. Ahlgrim et al. [2005]) et la structure de projection d'arbre (ou d'arborescence) d'autre part (cf. Kouwenberg [2001]). La différence principale entre ces deux structures de projection se situe au niveau de la nature de la dépendance entre les différentes trajectoires simulées. Alors que pour les structures de projection linéaire, une seule trajectoire est dérivée à partir de chaque nœud, les structures par arborescence supposent quant à elles que chaque nœud possède différents nœud-enfants et ainsi différentes trajectoires possibles sont déduites à partir de chaque nœud.

Par exemple, considérons le cas où la projection des rendements des actions se déroule sur deux périodes seulement et que pour les deux structures nous obtenons n_1 nœuds (ou rendements) à la fin de la première période. Si nous optons pour une structure linéaire de projection, il n'est possible d'obtenir que n_1 nœuds (ou rendements) à la fin de la deuxième période, chacun d'entre eux forme avec le nœud précédent une trajectoire distincte. Si par

contre nous optons pour une structure d'arbre, le nombre de scénarios à la fin de la deuxième période est m_1 , avec m_1 souvent supérieur à n_1 puisque différents scénarios de rendement à partir de chacun des n_1 nœuds, simulés fin de la première période, sont projetés.

Le graphique suivant illustre la différence entre ces deux structures :



Source Kouwenberg [2001]

Fig. 5 : Comparaison entre deux structures schématiques de projection des scénarios

La structure d'arbre des scénarios est la forme la plus récente et la plus complexe à utiliser. Notre intérêt sera focalisé dans la suite sur ce type de structure : il s'agit d'une structure plus adaptée que la structure linéaire pour l'application des techniques d'optimisation dynamique, en particulier dans le cas de la détermination de l'allocation stratégique d'actifs optimale (cf. Kouwenberg [2001]). En effet, chaque niveau dans l'arbre représente une date future (ou un moment de prise de décision) et les différents nœuds à chaque niveau représentent les réalisations possibles de la variable modélisée à cette date.

Le graphique suivant représente un exemple plus détaillé de cette structure :

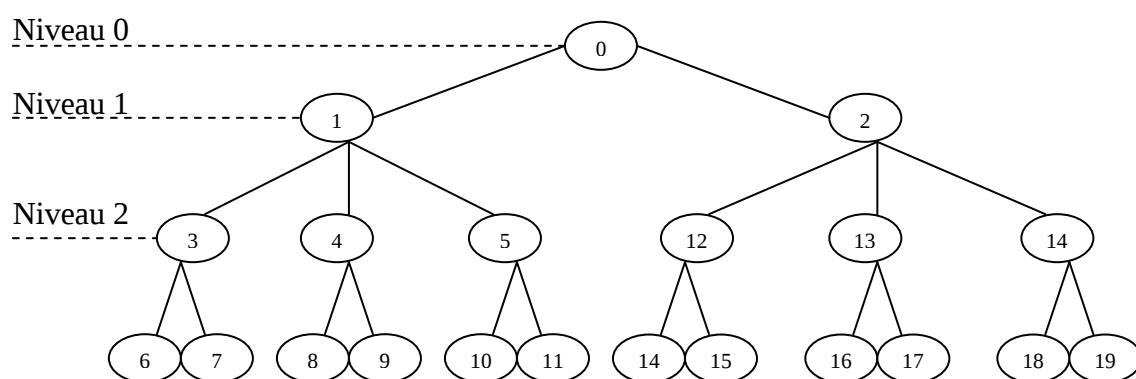


Fig. 6 : Exemple détaillé de la structure d'arbre des scénarios

Comme il est montré dans ce graphique, le nombre des nœuds-enfants à chaque niveau n'est pas nécessairement égal à celui du niveau suivant. Par exemple, le nœud 0 dans le graphique a deux nœuds-enfants alors que les nœuds 1 et 2 ont trois nœuds-enfants. Deux niveaux de l'arbre peuvent ne pas présenter la même période de temps. Par exemple, dans le graphique ci-dessus, le niveau 0 peut représenter le début de l'année 0, le niveau 1 la fin de la deuxième

année et le niveau 2 la fin de la dixième année. De même, dans certains arbres de scénarios complexes, tel que présenté ci-dessous, il pourrait avoir différents nombres de nœuds-enfants pour les nœuds d'un même niveau.

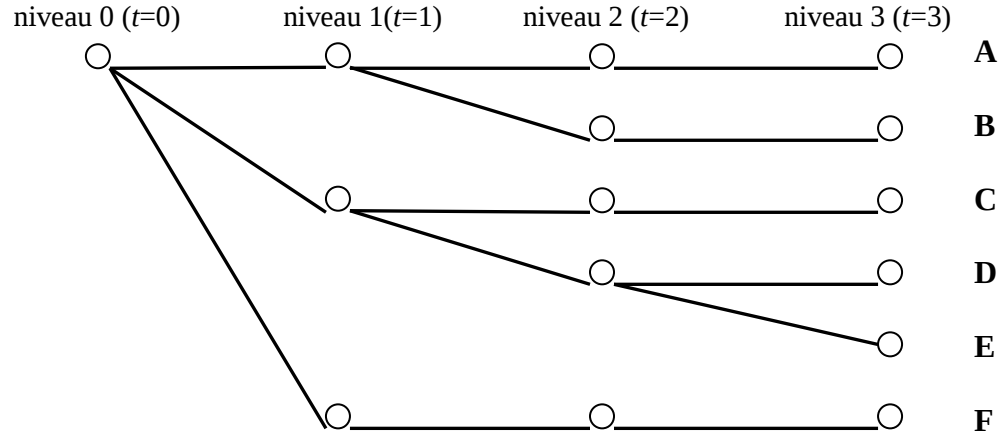


Fig.7 : Arbre de scénarios avec différents nombres de nœuds-enfants pour les nœuds d'un même niveau

Il existe différentes représentations mathématiques possibles de l'arbre des scénarios. Il est référé ici à la formulation de Hochreiter et al. [2002].

Considérons d'abord un processus stochastique $(\xi_t)_{t=0,1,\dots,T}$ discret dans le temps et continu dans l'espace et supposé représenter la dynamique des rendements des actions. L'analyse suivante peut être appliquée aux autres variables générées par le GSE. $\xi_0 = x_0$ représente la valeur d'aujourd'hui et elle est supposée être constante. La distribution de ce processus peut être le résultat d'une estimation, paramétrique ou non, basée sur des données historiques.

L'un des objectifs lors de la mise en place d'un GSE est de trouver un processus stochastique $\bar{\xi}_t$, qui prend seulement des valeurs finies et qui est aussi proche que possible du processus réel des rendements des actions : nous parlons dans ce cas de problème d'approximation. A titre d'hypothèse, le GSE est supposé avoir une structure de projection sous forme d'arbre pour ses différentes variables. Nous définissons, pour cela, l'espace d'état fini de $\bar{\xi}_t$ par S_t :

$$P\left\{\bar{\xi}_t \in S_t\right\} = 1$$

Soit $\text{card}\{S_t\}$ le cardinal de S_t . Si $x \in S_t$, nous appelons le facteur de branchement de x , le nombre des noeuds issus directement de x , c'est à dire la quantité :

$$b(x,t) = \text{card}\left\{y : P\left\{\bar{\xi}_{t+1} = y, \bar{\xi}_t = x\right\} > 0\right\}$$

Intuitivement, le processus $(\bar{\xi}_t)_{t=0,\dots,T}$ peut être représenté sous forme d'arbre, avec comme racine le noeud $(x_0, 0)$. Les noeuds (x, t) et $(y, t+1)$ sont connectés par un arc si $P\left\{\bar{\xi}_t = x, \bar{\xi}_{t+1} = y\right\} > 0$.

La collection de tous les facteurs de branchement $b(x, t)$ détermine la taille de l'arbre. Typiquement, le facteur de branchement est choisi avant et indépendamment de x . Dans ce cas, la structure de l'arbre est déterminée par le vecteur $[b(1), b(2), b(3), \dots, b(T)]$. Par exemple un arbre $[5, 3, 3, 2]$ a pour nombre de niveaux 5 (y compris le niveau 0) et pour nombre de noeuds $1 + 5 + 5.3 + 5.3.3 + 5.3.3.2 = 156$ noeuds. Le nombre des arcs est toujours égal au nombre des noeuds moins 1. Dans ce cadre d'analyse, il est possible de considérer que la structure linéaire constitue un cas particulier de la structure d'arbre avec un facteur de branchement égal à 1 à partir de la deuxième composante du vecteur de la structure ci-dessus (c'est-à-dire $[5, 1, 1, 1]$).

Selon Dupacova et al. [2000], le problème d'approximation principal est un problème d'optimisation de l'un des deux types suivants et il est souvent fonction de la méthode de génération de scénarios retenue :

- Le problème à structure donnée (*The given-structure problem*) : quel processus discret $(\bar{\xi}_t)_{t=0,\dots,T}$ avec une structure de branchement $[b(1), b(2), b(3), \dots, b(T)]$ est le plus proche d'un processus donné $(\xi_t)_{t=0,\dots,T}$? Bien évidemment, la notion de proximité est à définir de manière appropriée.
- Le problème à structure libre (*The free-structure problem*) : ici aussi le processus $(\xi_t)_{t=0,\dots,T}$ est à approximer par $(\bar{\xi}_t)_{t=0,\dots,T}$ mais sa structure de branchement est à définir librement excepté que le nombre total des noeuds est fixé à l'avance. Ce problème d'optimisation hybride et combinatoire est plus complexe que le problème à structure donnée.

Il est à noter que la difficulté majeure lors de l'utilisation des arbres de scénarios est l'augmentation exponentielle dans le nombre de scénarios. Si trois scénarios sont générés pour chaque noeud à n'importe quel niveau parmi 21 niveaux par exemple, le nombre de scénarios générés sera 3^{20} (presque 3,5 milliards de scénarios). Le recours à l'utilisation de telle structure dans le cadre de l'allocation stratégique d'actifs fera l'objet d'une étude plus approfondie à la fin de ce travail en particulier dans le cadre de la proposition d'un nouveau modèle d'ALM (cf. chapitre 3 de la partie II).

Au delà du choix de la structure de projection pour les variables d'un GSE, l'analyse des approches possibles pour la détermination des valeurs futures des variables projetées, appelée aussi méthodologies de génération des scénarios, constitue un problème souvent rencontré par le constructeur d'un GSE. Le choix d'une méthodologie particulière n'est pas sans impact sur les résultats obtenus *in fine* (cf. Ahlgrim et al. [2008]).

II-2 Méthodologies pour la génération de scénarios économiques

Si nous nous plaçons dans le même cadre d'analyse que celui de la sous-section précédente, nous pouvons dire que cette partie vise à répondre à la question suivante: comment pouvons-nous déterminer la valeur d'un nœud. Pour cela, différentes méthodologies, ayant pour finalité la génération de scénarios économiques, peuvent être trouvées dans la littérature (cf. Kaut et Wallace [2003] et Mitra [2006]). Il est proposé dans cette sous-section de les classer en quatre groupes: les approches basées sur l'échantillonnage, les approches basées sur le *matching* des propriétés statistiques, les approches basées sur les techniques de *Bootstrapping* et les approches basées sur l'Analyse en Composantes Principales. Ce dernier groupe n'est pas indépendant des autres comme nous allons le voir dans ce qui suit.

II-2-1 Les approches basées sur l'échantillonnage

Ces approches peuvent être classées en deux sous catégories : l'échantillonnage pur univarié (ou traditionnel) et l'échantillonnage à partir de marginales et de corrélations spécifiées. Cette dernière a le mérite de générer des scénarios dans lesquels la corrélation entre les variables converge vers celle ciblée par le modélisateur.

L'échantillonnage pur est la méthode de génération de scénarios la plus connue. A chaque nœud de l'arbre de scénarios, différentes valeurs sont tirés de façon aléatoire à partir du processus stochastique $\{\xi_t\}$. Cela se fait soit par un tirage direct à partir de la distribution de $\{\xi_t\}$, soit par l'évolution du processus selon une formule explicite: $\xi_{t+1} = z(\xi_t, \varepsilon_t)$.

Dans ce cadre, la dynamique de plusieurs variables financières peut être supposée suivre un processus stochastique de type mouvement brownien géométrique, ou bien l'une de ses variantes. Les scénarios d'évolution de ces variables sont ainsi simulés à partir des hypothèses sur la discrétisation d'un processus brownien géométrique définit par exemple par:

$$dS(t) = \mu S(t)dt + \sigma S(t)dB(t)$$

$S(t)$ est le prix de l'actif, μ et σ sont respectivement le *drift* et la volatilité. Le terme $dB(t)$ est un mouvement brownien, c'est-à-dire $B(t_2) - B(t_1) \sim N(0, |t_2 - t_1|)$. Il est ainsi possible de simuler des processus stochastiques sur un intervalle de temps donné, en attribuant des valeurs aléatoires au mouvement brownien et en calculant par la suite $S(t)$.

Les méthodes traditionnelles d'échantillonnage d'une variable aléatoire permettent de constituer des échantillons seulement à partir d'une variable aléatoire univariée ; lorsque nous voulons tirer un vecteur aléatoire (correspondant à différentes variables), on aura besoin de tirer chaque composante marginale (chaque variable) de façon séparée pour les rassembler ensuite. Le résultat obtenu sera un vecteur de variables aléatoires indépendantes.

Concernant la convergence vers les moments statistiques souhaités (moyenne, variance, etc.), il existe différentes méthodes pour améliorer l'algorithme de l'échantillonnage pur. Nous pouvons par exemple utiliser les méthodes de quadrature pour l'intégration ou les suites à discrétion faible (cf. Pennanen et al. [2002]). Pour les distributions symétriques il est possible d'utiliser les échantillonnages antithétiques. Une autre méthode pour améliorer la méthode d'échantillonnage pur est d'ajuster l'arbre obtenu de façon à avoir les valeurs cibles de la moyenne et de la variance (cf. Cariño et al. [1994]).

Comme mentionné ci-dessus, les méthodes d'échantillonnage traditionnel ont des limites au niveau de la génération des vecteurs multivariés, en particulier ceux avec une corrélation spécifiée. Cependant, il existe des méthodes qui résolvent ce problème en se basant sur des approches d'échantillonnage traditionnel pour ensuite ajuster la technique de contrôle de la corrélation entre les scénarios projetés.

Ces méthodes, que nous avons appelé « échantillonnage à partir de marginales et de corrélations spécifiées » constituent donc une extension des approches d'échantillonnage traditionnel. La différence principale se situe dans le fait qu'elles permettent à l'utilisateur de spécifier à l'avance les distributions marginales ainsi que la matrice de corrélation cible. En général, il n'y a aucune restriction sur les distributions marginales, elles peuvent même appartenir à différentes familles.

A titre d'exemple, Deler et al. [2001] proposent une méthode permettant de générer des variables aléatoires à partir des séries temporelles multi-variées $\{X_t; t=1,2,\dots\}$, où $X_t = (X_{1,t}, X_{2,t}, \dots, X_{k,t})$ est un vecteur aléatoire de dimension $(k \times 1)$ correspondant à l'observation à la date t de ces variables.

Pour cela, ils construisent un processus appelé processus de base Z_t (assimilé à un vecteur auto régressif gaussien standard) et le transforme, à travers un système de translation de Johnson [1949], en un processus ayant au moins les quatre premiers moments (moyenne, écart-type, le coefficient de dissymétrie ou *skewness* et coefficient d'aplatissement ou *kurtosis*) du processus X_t .

Il ajuste enfin la structure de corrélation de ce processus de base Z_t de façon à obtenir celle du processus X_t . L'approche utilisée dans ce cas est celle de Deler et al. [2001]. Elle se base sur la résolution d'un ensemble d'équations permettant de déterminer les corrélations à retenir entre les variables du vecteur Z_t afin de refléter la structure de corrélation cible observée entre les variables du vecteur X_t . D'autres exemples de ces méthodes se trouvent dans Dupacova et al. [2009].

Il est à noter finalement que ces approches basées sur l'échantillonnage sont utilisées dans le cas où nous avons des hypothèses sur les fonctions de distribution des composantes marginales (ou des différentes variables modélisées).

II-2-2 Les approches basées sur le *matching* des propriétés statistiques

Dans les situations où il n'y a pas d'hypothèses sur la distribution marginale du processus de génération de scénarios, les approches basées sur le *matching* des propriétés statistiques, en particulier les moments, sont les plus adaptées. Un processus de génération des scénarios par le *matching* des moments s'intéresse souvent aux trois ou aux quatre premiers moments de chacune des variables projetées (moyenne, variance, skewness, kurtosis) ainsi qu'à la matrice de corrélation.

Ces méthodes peuvent être étendues à d'autres propriétés statistiques (tel que les quantiles, etc.). Le générateur de scénarios par le *matching* des moments va ensuite construire une distribution discrète satisfaisant les propriétés statistiques sélectionnées. Par exemple, Hoyland et al. [2003] commencent par spécifier le nombre minimal de scénarios qu'il faut

générer pour ensuite obtenir l'arbre de scénarios par une optimisation non linéaire, où l'objectif est de minimiser l'erreur entre les moments théoriques de la variable et ceux fournis par l'arbre.

II-2-3 Les approches basées sur les techniques de *Bootstrapping*

Il s'agit des approches les plus simples pour générer des scénarios en utilisant seulement les données historiques disponibles sans aucune modélisation *a priori* de la dynamique d'évolution des variables du GSE (cf. Albeanu et al. [2008]). Le *Bootstrapping* se base essentiellement sur la constitution d'échantillons à partir des données observées. Dans ce cadre, les valeurs de chaque scénario représentent un échantillon de rendements d'actifs obtenu par un tirage aléatoire de certains rendements observés déjà dans le passé.

Par exemple, afin de générer des scénarios de rendement sur les dix prochaines années, un échantillon de 120 rendements mensuels tirés aléatoirement sur les 120 rendements des dix dernières années est utilisé. Ce processus est répété un certain nombre de fois afin de générer plusieurs scénarios possibles dans le futur. Autrement, dans le cas où nous avons k actifs à projeter, et en supposant qu'un historique de p périodes est disponible, nous tirons avec remise un entier entre 1 et p et nous prenons toutes les valeurs des k actifs à cette même date de façon à tenir compte de la corrélation historique entre ces derniers.

II-2-4 L'utilisation de l'Analyse en Composantes Principales

L'Analyse en Composantes Principales (ACP) est une méthode générique d'analyse de données ayant plusieurs dimensions (cf. Bouroche et al. [1980]). Elle permet l'identification des facteurs clés régissant les tendances de ces données et de réduire leur dimension tout en conservant le plus d'information possible.

Pour cela, l'ACP se base sur l'identification des vecteurs propres, des valeurs propres et des covariances. En effet, il s'agit d'une méthode descriptive qui dépend d'un modèle géométrique plutôt que d'un modèle probabiliste. L'ACP propose de réduire la dimension d'un ensemble des données (échantillon) en trouvant un nouvel ensemble de variables plus petit que l'ensemble original des variables, qui néanmoins contient la plupart de l'information de l'échantillon.

Autrement dit, pour un ensemble de données dans un espace à N dimensions, nous recherchons un sous-espace à k dimensions (défini par k variables) tel que la projection des données dans ce sous-espace minimise la perte d'information. Ces k variables seront appelés composantes principales et les axes qu'elles déterminent axes principaux. L'implémentation numérique de la méthode ACP est ainsi accessible. En pratique, l'ACP a pour objet de réduire le nombre des variables de départ du modèle pour ensuite appliquer une approche traditionnelle de génération de scénarios (parmi celles proposées ci-dessus) pour les composantes principales.

II-3 Probabilité risque-neutre Vs probabilité réelle

Les structures par terme dans les modèles d'absence d'arbitrage projettent des trajectoires de taux d'intérêt futurs qui émanent de la courbe des taux existante. L'application de tels modèles suppose que les prix sont déterminés à partir du principe d'absence d'opportunité

d'arbitrage dont découle « la probabilité risque-neutre » (cette technique est surtout intéressante lors du développement des techniques d'arbitrage).

Sous cette probabilité, le rendement de tous les actifs est supposé être le taux sans risque et la somme actualisée d'un ensemble de flux futurs est égale à la somme de ces flux actualisés au taux sans risque. Cette probabilité a pour objectif les opérations de « *pricing* » et reflète un consensus entre deux parties acceptant le transfert de risque. En fait, la méthodologie d'évaluation risque-neutre est utilisée par les banques d'investissement et les académiques pour l'évaluation des produits dérivés. Pour un modèle de retour à la moyenne, la tendance de long terme sera le taux sans risque pour toutes les variables. Nous estimons ainsi une moyenne des flux de trésorerie espérés du produit dérivé en utilisant les techniques de *Monte Carlo*. Ces flux sont actualisés au taux sans risque afin d'obtenir la valeur économique du produit dérivé.

La probabilité risque-neutre fournit des résultats cohérents par rapport à des marchés neutres face au risque mais elle n'est pas recommandée pour la mesure du risque ou pour la prévision (Adam, [2007]). Par ailleurs, les structures par terme dans les modèles d'absence d'opportunité d'arbitrage sont fréquemment plus difficiles à implémenter que leurs contreparties des modèles d'équilibre. Cette probabilité suppose un nombre d'hypothèses qui sont peu conforme au monde réel. Comme conséquence, les valeurs obtenues par cette méthodologie doivent être interprétées avec prudence. Elles peuvent constituer un benchmark intéressant, surtout lors de l'évaluation de produits dérivés par les techniques actuarielles.

La probabilité réelle a quant à elle un objectif de simulation (ou de projection) réaliste :

- Elle corrige les simulations risque-neutre et montre par exemple que sur le long terme l'investissement en obligations est plus intéressant que l'investissement en monétaire (cf. Campbell et al. [2001]).
- La simulation réelle intègre des primes de risque comme contrepartie du risque supplémentaire assumé.
- Elle corrige le calcul des indicateurs de risque (dans les scénarios de stress, analyse de sensibilité, etc.)
- Elle permet à la société de tenir compte de l'évolution de sa propre valeur de marché lors des simulations.

La probabilité réelle est utilisée lors de l'optimisation de l'exposition au risque par le gestionnaire actif-passif (Adam, [2007]). La prime de risque indique une différence entre la performance espérée de l'investisseur et le taux sans risque. Mathématiquement, le passage de la probabilité risque neutre à la probabilité réelle se fait par l'introduction d'une prime de risque μ :

$$dB_t^{réel} = dB_t^{RN} + \mu \cdot dt$$

II-4 L'exclusion de valeurs négatives pour les modèles de taux

Beaucoup de discussions existent sur la nécessité et la façon avec laquelle nous pouvons limiter la simulation des variables négatives (Bassetto [2004]). Les variables autour desquelles se concentrent ces discussions sont les taux réels et les taux nominaux, mais l'inflation est aussi concernée. La question se complique d'avantage lorsque le taux nominal est déduit à partir des deux autres taux. Parmi les approches possibles :

- soit l'utilisation de bornes inférieures pour les taux réels et d'inflation,
- soit la fixation de bornes inférieures uniquement pour les taux d'inflation (qui peuvent donc être négatifs) : les taux réels sont quant à eux déduits dans un deuxième temps de façon à avoir des taux nominaux positifs.

III- Mesure de qualité

La mesure de la qualité d'un GSE peut avoir deux caractères : qualitatif et quantitatif. Dans les deux cas, l'objectif sera de garantir des états de sortie du modèle qui permettent de prendre les meilleures décisions.

III-1 Mesure qualitative de la qualité d'un GSE

Pour Hibbert et al. [2001], les propriétés qu'un « bon » GSE doit avoir sont les suivantes : la représentativité, la plausibilité économique, la parcimonie, la transparence et l'évolution.

Le modèle doit « imiter » le comportement des actifs financiers dans le monde réel en captant leurs principales caractéristiques (représentativité). De même, les différents scénarios générés devront être plausibles et raisonnables pour les experts du marché financier. Cela passe par l'étude de la forme de la distribution et des corrélations des différentes variables du modèle.

A ce titre, le degré de correspondance des résultats aux données historiques constitue un champ d'investigation intéressant, malgré les critiques à l'encontre de l'hypothèse de reproduction des événements historiques dans le futur. L'utilité de la comparaison des résultats obtenus par rapport aux observations déjà réalisées sur le marché peut varier en fonction des attentes du gestionnaire en termes d'évolution future du contexte macro-économique. En fait, l'objectif de similitude entre les deux séries (projetées et historiques) peut être abandonné si on juge par exemple un changement profond du contexte économique par rapport au passé (exemple : modification importante des poids des économies à l'échelle internationale, explosion de l'endettement public,...).

Le comportement joint des différentes variables simulées dans les scénarios doit aussi présenter un niveau suffisant de plausibilité et de cohérence en respectant les principes économiques. Cependant, il est à noter qu'il n'y a pas de consensus sur certaines propriétés importantes des actifs financiers. L'exemple du principe de retour à la moyenne sur le marché des actions, qui stipule que les rendements sur le long terme des actions convergent vers une tendance donnée, en est une illustration (cf. Hibbert et al. [2001]).

La parcimonie fait référence, quant à elle, à la simplicité des modèles proposés par le GSE, leur permettant d'être d'une part compréhensibles par les utilisateurs et d'autre part implémentables sur des supports informatiques. Par ailleurs, la parcimonie permet une meilleure capacité prospective.

La transparence est enfin un élément déterminant pour le succès du modèle, fortement lié à son potentiel de communication : les résultats obtenus devront être présentés sous forme de graphiques clairs et détaillés. Finalement, Le potentiel du modèle à évoluer et à élargir le champ de son étude est un élément d'attractivité supplémentaire.

Pour Zenios [2007], les critères qu'un GSE doit satisfaire afin de garantir sa bonne qualité sont au nombre de trois : l'exactitude (*correctness*), la précision (*accuracy*) et la cohérence (*consistency*).

Les scénarios devront avoir des propriétés qui sont **justifiées et soutenues par différentes publications académiques**. Par exemple, la structure par terme devra refléter le phénomène de retour à la moyenne des taux d'intérêt constaté dans plusieurs études universitaires et techniques (cf. Vasicek [1977] et Hibbert et al. [2001]). De même, la structure par terme pourra être déduite à partir des changements dans le niveau, la pente et la courbure tels que mentionné dans certaines études économiques (cf. Heath et al. [1992]).

De même, afin de vérifier la condition d'exactitude, les scénarios devront couvrir les scénarios importants observés dans le passé ainsi que tenir compte des **événements qui ne sont pas observés avant**, mais qui ont de fortes chances d'être observés sous les conditions actuelles de marché.

Comme dans plusieurs cas, les scénarios représentent une discrétisation d'un scénario continu, l'accumulation d'un nombre d'erreurs dans la discrétisation est inévitable. Différentes approches peuvent être utilisées pour assurer que l'échantillon de scénarios représente la fonction de distribution continue sous-jacente.

La précision est assurée lorsque **le premier moment jusqu'au plus grand moment de la distribution des scénarios convergent vers ceux de la distribution théorique sous-jacente**. (le *matching* des moments et des propriétés statistiques sont souvent utilisés afin d'assurer que les scénarios gardent les moments théoriques de la distribution qu'ils représentent).

La demande de précision peut mener à la génération d'un nombre élevé de scénarios. Ceci dans le but de créer une discrétisation fine de la distribution continue et afin d'accomplir la précision qu'on considère appropriée et acceptable pour le problème en question. Cependant, compte tenu des erreurs de modèle et de calibrage il est recommandé d'éviter de consacrer trop d'énergie à l'obtention d'un niveau de précision, formellement satisfaisant mais illusoire de point de vue opérationnel.

Lorsque les scénarios sont générés pour différents instruments (par exemple, les obligations, la structure par terme, etc.), il est important de voir que les scénarios sont cohérents en interne. Par exemple, les scénarios dans lesquels une augmentation dans le taux d'intérêt tient lieu simultanément avec une hausse des prix des obligations sont théoriquement non cohérents (cf. Ahlgrim et al. [2008]), même si à une échelle individuelle chacun des deux scénarios, celui du taux et celui des prix, peut être avoir lieu étant donné certains facteurs exogènes (par exemple, une augmentation des taux par la banque centrale accompagnée par l'augmentation de la demande sur les obligations comme actifs refuges sur le marché). La **corrélation** entre les différents instruments peut être utilisée afin d'assurer la cohérence des scénarios.

L'étude de la qualité de certaines méthodologies de génération de scénarios est illustrée ci-dessus :

Concernant les approches basées sur l'échantillonnage : elles permettent de probabiliser un ensemble de scénarios futurs en tenant compte de la possibilité de la reproduction des scénarios passés. Si le choix de la dynamique des variables est justifié par des références

académiques, ces approches assurent la condition d'exactitude. Par ailleurs, les conditions de précision et de cohérence peuvent aussi être préservées : avec un choix de techniques de calibrage et de discrétisation adéquates (précision) et en tenant compte de la structure de corrélation entre les variables dans les modèles (cohérence).

Nous pouvons noter que le *matching des moments* assure par définition la condition de précision puisqu'elle garde les propriétés des moments. De même, le *matching* des matrices de covariance assure la cohérence des scénarios. Cependant, cette approche est assez générale et n'est pas validée par des études académiques ce qui est nécessaire pour que l'approche respecte la condition d'exactitude.

De son côté le *bootstrapping* des données historiques préserve la corrélation observée, mais ne satisfait pas la condition d'exactitude pour la génération de scénarios car il ne suggérera pas un rendement qui n'est jamais observé dans le passé (ce qui est contre-intuitif, comme la reproduction du passé ne constitue qu'une partie des scénarios possibles dans le futur et non la totalité des scénarios potentiels, confirmation donnée par la crise économique actuelle qui a donné lieu à des rendements records jamais observés au paravent). Lorsqu'ils sont tirés de façon appropriée, les scénarios de cette approche satisfont les conditions de précision et de cohérence puisqu'ils sont en parfaite cohérence avec les observations réelles du passé.

III-2 Mesure quantitative de la qualité d'un GSE

Dans cette sous-section, inspirée des travaux de Kaut et Wallace [2003], nous essayons de présenter certains critères permettant de tester de façon quantitative la qualité d'un GSE non pas en tant qu'outil de génération de la distribution réelle de ses variables mais plutôt en tant qu'outil d'aide à la prise de décision. L'intérêt dans ce travail sera porté sur la performance des scénarios comme un *input* au modèle de prise de décision et sur la mesure de la qualité des décisions subséquentes que proposent ces scénarios (stabilité, absence de biais, etc.). L'objectif n'est donc pas la mesure du degré de correspondance des scénarios projetés par rapport à la distribution réelle des variables du GSE, une distribution qui reste peu accessible quelque soit les outils d'estimation et d'approximation employés. Nous notons aussi que la connaissance de la distribution réelle enlève le besoin de passer par un GSE pour prévoir les valeurs futures des variables financières et macro-économiques.

III-2-1 Présentation du problème

Les notations utilisées dans la suite sont identiques à celles utilisées précédemment⁶. Nous nous intéressons à la fonction objectif d'un modèle d'optimisation synthétisé comme suit :

$$\min_{x \in X} F(x, \xi_t)$$

où ξ_t correspond donc à un processus stochastique discret dans le temps et continu dans l'espace.

La résolution directe du problème réel (théorique) étant supposée être impossible, le passage par un processus discret $\{\bar{\xi}_t\}$ dont l'évolution est représentée sous forme d'un arbre ne constitue qu'une solution d'approximation (empirique). Une fois effectué, la fonction objectif devient :

⁶ Le choix de la structure de projection n'influence pas l'analyse menée dans cette sous section.

$$\min_{x \in X} F(x, \bar{\xi}_t)$$

Comme mentionné ci-dessus, nous nous proposons dans ce travail de lier la performance du GSE étudié à la qualité des décisions qu'il permet d'obtenir et non pas à la performance statistique d'approximation et d'estimation du processus réel. D'autre part, il est à rappeler que l'erreur d'estimation des paramètres ne fait pas partie de l'objet de cet article ce qui ne réduit en aucun cas l'importance d'étudier la robustesse des décisions obtenues par le GSE par rapport aux choix effectués pour ces paramètres (cf. Meucci [2005]).

L'erreur, au niveau décisionnel, de l'approximation d'un processus stochastique $\{\xi_t\}$ par une discrétisation $\{\bar{\xi}_t\}$, peut être définie, soit comme la différence entre les valeurs optimales de la fonction objectif des deux problèmes : le problème réel (théorique) et le problème d'approximation (empirique), soit comme la différence entre les solutions proposées par ces deux mêmes problèmes.

Le premier type d'erreur, lié à la valeur optimale de la fonction objectif, peut être formulé comme suit :

$$\begin{aligned} e_f(\xi_t, \bar{\xi}_t) &= F\left(\arg \min_x F(x; \bar{\xi}_t); \xi_t\right) - F\left(\arg \min_x F(x; \xi_t); \xi_t\right) \\ &= F\left(\arg \min_x F(x; \bar{\xi}_t); \xi_t\right) - \min_x F(x; \xi_t) \end{aligned} \quad (1)$$

Nous notons que $e_f(\xi_t, \bar{\xi}_t) \geq 0$, étant donné que le deuxième élément de l'écart est le vrai minimum, alors que le premier est la valeur de la fonction objectif à une solution approximée. La comparaison est effectuée non pas au niveau des solutions optimales (x^*) mais plutôt au niveau des valeurs de la fonction objectif correspondantes. En fait, notre intérêt porte en premier lieu sur la valeur minimale et non pas sur les différentes solutions qui permettent *in-fine* d'avoir cette valeur (appelée aussi la performance).

Il y a deux problèmes rencontrés dans l'équation ci-dessus de $e_f(\xi_t, \bar{\xi}_t)$:

- trouver la valeur « réelle » de $F(x; \bar{\xi}_t)$ pour une solution donnée x .
- trouver la valeur « réelle » de $F(x; \xi_t)$ qui suppose la détermination de la solution réelle x^* .

Alors que le second problème est souvent difficile à résoudre, car il nécessite la résolution du problème d'optimisation avec le processus continu réel, le premier problème peut être résolu, à travers la simulation en temps discret par exemple. Dans la sous-section suivante, nous présentons différents indicateurs de mesure de l'erreur lié à la qualité des décisions obtenues par un GSE donné. Ces indicateurs se basent sur la définition de conditions à satisfaire par la solution empirique.

III-2-2 Le test de la qualité des décisions obtenues par le GSE

Il existe au moins deux conditions nécessaires à satisfaire par un GSE afin de garantir la fiabilité des décisions déduites à partir de ses projections. La première condition est la **stabilité** de la valeur de la fonction objectif optimale par rapport aux différents arbres de scénarios : c'est-à-dire si nous générons plusieurs arbres (avec les mêmes paramètres initiaux : tendance, volatilité, etc.) et que nous résolvons le problème d'optimisation pour chacun d'eux, nous devrions avoir les mêmes valeurs optimales de la fonction objectif (mêmes performances). L'autre condition est que l'arbre de scénarios n'introduit **aucun biais** au niveau de la vraie solution (x^*) et ainsi la solution proposée empiriquement doit correspondre à la solution optimale réelle. La première condition peut dans certaine mesure être testée alors qu'un test direct de la deuxième est souvent considéré comme impossible étant donné que la fonction réelle est inconnue.

III-2-2-1 La condition de stabilité

Deux types de stabilité peuvent être avancés : la stabilité interne de l'échantillon et la stabilité externe de l'échantillon. Ces deux types d'indicateurs interviennent au niveau de la valeur optimale de la fonction objectif et non au niveau de la solution optimale du problème.

En fait, la stabilité interne de l'échantillon signifie que si nous générons k arbres de scénarios avec la discrétisation $\{\bar{\xi}_t\}$ d'un processus donné $\{\xi_t\}$, et que si nous résolvons le problème d'optimisation pour chacun de ces arbres, nous devrions avoir (approximativement) la même valeur optimale de la fonction objectif avec comme solutions optimales x_k^* , $k = 1, \dots, K$.

$$\text{Autrement : } F\left(x_k^*; \bar{\xi}_{tk}\right) \approx F\left(x_l^*; \bar{\xi}_{tl}\right) \quad k, l \in 1 \dots K$$

La stabilité externe de l'échantillon concerne quant à elle les performances des arbres de la distribution réelle et elle est définie comme suit :

$$F\left(x_k^*; \xi_t\right) \approx F\left(x_l^*; \xi_t\right) \quad k, l \in 1 \dots K$$

Une autre formulation possible de ces deux indicateurs sera :

$$\text{La stabilité interne : } \min_x F\left(x; \bar{\xi}_{tk}\right) \approx \min_x F\left(x; \bar{\xi}_{tl}\right)$$

$$\text{La stabilité externe : } F\left(\argmin_x F\left(x; \bar{\xi}_{tk}\right); \xi_t\right) \approx F\left(\argmin_x F\left(x; \bar{\xi}_{tl}\right); \xi_t\right)$$

En utilisant l'équation (1) de la fonction d'erreur, nous déduisons que pour le cas de la stabilité externe: $e_f\left(\xi_t, \bar{\xi}_{tk}\right) = e_f\left(\xi_t, \bar{\xi}_{tl}\right)$

Il y a une différence importante entre les deux définitions : alors que pour la stabilité interne de l'échantillon nous avons besoin seulement de résoudre le problème d'optimisation basé sur les scénarios projetés par le GSE, la mesure de la stabilité externe passe par l'évaluation de la

fonction objectif réelle $F(x; \xi_t)$. En pratique, il n'est pas évident de tester précisément la stabilité externe comme nous n'avons pas la distribution réelle des variables projetées (rendements, taux, etc.).

Cela peut être démontré par l'exemple uni-périodique et uni-dimensionnel suivant⁷ illustrant la différence entre les deux types d'indicateurs :

$$\min_{x \in \mathfrak{R}} F(x; \xi) = E^{\xi} [(x - \xi)^2]$$

Ce problème peut être résolu théoriquement et avoir sa solution explicite pour n'importe quelle distribution de ξ :

$$\begin{aligned} F(x; \xi) &= E[(\xi - x)^2] \\ &= E[((\xi - E[\xi]) + (E[\xi] - x))^2] \\ &= E[(\xi - E[\xi])^2] + E[2(\xi - E[\xi])(E[\xi] - x)] + E[(E[\xi] - x)^2] \\ &= Var[\xi] + 0 + (x - E[\xi])^2 \end{aligned}$$

et la solution optimale est alors :

$$x^* = \underset{x \in \mathfrak{R}}{\operatorname{argmin}} F(x; \xi) = E[\xi]$$

$$F(x^*; \xi) = \min_{x \in \mathfrak{R}} F(x; \xi) = Var[\xi]$$

Cette solution est dite explicite parce que son obtention ne nécessite pas la simulation de la distribution de ξ et elle constitue du fait une solution générale et déterministe. Considérons le cas dans lequel nous générons des arbres de scénarios $\vec{\xi}_k$ ($k=1...K$) et que nous obtiendrons les solutions $x_k^* = E[\vec{\xi}_k]$.

Nous supposons dans un premier temps que la méthode de génération de scénarios est telle que tous les échantillons $\vec{\xi}_k$ ont d'une part la même moyenne réelle $E[\vec{\xi}_k] = E[\xi]$ et d'autre part des variances $Var[\vec{\xi}_k]$ différentes. Ainsi $F(x_k^*; \vec{\xi}_k) = Var[\vec{\xi}_k]$ sera différente d'un échantillon à l'autre d'où une non stabilité interne. En même temps, $x_k^* = x^*$, donc $F(x_k^*; \xi) = F(x^*; \xi) = Var[\xi]$, ce qui confirme une stabilité externe de l'échantillon.

Si par contre, nous supposons que la méthode de génération de scénarios produit des échantillons ayant tous la même variance réelle ($Var[\vec{\xi}_k] = Var[\xi]$), mais des moyennes

⁷ Pour simplifier, ξ désigne ξ_t dans cet exemple.

différentes, nous obtenons dans ce cas $F(x_k^*; \vec{\xi}_k) = Var[\vec{\xi}_k] = Var[\xi]$ et la stabilité interne de l'échantillon est ainsi vérifiée. De l'autre côté, la valeur de $F(x_k^*; \xi) = Var[\xi] + \left(E[\vec{\xi}_k] - E[\xi]\right)^2$ diffère entre les échantillons et la résolution du problème ne vérifie pas la stabilité externe de l'échantillon.

Une question peut être posée à ce stade : quel est le type de stabilité le plus important pour juger de la performance de la méthode de génération de scénarios ? Avoir la stabilité externe sans avoir la stabilité interne signifie que la performance obtenue par approximation peut ne pas correspondre à la performance réelle. Le cas opposé, c'est-à-dire stabilité interne sans la stabilité externe, est quant à lui plus gênant pour le processus décisionnel, puisque la performance réelle des solutions empiriques dépendra de l'arbre de scénarios réel sélectionné, et nous ne pouvons pas ainsi trancher au niveau de la meilleure solution obtenue par approximation alors que dans le premier cas nous avons au moins une solution empirique stable théoriquement mais dont la performance empirique reste non stable, ce qui paraît *a priori* moins gênant.

Une autre question possible est l'étude d'une stabilité interne de l'échantillon non pas au niveau de la performance mais au niveau des solutions elles même : ainsi nous cherchons à avoir la même solution optimale quel que soit l'arbre de scénarios généré.

Dans l'exemple ci-dessus, nous pouvons remarquer qu'il est possible d'avoir une non stabilité interne de la valeur optimale de la fonction objectif (différentes variances empiriques $Var[\vec{\xi}_k]$) tout en gardant une stabilité interne au niveau des solutions elles mêmes (les solutions $x_k^* = E[\xi]$ sont les mêmes pour tous les arbres). Ceci permet de déduire immédiatement la stabilité externe de l'échantillon. Ainsi, si nous détectons une stabilité interne de l'échantillon au niveau de la fonction objectif, nous devrions regarder aussi au niveau des solutions. Par contre l'autre sens d'analyse n'est pas toujours vrai et nous pouvons avoir une stabilité externe même si la stabilité interne au niveau des solutions n'est pas vérifiée.

Nous pouvons conclure *in-fine* que la stabilité est une condition minimale à respecter par la méthode de génération de scénarios. Ainsi, avant de commencer à travailler avec un nouveau modèle d'optimisation, ou une nouvelle méthode de génération de scénarios, il sera judicieux d'effectuer les tests de stabilité : le test de la stabilité interne de l'échantillon et dans la mesure du possible le test de la stabilité externe de l'échantillon.

III-2-2-2 La condition d'absence de biais

En plus de la condition de stabilité (à la fois interne et externe de l'échantillon), la méthode de génération de scénarios ne devra pas introduire de biais au niveau de la solution proposée. Autrement, la solution du problème d'approximation basé sur les scénarios fournis par notre GSE :

$$x^* = \underset{x}{argmin} F(x; \vec{\xi}_t)$$

devra être une solution optimale du problème original. Ainsi, la valeur de la vraie fonction objectif pour la solution obtenue, $F\left(\vec{x}^*; \vec{\xi}_t\right)$, devra être approximativement égale à la valeur optimale réelle $\min_x F(x; \xi_t)$:

$$F\left(\vec{x}^*; \vec{\xi}_t\right) = F\left(\arg \min_x F\left(x; \vec{\xi}_t\right); \vec{\xi}_t\right) \approx \min_x F\left(x; \vec{\xi}_t\right)$$

En utilisant l'équation (1) de la fonction d'erreur, cette condition s'exprime comme suit :

$$e_f\left(\vec{\xi}_t, \vec{\xi}_t\right) \approx 0.$$

Le problème est que le test de cette propriété n'est pas possible dans la plupart des cas, étant donné que cela nécessite la résolution du problème d'optimisation avec le processus continu réel, supposé jusque là comme inconnu.

III-3 Application numérique

Afin d'illustrer les points développés jusque là, nous présentons dans ce qui suit un exemple numérique simplifié prenant le cas d'un investisseur (compagnie d'assurance par exemple) intéressé par la répartition de son capital entre $n=3$ actifs financiers sur le marché européen : les actions, les obligations et le monétaire. Chacun de ces trois classes d'actifs sera représenté par un indice dont la série de données historiques mensuelles, entre le 01/01/1999 et le 31/10/2009, est récupérée à partir de la base de données Bloomberg : ainsi nous avons retenu l'indice EONIA pour le monétaire, l'indice obligataire EFFAS-Euro pour les obligations et l'indice DJ Eurostoxx50 pour les actions. Nous supposons aussi que ce problème d'optimisation de portefeuille est mono-périodique : la projection des valeurs de rendement de chacun des trois actifs se fait sur une seule période (par exemple une année).

Au niveau du choix de **la méthodologie de génération des scénarios** du GSE, un processus stochastique reflétant une hypothèse de normalité des rendements est associé à chacun des trois indices, soit :

$$R^i = \mu_i + \sigma_i \varepsilon_{t,i} \quad \text{avec } i = 1, 2, 3$$

R^i correspond au rendement de l'indice i . La tendance (ou la moyenne annualisée) μ_i et la volatilité σ_i des rendements des trois indices sont estimées sur une base historique (elles correspondent respectivement à la moyenne annualisée et à la volatilité empiriques de chacune des trois séries de données retenues). $\varepsilon_{t,i}$ suit une loi normale de moyenne nulle et de variance égale à 1.

La structure de dépendance entre les variables du GSE considéré est bien une structure par corrélation. Les trois indices sont simultanément dépendants l'un de l'autre lors de leur projection et le niveau des corrélations est déterminé à travers les corrélations fixées en amont entre leurs bruits $\varepsilon_{t,i}$ ($i = 1, 2, 3$).

La fonction objectif de notre investisseur sera défini comme suit :

$$\begin{aligned} \text{Min } F(R_p; \varepsilon_t) &= \sigma(R_p) \\ \text{tel que } p &\in U \\ E(R_p) &\geq m_0 \end{aligned}$$

Avec R_p le rendement total du portefeuille composé des trois actifs (monétaire, obligations et actions), $E(R_p)$ la valeur espérée du rendement du portefeuille, $\sigma(R_p)$ la volatilité du rendement du portefeuille et m_0 la valeur minimale du rendement espéré du portefeuille (égale à 4 % dans notre cas). En pratique, l'espérance sera estimée par la moyenne empirique.

Rappelons que pour déduire le rendement du portefeuille sur une période donnée, la formule suivante est utilisée :

$$R_p = \sum_{j=1}^{n=3} p_j R^j$$

Avec p_j et R^j correspondants respectivement au poids et au rendement de l'actif j sur la même période.

U correspond à l'univers des allocations (ou compositions de portefeuille) à tester. La détermination de cet univers passe par la fixation d'un pas, par exemple de 5 %, pour les poids des différents actifs. Autrement :

$$U = \left\{ (p_1, \dots, p_n) \mid \sum_{j=1}^n p_j = 1, 1 \geq p_j = c \times 0,05 \geq 0, \forall j \in \{1, \dots, n\}, c \in \{0,1, \dots, 20\} \right\}$$

Dans notre cas, 3 actifs et un pas de 5 %, nous obtenons 231 allocations à tester.

$F(R_p; \varepsilon_t)$ correspond à la fonction objectif réelle présentée dans la sous-section précédente comme $F(x; \xi_t)$ ⁸.

La structure de projection adoptée dans cette application est une structure linéaire et ce afin de simplifier la méthode de résolution de la fonction objectif. Cette structure peut être illustrée par exemple par le graphique 8 correspondant à la variable rendement des actions projetée sur une période de 20 ans avec 500 trajectoires. Dans ce graphique, $R_{i,t}^{actions}$ représente le rendement de l'indice des actions entre la date t et $t-1$ à la $i^{\text{ème}}$ trajectoire (ou scénario).

La taille de la structure de projection sera définie comme le nombre de trajectoires simulées dans cette structure. Un ensemble de scénarios désignera l'ensemble des trajectoires obtenues d'une variable dans le cas d'une structure de projection linéaire (équivalent de l'arbre de scénarios dans le cas de structure par arborescence).

Dans notre étude nous envisageons de procéder comme suit : nous fixons différentes tailles des structures de projection adoptées dans ce GSE : par exemple 50, 1000, 5000 et 10 000 trajectoires. Pour chacune de ces valeurs, nous faisons tourner ce GSE $m=30$ fois en

⁸ La caractérisation « réelle » vient du fait que dans notre cas la résolution du problème passe par la simulation et la projection de scénarios stochastiques pour les différents actifs: on parle dans ce cas d'une résolution empirique basée sur la méthode *Monte Carlo* et non d'une résolution réelle déterminée généralement à l'aide de formules explicites.

résolvant à chaque fois le problème d'optimisation présenté ci-dessus. 120 valeurs optimales empiriques de la fonction objectif sont ainsi obtenues (30 valeurs pour chacune des 4 tailles proposées). Les valeurs de ces différents paramètres (taille, nombre de fois qu'on fait tourner notre GSE) sont ici choisies de façon arbitraire. L'idée est de tester la stabilité et la distribution de la valeur optimale empirique de la fonction objectif ainsi que de la solution empirique optimale (c'est-à-dire l'allocation optimale d'actifs) par rapport à la taille de la structure. Cette étude essayera de juger *in fine* la qualité des solutions fournies par notre GSE pour le preneur de décision.

Avant d'exposer et d'analyser les résultats obtenus, il est à rappeler que la stabilité interne de l'échantillon nécessite que, pour une taille de structure donnée, les valeurs optimales de la fonction objectif ne dépendent pas de l'ensemble de scénarios projetés. L'idéal sera donc, dans ce cas, d'avoir, pour une taille de structure donnée, 30 valeurs optimales approximativement identiques de la fonction objectif. Pour la stabilité externe de l'échantillon, les valeurs réelles de la fonction objectif (ici la variance) doivent être approximativement les mêmes pour les différents ensembles de scénarios projetés.

De même, ces valeurs devront être approximativement égales à celles obtenues dans le test de stabilité interne. Comme mentionné précédemment, cette stabilité externe de l'échantillon peut être déduite à partir de la stabilité interne de l'échantillon au niveau de la solution empirique retrouvée (l'allocation optimale dans notre cas).

Dans le cadre du test de la stabilité interne de l'échantillon, le tableau 3 présente les statistiques relatives aux 30 valeurs optimales de la fonction objectif et ce pour les différentes tailles proposées.

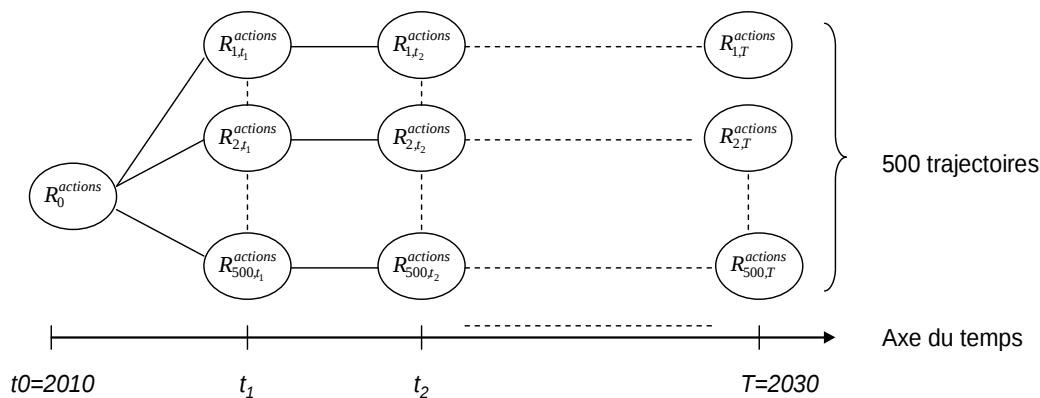


Fig. 8 : Projection du rendement des actions selon une structure schématique linéaire

Description du test			Nombre de scénarios			
Type de test	Fonction objectif	Valeur	50	1000	5000	10 000
Interne	$F\left(R_p; \vec{\varepsilon}_k\right)$	Moyenne	0,0288	0,0275	0,0275	0,0277
		Ecart-type	0,0050	0,0011	0,0007	0,0005

Tab. 3 : Propriétés statistiques des valeurs de la fonction objectif pour les différentes tailles retenues des structures de projection de scénarios

De même, le graphique ci-dessous illustre la distribution de ces valeurs optimales (médiane, quantile 25 % et 75 % ainsi que le minimum et le maximum sur les 30 arbres de chaque taille).

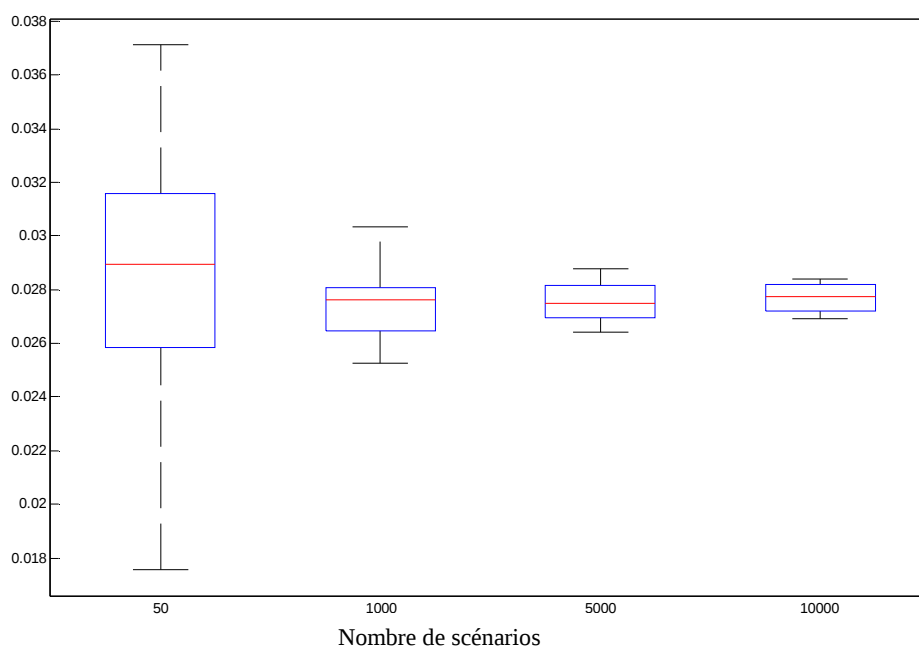


Fig. 9 : Box plot de la distribution des valeurs de la fonction objectif pour différentes tailles de scénarios

Comme nous pouvons le constater, plus le nombre de scénarios augmente plus la valeur optimale de la fonction objectif a tendance à baisser et à se stabiliser autour d'une valeur centrale (en particulier à partir de 5 000 scénarios). La condition de stabilité interne de l'échantillon peut ainsi être respectée par notre GSE à partir d'une certaine taille de la structure.

D'un autre côté, si la variation du poids optimal de chaque actif est relativement faible (bornée entre deux valeurs proches) cela garantit la stabilité externe de l'échantillon. Dans ce cadre, nous analysons les 30 vecteurs d'allocations optimales obtenus pour les ensembles de scénarios ayant pour taille 10 000. Les statistiques de ces solutions sont présentées dans le tableau suivant, ainsi que le graphique correspondant ci-dessous, et permettent de constater relativement une stabilité externe de l'échantillon (vu que les écarts types et/ou les étendus des poids des différents indices sont relativement faibles).

	Moyenne	Ecart-type	Etendue	Minimum	Maximum
Monétaire	0,37	0,0252	0,05	0,35	0,4
Obligations	0,55	0,0497	0,1	0,5	0,6
Actions	0,08	0,0254	0,05	0,05	0,1

Tab. 4: Propriétés statistiques des poids des 30 allocations optimales obtenues pour la taille de 10 000 scénarios

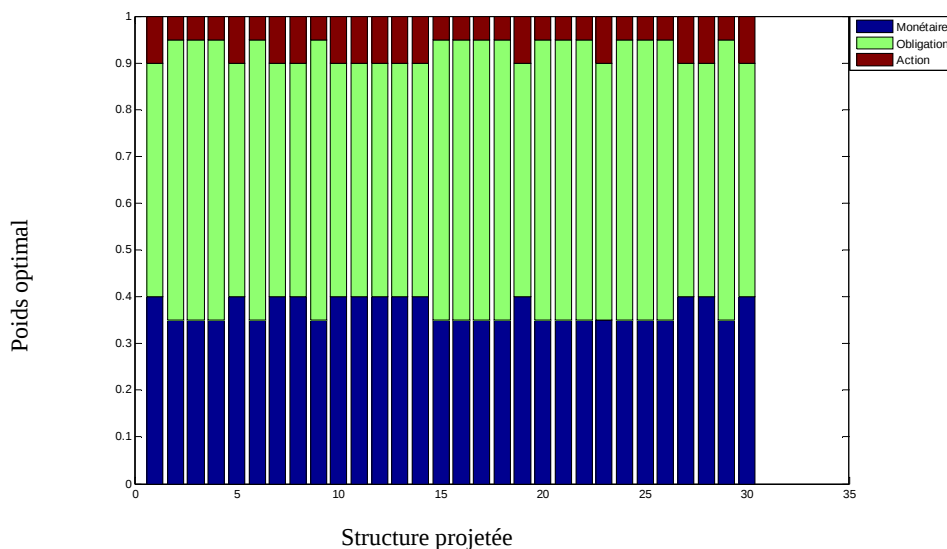


Fig. 10 : Les allocations optimales obtenues pour les différents ensembles de scénarios (30 ensembles) ayant chacun la taille de 10 000 scénarios

Conclusion du chapitre 1

Les résultats obtenus à travers cette application constituent des indicateurs que le décideur peut utiliser pour justifier la stabilité de sa décision finale par rapport à son générateur de scénarios : les graphiques et les tableaux présentés confirment la possibilité d'avoir une stabilité interne et une stabilité externe de l'échantillon. En pratique, les hypothèses de travail et les variables de risque prises en compte dans le processus décisionnel sont évidemment beaucoup plus nombreuses et plus compliquées mais cela ne limite pas l'intérêt de ces indicateurs comme des outils d'aide à la mesure de la performance de la décision.

La prochaine étape est une illustration plus globale du processus de construction d'un GSE. Le contexte de cette construction est celui de la projection sur le long terme de grandeurs réelles pour des fins de gestion actif-passif (cas des fonds de retraite et de certains contrats d'assurance). Le but du chapitre suivant est donc d'exposer la démarche suivie pour répondre à ce besoin ainsi que les points qui, en pratique, nécessitent une attention particulière de la part du modélisateur.

Chapitre 2 : Construction d'un générateur de scénarios économiques

I- Conception et composantes

Le générateur de scénarios économiques construit dans l'étude suivante tient compte des relations entre les différentes variables économiques et financières. La méthodologie de mise en place du modèle ainsi que les techniques de calibrage utilisées sont également exposées.

L'objet de l'utilisation de ce GSE sera la projection sur le long terme de différentes variables financières. Il s'agit en quelque sorte de préparer le terrain au fonctionnement des modèles de gestion actif-passif d'un régime de retraite comme nous allons le voir dans la partie II.

Nous supposons que les actifs constituant l'univers d'investissement sont les suivants : le monétaire, les obligations d'Etat à taux fixe (de différents niveaux de duration), les obligations d'Etat à taux indexé à l'inflation (de différents niveaux de duration), les obligations non gouvernementales à taux fixe (de duration moyenne 5 ans), les actions et l'immobilier.

Univers d'investissement	
Monétaire	-
ZC 5 ans (nominal)	obligations d'Etat à taux fixe de maturité 5 ans
ZC 10 ans (nominal)	obligations d'Etat à taux fixe de maturité 10 ans
ZC 15 ans (nominal)	obligations d'Etat à taux fixe de maturité 15 ans
ZC 8 ans (réel)	obligations d'Etat à taux indexé à l'inflation de maturité 8 ans
ZC 15 ans (réel)	obligations d'Etat à taux indexé à l'inflation de maturité 15 ans
Crédit 5 ans	obligations non gouvernementales à taux fixe de maturité 5 ans
Actions	-
Immobilier	-

Tab.5 : Univers d'investissement considéré dans l'étude

I-1 Conception théorique

Une première étape dans la conception théorique du GSE objet de l'illustration vise à arbitrer entre un modèle par cascade et un modèle basé sur les corrélations.

Comme il est mentionné dans la sous-section I-2 du chapitre précédent, les modèles basés sur les corrélations s'appuient sur des descriptions *ad hoc* de chaque classe d'actif et les agrègent ensuite pour proposer une description globale de l'actif. Dans le cadre d'un modèle basé sur les corrélations, chaque classe est ainsi modélisée finement ce qui, dans le cadre d'une démarche de gestion actif-passif (ALM), permet de décrire de manière précise l'allocation effectuée.

Les modèles par cascade proposent une description structurée de plusieurs classes d'actifs à partir en général d'une variable explicative de référence (comme par exemple le modèle de Wilkie [1995] qui s'appuie sur une modélisation de l'inflation dont découlent les taux et les

prix des actions). Les classes d'actifs tels que les taux d'intérêt peuvent apparaître modélisées de manière relativement sommaire en comparaison à la précision des modèles basés sur les corrélations.

Toutes ces observations, ainsi que celles présentées dans la sous-section II-2 du chapitre précédent, nous ont conduits à retenir un modèle proche de celui d'Ahlgrim et al. [2005] plutôt qu'un modèle du type Wilkie [1995] dans le cadre de la présente étude. Le modèle GSE développé ne supposera pas l'existence d'une structure par cascade entre ces variables.

Concernant le choix d'un modèle de structure par terme des taux, ce dernier doit tenir compte de l'objectif de planification stratégique de long terme qui diffère d'un objectif d'application de court terme (basé sur la comparaison des prix obtenus avec ceux des produits dérivés observés sur le marché). Dans cette étude, nous n'envisageons pas d'utiliser le modèle de GSE, objet de l'illustration, pour des opérations de « *trading* ». Nous cherchons plutôt à fournir aux fonds de retraite un ensemble de scénarios de taux d'intérêt et de rendement d'actifs qui soit plausible par rapport aux tendances futures et aux réalisations passées. Un cadre de modélisation avec des hypothèses **d'équilibre général** sera favorisé. Les différents actifs n'auront pas le même rendement espéré et une prime de risque sera accordée à chacun d'entre eux. En choisissant la structure par terme des taux, nous essayeront également de respecter les points mentionnés ci-dessous relatives à l'étude historique des mouvements de la courbe des taux.

L'étude historique des mouvements de la courbe des taux met en relief les points suivants (Martellini et al. [2003]) :

- les taux d'intérêt nominaux ne sont pas négatifs.
- les taux d'intérêt sont affectés par des effets de retour à la moyenne.
- les taux longs et les taux courts ne sont pas parfaitement corrélés.
- les taux à court terme sont plus volatiles que les taux à long terme.
- les courbes de taux d'intérêt peuvent avoir différentes formes (quasi-plate, croissante, décroissante, avec une bosse, avec un creux).

Dans le modèle développé, la structure par terme des taux nominaux est déduite à partir de la structure par terme des taux réels et de la structure par terme des taux d'inflation anticipés. Nous avons choisi de modéliser ces deux dernières par le modèle de Hull et White à 2 facteurs [1994] appliqué dans un cadre d'équilibre général, avec un principe de retour à la moyenne et avec un *drift* dépendant du temps. Il se base sur le principe de retour à la moyenne des taux d'intérêt. Ce modèle a le mérite d'être utilisé pour des fins de projection de long terme et dans un cadre d'analyse d'un « monde réel » (Ahlgrim et al. [2005]). Il fournit aussi des formules explicites pour les prix des obligations zéro-coupon ce qui permet à son utilisateur d'obtenir les taux de rendements de différentes maturités, à n'importe quelle date durant la période de simulation (courbe des taux).

Le choix d'un modèle multi-factoriel garantit une flexibilité au niveau de la forme de la courbe des taux. Il est montré d'un autre côté que le pouvoir explicatif d'un modèle à deux paramètres est suffisant pour la prévision des changements de taux d'intérêt (Date et al. [2009]). Ces deux facteurs retenus seront le taux long (sur un horizon long, 10 ans par exemple) et le taux court ou le taux instantané (sur un horizon court, 3 mois par exemple). Nous pouvons ainsi obtenir différentes formes de la courbe des taux en fonction de l'évolution de sa partie de long terme par rapport à sa position initiale. De même, en considérant deux facteurs, nous évitons la limite du modèle à un seul facteur qui est celle

d'une corrélation parfaite, égale à 1, entre les taux de rendement exigés sur différentes maturités. Ainsi, les taux courts et les taux longs n'auront pas à évoluer toujours dans un « *lock-step* ».

En ce qui concerne les actions, différents auteurs ont soulevé la problématique des régimes qui caractérisent le processus d'évolution de la rentabilité des actifs financiers. Le GSE que nous allons mettre en place ne suppose pas le retour à la moyenne pour les rendements des actions. La dynamique des rendements des actions se base sur le modèle de changement de régime tel que présenté par Hardy [2001]. Les actions peuvent évoluer dans le cadre de l'un de deux régimes suivants : un régime à espérance de rendement élevée et à volatilité faible des rendements des actions et un régime à espérance de rendement relativement faible (voire même négative) et à volatilité relativement élevée. Ce modèle permet de générer un *skewness* (coefficient d'asymétrie) et un *kurtosis* (coefficient d'aplatissement) plus proche de ceux de la distribution empirique des rendements des actions (Hibbert et al. [2001]). Le rendement des actions sous chaque régime est supposé suivre un mouvement brownien géométrique tel que celui de Black et Scholes [1973].

L'un des avantages de l'application des modèles à changement de régimes est la prise en compte de la non-normalité des distributions des processus à modéliser. De même, ce modèle permet de tenir compte de l'instabilité temporelle de la relation entre les taux (en particulier l'inflation) et les rendements des actions dont l'origine peut être attribuée à l'émergence de bulles spéculatives et/ou à des fluctuations non dictées par les fondamentaux (cf. Fama [1981]).

A ce niveau, notre étude se distingue par l'adoption d'une démarche novatrice dont l'objectif est de tenir compte de l'instabilité temporelle de la corrélation entre les rentabilités des actifs. Autrement dit, nous tenons compte dans la modélisation des possibilités d'évolution de cette corrélation selon différents régimes.

L'idée principale est de partir de l'approche classique de changement de régime, qui concerne souvent les rendements moyens ainsi que les volatilités, pour supposer que les corrélations basculent elles aussi d'un régime à l'autre. Nous introduisons un nouveau concept qui est celui de la « **corrélation à risque** » correspondant à la matrice de corrélation dans le régime de crise (régime à forte volatilité des marchés). Elle reflètera *a priori* la perception subjective de l'investisseur en fonction des risques auxquels il est exposé. La matrice de corrélation dans le régime à faible volatilité sera considérée comme celle reflétant la stabilité des marchés (estimée par exemple à partir de données historiques hors périodes de crises).

Concernant les valeurs des obligations non gouvernementales (à risque de crédit), ces dernières sont déterminées à travers l'approche des coefficients d'abattement : approche développée par les assureurs dans le cadre de la projection de leurs évaluations de provisions économiques (MCEV⁹ et Solvabilité II). Le monétaire est modélisé selon un modèle log-normal à volatilité faible. L'immobilier est modélisé selon un modèle de Vasicek [1977] (l'équivalent de Hull et White [1990] avec un seul facteur : le taux court).

La présentation du GSE, objet de l'illustration, sera accompagnée par la comparaison de ce dernier avec le modèle d'Ahlgrim et al. [2005] notamment au niveau du calibrage des paramètres.

⁹ *Market Consistent Embedded Value* est le standard de valorisation des compagnies d'assurance en juste valeur.

I-2 Composantes

En matière de gestion actif-passif, il s'avère utile de s'appuyer sur une modélisation des principales classes d'actifs disponibles dans le but de construire un modèle global dans lequel ces classes sont représentées de manière cohérentes dans une perspective de long terme en phase avec les besoins d'un régime de retraite.

I-2-1 Structure de corrélation

Deux principales questions s'imposent lors de l'étude des modèles de corrélation. D'une part, quelle est l'importance de l'estimation de la corrélation dans le contexte de l'application considérée? D'autre part, quel est le meilleur modèle d'estimation de cette corrélation ?

Concernant la première question sur l'importance de l'estimation des corrélations, nous nous référons à l'étude de Skintzi et al. [2007]. Ces derniers explorent l'importance de l'estimation de la corrélation dans le contexte de la gestion des risques. En particulier, ils étudient l'effet de l'erreur d'estimation de la matrice de corrélation sur le niveau estimé de la Valeur à Risque¹⁰ (VaR ou *Value-at-Risk*) d'un portefeuille d'actifs donné. Ils constatent que l'erreur d'estimation de la matrice de corrélation influence significativement le calcul de la VaR et mettent ainsi en évidence l'importance de minimiser cette erreur pour les gestionnaires de risque (*risk managers*).

Plus précisément, plus l'erreur d'estimation de la matrice de corrélation augmente plus l'erreur au niveau de la VaR augmente. Ceci est vrai que ce soit avec des portefeuilles linéaires (composés d'actifs classiques) ou avec des portefeuilles composés d'options. Pour ces deux portefeuilles, nous supposons que la VaR est calculée respectivement avec la matrice de variances-covariances et les méthodes de simulation *Monte Carlo*. Ainsi le choix d'un modèle qui minimise l'erreur d'estimation des corrélations est important pour le gestionnaire des risques.

Un résultat additionnel obtenu par les auteurs consiste dans le fait que plus le « vrai » niveau de la corrélation est faible plus la sensibilité du pourcentage d'erreur de calcul de la VaR est élevé. Ceci implique que, dans le cas d'un portefeuille linéaire, la sensibilité est d'autant plus importante que le portefeuille est bien plus diversifié. De même, dans le cas d'un portefeuille d'options, la sensibilité de l'erreur de calcul de la VaR par rapport à l'erreur d'estimation des corrélations est d'autant plus forte que le niveau des corrélations entre les actifs sous-jacents est faible.

Par ailleurs, l'étude de Skintzi et al. [2007] montre que, dans le cas d'un portefeuille d'options, le biais induit par l'erreur d'estimation des corrélations dépend du niveau par rapport à la monnaie et de la maturité. Les résultats montrent à ce niveau que la VaR relative à un portefeuille d'options avec des maturités courtes et *dans la monnaie* est la plus sensible à l'erreur de corrélation. Bien évidemment, ce résultat n'est pas le bienvenu chez les gestionnaires de risque, du fait que la plupart de l'activité de transaction des options est concentré sur des options à maturité courte et proche de la monnaie.

¹⁰ La VaR d'un portefeuille d'actifs financiers correspond au montant de pertes maximales sur un horizon de temps donné et avec un niveau de confiance prédéfini. Il s'agit la mesure de risque de référence utilisée dans Solvabilité II.

Finalement, les résultats de Skintzi et al. [2007] indiquent que la sensibilité de la VaR à l'erreur d'estimation de la matrice de corrélation dépend également de la méthode de calcul de la VaR. Ainsi, la pertinence des résultats obtenus par les simulations *Monte Carlo* doit être vérifiée à travers différents tests de *backtesting*. De même, il est souvent admis que, pour les portefeuilles linéaires, le gestionnaire de risque est indifférent dans le choix entre la matrice de variances-covariances et les simulations *Monte Carlo* : la VaR calculée à travers les méthodes *Monte Carlo* devra converger en principe vers celle obtenue par la matrice de variances-covariances. Cependant, comme le biais provenant de l'erreur d'estimation des corrélations dépend de la méthode employée, l'utilisation des simulations *Monte Carlo* est de préférence à éviter même dans le cas de petits portefeuilles linéaires.

Skintzi et al. [2007] confirment également que quelle que soit l'approche d'estimation retenue, l'erreur est toujours inévitable. En particulier, dans les périodes de stress des marchés, où la VaR est supposé être utile, les corrélations dévient significativement de leurs valeurs estimés initialement (voir Bhansali et Wise [2001] qui mettent en évidence les limites des techniques statistiques dans ce cas). A ce niveau, la particularité de l'étude menée par Skintzi et al. [2007] consiste dans le fait que ces derniers se sont intéressés à la relation entre l'erreur d'estimation de la matrice de corrélation et l'erreur de calcul de la VaR plutôt que de s'intéresser au modèle qui minimise l'erreur d'estimation de la matrice de corrélation.

Autrement dit, ils ont supposé que ce qui compte finalement pour les gestionnaires de risque est de minimiser l'erreur de calcul de la VaR, quel que soit le niveau d'erreur obtenu dans l'estimation de la matrice de corrélation. Ainsi, ils partent de la question : est ce que le choix d'une approche d'estimation de la matrice de corrélation, plutôt qu'une autre approche, impacte le niveau de la VaR obtenu *in fine* et à quel niveau ? Si de telles relations systématiques n'existent pas, le gestionnaire de risque se trouve libre dans le choix de son modèle d'estimation de la matrice de corrélation. Il pourra ensuite, par exemple, chercher à identifier d'autres facteurs menant à des erreurs dans son modèle de calcul de la VaR.

Même si l'importance des hypothèses de corrélation est de premier plan dans certaines applications, du fait de leur influence sur les décisions finales du gestionnaire (cf. Wainwright [1990]), certains chercheurs ont tendance à associer un degré plus faible à cette importance. C'est le cas de Black&Litterman [1992] qui affirment, que dans le cas d'une allocation d'actifs mono-périodique type Markowitz, le modèle est significativement plus sensible aux hypothèses émises sur les espérances et les volatilités des différentes classes d'actifs plutôt qu'aux hypothèses de corrélations adoptées dans le modèle. L'importance des hypothèses de corrélation est donc fonction du contexte d'application et dépend des spécificités de l'étude envisagée.

Concernant la deuxième question sur la performance de la méthode d'estimation des corrélations, de nombreux travaux ont essayé de trouver la réponse en évaluant la performance de la prévision de différents modèles de corrélations dans le cas de multiples applications financières. Citons par exemple, Elton et Gruber [1973] et Chan et al. [1999] dans le contexte de l'allocation d'actifs, Beder [1995] et Alexander et Leigh [1997] dans le contexte de la gestion des risques, Gibson et Boyer [1998] et Byström [2002] dans le contexte du *pricing* des options. Cependant, l'évidence en matière de la sélection du meilleur modèle de corrélation est loin d'être tranchée. De même, du fait que la plupart des applications nécessite la prévision simultanée de la volatilité et de la corrélation, l'évaluation de la performance de la prévision de la corrélation indépendamment de celle de la volatilité apparaît comme peu cohérente.

La corrélation entre les taux (en particulier l'inflation) et les rendements des actions sera étudiée dans la suite. Ce sujet se caractérise par le fait d'être controversé dans la littérature. L'équation de Fisher [1930] stipule que le taux nominal de rentabilité d'un actif financier, tel que les actions, est égal à la somme de l'inflation anticipée et du taux réel de rentabilité de l'action. Cette identité est basée sur deux hypothèses. La première est relative à l'efficience du marché des actions; tandis que la seconde stipule que le taux réel de rentabilité est déterminé par des facteurs réels, et qu'ainsi il est indépendant des anticipations inflationnistes.

Or, de multiples travaux empiriques (dont Fama [1981] et Kim [2003]) destinés à tester l'équation de Fisher révèlent une relation négative entre la rentabilité des valeurs et l'inflation. Cette nouvelle relation est qualifiée dans la littérature économique de "*stock return-inflation puzzle*". Une multitude d'études fut consacrée à l'analyse de cette relation inverse entre la rentabilité des actifs et une variété de mesures de l'inflation ou de l'inflation anticipée.

L'une des hypothèses implicites à la relation de Fisher est que les marchés financiers sont efficients. L'adoption d'une telle hypothèse est synonyme d'une relation stable entre les prix des actifs financiers et leurs déterminants fondamentaux. Or, les travaux de Shiller [1981] dont les résultats révélaient que la variance des cours boursiers est trop élevée par rapport à celles des fondamentaux vont à l'encontre de l'hypothèse d'efficience des marchés. Artus [1998] affirmait que la relation entre l'inflation et le prix des actifs peut être instable.

L'absence d'une relation stable, au moins à court terme, entre les cours boursiers et leurs déterminants fondamentaux peut être expliquée au niveau macroéconomique, soit par l'irrationalité ou la myopie des agents économiques, soit par l'existence de bulles rationnelles. Autrement dit, étudier la relation entre les prix des actifs financiers et l'inflation sans tenir compte de ces imperfections et anomalies conduit nécessairement à des relations erronées. Le graphique 11 illustre cette instabilité au niveau des corrélations glissantes entre l'inflation et les rendements des actions aux Etats-Unis : par exemple, pour deux crises différentes (celle de 2001 et de celle de 2008) ce niveau de corrélation passe de $-0,2$ à $0,2$ ce qui s'explique par la différence de la source de la crise et de ses effets sur le comportement des investisseurs.

En adoptant un modèle issu de la combinaison de la théorie de la demande de la monnaie et de la théorie quantitative de la monnaie, Fama [1981] affirmait que la relation négative entre les taux nominaux de rentabilité des actions et l'inflation n'est que le reflet du lien négatif entre cette dernière et l'activité économique réelle. Il explique que dans la mesure où l'activité économique réelle est négativement corrélée à l'inflation et puisque la rentabilité des actions est corrélée positivement à l'activité économique, la corrélation négative entre l'inflation et les taux nominaux de la rentabilité des actions est fausse. Elle ne représente qu'une relation proxy du lien entre les évolutions des prix des actions et de la production.

Ainsi, la thèse dite de "*proxy effect hypothesis*" de Fama est souvent avancée pour justifier les résultats empiriques opposés à l'identité de Fisher. Il faut souligner que la plupart des travaux empiriques relatifs à l'analyse de la relation de l'inflation et la rentabilité des actions ne tiennent pas compte de l'instabilité temporelle de cette relation dont l'origine peut être attribuée à l'émergence de bulles spéculatives et/ou à des fluctuations non dictées par les fondamentaux.

L'instabilité dans le temps de la corrélation entre les différents facteurs de risque, quel que soit l'horizon de calcul de cette corrélation (5 ans, 10 ans, etc.), constitue un élément d'étude important au niveau du choix du niveau de finesse d'un modèle d'actifs (cas de l'allocation

d'actifs). Cette importance provient de l'impact significatif des hypothèses de corrélation retenues sur les résultats finaux obtenus et ainsi sur la décision finale du *manager*.

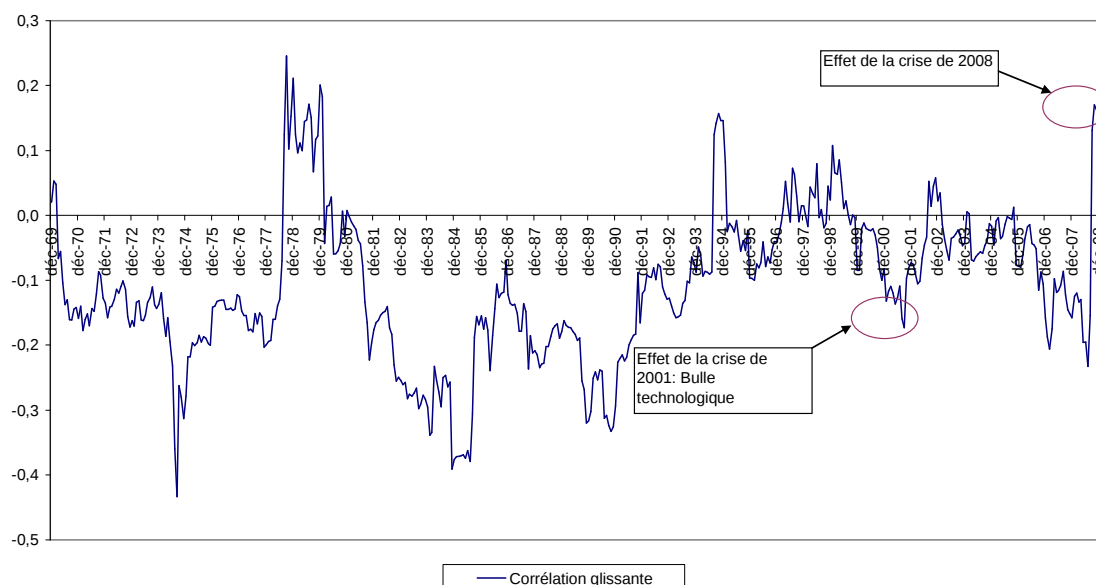


Fig.11 : Evolution de la corrélation (glissante) sur 5 ans entre l'inflation et le rendement des actions aux Etats-Unis depuis 1970

L'omission des régimes qui caractérisent les processus générateurs de la rentabilité des actions dans l'analyse de la relation " inflation-rentabilité " peut introduire des biais et suggérer la présence d'une corrélation négative. Si l'introduction des régimes génère une relation conforme à l'identité de Fisher (corrélation positive entre l'inflation anticipée et la rentabilité des actions), nous pourrions mettre en doute les arguments avancés par Fama. Et le rendement des actions serait, alors, un indicateur avancé de l'inflation future.

A ce niveau, notre étude se distingue par l'adoption d'une démarche novatrice dont l'objectif est de tenir compte de l'instabilité temporelle de la corrélation entre les rentabilités des actifs. Autrement dit, nous tenons compte dans la modélisation des possibilités d'évolution de cette corrélation selon différents régimes.

L'idée principale est de partir de l'approche classique de changement de régime, qui concerne souvent les rendements moyens ainsi que les volatilités, pour supposer que les corrélations basculent elles aussi d'un régime à l'autre. Un nouveau concept est ainsi introduit et sera appelé la « **corrélation à risque** ». Il s'agit de la matrice de corrélation dans le régime de crise (régime à forte volatilité des marchés). Elle reflètera *a priori* la perception subjective de l'investisseur en fonction des risques auxquels il est exposé. La matrice de corrélation dans le régime à faible volatilité sera considérée comme celle reflétant la stabilité des marchés (estimée par exemple à partir de données historiques hors périodes de crise).

En fait, pour le régime de faible volatilité (régime 1 pour le rendement des actions), nous avons supposé une matrice de corrélation, entre les variables économiques et financières (le rendement des actions, le rendement de l'immobilier, le monétaire, les taux d'inflation anticipés longs et courts, les taux réels longs et courts) qui est estimée à partir des données historiques après élimination des périodes aberrantes ou de crises. Elle correspond alors à la

matrice de corrélation en situation de stabilité économique. Cependant, nous avons choisi de fixer des corrélations dans le cas de situations de crises (régime 2 pour les rendements des actions), qui correspondent aux corrélations « à risque » pour l'investisseur : il s'agit de la relation entre les variables qui reflètera les craintes de l'investisseur. Ainsi, nous mettons en relief l'exposition au risque extrême en cas de crise.

Si nous considérons par exemple que notre investisseur craint une baisse des actions accompagnée par une hausse des taux (baisse de la valeur des obligations) due à une hausse des taux d'inflation (cas d'un passif indexé à l'inflation) et une baisse des taux réels (cas de ralentissement de l'économie), alors les corrélations sont fixées en situation de crise (régime 2 pour les actions) de façon à avoir une corrélation négative élevée entre les actions et les taux d'inflation, une corrélation positive entre les actions et les taux réels et implicitement une corrélation négative entre les taux d'inflation et les taux réels.

Dans un cadre d'évaluation *Mark-to Market* des engagements, comme la durée du passif est élevée par rapport à celle de l'actif, ce qui est le cas en général, du fait que les maturités des produits obligataires sont relativement courtes, la crainte de l'investisseur sera la baisse des actions accompagnée par une baisse des taux d'inflation anticipés et des taux réels : ce qui ramène le fonds de retraite à une situation de sous financement. Une fois que ce scénario est choisi comme scénario catastrophe, nous fixerons des corrélations positives élevées entre ces variables pour le cas d'un régime à volatilité élevée des actions.

Cette démarche intègre ainsi la méthodologie déterministe de « *stress scenario* » dans un contexte de simulation stochastique. Cela permet, même si la démarche proposée ne présente pas assez d'éléments quantitatifs, de tenir compte dans les simulations de la perception subjective du risque extrême par l'investisseur.

La corrélation entre les actifs est présentée ici dans un cadre « **non linéaire** » (avec deux régimes de corrélation). Au niveau de chaque régime, les corrélations entre les différents actifs du portefeuille peuvent être illustrées par exemple à travers la matrice ci-dessous :

	$dz_t^{OBLIGATION}$	dz_t^{ACTION}	$dz_t^{IMMOBILIER}$	$dz_t^{MONETAIRE}$
$dz_t^{OBLIGATION}$	dt	$\rho_{12}dt$	$\rho_{13}dt$	$\rho_{14}dt$
dz_t^{ACTION}	$\rho_{12}dt$	dt	$\rho_{23}dt$	$\rho_{24}dt$
$dz_t^{IMMOBILIER}$	$\rho_{13}dt$	$\rho_{23}dt$	dt	$\rho_{34}dt$
$dz_t^{MONETAIRE}$	$\rho_{14}dt$	$\rho_{24}dt$	$\rho_{34}dt$	dt

Tab. 6 : Illustration de la matrice de corrélation entre les différents actifs du portefeuille

où ρ_{12} : la corrélation entre les obligations et les actions.

Notons, cependant, que les difficultés liées à l'agrégation des risques peuvent être également résolues par l'introduction de copules : ces fonctions permettent de modéliser l'intégralité de la structure de dépendance entre plusieurs variables et donc de prendre en compte les dépendances non linéaires (cf. Horta et al. [2008], Armel et al. [2010]).

I-2-2 Choix des modèles de diffusions

Dans cette étude, le choix d'un modèle de diffusion concerne les variables financières suivantes : la structure par terme des taux d'intérêt (y compris les taux nominaux, réels et d'inflation), le rendement des actions, le rendement des obligations crédits et le rendement de l'immobilier. Ces variables sont retenues compte tenu de leur importance dans le cadre d'une allocation stratégique de long terme d'un fonds de retraite (cf. Campbell et al. [2001]).

I-2-2-1 *Le modèle de structure par terme des taux d'intérêt*

Une attention particulière doit être portée à l'inflation, qui est utilisée pour revaloriser les cotisations et les prestations de tels régimes. Par ailleurs, un modèle d'actifs dans un contexte de retraite se doit d'une part d'intégrer des fluctuations de court terme sur la valeur des actifs et d'autre part d'être cohérent à long terme avec les équilibres macro-économiques reliant l'inflation, les taux d'intérêt et le taux de croissance de l'économie.

En effet, à long terme, il est généralement considéré (cf. Planchet et Thérond [2007]) qu'il existe un lien étroit entre le taux d'intérêt nominal prévalant sur les marchés financiers, le taux d'inflation, c'est-à-dire l'évolution de l'indice des prix, et le taux d'intérêt réel. Plus précisément, le taux d'intérêt nominal est égal (ou converge) à la somme du taux d'intérêt réel et du taux d'inflation :

$$\text{Taux d'intérêt nominal} \approx \text{Taux d'inflation} + \text{Taux d'intérêt réel}$$

Simultanément, il existe un lien étroit à long terme entre le taux d'intérêt réel et le taux de croissance réel de l'économie. Nous considérons ainsi que le taux d'intérêt réel ne peut être longuement très différent du taux de croissance de l'économie sauf à provoquer des arbitrages entre activité financière et activité réelle d'une part, entre investissement dans un pays et investissement dans d'autres pays d'autre part.

Nous pouvons donc écrire :

$$\text{Taux d'intérêt réel à long terme} \approx \text{Taux de croissance de l'économie à long terme}$$

Ces contraintes ont donc été prises en compte dans le choix du modèle. Le graphique 12 montre une corrélation forte entre l'inflation et les taux d'intérêt nominaux de court terme au Royaume-Uni à partir de 1960.

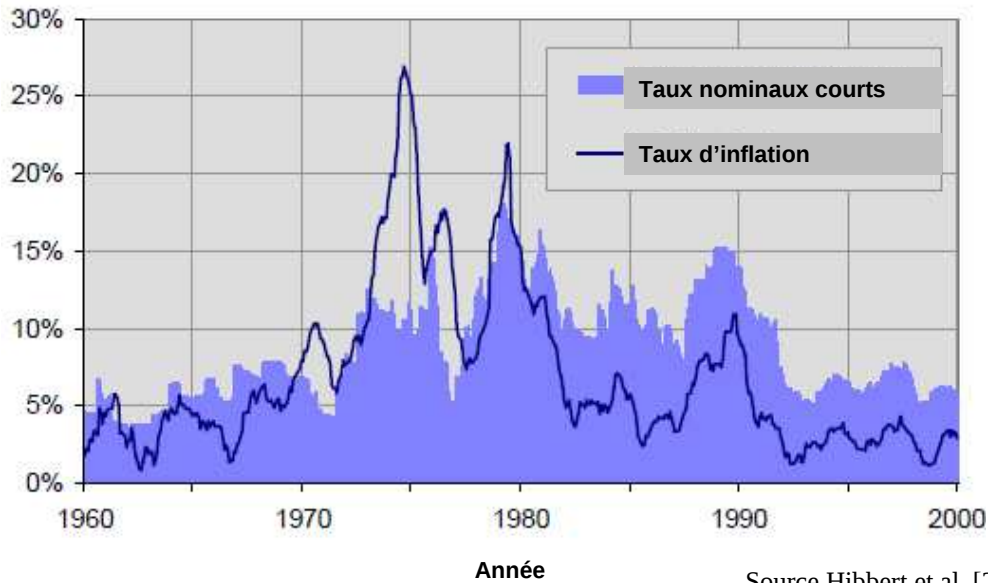
L'idée de base dans le modèle construit : la structure par terme des taux d'intérêt nominaux est établie à partir de deux composantes séparées :

- La structure par terme des taux réels : nous supposons qu'il s'agit de la courbe des taux des obligations indexées à l'inflation.
- Un modèle de l'inflation anticipée pour différents horizons. Il est implicitement supposé que les investisseurs comprennent le processus générant l'inflation et ajustent leurs anticipations de façon cohérente avec la trajectoire déjà déterminée.

Les deux structures par terme sont combinées afin d'obtenir celle des taux nominaux en tenant compte de :

- La corrélation entre ces deux composantes,

- La prime de risque du temps (prime de terme) et/ou la prime de risque inflation.



Source Hibbert et al. [2001]

Fig. 12 : La corrélation entre les taux d'inflation et les taux nominaux courts au Royaume-Uni

i. Taux d'intérêt réel : Modèle de Hull-White à 2 facteurs

Nous utilisons un modèle qui est l'extension de l'un des premiers modèles stochastiques de taux d'intérêt : le modèle de Vasicek [1977]. Ce dernier spécifie un processus stochastique en temps continu avec une propriété de retour vers la moyenne pour les taux de court terme (ou les taux instantanés). Nous pouvons ajouter une prime de risque demandée par l'investisseur pour avoir détenu l'obligation plus longtemps que le monétaire. Ce modèle est élargi par l'addition d'un deuxième facteur stochastique. Ce facteur prévoit que le niveau des taux d'intérêt de long terme suit lui aussi un processus stochastique de retour à la moyenne. Ce modèle est inspiré d'un modèle général de structure par terme à 2 facteurs décrit par Hull et White [1994]. Les équations régissant ce processus sont les suivantes :

$$\begin{aligned} dl_r(t) &= k_1(\mu_{lr} - l_r(t))dt + \sigma_1 dZ_1(t) \\ dr(t) &= k_2(l_r(t) - r(t))dt + \sigma_2 dZ_2(t) \end{aligned}$$

Avec :

$r(t)$ = le taux d'intérêt réel à court terme (à l'instant t) ou le taux réel instantané.

$l_r(t)$ = le taux d'intérêt réel à long terme (à l'instant t).

k_2 = le paramètre d'auto régression du processus des taux réels à court terme (vitesse de retour à la moyenne).

k_1 = le paramètre d'auto régression du processus des taux réels à long terme.

σ_2 = la volatilité annualisée (écart-type) des taux d'intérêt réels à court terme.

σ_1 = la volatilité annualisée (écart-type) des taux d'intérêt réels à long terme.

μ_{lr} = le taux d'intérêt réel à long terme moyen (ou le niveau de retour à la moyenne de $l_r(t)$).

$dZ_2(t)$ = le choc au processus de taux réels à court terme avec une distribution $N(g_r dt, dt)$.

$dZ_1(t)$ = le choc au processus de retour à la moyenne des taux réels à court terme avec une distribution $N(g_r dt, dt)$.

g_r = un paramètre qui reflète l'excès de rendement entre les obligations indexées à l'inflation de court et de long terme.

lb_r = la borne inférieure du taux réel à court terme $r(t)$.

hb_r = la borne supérieure du taux réel à court terme $r(t)$.

Le graphique suivant donne une idée sur les résultats pouvant être obtenus par ce modèle. Ici $r(t)$ représente le taux *forward* instantané (court terme). Nous pouvons ensuite déduire la courbe des taux spots en combinant les taux *forwards* appropriés.

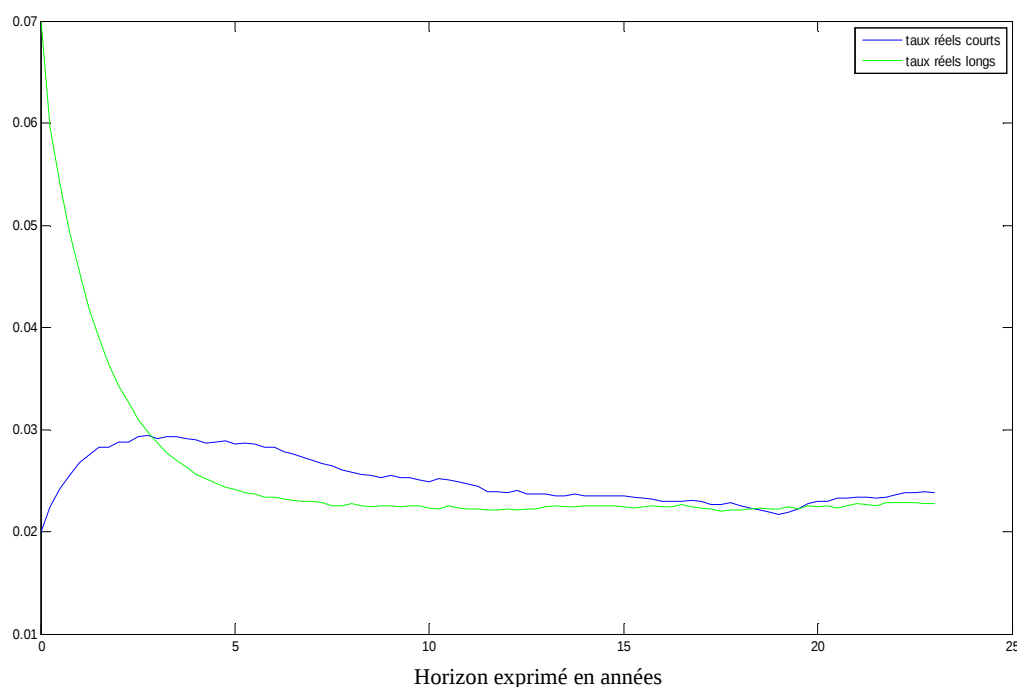


Fig.13 : Illustration du phénomène de retour à la moyenne des taux réels dans le cadre du modèle de Hull et White à deux facteurs

Le prix, à l'instant t , d'une obligation zéro-coupon qui paie une unité en terme réel à l'instant T est donné par les équations suivantes (cf. Hibbert et al. [2001]) :

$$P_{réel}(t, T) = \exp[A(T - t) - B_2(T - t)r(t) - B_1(T - t)l_r(t)]$$

Avec :

$$B_2(s) = \left[\frac{1 - e^{-k_2 s}}{k_2} \right]$$

$$B_1(s) = \frac{k_2}{k_2 - k_1} \left[\frac{1 - e^{-k_1 s}}{k_1} - \frac{1 - e^{-k_2 s}}{k_2} \right]$$

$$A(s) = (B_2(s) - s) \left(\mu_{lr} - \frac{\sigma_2^2}{2k_2} \right) + B_1(s) \mu_{lr} - \frac{\sigma_2^2 B_2(s)^2}{4k_2} + \frac{\sigma_1^2}{2} \left[\frac{s}{k_1^2} - 2 \frac{(B_1(s) + B_2(s))}{k_1^2} + \frac{1}{(k_2 - k_1)^2} \frac{(1 - e^{-2k_2 s})}{2k_2} - \frac{2k_2}{k_1(k_2 - k_1)} \frac{(1 - e^{-(k_1 + k_2)s})}{(k_1 + k_2)} + \frac{k_2^2}{k_1^2(k_2 - k_1)^2} \frac{(1 - e^{-2k_1 s})}{2k_1} \right]$$

Nous pouvons aussi calculer le rendement continu composé à l'instant t pour une maturité T :

$$R_r(t, T) = -\log\{P_{réel}(t, T)\} / (T - t)$$

ii. Le modèle d'inflation : Modèle de Hull-White à 2 facteurs

La même structure par terme est utilisée :

$$\begin{aligned} dl_i(t) &= k_3(\mu_{li} - l_i(t))dt + \sigma_3 dZ_3(t) \\ di(t) &= k_4(l_i(t) - i(t))dt + \sigma_4 dZ_4(t) \end{aligned}$$

Avec :

$i(t)$ = le taux d'inflation à court terme (à l'instant t) ou le taux d'inflation instantané.

$l_i(t)$ = le taux d'inflation à long terme (à l'instant t).

k_4 = le paramètre d'auto régression du processus des taux d'inflation à court terme (vitesse de retour à la moyenne).

k_3 = le paramètre d'auto régression du processus des taux d'inflation à long terme.

σ_4 = la volatilité annualisée (écart-type) des taux d'inflation à court terme.

σ_3 = la volatilité annualisée (écart-type) des taux d'inflation à long terme.

μ_{li} = le taux d'inflation à long terme moyen (ou le niveau de retour à la moyenne de $l_i(t)$).

$dZ_4(t)$ = le choc au processus de taux d'inflation à court terme avec une distribution $N(g_i dt, dt)$.

$dZ_3(t)$ = le choc au processus de taux d'inflation à long terme avec une distribution $N(g_i dt, dt)$.

g_i = un paramètre qui reflète la prime de risque inflation entre les obligations nominales et les obligations indexées à l'inflation.

lb_i = la borne inférieure du taux d'inflation à court terme $i(t)$.

hb_i = la borne supérieure du taux d'inflation à court terme $i(t)$.

Nous pouvons de la même façon ensuite déduire la structure par terme des anticipations d'inflation à partir du taux d'inflation à court terme $i(t)$ et de la valeur de $l_i(t)$.

Soit alors : $P(t, T)$ = la valeur actualisée d'une obligation à l'instant t qui paie 1 unité à l'instant T et dont l'actualisation se fait par rapport à l'anticipation d'inflation.

Le rendement équivalent sera :

$$R_i(t, T) = -\log\{P_{inflation}(t, T)\} / (T - t)$$

Nous pouvons aussi inclure la prime de risque d'inflation dans la structure par terme des anticipations d'inflation. Comme pour les taux réels, deux bornes (supérieure et inférieure) peuvent être fixées par l'investisseur.

iii. La structure par terme des taux nominaux

Hibbert et al. [2001] explicitent la valeur du taux nominal et le prix d'un zéro coupon nominal en fonction de la relation entre les taux réels et les taux d'inflation anticipés dans le cadre d'un modèle de Hull et White à 2 facteurs.

- Cas de l'indépendance entre les taux réels et les taux d'inflation anticipés :

$$\text{Taux nominal} = i + r$$

où :

i est l'inflation (de court terme ou de long terme suivant le cas) ;

r est le taux d'intérêt réel (de court terme ou de long terme suivant le cas) .

L'origine de cette formule, qui n'est qu'une approximation, provient de la proposition de Fisher [1930] qui stipule que le taux d'intérêt nominal i est déduit de la relation suivante :

$$\text{Taux nominal} = [(1+i) \times (1+r)] - 1$$

Comme conséquence nous obtenons (Hibbert et al. [2001]) :

$$P_{nom}(t, T) = P_{réel}(t, T) \times P_{inflation}(t, T)$$

- Cas de la dépendance entre les taux réels et les taux d'inflation anticipés :

Nous ajoutons une quantité reflétant la covariance entre les deux variables (Hibbert et al. [2001]) :

$$P_{nom}(t, T) = P_{réel}(t, T) \times P_{inflation}(t, T) + \rho \sqrt{[Var(\exp\{-R_r(t, T)\})Var(\exp\{-R_i(t, T)\})]}$$

Avec :

ρ = la corrélation entre le choc du taux réel instantané et le taux d'inflation anticipé instantané $dZ_2(t)$ et $dZ_4(t)$.

$$R_r(t, T) = \int_t^T r(s) ds$$

$$R_i(t, T) = \int_t^T i(s) ds$$

Afin de calculer cette covariance, nous partons de l'expression de la variance suivante correspondant aux taux réels:

$$Var(\exp\{-R_r(t, T)\}) = \exp\{-2E(R_r(t, T)) + Var(R_r(t, T))\} \cdot (\exp\{Var(R_r(t, T))\} - 1)$$

Avec :

$$E(R_r(t, T)) = \mu_{lr} \cdot (T - t) + x_1 + x_2$$

$$Var(R_r(t, T)) = y_1 \cdot (T - t) - 2y_2y_3 - 2y_4y_5 + (1/k_2)y_6y_7 - 2y_8y_9 + y_{10}y_{11}$$

$$x_1 = (r(t) - \mu_{lr}) \frac{(1 - \exp\{-k_2(T - t)\})}{k_2}$$

$$x_2 = \frac{k_2}{k_1 - k_2} \frac{(1 - \exp\{-k_1(T - t)\})}{k_1} - (l_r(t) - \mu_{lr}) \frac{(1 - \exp\{-k_2(T - t)\})}{k_2}$$

$$y_1 = \left(\frac{\sigma_2}{k_2}\right)^2 + \left(\frac{\sigma_1}{k_1}\right)^2$$

$$y_2 = \left(\frac{\sigma_2}{k_2}\right)^2 - \frac{\sigma_{k_1}^2}{k_1(k_2 - k_1)}$$

$$y_3 = \frac{1 - \exp\{-k_2(T - t)\}}{k_2}$$

$$y_4 = \frac{k_2 \sigma_1^2}{k_1^2(k_2 - k_1)}$$

$$y_5 = \frac{(1 - \exp\{-k_1(T - t)\})}{k_1}$$

$$y_6 = \sigma_2^2 + \frac{k_2^2 \sigma_1^2}{(k_2 - k_1)^2}$$

$$y_7 = \frac{(1 - \exp\{-2k_2(T - t)\})}{2k_2}$$

$$y_8 = \frac{k_2 \sigma_1^2}{k_1(k_2 - k_1)^2}$$

$$y_9 = \frac{(1 - \exp\{-(k_2 + k_1)(T - t)\})}{k_2 + k_1}$$

$$y_{10} = \frac{k_2^2 \sigma_2^2}{k_1^2(k_2 - k_1)^2}$$

$$y_{11} = \frac{(1 - \exp\{-2k_1(T-t)\})}{2k_1}$$

Une expression analogue est utilisée pour la variance correspondant à la structure des taux d'inflation anticipés. Ainsi, à travers la prise en compte du produit de ces deux variances et de la corrélation, nous obtiendrons le terme de covariance qui permettra le calcul du prix d'un zéro-coupon nominal dans le cadre de la dépendance des taux.

iiii. Prise en compte de la prime de terme dans la STTI

La prime de terme (ou de l'horizon) a pour effet d'ajuster la trajectoire du taux d'intérêt de court terme, de façon à ce que la courbe des taux ne soit pas un estimateur biaisé de l'évolution future effective des taux courts (cf. Hibbert et al. [2001]).

Etant données les valeurs actuelles à l'instant t de $r(t)$ et de $l_r(t)$ (supposées ici être respectivement le taux nominal court et long terme), nous pouvons calculer les valeurs espérées et les variances de ces deux variables à la date $T > t$. Elles sont obtenues par les expressions suivantes :

$$\begin{aligned} E[l_r(T) | l_r(t)] &= \mu_{lr} + e^{-k_1(T-t)}(l_r(t) - \mu_{lr}) \\ E[r(T) | r(t), l_r(t)] &= \mu_{lr} + e^{-k_2(T-t)}(r(t) - \mu_{lr}) + \frac{k_2}{k_2 - k_1}(e^{-k_1(T-t)} - e^{-k_2(T-t)})(l_r(t) - \mu_{lr}) \\ \text{Var}[l_r(T) | l_r(t)] &= \frac{\sigma_1^2}{2k_1}(1 - e^{-2k_1(T-t)}) \\ \text{Var}[r(T) | r(t), l_r(t)] &= \frac{\sigma_2^2}{2k_2}(1 - e^{-2k_2(T-t)}) + \frac{\sigma_2^2}{2k_2}(1 - e^{-2k_2(T-t)}) + \dots \\ &\quad \dots \left(\frac{k_2 \sigma_1}{k_2 - k_1} \right)^2 \left[\frac{1}{2k_2}(1 - e^{-2k_2(T-t)}) + \frac{1}{2k_1}(1 - e^{-2k_1(T-t)}) - \frac{2}{k_2 - k_1}(1 - e^{-(k_2+k_1)(T-t)}) \right] \end{aligned}$$

Comme $r(t)$ et $l_r(t)$ sont normalement distribuées, nous aurons besoin seulement des moments ci-dessus afin de simuler ces distributions et intégrer à la STTI la prime de terme.

Ainsi :

$$\begin{cases} l_r(T) = E[l_r(T) | l_r(t)] + \sqrt{\text{Var}[l_r(T) | l_r(t)]} Z_1(T) \\ r(T) = E[r(T) | r(t), l_r(t)] + \sqrt{\text{Var}[r(T) | r(t), l_r(t)]} Z_2(T) \end{cases}$$

Avec : $Z_1(T)$ et $Z_2(T)$ sont deux mouvement browniens standards indépendants.

Cependant, si nous avons une prime de terme non nulle, g_r , ces équations sont ajustées de la façon suivante :

$$\begin{cases} l_r(T) = E[l_r(T) | l_r(t)] + \sqrt{\text{Var}[l_r(T) | l_r(t)]} \cdot (Z_1(T) + g_r \cdot \sqrt{(T-t)}) \\ r(T) = E[r(T) | r(t), l_r(t)] + \sqrt{\text{Var}[r(T) | r(t), l_r(t)]} \cdot (Z_2(T) + g_r \cdot \sqrt{(T-t)}) \end{cases}$$

La nouvelle courbe obtenue reflètera l'attitude de l'investisseur vis-à-vis du risque. Par exemple, si l'investisseur exige un rendement supplémentaire sur une obligation de long terme, alors g_r est négative et la valeur espérée finale d'une position dynamique dans le monétaire sera inférieure à la valeur espérée finale d'un investissement initial en obligation de maturité par exemple de 10 ans (avec l'hypothèse d'un rebalancement continu permettant de garder la même maturité dans le portefeuille).

La prime de l'horizon peut être exprimée en termes de rendement espéré sur une obligation zéro-coupon (*Yields*), ou en termes de rentabilité espérée (*Returns*). En se référant à Hibbert et al. [2001], nous avons :

$$\text{Prime de Terme (Rendement)} = -g_r \left(\frac{\sigma_2}{k_2} + \frac{\sigma_1}{k_1} \right) - \frac{1}{2} \left(\frac{\sigma_2^2}{k_2^2} + \frac{\sigma_1^2}{k_1^2} \right)$$

$$\text{Prime de Terme (Rentabilité)} = -g_r \left(\frac{\sigma_2}{k_2} + \frac{\sigma_1}{k_1} \right)$$

Notons qu'une prime de terme positive nécessite une valeur de g_r négative.

Dans le modèle de l'inflation, le paramètre de la prime de risque reflète le rendement supplémentaire exigé sur une obligation nominale par rapport à une obligation indexée à l'inflation, et cette « prime de risque inflation » est intégrée de la même façon dans le modèle de taux d'intérêt réel.

I-2-2-2 Le modèle des actions

Face aux limites des modèles classiques (hypothèses de normalité, etc.), nous avons opté pour l'utilisation du modèle de changement de régime (RSLN : *Regime Switching Log Normal*, cf. Hardy [2001]) pour le cas du GSE objet de l'illustration.

En effet, une méthode simple pour tenir compte de la volatilité stochastique est de supposer que la volatilité peut prendre une des K valeurs discrètes, en passant d'une valeur à l'autre de façon à définir par le modélisateur. Ceci est la base du modèle lognormal de changement de régime. Ce modèle maintient le caractère de la simplicité fourni par le modèle lognormal classique. De même, il permet de refléter les mouvements extrêmes observés sur les marchés.

Le principe de changement de régime a été introduit par Hamilton [1989] qui a décrit un processus de changement de régime autorégressif. Dans Hamilton et Susmel [1994] différents modèles de changement de régime sont analysés, en testant un nombre différent de régime et des modélisations différentes. L'objectif est de modéliser des séries économétriques hebdomadaires. Les deux chercheurs concluent qu'un modèle à volatilité autorégressive conditionnellement hétéroscédastique de type ARCH est le plus adéquat.

Le modèle de changement de régime a été proposé la première fois comme un modèle traduisant la dynamique des actions par Hardy [2001]. L'idée sous jacente à ce modèle est que les marchés peuvent passer d'un instant à l'autre d'un régime à faible volatilité à un régime à

volatilité relativement élevée. Les périodes de volatilité élevée peuvent être induites par exemple par une politique monétaire de court terme ou par une incertitude économique élevée. Ceci permet de générer des distributions de rendements qui sont plus pertinentes par rapport aux distributions empiriques.

La distribution des rendements des actions est ainsi supposée transiter entre différents états de la nature (régimes). Une matrice de transition contrôle les probabilités de passage entre les régimes. En fait, le processus proposé est markovien : la probabilité de changement de régime dépend seulement du régime courant et non pas de la totalité de la trajectoire historique du processus.

I-2-2-2-1 Le modèle initial de changement de régime

Les rendements des actions ont été décomposés en deux parties : rendement des actions hors dividendes et taux de dividende (Hardy [2001] et Ahlgrim et al. [2005]).

i. Rendement des actions hors dividendes

Les rendements (hors dividendes) des actions, notés s_t , sont égaux au taux d'intérêt nominal (taux sans risque), soit $i_t + r_t$, augmenté d'un excès de rendement attribuable à l'appréciation du capital et noté x_t :

$$s_t = i_t + r_t + x_t$$

Hardy [2001] applique le modèle à changement de régime à l'excès de rendement des actions. En outre, Ahlgrim et al. [2005] reprennent la même approche et estiment les six paramètres de ce modèle à partir de deux jeux de données distincts : un premier jeu (*large stocks*) comprenant des séries observées entre 1871 et 2002 et un deuxième jeu (*small stocks*) comprenant des séries observées entre 1926 et 1999.

ii. Taux de dividende

Ahlgrim et al. [2005] proposent de modéliser le taux de dividende, noté y_t , par :

$$d(\log(y_t)) = \kappa_y (\mu_y - \log(y_t))dt + \sigma_y dB_{y,t}$$

où :

κ_y est la vitesse de retour à la moyenne ;

μ_y est le logarithme du taux de dividende moyen.

Le modèle sur le logarithme du taux de dividende est comparable à celui relatif à l'inflation (avec un seul facteur).

Toutefois, une des principales difficultés du sous-modèle sur le taux de dividende est la disponibilité des données. En effet, les séries de données sur les taux de dividende avec un historique suffisamment long ne sont pas disponibles publiquement. Aussi, pour estimer les paramètres de ce sous-modèle, Ahlgrim et al. [2005] utilisent des données privées. Une autre alternative proposée par Ahlgrim et al. [2005] serait de retenir les taux de dividendes sur un nombre limité d'actions.

De ce fait, nous préférons travailler avec des rendements d'actions déterminés avec une hypothèse de réinvestissement des dividendes.

I-2-2-2 Rendement des actions avec dividendes réinvestis : approche alternative

Au regard des difficultés susceptibles d'être rencontrées dans le cadre des spécifications du modèle d'Ahlgrim et al. [2005] (modélisation à partir de l'excès de rendement sur les actions et des taux de dividendes), une alternative est proposée pour modéliser le rendement des actions.

Cette alternative repose sur une modélisation des rendements totaux des actions (et non uniquement l'excès de rendement) avec dividendes réinvestis. En termes de choix de modèle, l'approche à changement de régime d'Hardy [2001] est maintenue.

Soit S_t le cours des actions (dividende réinvesti) à la date t et $Y_t = \log(S_{t+1} / S_t)$ le logarithme du rendement des actions pour la $t+1^{\text{ième}}$ période. Dans le modèle à changement de régime, nous considérons :

$$\begin{aligned} Y_t \mid \rho_t &\sim N(\mu_{\rho_t}, \sigma_{\rho_t}^2) \text{ où } \rho_t \text{ représente le régime appliqué à l'intervalle } [t, t+1] \text{ (avec } \\ &\rho_t = 1 \text{ ou } 2) ; \\ p_{i,j} &= \Pr[\rho_{t+1} = j \mid \rho_t = i] \text{ représente la probabilité de changement de régime (avec } i=1 \\ &\text{ ou } 2, j=1 \text{ ou } 2) ; \\ \Theta &= \{\mu_1, \mu_2, \sigma_1, \sigma_2, p_{1,2}, p_{2,1}\} \text{ représente les paramètres à estimer.} \end{aligned}$$

Par ailleurs, nous notons que Hardy a testé dans ce cadre le cas des modèles à deux et à trois régimes ($K=2,3$) et elle constate qu'en ajoutant un troisième régime la modélisation de S_t n'est pas améliorée de façon significative. Garder deux régimes seulement permet en plus la simplification du modèle et une pertinence meilleure lors de l'estimation des paramètres.

Ainsi, dans le modèle de GSE objet de l'illustration, nous avons également supposé l'existence unique de deux régimes. La projection de la trajectoire des rendements entre ces deux régimes se base sur un algorithme de simulation d'une distribution uniforme cumulée. Cette dernière permet de déterminer le régime de volatilité de la période suivante. Les corrélations des actions (au niveau des browniens) avec les autres variables du modèle diffèrent selon les régimes. En cas de régime à volatilité élevée (crise), nous supposons une corrélation « à risque » pour l'investisseur comme expliqué ci-dessus.

I-2-2-3 Le modèle des obligations "crédit"

Portrait et Poncet [2008] précisent que les titres de créance sont généralement affectés par un risque de crédit, qui englobe un risque de défaut (défaillance de l'emprunteur) et un risque de signature (détérioration du crédit de l'emprunteur). Le *spread* de taux traduit ce risque de crédit : il est égal à la différence entre le taux promis (*a priori* observable pour chaque titre) et le taux sans risque.

Dans le détail, le *spread* est composé de l'espérance de baisse du taux dû à la prise en compte du défaut, d'une prime de risque reflétant l'aversion au risque dans le monde réel et d'une prime de liquidité (cf. Portrait et Poncet [2008]). Deux approches peuvent être envisagées *a priori* pour intégrer le risque de crédit au modèle décrit ici :

- la modélisation des probabilités de défaut historiques et leur lien avec les *spreads* observés pour pouvoir calculer le prix des titres présentant une possibilité de défaut ;
- une démarche simplifiée utilisée couramment dans les MCEV (*Market Consistent Embedded Value*) et les états Solvabilité II (QIS ou *Quantitative Impact Studies*, bilan économique) des assureurs vie et basée sur le calibrage d'abattements sur le nominal du titre obligataire.

Ces deux approches sont décrites ci-après, ainsi qu'une troisième approche d'un modèle « mélange » (qui fera l'objet d'une étude comparative avec la démarche des coefficients d'abattement).

I-2-2-3-1 Modélisation des probabilités de défaut

Deux éléments doivent être décrits ici :

- l'impact du défaut sur la détermination du prix du titre ;
 - la dynamique des probabilités historiques de défaut.
- i. L'impact du défaut sur le prix du titre :

À l'image de l'illustration de Portrait et Poncet [2008], considérons un titre zéro-coupon de durée T quelconque, qui vaut 1 en date 0 et qui promet un flux $X_{max} = e^{(r_T^{max} T)}$ en date T , avec r^{max} le taux servi par le titre. L'espérance du flux aléatoire X réellement payé par le titre en T s'écrit alors :

$$E^*(X) = (1 - \phi_T^*) e^{(r_T^{max} T)} + \phi_T^* \alpha e^{(r_T^{max} T)}$$

où :

ϕ_T^* représente la probabilité cumulée, en univers risque neutre, que l'émetteur soit en défaut de paiement en date T ;

α représente la fraction de remboursement de sa dette en cas de défaut de paiement en date T (taux de recouvrement).

Par ailleurs, dans un univers risque-neutre la prime de risque est nulle, et en négligeant la prime de liquidité, nous avons :

$$E^*(X) = e^{(rT)}$$

où r représente le taux sans risque. De ces deux dernières égalités, nous déduisons que le *spread* s est égal à :

$$s = r_T^{max} - r = -\frac{1}{T} \log(1 - \phi_T^*(1 - \alpha))$$

d'où nous déduisons que la probabilité de défaut cumulée entre 0 et T est :

$$\phi_T^* = \frac{1}{1-\alpha} \left(1 - \frac{1}{e^{(sT)}} \right)$$

Cette approche établie donc un lien formel entre la possibilité de défaut et l'existence d'un *spread* par rapport au taux sans risque. Le prix d'une obligation présentant un risque de défaut $P_D(0,T)$ est alors obtenu à partir du prix sans risque de défaut $P(0,T)$ via la formule :

$$P_D(0,T) = (1-\phi_T^*)P(0,T) + \phi_T^*\alpha P(0,T).$$

Les *spreads* sont utilisés pour calculer les probabilités de défaut "risque neutre" intervenant dans la formule ci-dessus.

ii. Les probabilités historiques de défaut

Dans le cadre de l'univers risque-neutre utilisé ci-dessus pour la détermination de l'impact sur le prix de la possibilité de défaut, la prime de risque caractérisant l'aversion au risque est nulle. Au niveau de la probabilité de défaut, cette absence de primes de risque est compensée par une surpondération des événements défavorables et conduit ainsi à une majoration significative de cette probabilité. En d'autres termes, les probabilités de défaut en univers risque-neutre estimées à partir des *spreads* sont nettement supérieures aux probabilités de défaut historiques (cf. Portrait et Poncet [2008] pour une illustration).

Mais dans la modélisation des scénarios économiques il est nécessaire de disposer non seulement du prix des actifs obligataires à chaque date, mais également de tenir compte des probabilités de défaut historiques observées sur le marché. Ces probabilités de défaut devront être cohérentes avec le rating des contreparties et avec la maturité des titres.

Usuellement ces probabilités sont présentées en fonction de la notation de l'émetteur pour un horizon d'un an. La question de la dynamique temporelle de ces probabilités est délicate et fait l'objet d'une abondante littérature (voir par exemple Bluhm et Overbeck [2007]). Afin de ne pas complexifier trop le modèle, il est retenu ici de considérer les probabilités de défaut comme fixes pour un niveau de notation déterminé.

En pratique, nous utiliserons les probabilités de défaut cumulées communiquées par des agences de notations telles que Moody's ou Standard & Poor's. Ces probabilités historiques permettent de simuler des événements de défaut lors de la projection des flux.

I-2-2-3-2 Démarche des coefficients d'abattement

L'approche ci-dessus présente cependant trois défauts importants en pratique :

- l'estimation des paramètres, notamment dans une perspective de long terme, n'est pas simple et de ce fait les modèles ainsi calibrés s'avèrent peu robustes ;
- la manipulation d'une structure par terme des défauts est contraignante en termes de stockage comme de simulation ;
- la simulation du défaut dans le cadre de la projection de l'actif alourdit l'algorithme et s'avère coûteuse en temps de calcul.

C'est pourquoi il est proposé d'utiliser une méthode alternative développée par les assureurs dans le cadre de leurs évaluations de provisions économiques (MCEV et Solvabilité II). Cette méthode consiste simplement à utiliser pour les obligations avec défaut le modèle d'évaluation des obligations sans risque de défaut en abattant le nominal de manière à recalculer le prix théorique issu du modèle avec le prix de marché initial.

Le calcul du rendement du crédit repose sur le coefficient d'abattement qui est le rapport entre le prix du marché de l'obligation crédit et le prix théorique de l'obligation zéro coupon nominale d'une même maturité.

$$\text{coeff_abtt} = P_{\text{crédit}} / P_{\text{nominal}}$$

Où coeff_abtt est le coefficient d'abattement recherché. De même, par analogie au calcul du prix de l'obligation nominale, le calcul du prix du marché de l'obligation crédit versant un euro à sa maturité est réalisé à partir de l'actualisation des flux au taux actuariel de marché de l'obligation, soit :

$$P_{\text{crédit}} = \frac{1}{(Tx_{\text{crédit}} + 1)^M}$$

où :

$Tx_{\text{crédit}}$: le taux actuariel observé sur le marché des obligations *corporates* ;

$P_{\text{crédit}}$: le prix de l'obligation à l'instant du calcul ;

M : la maturité restante.

Cette modélisation du défaut est déterministe (le coefficient calculé à la date initiale est supposé être constant sur toute la période de projection), mais présente de nombreux avantages pratiques :

- calibrage simple et cohérence avec les prix de marché ;
- pas de simulation du défaut lors des projections, la valeur des titres reflète déjà ce risque ;
- approche légitimée par son statut de « standard de place ».

Elle présente l'inconvénient de ne pas impacter le rendement des titres par l'introduction du défaut, qui se traduit uniquement par un effet de minoration de la valeur.

I-2-2-3-3 Modèle « mélange » de crédit

Ce modèle s'appuie sur la modélisation des *spreads* avec un modèle à changement de régime, selon la logique suivante :

- le *spread* est projeté avec un modèle RSLN2;
- le prix des zéro-coupons est ensuite calculé à partir du taux d'intérêt (projeté par un modèle de Vasicek) majorés de ce *spread* en remplaçant dans la formule fermée du prix du zéro-coupon sans risque de défaut le taux court par le taux court majoré du *spread* ;
- les prix projetés sont ensuite utilisés pour déterminer le taux de rendement ρ_0 et sur ce rendement est appliqué le taux de défaut et de recouvrement pour obtenir finalement :

$$\rho = (1 + \rho_0) \times (1 - \pi \times (1 - d)) - 1$$

avec π la probabilité de défaut historique et d le taux de recouvrement.

Autrement, le « spread » de crédit ou l'écart de taux de rendement exigé entre les obligations gouvernementales et les obligations non gouvernementales est modélisé selon un modèle mixte RSLN-B&S tel que décrit ci-dessus. Le « spread » simulé est intégré dans le modèle de taux d'intérêt afin de déterminer le prix d'un zéro-coupon avec risque de crédit. Ce prix ne tient pas compte de la probabilité de défaut : cette dernière est présente lors du calcul de la performance de l'actif et elle est considérée comme constante. De même, le taux de recouvrement, qui reflète la part récupérée en cas de défaut de la créance initiale, est supposé comme étant une constante.

Outre les difficultés pratiques de calibrage des *spreads* à partir de données historiques, ce modèle s'appuie sur deux approximations fortes et délicates à justifier :

- Déterminer le prix d'un titre présentant un risque de défaut utilisant la formule fermée du modèle de base dans laquelle le taux court r_t est remplacé par le taux court majoré du *spread* implique que $x_t = r_t + s_t$ devrait lui-même suivre un modèle de Vasicek, ce qui n'est pas compatible avec l'hypothèse RSLN2 faite sur s_t ;
- La formule $\rho = (1 + \rho_0) \times (1 - \pi \times (1 - d)) - 1$ doit être calculée avec des estimations « risque neutre » de π et de d , ce qui n'est pas l'usage qui en est fait actuellement.

Au surplus, le défaut est introduit deux fois dans ce modèle : la première *via* le *spread* dans la formule de prix et une seconde fois dans le rendement, alors que la justification de l'existence du *spread* est précisément l'existence du défaut, comme nous l'avons discuté au début de cette partie.

I-2-2-4 Le modèle de l'immobilier

Ahlgrim et al. [2005] admettent que les rendements issus de l'immobilier, notés $(re)_t$, sont modélisés par Hulle et White à un facteur (l'équivalent de Vasicek [1977]) :

$$d(re)_t = \kappa_{re} (\mu_{re} - (re)_t) dt + \sigma_{re} dB_{re,t}$$

où :

- κ_{re} est la vitesse de retour à la moyenne ;
- μ_{re} est le rendement moyen issu de l'immobilier.

I-2-2-5 Le monétaire

Le monétaire (M) est modélisé selon un modèle lognormal avec un brownien corrélé avec les autres actifs. Il représente l'actif sans risque de court terme.

$$dM_t = \mu_M \times dt + \sigma_M \times dZ_M$$

Avec :

M est le taux de rendement du monétaire

μ_M est le rendement annuel moyen du monétaire

σ_M est la volatilité du rendement du monétaire

dZ_M est un mouvement brownien corrélé avec les autres variables du modèle

Le graphique 14 illustre la structure adoptée dans le GSE proposé et met en évidence le lien entre les différentes variables.

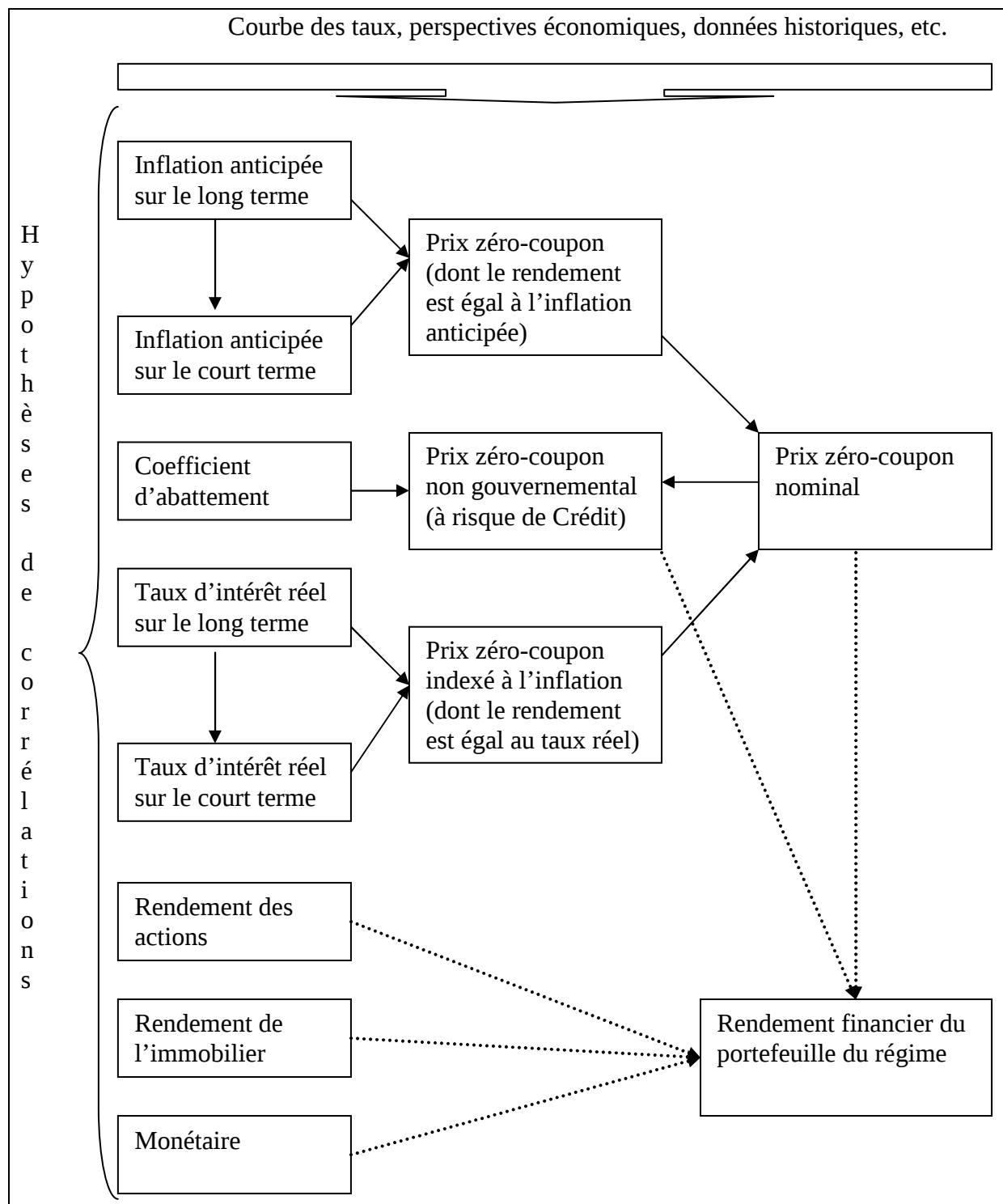


Fig. 14 : Illustration de la structure du GSE construit

II- Calibrage

II-1 Présentation du contexte de l'étude

Dans cette sous-section, nous présentons de façon illustrative certaines approches possibles pour le calibrage et l'estimation des paramètres des différents modèles de diffusion. Pour cela il convient de présenter les caractéristiques des données retenues ainsi que les approches générales retenues pour le calibrage.

III- Projection des scénarios

La présente section est relative à l'étape du *backtesting*. Plus précisément, elle présente et commente les résultats des projections du rendement, de la volatilité et des corrélations réalisées dans le cadre de la validation du modèle et de son paramétrage. Il est rappelé que le calibrage des modèles est utilisé à titre illustratif.

III-1 Éléments sur la mise en œuvre

Au niveau de la structure schématique de projection des scénarios, nous avons opté pour la structure linéaire (cf. sous-section II-1 du chapitre 1 de cette partie). Pour rappel, ci-dessous un exemple de cette structure lors d'une projection de 500 trajectoires de rendement annuel des actions entre 2010 et 2030 :

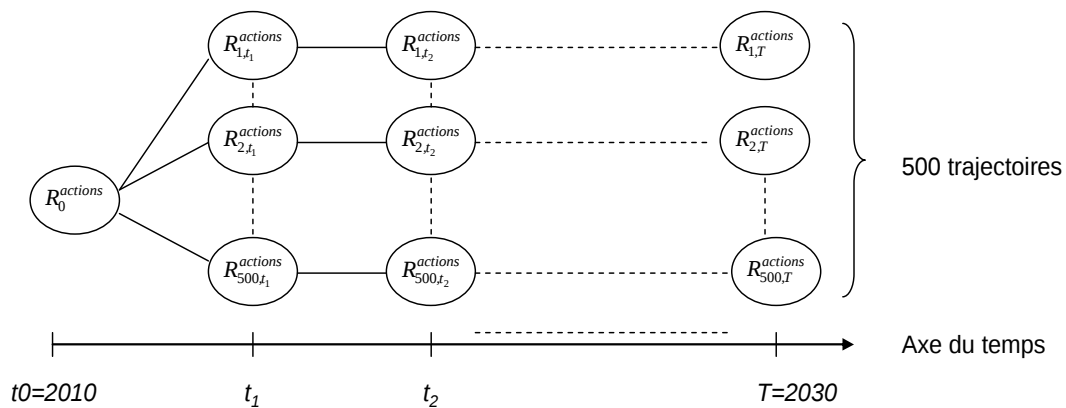


Fig. 15: Illustration de la structure schématique linéaire dans le cas de la projection des rendements des actions

En vue de simuler des browniens ayant une structure de corrélation prédéfinie, nous utilisons la décomposition de Choleski de leur matrice de variance covariance (cf. tableau 6).

En fait, si nous considérons un vecteur aléatoire gaussien $X = (X_1, X_2, \dots, X_n)$ dont la distribution est multi-normale non dégénérée, sa densité est de la forme :

$$f_X(X) = \frac{1}{(2\pi)^{\frac{n}{2}} \times \det(\Sigma)^{\frac{1}{2}}} \times \exp\left[-\frac{1}{2}(x - \mu) \cdot \Sigma^{-1} \cdot (x - \mu)\right]$$

Avec :

$\mu = (\mu_1, \mu_2, \dots, \mu_n)$: Vecteur espérance des lois marginales

Σ : Matrice de variance-covariance

Plus précisément, cette dernière matrice s'écrit dans le cas de $n=4$:

$$\Sigma = \begin{pmatrix} 1 & \rho_{12} & \rho_{13} & \rho_{14} \\ \rho_{12} & 1 & \rho_{23} & \rho_{24} \\ \rho_{13} & \rho_{23} & 1 & \rho_{34} \\ \rho_{14} & \rho_{24} & \rho_{34} & 1 \end{pmatrix}$$

Comme Σ est symétrique et définie positive ($Y'\Sigma Y > 0, \forall Y \neq 0$), il est possible de trouver

une matrice T unique, définie comme suit : $T = \begin{pmatrix} b_{11} & 0 & 0 & 0 \\ b_{21} & b_{22} & 0 & 0 \\ b_{31} & b_{32} & b_{33} & 0 \\ b_{41} & b_{42} & b_{43} & b_{44} \end{pmatrix}$ tel que : $\Sigma = TT'$.

Cette dernière expression correspond à la décomposition de Choleski de Σ . Elle est utilisée ensuite pour la simulation de browniens qui reflètent la même structure de dépendance existante entre les composantes du vecteur X puisque nous savons que si Z est un vecteur indépendant identiquement distribué (i.i.d) de loi commune normale centrée réduite $N(0,1)$, alors le vecteur $X = TZ + \mu$ a pour matrice de variance-covariance Σ et pour moyenne μ .

Nous utilisons cette relation pour retrouver de proche en proche les valeurs des coefficients de la matrice triangulaire T :

- $X_1 = b_{11}.Z_1 + \mu_1$

En considérant la variance de chacun des termes de cette égalité, nous obtenons $b_{11}=1$

- $X_2 = b_{21}.Z_1 + b_{22}.Z_2 + \mu_2$

En considérant la variance de chacun des termes de cette égalité, nous obtenons $b_{11} + b_{22}=1$

Si nous nous intéressons à la covariance, nous tombons sur l'égalité suivante :

$$\text{cov}(X_1, X_2) = \rho_{12} = E[b_{11}.Z_1(b_{12}.Z_1 + b_{22}.Z_2)]$$

Finalement, $b_{21} = \frac{\rho_{12}}{b_{11}} = \rho_{12}$ et $b_{22} = (1 - \rho_{12}^2)^{\frac{1}{2}}$

- Etc.

Le graphique suivant illustre le cas où nous souhaitons simuler des variables négativement corrélées. Cela passe par la génération dans un premier temps de browniens indépendants (composants le vecteur Z) i.i.d et de loi $N(0,1)$.

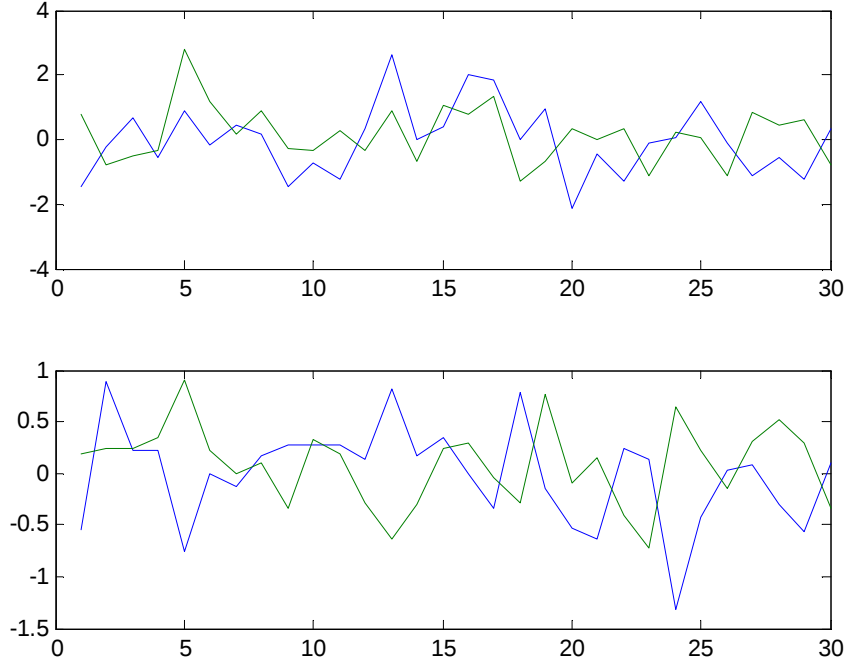


Fig. 16 : Illustration du passage de browniens indépendants à des browniens corrélés négativement

Notons que ce modèle suppose que les taux d'intérêts suivent la loi normale, donc il y a une possibilité d'obtenir des taux réels négatifs. Deux bornes, une borne supérieure et une autre inférieure, peuvent être fixée par l'investisseur pour les taux instantanés et les taux longs.

Les modèles développés (un modèle pour les taux courts, notés r_t , et un modèle pour les taux longs, notés l_t) reprennent l'approche de retour à la moyenne du modèle de Hull et White [1994] à deux facteurs :

$$\begin{aligned} dr_t &= \kappa_r (l_t - r_t) dt + \sigma_r dB_{r,t} \\ dl_t &= \kappa_l (\mu_l - l_t) dt + \sigma_l dB_{l,t} \end{aligned}$$

Nous rappelons que la discrétisation retenue est la discrétisation exacte (cf. Planchet et al. [2005])¹¹ :

$$\begin{aligned} r_{t+1} &= r_t \times e^{(-\kappa_r \Delta t)} + l_t (1 - e^{(-\kappa_r \Delta t)}) + \varepsilon_{r,t} \times \sigma_r \times \sqrt{\frac{1 - e^{(-2\kappa_r \Delta t)}}{2\kappa_r}} \\ l_{t+1} &= l_t \times e^{(-\kappa_l \Delta t)} + \mu_l (1 - e^{(-\kappa_l \Delta t)}) + \varepsilon_{l,t} \times \sigma_l \times \sqrt{\frac{1 - e^{(-2\kappa_l \Delta t)}}{2\kappa_l}} \end{aligned}$$

¹¹ La discrétisation exacte doit être utilisée lorsqu'elle est disponible, cette discrétisation étant la seule à ne pas introduire de biais sur la loi du facteur discrétisé.

où :

κ_r (resp. κ_l) est la vitesse de retour à la moyenne ;

μ_l est le taux d'intérêt à long terme moyen.

D'un autre côté et compte tenu de l'absence de données fiables pour la partie long terme, il a été retenu pour la moyenne de long terme l'objectif BCE de 2 % et pour les deux autres paramètres des niveaux permettant d'avoir des résultats cohérents, sans aucun calibrage statistique.

Mise en œuvre de la démarche crédit

La mise en œuvre de la démarche repose sur la spécification du prix des obligations sans risque de défaut et sur le prix observé sur le marché des obligations présentant un risque de défaut.

Les prix $O(T)$ des obligations sans risque de défaut de nominal N , de maturité T et de taux facial γ se déduisent via la formule suivante (cf. Planchet et al. [2005]) :

$$O(T) = N \times \left(\gamma \sum_{i=1}^T P_i(0, i) + P_i(0, T) \right)$$

où $P_i(t, T)$ représente le prix à la date t des zéro-coupon de maturité T . La démarche présentée consiste à déterminer l'abattement α tel que :

$$S(T) = \alpha \times O(T)$$

où $S(T)$ représente le prix de marché observé pour les obligations de maturité T avec risque de défaut. En pratique, nous noterons que α dépend de la maturité et du rating des obligations présentant un risque de défaut.

Ci-après des paramètres relatifs à l'environnement du calcul. Notons qu'avec 2 000 simulations nous obtiendrons des résultats qui sont égaux aux résultats observés avec un nombre de simulations supérieur.

Horizon de projection (en années)	50
Nombre de pas	200
Nombre de simulations	2 000

Tab. 17 : Paramètres utilisés dans le GSE construit

III-2 Tests du modèle : paramètres et matrice de corrélation

Cette sous-section est relative aux tests du modèle, et en particulier aux tests sur les paramètres et la matrice de corrélation. L'analyse et l'interprétation de ces tests devra être traitée avec prudence (existence de biais dus à l'absence de linéarité dans les transformations effectuées, prise en compte de la périodicité des projections, etc.).

Face aux limites de l'estimation des paramètres des modèles de GSE, en particulier à cause de la non évidence de l'hypothèse de la reproduction de l'historique dans le futur, nous supposons que le seul garant de la fiabilité du GSE considéré par la suite sera le degré de correspondance des résultats obtenus (projections) par rapport à ceux prévus lors du calibrage (avis des experts, paramètres historiques, anticipations à partir des valeurs actuelles de marché des différents produits et variables financières, etc.). Le respect de cette correspondance sera considéré comme une condition minimale à vérifier.

III-2-1 Tests sur les paramètres et sur les développements informatiques

Les tests sur les paramètres peuvent être organisés en deux parties.

› Comparaison de valeurs théoriques et de valeurs empiriques

Le premier point est relatif à la comparaison de valeurs théoriques et de valeurs empiriques estimées sur les trajectoires simulées. À cet effet, il convient de distinguer différentes classes d'actifs.

Cas des actions et de l'immobilier

Pour les actions et l'immobilier, nous pouvons comparer les valeurs théoriques et empiriques des rendements ou des volatilités. Soit $x_{i,t}, x_{i,t+1}, \dots, x_{i,t+h}$ la trajectoire des rendements empiriques de la simulation i jusqu'à l'horizon $t+h$:

- le rendement moyen empirique est égal à $\hat{\mu} = \frac{1}{I} \times \frac{1}{H+1} \sum_{i=1}^I \sum_{h=0}^H x_{i,t+h}$ (où I représente le nombre total de simulations et H l'horizon de la projection) ;

- la volatilité moyenne empirique est calculée par $\hat{\sigma} = \sqrt{\frac{1}{I+H+1} \times \sum_{i,h}^{I,H} (x_{i,t+h} - \hat{\mu})^2}$.

Cas des produits monétaires

Dans le cas des produits monétaires, la comparaison présentée ci-dessus peut également être appliquée.

Cas des produits obligataires

Concernant les produits obligataires, la comparaison en lecture directe est plus difficile car dans les sorties du modèle nous disposons du prix des obligations, desquels sont déduits les rendements et les volatilités associés, alors que les paramètres retenus pour ces classes d'actifs portent sur la moyenne des taux d'intérêt réel et des taux d'inflation.

La comparaison des valeurs théoriques et empiriques pourrait toutefois être effectuée à partir des formules de prix des obligations zéro-coupon, sans tenir compte du risque de défaut. En effet, par exemple dans les modèles de type Vasicek à un-facteur les prix à la date t des zéro-coupon de maturité T , notés $P_t(t, T)$, sont déterminés à partir des taux r (κ représentant la vitesse de retour à la moyenne et μ représentant le taux moyen) par la relation suivante (cf. Hull [2000]) :

$$P(t, T) = A(t, T) e^{(-B(t, T) \times r)}$$

Avec :

$$B(t, T) = \frac{1 - e^{(-\kappa(T-t))}}{\kappa}$$

$$A(t, T) = \exp\left(\frac{(B(t, T) - T + t)(\kappa^2 \times \mu - \sigma^2 / 2)}{\kappa^2} - \frac{\sigma^2 B(t, T)^2}{4\kappa}\right)^{12}$$

Il convient ainsi de comparer le prix théorique $P(t, T)$ obtenu avec le paramètre de taux moyen au prix empirique $\hat{P}_l(t, T)$ obtenu avec la moyenne (sur l'ensemble des simulations et l'ensemble de la durée de projection) des taux projetés.

› Tests aux limites (cas du modèle des actions)

Le deuxième point est relatif aux tests aux limites. À cet effet, il s'agit par exemple de retenir un seul régime pour le RSLN2 en mettant une probabilité de transition à 0 ou de considérer des versions déterministes en mettant les volatilités à 0. Nous pouvons également vérifier pour le modèle de taux que nous retrouvons le modèle à un facteur en mettant les paramètres du taux long à 0. Dans le même esprit, nous pouvons vérifier que les sens de variation des grandeurs modélisées sont conformes à la théorie (observer ce qui se passe lorsque la volatilité augmente sur une option, une obligation, etc.).

III-2-2 Tests sur les matrices de corrélation

Une fois les tests sur les paramètres et les développements informatiques réalisés, il convient de décrire un test de cohérence du paramétrage consistant à estimer empiriquement la matrice de corrélation des rendements et à la comparer à celle que nous observons sur le marché (si elle est disponible). En fait la matrice de corrélation théorique (*input* du GSE) et par la suite la matrice de corrélation empirique (*output* du GSE) peut différer de celle observée réellement sur le marché (avec des données historiques) en fonction des prévisions de l'investisseur. Cela ne remet pas en cause l'intérêt de la comparaison entre ces deux dernières matrices par exemple pour effectuer des tests de sensibilité.

La matrice que nous observons sur le marché prend également la forme suivante :

$$\begin{pmatrix} \rho_{1,1} & \rho_{1,2} & \dots & \dots & \rho_{1,K} \\ \rho_{2,1} & \rho_{2,2} & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots \\ \rho_{K,1} & \dots & \dots & \dots & \rho_{K,K} \end{pmatrix}$$

où $\rho_{i,j}$ représente la corrélation entre les classes d'actifs numéros k_1 et k_2 , et est déterminée par :

¹² *exp* désigne ici la fonction exponentielle.

$$\rho_{k_1, k_2} = \frac{\text{cov}(k_1, k_2)}{\sigma(k_1) \times \sigma(k_2)} = \frac{\sum_{i,h} \left((x_{k_1})_{i,h} - \overline{(x_{k_1})} \right) \left((x_{k_2})_{i,h} - \overline{(x_{k_2})} \right)}{\sqrt{\sum_{i,h} \left((x_{k_1})_{i,h} - \overline{(x_{k_1})} \right)^2} \cdot \sqrt{\sum_{i,h} \left((x_{k_2})_{i,h} - \overline{(x_{k_2})} \right)^2}}$$

avec :

(x_{k_1}) représentant le rendement des actifs de la classe k_1 ;

(x_{k_2}) représentant le rendement des actifs de la classe k_2 ;

i représentant le nombre de simulations (I simulations au total) ;

h représentant l'année de projection (H années de projection au total).

III-3 Résultats de projection

Nous présentons ci-après des graphiques mettant en évidence quelques propriétés théoriques du modèle de GSE objet de l'étude (exemple : retour à la moyenne pour les taux, changement de régime pour les actions, etc.) ainsi que certains résultats sur les niveaux des rendements.

Rendements moyens et volatilités sur toute la période de projection

Le tableau suivant illustre la volatilité et le rendement moyen obtenus sur toute la période de projection. En particulier, nous observons que la volatilité est cohérente avec le rendement.

Pour le monétaire, les actions et l'immobilier nous observons que les rendements moyens et les volatilités sont du même ordre de grandeur que les paramètres du calibrage ce qui est donc cohérent avec les approches retenues pour modéliser ces actifs.

Rendements moyens et volatilités	Rendement	Volatilité
Monétaire	3,51 %	1,52 %
ZC 5 ans (nominal)	4,25 %	4,42 %
ZC 10 ans (nominal)	4,32 %	6,67 %
ZC 15 ans (nominal)	4,34 %	7,46 %
ZC 8 ans (réel)	3,59 %	3,96 %
ZC 15 ans (réel)	3,61 %	4,49 %
Crédit 5 ans	4,27 %	4,42 %
Actions	7,22 %	16,53 %
Immobilier	5,91 %	8,81 %

Tab. 18 : *Exemple des projections de rendement obtenues par le GSE construit*

La prime de risque entre les obligations nominales et les obligations indexées est égale à près de 75 points de base pour les ZC 15 ans et apparaît *a priori* importante. Nous noterons toutefois qu'entre 2004 et 2009 certains fonds présentent des primes de risques pour ces

classes d'actifs de près de 100 points de base¹³. Sauf précision contraire, tous les chiffres présentés ci-après sont établis à partir des résultats ci-dessus obtenus dans le cadre du scénario central.

Courbe des taux des obligations indexées sur l'inflation d'une maturité de 15 ans

La courbe des taux présentée ci-après en bleu (cf. graphique 17) correspond à la courbe des taux déduite des valeurs des taux de rendement annualisés et projetés de l'obligation indexées d'une maturité de 15 ans avec un pas trimestriel (ce qui correspond à un nombre de 60 pas au total). En rouge une approximation par une fonction linéaire par morceau est tracée. Cette fonction correspond aux quatre droites dont les pentes sont calculées respectivement sur les quatre périodes 0-2 ans, 2-5 ans, 5-10 ans, 10-15 ans. Il est précisé que l'unité de temps considéré dans le graphique est trimestrielle. Ces pentes ont les valeurs suivantes :

Périodes	Période 0-2	Période 2-5	Période 5-10	Période 10-15
Unité de temps trimestrielle	0,15 %	0,04 %	0,01 %	0,0046 %
Unité de temps annuelle	0,59 %	0,15 %	0,06 %	0,02 %

Tab. 19 : *Pentes de la courbe des taux obtenues via le GSE construit*

Rappelons que la pente sur une période s'étalant entre t et $t+h$ est la différence entre le taux à l'instant $t+h$ et le taux à l'instant t le tout divisé par h .

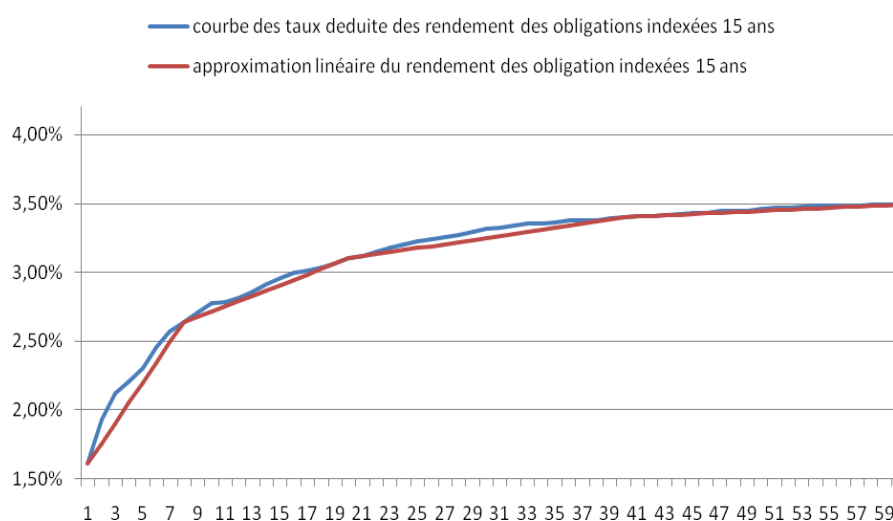


Fig.17 : *Courbe des taux des obligations indexées sur l'inflation d'une maturité de 15 ans*

¹³ cf. par exemple les fonds suivants :

<http://www.sgam.fr/portal/site/sgamfr/menuitem.8d6c6203bb02b5bf1398c6e5100000f7/?vgnextoid=5ce99634dd134010VgnVCM1000000100007fRCRD&countryiso=FR&spaceid=1&lang=fr&isin=LU0123756628>

<http://www.sgam.fr/portal/site/sgamfr/menuitem.8d6c6203bb02b5bf1398c6e5100000f7/?vgnextoid=5ce99634dd134010VgnVCM1000000100007fRCRD&countryiso=FR&spaceid=1&lang=fr&isin=LU0197574634>.

Une autre façon de construire la courbe des taux sera de partir des taux courts simulés selon les modèles de taux proposés (Hull et White à 2 facteurs), pour ensuite déterminer les taux sur différents horizons (ou maturités) en utilisant la formule présentée dans le paragraphe I-2-2-1 du chapitre 2, à savoir : $R(t,T) = -\log\{P(t,T)\}/(T-t)$.

Dans ce cadre, nous avons considéré le cas où le taux réel long observé initialement $l_r(0)$ est significativement plus élevé que le niveau moyen de long terme μ_{lr} . Cela donne lieu à une courbe des taux réels avec une inflexion. La STTI des taux nominaux aura également cette même forme puisque elle est déduite à la fois à partir des taux réels et des taux d'inflation (anticipés). Nous confirmons ainsi l'avantage, précédemment cité, du choix d'un modèle de taux à deux facteurs : ce dernier permet d'obtenir différentes formes de la courbe des taux en conformité avec ce qui est observé réellement (cf. Martellini et al. [2003]).

Le phénomène de retour à la moyenne, qui par hypothèse concerne les variables de taux du GSE construit, peut être mis en évidence au niveau des courbes de taux elles-mêmes. Le graphique 19 illustre ce cas de figure. En construisant les STTI à six dates différentes, à savoir (la date initiale et cinq dates futures : fin de la deuxième année, fin de la cinquième année, fin de la dixième année, fin de la quinzième année et fin de la vingtième année), nous constatons à partir d'une certaine date une certaine stabilité de la courbe. La maturité maximale représentée par la STTI est ici 30 ans.

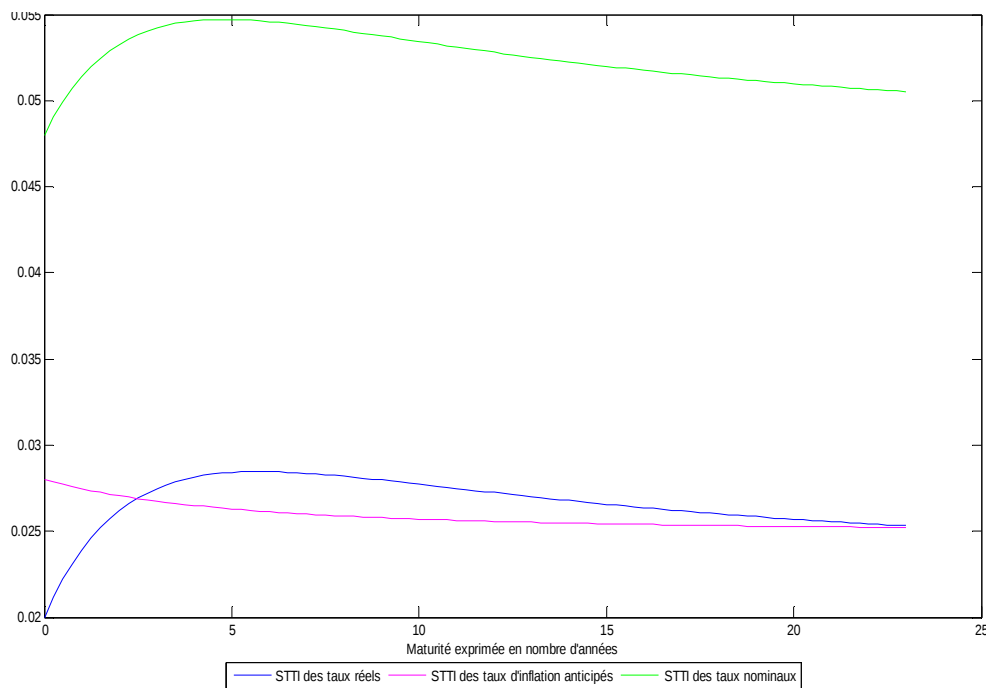


Fig. 18 : Structure par terme des taux nominaux, des taux réels et des taux d'inflation en cas de taux réels longs élevés par rapport à ceux de l'équilibre : Mise en évidence de la courbure de la STTI

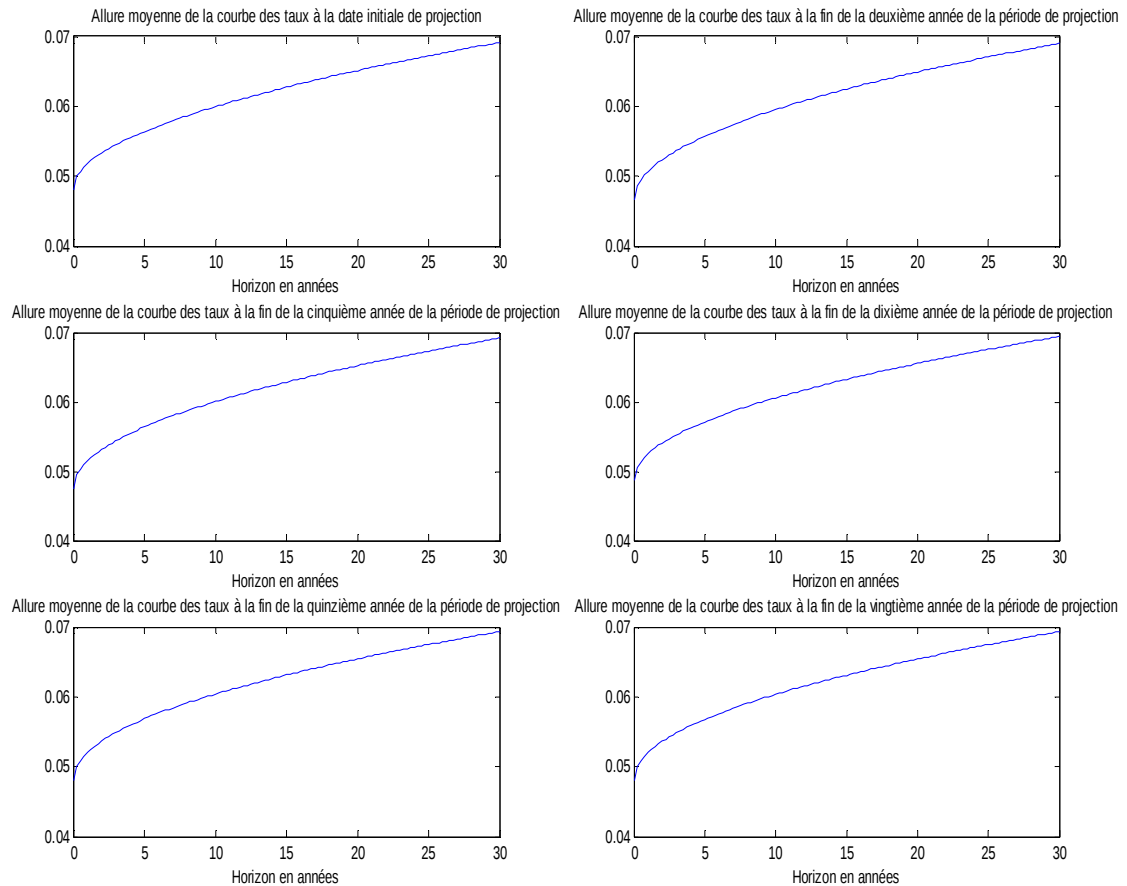


Fig.19 : Evolution dans le temps de la STTI : Mise en évidence du phénomène de retour à la moyenne

Ce même phénomène est confirmé pour ce qui concerne les écarts absolus moyens, cette fois-ci, entre les taux de deux maturités différentes (cf. graphique 20). Nous appellerons « pentes moyennes de taux » les écarts absolus moyens entre deux taux de maturités différentes évalués à la même date, ils sont calculés respectivement sur les quatre périodes 0-2 ans, 2-5 ans, 5-10 ans, 10-15 ans. Ces écarts, exprimés en point de base, se stabilisent après cinq années autour de leurs niveaux moyens. De même, nous constatons une décroissance de cet écart en augmentant les maturités considérées (passage d'un écart moyen de 57 points de base (pb) entre les taux de maturité 2 ans et ceux instantanés à 28 pb entre les taux de maturité 15 ans et 10 ans). Cela est cohérent avec les hypothèses considérées d'une faible volatilité sur les niveaux des taux longs.

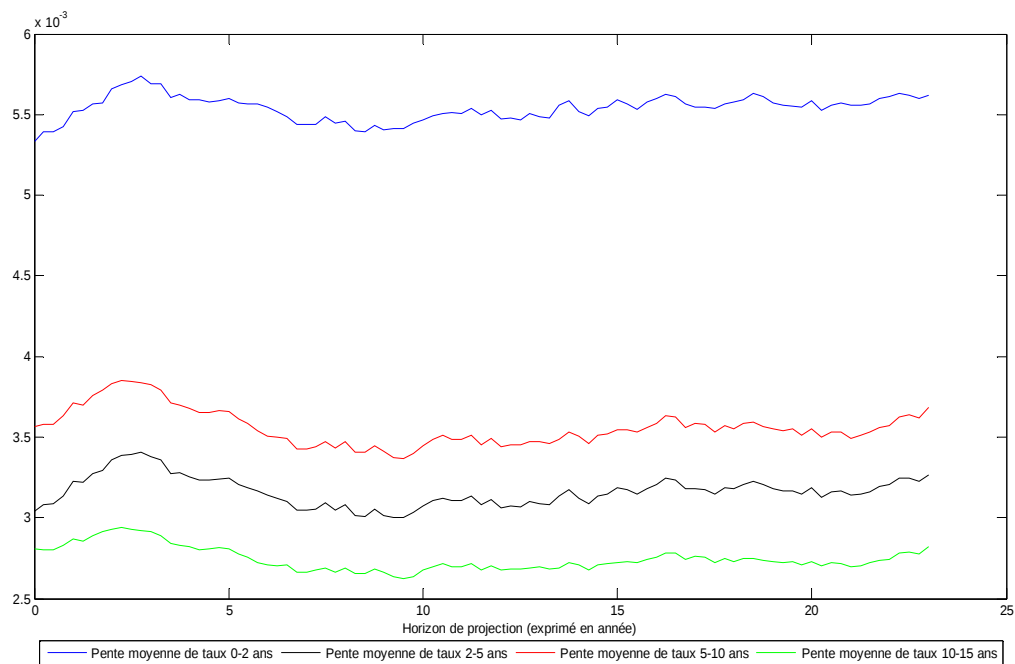


Fig. 20 : Evolution des pentes moyennes de taux (écarts absolus moyens entre les taux de deux maturités différentes) tout au long de la période de projection

	Pente 0-2 ans	Pente 2-5 ans	Pente 5-10 ans	Pente 10-15 ans
Moyenne projetée*	57	33	37	28

* Moyenne exprimée en point de base (et sur toute la période de projection)

Tab. 20 : Moyennes des différentes pentes de taux projetées

Les graphiques ci-dessous représentent, sous forme d'histogrammes, les distributions obtenues de différentes variables de taux et/ou de rendement. Les histogrammes permettent d'avoir une vision globale sur la fréquence des niveaux de taux et/ou de rendement simulés. Certes, il existe plusieurs tests statistiques pour juger de la conformité des valeurs projetées avec les hypothèses de départ retenues dans le modèle (normalité, etc.). Cependant l'étude des modèles pris individuellement ne constitue pas notre principal centre d'intérêt et c'est plutôt la structuration de ces différents modèles entre eux, dans l'objectif de la construction d'un GSE, qui suscite le plus notre attention.

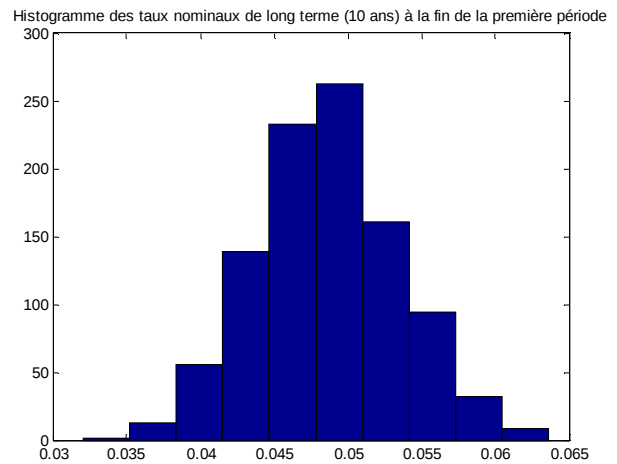
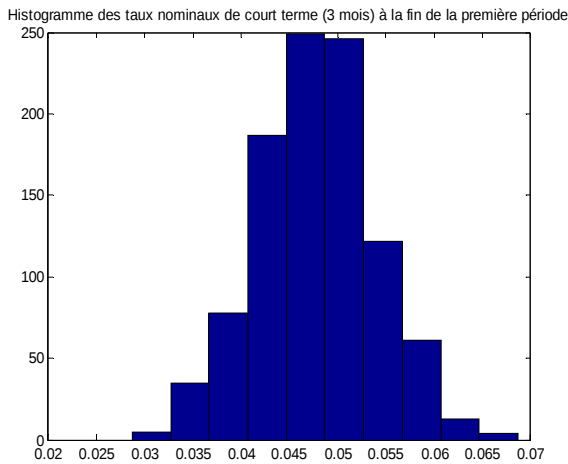


Fig. 21 : *Histogramme des taux nominaux de court terme (3 mois) à la fin de la première période*

Fig. 22 : *Histogramme des taux nominaux de long terme (10 ans) à la fin de la première période*

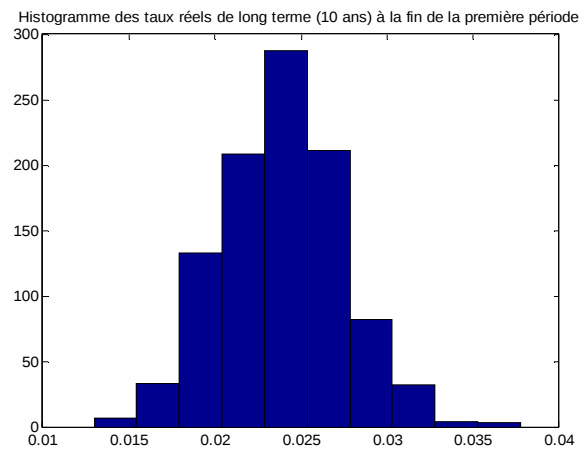
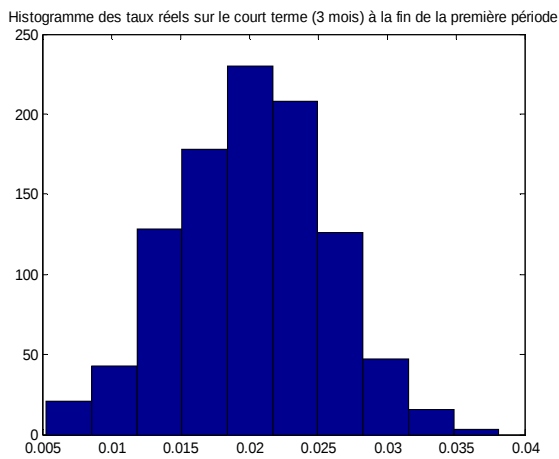


Fig. 23 : *Histogramme des taux réels sur le court terme (3 mois) à la fin de la première période*

Fig. 24 : *Histogramme des taux réels de long terme (10 ans) à la fin de la première période*

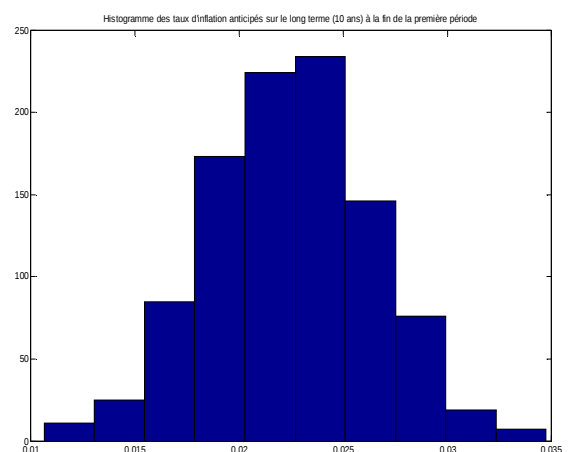
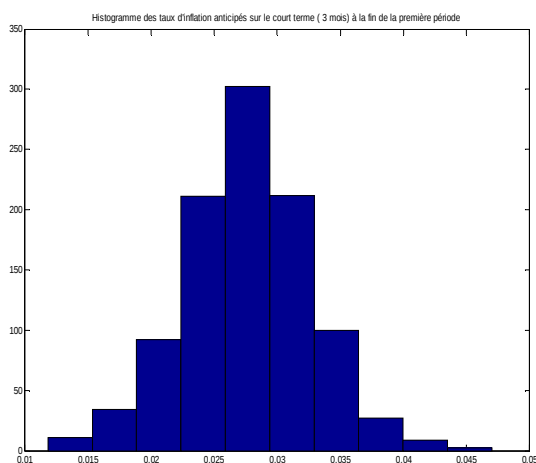


Fig. 25 : *Histogramme des taux d'inflation anticipés de court terme (3 mois) à la fin de la première période*

Fig. 26 : *Histogramme des taux d'inflation anticipés de long terme (10 ans) à la fin de la première période*

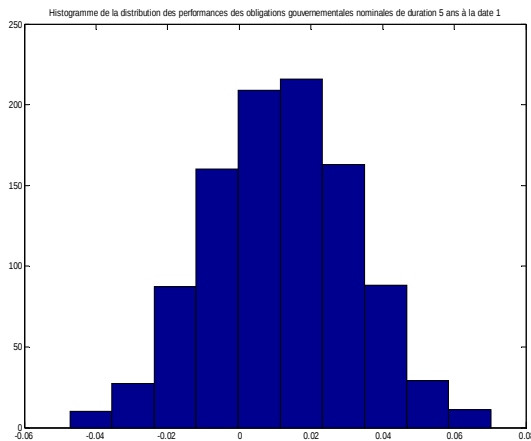


Fig. 27 : *Histogramme de la distribution des obligations nominales gouvernementales de duration 5 ans sur la première période*

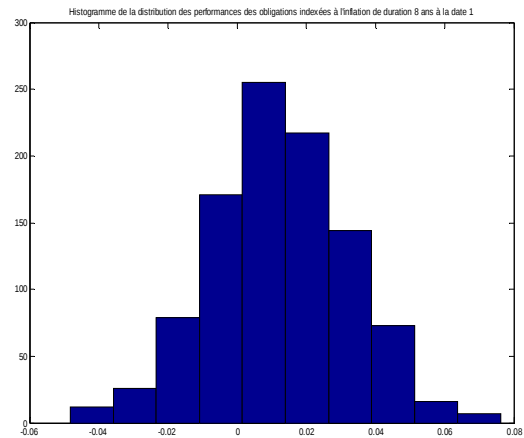


Fig. 28 : *Histogramme de la distribution des obligations indexées inflation de duration 8 ans sur la première période*

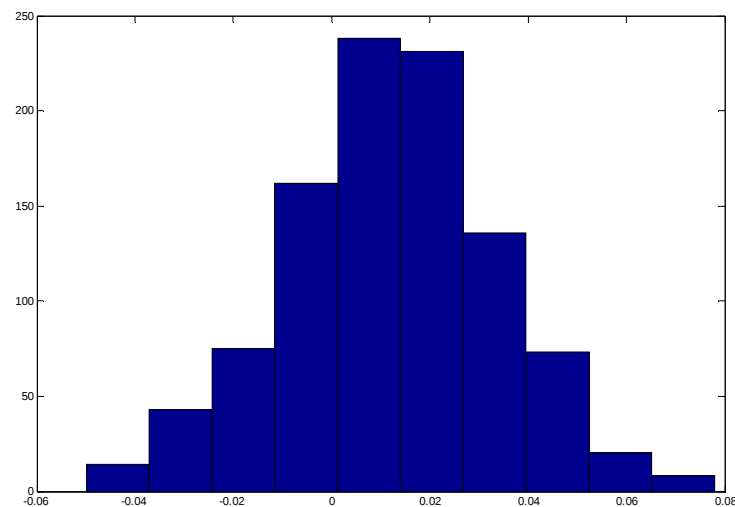


Fig. 29 : *Histogramme de la distribution des rendements des obligations nominales non gouvernementales de duration 5 ans sur la première période*

Matrice de corrélation des browniens sur la période de projection

Il est important de préciser que les valeurs retenues dans la matrice de corrélation des browniens, utilisée en *input* du GSE, se basent sur des ajustements : certains niveaux de corrélations sont en effet attribués de façon subjective sans nécessairement s'aligner sur les résultats historiques. Cela nous a conduits à proposer la matrice suivante à titre illustratif :

	Action 1	Inflation longue	Inflation	Taux long réel	Taux court réel	Monétaire	Immobilier	Action 2
Action 1	100,00 %	-20,00 %	0,00 %	-10,00 %	0,00 %	1,00 %	75,00 %	0,00 %
Inflation longue	-20,00 %	100,00 %	10,00 %	-20,00 %	0,00 %	0,00 %	-20,00 %	40,00 %
Inflation	0,00 %	10,00 %	100,00 %	0,00 %	-30,00 %	20,00 %	30,00 %	-40,00 %
Taux long réel	-10,00 %	-20,00 %	0,00 %	100,00 %	10,00 %	0,00 %	-20,00 %	30,00 %
Taux court réel	0,00 %	0,00 %	-30,00 %	10,00 %	100,00 %	-1,00 %	-40,00 %	40,00 %
Monétaire	1,00 %	0,00 %	20,00 %	0,00 %	-1,00 %	100,00 %	6,00 %	0,00 %
Immobilier	75,00 %	-20,00 %	30,00 %	-20,00 %	-40,00 %	6,00 %	100,00 %	-37,00 %
Action 2	0,00 %	40,00 %	-40,00 %	30,00 %	40,00 %	0,00 %	-37,00 %	100,00 %

Tab. 21 : *Corrélations théoriques entre les différentes variables (input du GSE)*

La matrice ci-après illustre les corrélations entre les browniens projetés (obtenues après projection de 1000 scénarios). Nous observons que les valeurs affichées sont très proches des valeurs présentées dans la matrice de corrélation des browniens en *input*. Ces valeurs correspondent aux valeurs attendues et sont donc cohérentes.

Corrélations entre les browniens (modèle BS 2 régimes)	Action 1	Inflation longue	Inflation	Taux long réel	Taux court réel	Monétaire	Immobilier	Action 2
Action 1	100,00 %	-19,58 %	-0,39 %	-10,06 %	-0,21 %	1,03 %	74,94 %	0,00 %
Inflation longue	-19,58 %	100,00 %	9,97 %	-19,72 %	0,17 %	0,04 %	-19,81 %	40,10 %
Inflation	-0,39 %	9,97 %	100,00 %	0,37 %	-29,77 %	19,94 %	29,64 %	-39,98 %
Taux long réel	-10,06 %	-19,72 %	0,37 %	100,00 %	10,02 %	0,22 %	-19,94 %	29,84 %
Taux court réel	-0,21 %	0,17 %	-29,77 %	10,02 %	100,00 %	-0,94 %	-40,07 %	39,77 %
Monétaire	1,03 %	0,04 %	19,94 %	0,22 %	-0,94 %	100,00 %	5,95 %	0,09 %
Immobilier	74,94 %	-19,81 %	29,64 %	-19,94 %	-40,07 %	5,95 %	100,00 %	-36,95 %
Action 2	0,00 %	40,10 %	-39,98 %	29,84 %	39,77 %	0,09 %	-36,95 %	100,00 %

Tab. 22 : *Corrélations obtenues lors des projections entre les différentes variables*

Matrice de corrélation des rendements sur la période de projection

Nous nous intéressons ici à la dépendance entre eux des rendements des indices projetés. Cette dépendance est reflétée par les corrélations entre les browniens qui interviennent dans la dynamique des variables modélisées. Les valeurs cibles sont en pratique fournies sur la base d'avis d'experts, d'une étude historique ou encore déduits de manière implicite de l'observation du marché. Il s'agit en fait d'une matrice de corrélations jugée « raisonnable » par la place.

Globalement les valeurs obtenues dans le cadre du présent *backtesting* sont cohérentes avec les valeurs théoriques utilisées initialement en *input*. Nous notons que seules les corrélations entre l'immobilier et les obligations présentent des différences entre les valeurs cibles et celles obtenues par *backtesting*.

Plus particulièrement nous remarquons que :

- Les corrélations entre les zéro-coupons (ZC) nominaux et les ZC réels sont de l'ordre de 45 % à 55 % et donc sont cohérentes avec les valeurs cibles (pour mémoire ces valeurs sont de l'ordre de 60 %).
- La corrélation entre les actions et l'immobilier est de plus de 20 %, ce qui est cohérent avec la valeur cible (26 %).

- Les corrélations entre les ZC réels sont très fortes (de l'ordre de 99 %). Ces fortes corrélations sont d'un côté cohérentes avec l'approche du calcul des prix de ces obligations et représentent, d'un autre côté des niveaux proches de ceux de la cible. Elles traduisent le caractère très contraint de la courbe des taux, déterminée dans le cadre d'un modèle à deux facteurs.

Corrélations entre les indices projetés	Monétaire	ZC 5 ans (nominal)	ZC 10 ans (nominal)	ZC 15 ans (nominal)	ZC 8 ans (réel)	ZC 15 ans (réel)	Obligation crédit	Actions	Immobilier
Monétaire	100,00 %	-7,92 %	-8,79 %	-9,80 %	0,46 %	0,40 %	-7,92 %	1,06 %	1,45 %
ZC 5 ans (nominal)	-7,92 %	100,00 %	98,56 %	97,32 %	56,49 %	55,44 %	99,99 %	14,38 %	1,23 %
ZC 10 ans (nominal)	-8,79 %	98,56 %	100,00 %	99,73 %	46,91 %	46,80 %	98,56 %	15,85 %	1,45 %
ZC 15 ans (nominal)	-9,80 %	97,32 %	99,73 %	100,00 %	42,78 %	42,94 %	97,32 %	15,80 %	1,10 %
ZC 8 ans (réel)	0,46 %	56,49 %	46,91 %	42,78 %	100,00 %	99,68 %	56,49 %	5,28 %	5,66 %
ZC 15 ans (réel)	0,40 %	55,44 %	46,80 %	42,94 %	99,68 %	100,00 %	55,44 %	5,83 %	5,77 %
Obligation crédit	-7,92 %	99,99 %	98,56 %	97,32 %	56,49 %	55,44 %	100,00 %	14,38 %	1,23 %
Actions	1,06 %	14,38 %	15,85 %	15,80 %	5,28 %	5,83 %	14,38 %	100,00 %	20,72 %
Immobilier	1,45 %	1,23 %	1,45 %	1,10 %	5,66 %	5,77 %	1,23 %	20,72 %	100,00 %

Tab. 23 : Corrélations obtenues lors des projections entre les différents rendements

Corrélation des actifs dans un régime à faible volatilité

La matrice suivante présente la matrice de corrélation des actifs lorsque les actions évoluent dans un régime à faible volatilité. Nous remarquons que dans ce régime l'immobilier par exemple est positivement corrélé avec les actions.

Corrélation régime 1 (faible volatilité)	Monétaire	ZC 5 ans (nominal)	ZC 10 ans (nominal)	ZC 15 ans (nominal)	ZC 8 ans (réel)	ZC 15 ans (réel)	Obligation crédit	Actions	immobilier
Monétaire	100,00 %	-7,92 %	-8,77 %	-9,78 %	0,72 %	0,71 %	-7,92 %	1,03 %	1,62 %
ZC 5 ans (nominal)	-7,92 %	100,00 %	98,56 %	97,30 %	56,35 %	55,29 %	100,00 %	15,15 %	1,23 %
ZC 10 ans (nominal)	-8,77 %	98,56 %	100,00 %	99,73 %	46,70 %	46,57 %	98,56 %	16,63 %	1,54 %
ZC 15 ans (nominal)	-9,78 %	97,30 %	99,73 %	100,00 %	42,53 %	42,68 %	97,30 %	16,50 %	1,23 %
ZC 8 ans (réel)	0,72 %	56,35 %	46,70 %	42,53 %	100,00 %	99,68 %	56,35 %	5,92 %	5,37 %
ZC 15 ans (réel)	0,71 %	55,29 %	46,57 %	42,68 %	99,68 %	100,00 %	55,29 %	6,49 %	5,49 %
Obligation crédit	-7,92 %	100,00 %	98,56 %	97,30 %	56,35 %	55,29 %	100,00 %	15,15 %	1,23 %
Actions	1,03 %	15,15 %	16,63 %	16,50 %	5,92 %	6,49 %	15,15 %	100,00 %	21,88 %
Immobilier	1,62 %	1,23 %	1,54 %	1,23 %	5,37 %	5,49 %	1,23 %	21,88 %	100,00 %

Tab. 24 : Corrélations obtenues lors des projections entre les différents rendements dans un régime à faible volatilité

Corrélation des actifs dans un régime à forte volatilité

La matrice suivante présente la matrice de corrélation des actifs lorsque les actions évoluent dans un régime à forte volatilité. Nous remarquons que dans ce régime l'immobilier est négativement corrélé avec les actions. Cela peut s'expliquer par l'hypothèse de départ considérée : l'immobilier est considéré comme une valeur refuge en cas de crise.

Corrélation en régime 2 (forte volatilité)	Monétaire	ZC 5 ans (nominal)	ZC 10 ans (nominal)	ZC 15 ans (nominal)	ZC 8 ans (réel)	ZC 15 ans (réel)	Obligation crédit	Actions	immobilier
Monétaire	100,00 %	-7,66 %	-8,55 %	-9,57 %	0,84 %	0,78 %	-7,66 %	-0,21 %	1,82 %
ZC 5 ans (nominal)	-7,66 %	100,00 %	98,56 %	97,31 %	56,54 %	55,46 %	100,00 %	-38,98 %	1,01 %
ZC 10 ans (nominal)	-8,55 %	98,56 %	100,00 %	99,73 %	46,92 %	46,79 %	98,56 %	-35,82 %	1,31 %
ZC 15 ans (nominal)	-9,57 %	97,31 %	99,73 %	100,00 %	42,77 %	42,91 %	97,31 %	-32,13 %	0,97 %
ZC 8 ans (réel)	0,84 %	56,54 %	46,92 %	42,77 %	100,00 %	99,68 %	56,54 %	-44,60 %	5,21 %
ZC 15 ans (réel)	0,78 %	55,46 %	46,79 %	42,91 %	99,68 %	100,00 %	55,46 %	-44,22 %	5,32 %
Obligation crédit	-7,66 %	100,00 %	98,56 %	97,31 %	56,54 %	55,46 %	100,00 %	-38,98 %	1,01 %
Actions	-0,21 %	-38,98 %	-35,82 %	-32,13 %	-44,60 %	-44,22 %	-38,98 %	100,00 %	-10,57 %
Immobilier	1,82 %	1,01 %	1,31 %	0,97 %	5,21 %	5,32 %	1,01 %	-10,57 %	100,00 %

Tab. 25 : *Corrélations obtenues lors des projections entre les différents rendements dans un régime à forte volatilité*

Les graphiques suivants présentent l'évolution des rendements des actions entre les deux régimes sur une trajectoire simulée (choisie à titre illustratif). Le nombre de fois où ces rendements se trouvent sous le régime 2 (régime de « crise ») est nettement plus faible que celui sous le régime 1 (régime normal). Cela reflète une probabilité plus faible de se retrouver dans le régime de crise, comme supposé dans les hypothèses de départ dans le modèle des actions.

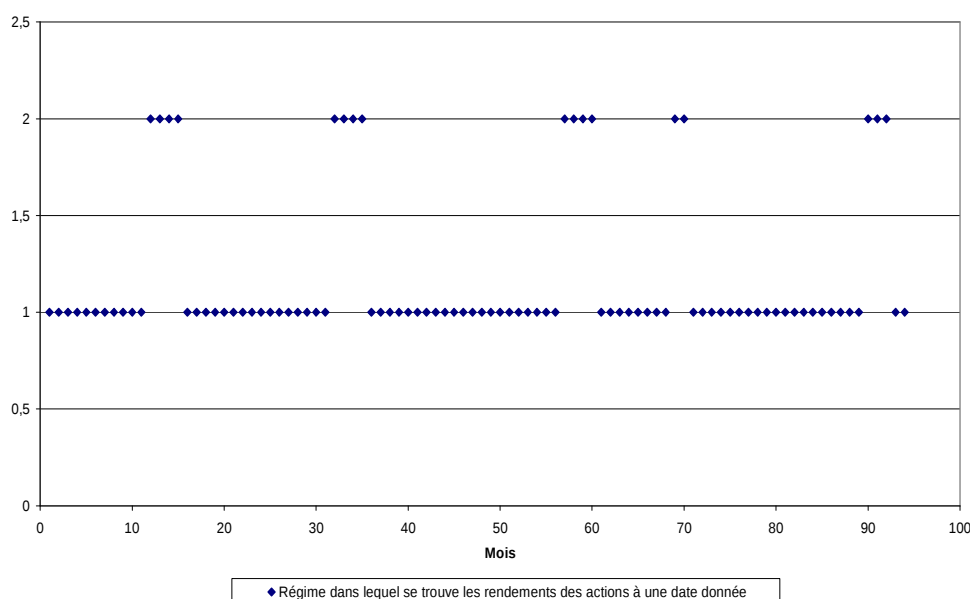


Fig. 30 : *Evolution du marché des actions entre le régime 1 (actions à faible volatilité) et le régime 2 (actions à forte volatilité) sur une trajectoire simulée*

Comparaison de l'approche du coefficient d'abattement avec l'approche « mélange » pour le crédit

Nous avons montré que les approximations retenues pour la méthode « mélange » sont délicates à justifier. Au-delà des descriptions fonctionnelles de cette approche, il s'agit ici de présenter quelques résultats obtenus en l'utilisant.

L'approche du coefficient d'abattement consiste à calculer le prix de l'obligation crédit en multipliant le prix des obligations nominales par un coefficient d'abattement dépendant de la maturité.

Aussi, dans un premier temps le tableau ci-après compare les rendements et les volatilités des obligations du crédit correspondant aux deux approches : l'approche du coefficient d'abattement et l'approche « mélange ». Notons que la discrétisation des formules, utilisée lors du calcul, est exacte dans les deux approches. Le calibrage retenu pour cette comparaison est le même que celui présenté dans la sous-section II-6 du chapitre 2 à l'exception de la matrice de corrélation des browniens qui intègre en plus une ligne et une colonne correspondant aux corrélations avec les *spreads* de crédit.

Obligation crédit 5 ans	Coefficient d'abattement	approche « mélange »
Rendements	4,27 %	3,98 %
Volatilités	4,42 %	4,27 %

Tab. 26 : *Discrétisation exacte pour Coefficient d'abattement et l'approche « mélange »*

Nous observons que les rendements et les volatilités sont du même ordre de grandeur. Cela signifie que les approches, bien que différentes, conduisent à des résultats proches. Ainsi, l'approche du coefficient d'abattement donne des résultats comparables à ceux de l'approche « mélange » en termes de rendements moyens et de volatilités.

Effet de la discrétisation

À titre illustratif, le tableau suivant présente les rendements et les volatilités des deux approches mais en retenant la discrétisation d'Euler pour toutes les variables modélisées dans le cas de l'approche « mélange » (nous rappelons que toutes les variables sont implémentées en tenant compte d'une discrétisation exacte dans le cas de l'autre approche des coefficients d'abattement). En termes de résultat, nous constatons que les rendements moyens dans les deux approches sont proches. Cependant les volatilités sont significativement différentes. Donc le choix de la discrétisation a un impact significatif sur le niveau de la volatilité des obligations crédit lorsqu'elles sont modélisées avec ces deux approches différentes.

Obligation crédit 5 ans	Discrétisation exacte	Discrétisation d'Euler
Rendements	4,27 %	3,96 %
Volatilités	4,42 %	7,43 %

Tab. 27 : *Discrétisation exacte pour Coefficient d'abattement et d'Euler pour l'approche « mélange »*

Conclusion du chapitre 2

L'objet du chapitre 2 de cette partie était de présenter un modèle stochastique de génération des scénarios économiques dont la vocation finale est plus proche du domaine de la gestion des risques que de celui du *pricing*. Nous avons également essayé d'étudier différents aspects liés au processus de calibrage du GSE construit et de mettre en avant les éléments pertinents relatifs à ce processus.

Nous rappelons qu'il existe, dans la littérature, d'autres modèles candidats (cf. Ahlgrim et al. [2008]). Nous ne prétendons pas que le modèle présenté est le « meilleur » modèle (d'ailleurs aucun modèle ne l'est) mais nous avons la garantie que celui-ci respecte certaines conditions minimales telles que la parcimonie et la plausibilité économique dans le sens où le GSE utilisé est capable d'imiter certains phénomènes financiers importants observés sur les marchés financiers. Ces phénomènes attirent souvent l'attention des analystes qui s'intéressent aux comportements joints de l'inflation, des taux d'intérêts et des marchés des actions. Les états de sortie (ou l'*output*) du modèle proposé ont été ainsi illustrés en partant initialement des paramètres obtenus via l'approche de calibrage décrite.

Conclusion de la partie I

Le glissement sémantique du « modèle d'actifs » vers le « générateur de scénarios économiques » matérialise la profonde évolution de ce sujet depuis une dizaine d'années maintenant.

Longtemps utilisée dans le cadre d'études spécifiques et ponctuelles, la construction d'une représentation des actifs dans lequel l'assureur peut investir ses avoirs et du contexte économique dans lequel ils évoluent est devenue un élément essentiel de la description des risques que celui-ci doit gérer.

Que ce soit dans le cadre prudentiel, pour la détermination des provisions et du capital de solvabilité, pour sa communication vers les tiers (MCEV et états comptables) ou pour ses besoins de pilotage techniques (choix d'allocations stratégique et tactique, tests de rentabilité, etc.), l'organisme assureur doit disposer d'un cadre rigoureux et cohérent prenant en compte l'ensemble des actifs de son bilan et les risques associés.

Si l'identification de ces risques peut être considérée comme relativement aboutie, la crise financière a mis en évidence certaines faiblesses dans leur modélisation. Deux éléments sont ainsi mis en évidence et vont sans doute donner lieu à de nombreux développements dans les prochaines années (cf. Planchet et al. [2009]) :

- la structure de dépendance entre les actifs ;
- le risque de liquidité, intimement associé à la gestion efficace des couvertures financières.

La quasi-totalité des modèles actuels s'appuient sur des structures de dépendance dans lesquels la corrélation tient une place centrale ; de nombreux travaux, dont nous pourrions trouver une synthèse dans Kharoubi-Rakotomalala [2008], mettent en évidence le caractère dynamique de l'intensité de la dépendance. En pratique, l'intensité de la dépendance augmente dans les situations défavorables, ce qui limite l'efficacité des mesures de

diversifications calibrées avec des structures ne prenant pas en compte cet effet. L'introduction de structures de dépendance non linéaires intégrant de la dépendance de queue apparaît ainsi comme un élément incontournable de l'évolution des générateurs de scénarios économiques.

Le risque de liquidité est également apparu comme un élément majeur de la crise des *subprimes* et, plus généralement, de la crise financière qu'elle a engendrée. Au moment où la généralisation des approches *market consistent* (basées sur les valeurs de marché) impose d'évaluer les options et garanties financières des portefeuilles dans la logique de détermination du coût de leur couverture, encore faut-il pour que le montant obtenu ait du sens que la couverture puisse être réajustée régulièrement, ce qui n'est possible qu'avec des actifs liquides. Cela impose par conséquent si ce n'est une adaptation de l'approche dans le cas d'actifs peu liquides, à tout le moins la prise en compte d'une prime de liquidité pour refléter dans le montant affiché ce risque d'impossibilité de gérer idéalement la couverture.

Au-delà de ces deux éléments structurants, l'efficacité opérationnelle des modèles mis en œuvre dépend, nous l'avons vu, dans une très large mesure de la pertinence des paramètres retenus pour les alimenter. La détermination de ses paramètres est complexe et fait appel à la fois à des considérations d'ordre statistique (exploitation des historiques), économique (cohérence des valeurs de long terme prédites par le modèle avec les relations économiques fondamentales), financières enfin (cohérence avec les prix observés sur le marché). La prise en compte rationnelle de ces différentes composantes nécessite une réflexion spécifique et fait partie intégrante des choix structurants en termes de gestion des risques que peut effectuer l'organisme assureur.

Enfin, la définition d'indicateurs de performance du modèle contribue à une meilleure alimentation du processus de prise de décision. L'exploitation de ces indicateurs dans le cas de modèles plus compliqués et plus concrets, tel que les modèles de gestion actif-passif, constitue notre prochain champ d'investigation.

PARTIE II :

**L'ALLOCATION STRATÉGIQUE D'ACTIFS DANS LE
CADRE DE LA GESTION ACTIF-PASSIF**

Introduction

Cette seconde partie est consacrée à l'élaboration de l'allocation d'actif elle-même. Rappelons que concernant les régimes de retraite par répartition, trois situations peuvent être recensées : régime provisionné (cas où le régime est tenu à la couverture totale des engagements souscrits par ses cotisants actuels et futurs), régime partiellement provisionné (couverture d'une partie seulement des engagements souscrits par ses cotisants actuels et futurs) et enfin régime non provisionné.

Le fonctionnement financier des régimes de retraite en France (objet de la plupart des applications illustratives dans ce travail) peut être schématisé comme suit : les flux de cotisations permettent de régler les flux de prestations. Le surplus, le cas échéant, permet d'alimenter une réserve destinée à régler une partie des prestations futures. Cette même réserve peut se voir prélever, le cas échéant, le solde technique débiteur (cotisations insuffisantes pour régler les prestations). Entre temps, la réserve est placée sur le marché financier et répartie entre différentes classes d'actifs.

La gestion actif-passif d'un régime par répartition partiellement provisionné peut être basée sur l'optimisation de la valeur de la réserve, compte tenu des contraintes liées au passif qu'il doit respecter. C'est dans ce cadre que le choix d'une « bonne » allocation stratégique joue un rôle essentiel dans le pilotage actif-passif d'un régime de retraite, étant donné que les réserves contribuent à part entière à la solidité du régime.

Cependant une difficulté consiste à définir la « meilleure » stratégie de placement de ces réserves sur les marchés financiers, notamment dans un contexte de fortes incertitudes économiques, comme c'est actuellement le cas du fait de la récente crise financière.

Il est à rappeler que le modèle d'allocation d'actifs le plus répandue est celui de Markowitz [1952]. L'objectif de ce modèle est de résoudre le problème suivant : pour une rentabilité espérée fixée par l'investisseur, trouver les proportions à investir dans chaque titre, conduisant à un portefeuille de risque minimum et sous un ensemble de contraintes. L'ensemble des portefeuilles vérifiant le programme d'optimisation de Markowitz pour différents niveaux de rentabilité espérée, décrit une hyperbole dans le repère (risque mesuré par l'écart type de la rentabilité / rentabilité espérée).

Le modèle de Markowitz n'est pas un modèle de gestion au jour le jour mais un modèle de structure à long terme qui a conduit à de nombreux développements sur la théorie de la gestion de portefeuille. La limite de ce modèle réside dans le fait d'utiliser la variance pour mesurer le risque du portefeuille. Quoique simple et parfois efficace, la variance est contre intuitive : Les investisseurs n'ont aucune aversion face à des rendements supérieurs à ceux prévus, ils s'intéressent seulement à ce qui est en dessous de leurs prévisions où de leurs objectifs personnels. Or, la variance est une mesure de dispersion qui ne fait aucune différence entre ce qui est en dessous et au-dessus de l'espérance de rendement. En plus il ne tient pas compte de l'existence d'un passif.

Dans la littérature, les modèles d'ALM sont généralement classés en trois groupes présentés chronologiquement comme suit.

Le premier groupe contient les modèles d'adossment (ou *matching*) et d'immunisation par la duration (cf. Macaulay [1938], Redington [1952]). Ces modèles se basent sur le fait que les

investissements sont essentiellement effectués dans des obligations. Ceci nous permet d'obtenir, soit un adossement des flux de trésorerie des actifs financiers à ceux du passif (*matching*), soit un adossement de la durée de l'actif à celle du passif (immunisation par la durée).

Ces techniques étaient utilisées jusqu'au milieu des années 80 et avaient comme inconvénients principaux la considération du risque de taux comme seule source de risque pour le fonds, ainsi que la nécessité d'un rebalancement périodique du portefeuille en ré-estimant à chaque fois la durée du passif, qui change continûment du fait du changement des taux d'intérêts et de l'écoulement du temps. Pierre [2009] présente une approche de la couverture de passif qui vise à résoudre les problèmes mentionnés ci-dessus. L'actif sera constitué dans ce cas d'un portefeuille de couverture de type taux (obligations ou dérivés) couplé à un portefeuille de rendement. Selon Pierre [2009], cette architecture permet une flexibilité suffisante pour permettre à l'actif de s'adapter aux mise à jour de la valeur des engagements lors de la revue des hypothèses actuarielles ayant permis leur évaluation.

Le deuxième groupe contient les modèles basés sur la simulation de scénarios déterministes et sur la notion de surplus (Kim et Santomero [1988], Sharpe et Tint [1990], Leibowitz et al. [1992]). Les modèles de surplus ont pour objet la minimisation du risque de perte du surplus (mesuré par la variance de la rentabilité du surplus) sous contraintes de rentabilité et de poids des actifs. Ils sont des modèles mono-périodiques, ce qui limite leur utilité en pratique pour des problèmes d'allocation sur le long terme.

Le troisième groupe de modèles utilise les techniques de simulation stochastique (*Monte Carlo*) pour modéliser l'évolution des différents éléments, que ce soit au niveau des actifs financiers et des engagements, ou au niveau des variables de marché et des variables démographiques (cf. Frauendorfer [2007], Munk et al. [2004], Waring [2004], Martellini [2006]). Ainsi les lois de probabilité associées aux résultats du fonds de retraite sur le long terme peuvent être estimées. A ce niveau, nous nous proposons de distinguer deux sous-groupes de modèles d'ALM basés sur les techniques stochastiques. L'élément clé de distinction sera si oui ou non les poids des différents actifs reviennent périodiquement à ceux de l'allocation stratégique définie initialement (si oui, les modèles seront appelés modèles à poids constants ou stratégie *Fixed-Mix*).

- Pour le premier sous-groupe de modèles à poids constants et malgré les avancées réalisées avec ces techniques (surtout au niveau de l'implémentation informatique), l'aspect dynamique de l'allocation stratégique reste encore marginalisé. En fait, ces modèles permettent de comparer des allocations constantes dans le temps (statiques) indépendamment des opportunités liées aux évolutions inter-temporelles des marchés (cf. Merton [1990], Kouwenberg [2001], Infanger [2002], Dempster et al. [2003]).
- Le deuxième sous-groupe de modèles, et le plus récent, est principalement inspiré de la théorie du choix de la consommation et de portefeuille développée par Merton [1971]. Il s'agit des modèles d'allocation dynamique ou inter temporels. Par exemple, à partir de la définition de la fonction objectif pour l'investisseur, ces modèles permettent la détermination d'une trajectoire des poids des différents actifs jusqu'à la date d'échéance (l'ajustement des poids est fonction des évolutions projetées du marché et de la règle de gestion prédéfinie). L'allocation stratégique retenue sera l'allocation optimale d'actifs

à la date initiale t_0 . Le cadre d'utilisation de ces modèles récents se heurte au problème d'implémentation vu la complexité des outils mathématiques employés (cf. Hainaut et al. [2005], Rudolf et al. [2004] et Yen et al. [2003]).

Concernant les modèles d'allocation d'actifs, notre étude est axée sur la comparaison des modèles disponibles selon les hypothèses sous-jacentes de « rebalancement » des portefeuilles : nous faisons la distinction entre une gestion « statique » (pour laquelle les poids reviennent périodiquement à ceux de l'allocation stratégique de long-terme - cas par exemple des modèles à poids constants *Fixed-Mix*) et une gestion dite « dynamique », pour laquelle les poids peuvent s'écarter définitivement de l'allocation stratégique initiale selon des règles de gestion prédéfinies.

Nous verrons que l'approche dynamique présente l'avantage théorique de la robustesse face aux changements de régime des marchés. L'autorisation du changement des poids des différentes classes d'actifs, sur la base d'une règle de gestion bien définie, constitue *a priori* un élément intéressant. Cela en effet permet l'ajustement des expositions aux différentes classes d'actifs suite à l'évolution des conditions de marché.

Cette réflexion sur la construction de modèles d'allocation d'actifs applicables dans une optique prévisionnelle à long-terme nous conduira à étudier une approche innovante fondée sur les techniques de « programmation stochastique » (cf. Birge et Louveau [1997]). Il s'agit d'une version adaptée d'une technique déjà utilisée dans le domaine de l'ingénierie pour la planification de la production (cf. Dantzig et al. [1990], Escudero et al. [1993]). L'objectif sera la mise en évidence et l'étude des caractéristiques de cette approche.

Comme déjà mentionné, cette thèse accorde également une attention toute particulière aux techniques numériques de recherche de l'optimum, qui demeurent des questions essentielles pour la mise en place d'un modèle d'allocation. Le point de départ sera notre constat d'un temps de calcul significatif dû simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Dans ce cadre, nous présentons un algorithme d'optimisation globale d'une fonction non convexe et bruitée. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement.

Le plan de cette deuxième partie sera comme suit :

- Premier chapitre : Il est question dans ce chapitre d'analyser les approches classiques (ou déterministes) de gestion actif-passif à savoir les modèles basés sur la notion de « duration » et de « surplus ». Nous détaillons différentes approches tout en mettant en évidence la différence entre elles. L'objectif de ce chapitre est de revoir l'état de l'art en matière de modèles d'ALM déterministe. Après avoir mis en évidence les limites de ces modèles, nous passons à l'étude de modèles plus élaborés et plus sophistiqués (objet des chapitres ultérieurs).

- Deuxième chapitre : Nous considérons ici une approche récente d'allocation stratégique d'actifs basée sur la stratégie dite « à poids constants » ou *Fixed-Mix* (cf. Merton [1990], Kouwenberg [2001], Infanger [2002], Dempster et al. [2003]).

Nous proposons une modélisation du régime-type de retraite et étudions certains critères d'allocation stratégique d'actifs en fonction du type du régime (provisionné, partiellement provisionné, etc.) sont étudiés. Nous passons ensuite à l'illustration de la stratégie *Fixed-Mix* avec une application sur les réserves d'un régime de retraite partiellement provisionné. Les résultats obtenus sont discutés et différents tests de sensibilité sont mises en place : ces tests sont liés principalement à l'impact des hypothèses de rendement ou de corrélation retenues.

Les difficultés rencontrées lors de la mise en place de la stratégie *Fixed-Mix*, dues essentiellement à la multiplicité du nombre de classes d'actifs considérées et au nombre de scénarios économiques simulés, nous ont mené à nous pencher sur les aspects d'optimisation numérique. Le point de départ est notre constat d'un temps de calcul significatif dû simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Dans ce cadre, une présentation détaillée des travaux menés lors de la rédaction de l'article de Rulhière et al. [2010] est effectuée. Un algorithme d'optimisation globale d'une fonction non convexe et bruitée est présenté. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement et compare les résultats obtenus avec l'algorithme de Kiefer-Wolfowitz- Blum (cf. Blum [1954], Kiefer et Wolfowitz [1952]).

- Troisième chapitre : Les modèles classiques d'ALM dynamiques sont explorés, notamment les techniques d'assurance de portefeuille (cf. Perold et Sharpe [1988]) et les techniques de programmation dynamique (cf. Cox et Huang [1989], Merton [1971]). A ce niveau, les techniques d'assurance de portefeuille basées sur la notion de CPPI ou *Constant Proportion Portfolio Insurance* (cf. Perold et Sharpe [1988]) sont mises en place et certains résultats relatifs à ce modèle sont étudiés. De même, les principes des techniques de programmation dynamique et leurs limites sont également mises en évidence.

Nous nous penchons par la suite sur une approche innovante fondée sur les techniques de « programmation stochastique » (cf. Birge et Louveaux [1997]). Il s'agit d'une version adaptée d'une technique déjà utilisée dans le domaine de l'ingénierie pour la planification de la production (cf. Dantzig et al. [1990], Escudero et al. [1993]). Dans ce cadre, nous mettons en place un modèle d'ALM dynamique basé sur les techniques de programmation stochastique.

Nous proposons, au cours d'une illustration numérique, une nouvelle méthodologie de génération de scénarios économiques que nous appelons méthodologie « des quantiles de référence ». Cette dernière permet de partir d'une structure linéaire de génération de scénarios (telle que décrite dans le chapitre 2 de la partie I) pour réduire la dimension du problème rencontré avec la stratégie *Fixed-Mix* tout en tenant compte de la corrélation entre les distributions des différentes variables projetées.

A travers la même application numérique, nous comparons certains résultats relatifs aux deux approches d'allocation stratégique d'actifs : celle basée sur la stratégie *Fixed-Mix* et celle basée sur les techniques de programmation stochastique. Nous testons également la sensibilité de cette deuxième approche par rapport au changement de certains de ses paramètres, toutes choses étant égales par ailleurs.

Chapitre 1 : Les modèles d'ALM classiques (déterministes)

Il est question dans ce qui suit d'analyser les approches classiques (ou déterministes) de gestion actif-passif à savoir les modèles basés sur la notion de « duration » et de « surplus ». Nous détaillons différentes approches tout en mettant en évidence la différence entre elles. L'objectif de ce chapitre est de revoir l'état de l'art en matière de modèles d'ALM déterministe.

I- Immunisation du portefeuille

Redington [1952] définit l'immunisation du portefeuille comme « l'investissement de l'actif d'une telle manière que le portefeuille soit protégé contre un changement des taux d'intérêt ». Autrement, c'est une stratégie d'investissement qui, dans le cas de l'assurance vie ou de celui des régimes de retraite, produit des flux exactement adossés en maturité et en valeur à ceux que doit payer l'entreprise. Cette technique suppose que le risque principal des portefeuilles financiers est celui du changement des taux d'intérêt.

Nous pouvons énumérer deux approches classiques appartenant à la catégorie de l'immunisation :

- l'adossement des flux de trésorerie (*cash flows*) à ceux du passif.
- l'adossement des durations de l'actif et du passif.

Nous rappelons cependant quelques notions qui seront reprises ultérieurement dans cette étude (cf. Le Vallois et al. [2003]).

Taux de rendement actuariel

En suivant à une démarche inverse à celle du calcul de la valeur actuelle d'un actif (ou de toutes séquences de flux fixes), il est possible de calculer, à partir du prix de marché, un taux d'actualisation des flux de trésorerie correspondants dit taux de rendement actuariel r_a . Ce dernier vérifie :

$$\text{Prix de marché} = \text{Valeur actuelle} = \sum_{i=1}^n \frac{F_{t_i}}{(1 + r_a)^{t_i}}$$

Avec F_{t_i} est le flux de l'actif financier (coupon, dividende, amortissement...) à l'époque t_i .

Courbe des taux zéro-coupon

Le calcul du taux de rendement actuariel des actifs financiers, et des obligations en particulier, présente l'inconvénient de ne pas donner la même valeur de taux pour les différents actifs. Dans le cas des obligations, cela est dû principalement à deux facteurs.

- Le premier facteur est que les investisseurs demandent aux émetteurs privés des primes de risque appelées *spread* de signature qui compense le risque de défaillance propre à chaque émetteur (ce *spread* diffère donc d'une obligation à l'autre).
- Le deuxième facteur est la dépendance des taux de rendement actuariels des échéances des obligations correspondantes. En fait même si nous limitons l'analyse aux emprunts d'Etat, tous identiques en terme de risque, nous

constatons encore des différences entre les taux de rendement actuariels de différents titres. Il faudrait utiliser un taux d'actualisation différent pour chaque échéance. Le seul cas où le taux de rendement actuariel observé correspond sans ambiguïté à une seule maturité est le cas où il n'y a qu'un seul flux de trésorerie. Ce cas existe puisqu'il correspond aux titres dits obligations zéro-coupon pour lesquels les intérêts sont versés en une seule fois au moment de l'échéance finale et unique.

L'observation des prix des zéro-coupon permettrait donc de bâtir une courbe des taux zéro-coupon. Cette courbe existe, mais en pratique elle est obtenue par des moyens différents compte tenu de la faible liquidité des zéro-coupon existants sur le marché. La courbe zéro-coupon est également appelée structure par terme des taux, ou encore courbe des taux spot.

Valeur actuelle nette des actifs obligataires

En pratique il n'est pas nécessaire de calculer la valeur actuelle des actifs obligataires, puisqu'il est plus simple d'observer leur valeur de marché. Cependant, dans le cas des obligations à taux fixe, la séquence des flux de trésorerie futurs associés à un titre est parfaitement connue.

Nous considérons même généralement que dans le cas des emprunts d'Etat, cette séquence est certaine et ne présente aucun risque de défaillance de l'émetteur. A l'aide de la structure par terme des taux (notée r_{t_i}) il est possible de reconstituer le prix d'équilibre de l'obligation sans

risque :
$$\text{Prix de marché} \approx \text{Prix estimé} = \sum_{i=1}^n \frac{F_{t_i}}{(1 + r_{t_i})^{t_i}}$$

Valeur actuelle des passifs

Par analogie avec la valorisation des actifs, la valorisation des flux du passif se fait en se basant sur la courbe des taux zéro-coupon. Cette dernière, permet de tenir compte de la maturité de chacun des flux pour lui attribuer un taux d'actualisation précis. De même, il faudrait ajouter à cette courbe des taux sans risque une prime de risque (*spread*). Elle permettra de retomber sur une « valeur de marché » des passifs. Le problème c'est que le niveau réel de cette prime est inobservable sur le marché dans la mesure où les engagements du passif ne s'échangent pas régulièrement sur des marchés liquides et organisés. En pratique, le calcul de la valeur actuelle du passif peut être réalisé avec une prime de risque arbitrairement choisie, éventuellement nulle.

I-1 Adossement des flux de trésorerie

Il s'agit de la procédure d'immunisation la plus simple et la plus ancienne. Elle consiste à investir la richesse initiale dans un portefeuille de titres (le plus souvent des zéro-coupons) qui produisent exactement, et aux échéances prévues, les flux du passif.

Autrement, lorsque sur chaque période les flux nets obtenus (total des entrées de flux – total des sorties de flux) sont toujours positifs ou nuls, l'actif est dit adossé au passif. Il est dit exactement ou parfaitement adossé si les flux nets sont nuls. Les anglo-saxons parlent dans ce cas de méthode de *cash flow matching*. Une fois que tous les passifs sont couverts par cette méthode, les actifs excédentaires sont alors considérés comme libres et représentatifs de la

situation nette réelle de la société (le *shareholder surplus*). Le même raisonnement peut être adopté dans le cas des régimes de retraite par répartition.

En pratique, l'adossement par les flux de trésorerie est toujours utilisé par les sociétés avec la différence que les passifs peuvent être regroupés sur d'autres critères que le seul « rendement ». En fait, les engagements au passif des assureurs par exemple sont généralement répartis par famille de contrats (en fonction des prix de revient, des taux de rendement actuariels, de la durée des titres...) appelés contons, puis les provisions correspondantes et les flux de trésorerie associés sont calculés. Cela permet outre l'adossement des flux de trésorerie, la bonne adéquation des actifs en termes de taux de rendement financier et comptable comparé au taux garanti moyen du passif correspondant.

L'adossement des flux de trésorerie est valable à très court terme, mais malheureusement, elle ne trouve pas d'applications pratiques à long terme. En effet, les flux de l'actif et ceux du passif sont très influencés par des facteurs externes, notamment les taux d'intérêt. Ils sont donc eux-mêmes sujet à des variations dans des sens différents, il importe donc de réitérer souvent cette immunisation.

I-2 Adossement par la duration

Cette technique consiste à appairer les sensibilités de l'actif et du passif vis-à-vis de la variation des taux d'intérêt. En d'autres termes, elle consiste à définir une stratégie d'investissement qui fait que la valeur de marché des actifs suit tout mouvement de la valeur actuelle des engagements. L'immunisation par la duration définit un portefeuille dont la valeur, au premier ordre, évolue comme la valeur actuelle des engagements. La règle de décision dans cette technique est basée sur l'indice de « sensibilité », défini par Macauley [1938]. Il est obtenu à partir de la formule de développement limite de Taylor du prix en fonction du taux d'intérêt.

Pour un titre, dont les caractéristiques contractuelles sont déterminées indépendamment du mouvement des taux, la sensibilité est exprimée comme la variation relative de cette valeur induite par une variation infinitésimale de taux d'intérêt.

- La sensibilité de l'actif :

Nous reprenons l'expression simplifiée de la valeur actuelle d'une obligation à taux fixe en fonction du taux d'actualisation actuariel r_a :

$$VA(r_a) = \sum_{i=1}^n \frac{F_{t_i}}{(1+r_a)^{t_i}}$$

La variation de la valeur actuelle pour une petite variation du taux d'actualisation est donnée par le calcul de la dérivée première :

$$\frac{dVA(r_a)}{dr} = VA'(r_a) = \sum_{i=1}^n -t_i \frac{F_{t_i}}{(1+r_a)^{t_i+1}}$$

Cette dérivation n'a de sens que si les flux F_{t_i} sont fixes par rapport à r .

La sensibilité est ainsi donnée par l'expression suivante :

$$sensibilité = \frac{dVA(r_a)}{VA \times dr} = -\frac{1}{VA} \times \sum_{i=1}^n t_i \frac{F_{t_i}}{(1+r_a)^{t_i+1}}$$

Lorsque les flux de trésorerie sont tous positifs, la sensibilité de la valeur actuelle aux variations du taux d'actualisation est nécessairement négative. Cela est effectivement le cas pour les obligations dont la valeur de marché croît quand les taux baissent et réciproquement. La sensibilité est appelée aussi duration modifiée.

- La duration de l'actif

La duration telle que définie par Macaulay [1938] est donnée par l'expression suivante :

$$duration = \frac{\sum_{i=1}^n t_i \frac{F_{t_i}}{(1+r_a)^{t_i}}}{\sum_{i=1}^n \frac{F_{t_i}}{(1+r_a)^{t_i}}} = \frac{\sum_{i=1}^n t_i \frac{F_{t_i}}{(1+r_a)^{t_i}}}{VA(r_a)}$$

Avec F_{t_i} une série de flux fixes.

La duration peut s'interpréter comme la durée de vie moyenne de l'obligation. En fait, chaque durée t_i étant pondérée par la valeur actuelle du flux correspondant $\frac{F_{t_i}}{(1+r_a)^{t_i}}$. Elle peut s'interpréter aussi comme l'élasticité du prix de l'obligation par rapport aux variations des taux. C'est la définition de Hicks [1946] qui relie la sensibilité à la duration :

$$duration = -\frac{\frac{dVA(r_a)}{VA(r_a)}}{\frac{dr_a}{1+r_a}} = -(1+r_a) \times \frac{\frac{dVA(r_a)}{VA(r_a)}}{dr_a} = -(1+r_a) \times sensibilité$$

La duration s'interprète aussi comme la date à la quelle les deux effets de la variation des taux, effet sur la valeur et effet sur le revenu, se compensent.

Elle est compatible avec une structure par terme des taux à condition de se limiter aux variations parallèles des courbes de taux. Autrement, la variation de taux dr_t doit s'appliquer d'une façon uniforme à la courbe des taux (chaque r_{t_i} devient $r_{t_i} + dr_t$) ce qui correspond bien à une translation uniforme de la courbe des taux. La duration s'exprime dans ce cas comme suit :

$$duration = \frac{\sum_{i=1}^n t_i \frac{F_{t_i}}{(1+r_{t_i})^{t_i}}}{\sum_{i=1}^n \frac{F_{t_i}}{(1+r_{t_i})^{t_i}}} = \frac{\sum_{i=1}^n t_i \frac{F_{t_i}}{(1+r_{t_i})^{t_i}}}{VA(r_{t_i})}$$

Si nous étendons la notion de la duration à des flux variables (obligations à taux variables par exemple), nous pourrions dans certains cas obtenir la relation ci dessus entre la duration et la sensibilité (tant que les dérivées $F'_{t_i}(r_a)$ sont connues), mais la différence c'est que la duration n'est plus assimilable à une durée moyenne.

La duration peut être négative si certains des flux sont négatifs. La sensibilité dans ce cas est positive : la valeur actualisée de la séquence examinée augmente quand les taux augmentent et inversement. La duration d'une obligation zéro coupon est égale à la durée de cette obligation.

Le calcul de la duration d'un portefeuille s'effectue par trois méthodes :

- La première consiste à déterminer la moyenne des durations pondérées par la valeur de marché des titres correspondant (coupons courus inclus). L'inconvénient de cette méthode c'est qu'elle ne peut avoir de sens que si les taux actuariels de tous les titres étaient identiques, ce qui bien entendu n'est jamais le cas.
- La deuxième méthode consiste à cumuler tous les flux de trésorerie du portefeuille et à calculer un taux de rendement actuariel unique en fonction de la valeur de marché totale, coupon courus inclus. Il est ensuite possible de calculer une duration globale puisque la formule de calcul est valable à toute séquence de flux fixes.
- La troisième méthode est la plus rigoureuse. Elle consiste à calculer la valeur actuelle et la duration de chaque titre avec une courbe des taux commune, puis à calculer la moyenne pondérée par les valeurs actuelles et non pas par les valeurs de marché. C'est théoriquement la meilleure solution, car chaque flux de trésorerie est actualisé avec le taux qui correspond à sa maturité. En pratique, la différence obtenue entre les différentes méthodes de calcul de la duration d'un portefeuille n'est pas significative.

Par ailleurs, nous pouvons étendre la notion de duration à d'autres actifs non obligataires malgré qu'ils sont plus ou moins sensibles à la variation des taux notamment les actions. Pour ces derniers le calcul de la duration peut être effectué dans le cadre du modèle Gordon-Shapiro [1956] (qui permet d'analyser le prix d'une action comme étant l'actualisation de la série des dividendes projetés) ou tout simplement dans le cadre de la régression des prix de marché par rapport aux variations de taux. Les actifs et les passifs conditionnels ne doivent donc pas faire l'objet de calculs élémentaires d'adossement en duration.

Pour accomplir une immunisation par la duration, l'investisseur doit acquérir des titres dont la duration moyenne est égale à la duration des flux du passif.

La notion de sensibilité est un indicateur de l'exposition au risque de taux par la Valeur Actuelle Nette (VAN) de la société.

En fait, à partir de la relation suivante : $VA_{nette} = VA_{actif} - VA_{passif}$ nous pouvons exprimer l'exposition de la VAN au risque de taux comme suit :

$$\frac{dVA_{nette}}{dr} = \frac{dVA_{actif}}{dr} - \frac{dVA_{passif}}{dr}$$

Si nous considérons le cas où la VAN initiale est égale à zéro, et où la sensibilité de l'actif et du passif sont identiques, alors la variation de taux est sans influence sur la VAN :

$$\frac{dVANette}{dr} = 0$$

Nous parlons dans ce cas de l'immunisation du risque de taux. En effet, la VAN devient insensible, sinon à toute variation de taux, du moins à une petite variation parallèle de la courbe des taux. Cet effet est obtenu en alignant les sensibilités de l'actif et du passif.

Nous notons également que :

- Plus la duration est élevée plus le risque est grand et plus son prix est sensible aux variations du taux de marché.
- Plus le coupon est faible, plus la sensibilité est élevée.
- Nous couvrons systématiquement le passif par une obligation de durée plus longue (car la duration est inférieure à la durée sauf pour le cas des zéro-coupon).
- En cas de baisse des taux, le gain réalisé sur la valeur de l'obligation est supérieur à la perte encourue sur le passif.
- L'immunisation en duration d'un portefeuille n'est parfaite que si elle est réalisée avec des instruments zéro-coupon. Elle ne tient pas compte de la forme de la courbe des taux, ni de ses déformations, ni de sa dynamique.

- Limites de l'adossement par la duration

Ils existent deux limites principales des indicateurs ci-dessus :

- Le domaine d'utilisation de ces concepts est limité aux variations parallèles de la courbe des taux.
- Les calculs de la duration et de la convexité ne gardent leurs significativité que pour des flux fixes, indépendants des taux de marché.

La première limite n'est pas sans conséquence mais elle n'invalide pas totalement l'analyse du risque de taux. En effet, les variations parallèles de la courbe des taux représentent la principale source de risque pour la plupart des portefeuilles obligataires (résultats obtenus par des analyses statistiques de la variance des taux en composante principale).

En revanche, la deuxième limite (flux fixes) est très contraignante pour l'analyse des passifs comprenant des options complexes exercées par les clients (tel que l'option de rachat de l'épargne). L'exercice de certains de ces options dépend fortement de l'évolution des taux sur le marché. Pour le cas d'option de rachat, une augmentation des niveaux de taux a pour conséquence évidente la hausse des niveaux de rachats afin de profiter de placement plus rentable sur le marché. En pratique, nous faisons souvent l'hypothèse que pour de petites variations des taux de marché, les flux de trésorerie liés au passif restent fixes. Les calculs de

sensibilité et de duration ont alors une validité « locale » dont l'intérêt est cependant limité pour l'adéquation actif-passif.

Par ailleurs, l'immunisation par la duration requiert un rebalancement périodique du portefeuille en réestimant à chaque fois la duration du passif, celle-ci change continûment du fait du changement des taux d'intérêt et de l'écoulement du temps. Lorsque nous utilisons une stratégie avec une duration constante à chaque rebalancement, nous risquons de se trouver dans des situations de redondance ou de déficit.

Le principal inconvénient reste l'hypothèse de changements parallèles dans la structure des taux. Pour lever partiellement cet inconvénient, nous pouvons utiliser le deuxième terme du développement limite de la fonction $P = f(r_t)$. Nous introduisons alors la notion de convexité.

- La convexité de l'actif

La sensibilité permet de traiter de la variation du cours des obligations par rapport aux petites variations parallèles de la courbe des taux. Mais lorsque la variation des taux est forte, l'erreur obtenue en mesurant la variation relative du cours par la droite tangente devient importante. Nous faisons donc recours au développement de Taylor du second ordre pour améliorer notre estimation de la variation de la valeur VA :

$$\Delta VA \approx VA' \times \Delta r + \frac{1}{2!} VA'' \times \Delta r^2$$

Formule où VA' et VA'' représentent respectivement la dérivée première et la dérivée seconde de $VA(r)$. En divisant par VA , Nous obtenons :

$$\frac{\Delta VA}{VA} \approx \frac{VA'}{VA} \times \Delta r + \frac{1}{2!} \frac{VA''}{VA} \times \Delta r^2$$

L'expression de la convexité est la suivante :

$$convexité = \frac{VA''}{VA} = \frac{1}{VA \times (1+r)^2} \times \left(\sum_{i=1}^n t_i \times (t_i + 1) \times \frac{F_{t_i}}{(1+r)^{t_i}} \right)$$

Avec cette définition, l'approximation de la valeur actuelle du cours de l'obligation est la suivante:

$$\Delta VA \approx VA \times \left[sensibilité \times \Delta r + \frac{convexité}{2} \times (\Delta r)^2 \right]$$

Comme la duration, la convexité est compatible avec l'analyse de la structure par terme des taux. Nous pouvons écrire :

$$convexité = \frac{1}{VA} \times \left(\sum_{i=1}^n t_i \times (t_i + 1) \times \frac{F_{t_i}}{(1+r_{t_i})^{t_i+2}} \right)$$

Comme pour la duration cette extension n'est valide que pour des variations parallèles de la courbe des taux.

Lorsque les flux de trésorerie sont tous fixes et positifs, la convexité est positive. Quel que soit le sens de la variation du taux d'actualisation, le terme $\frac{convexité}{2} \times (\Delta r)^2$ intervient positivement dans la variation du cours. Il est donc, habituel de considérer qu'une convexité supplémentaire liée à la dispersion des flux de trésorerie est « la bienvenue ».

L'insuffisance de l'indice de « duration », comme mesure de sensibilité, dans le cas systématique où les flux sont sensibles aux variations de taux a été la motivation majeure du développement d'autres approches (dont les modèles basés sur le surplus et les techniques de simulations). Le but est de fournir un indice de « sensibilité » plus adapté au cas des flux sensibles aux variations de taux.

II- Modèles basés sur le surplus

Ces modèles sont basés sur la notion du surplus S qui est souvent défini comme étant la différence entre la valeur de marché des actifs et la valeur actuelle du passif. Dans la suite, nous illustrons ces modèles via trois approches : celle de Kim et Santomero [1988], Sharpe et Tint [1990] et Leibowitz [1992]. Une quatrième approche permettant de lier simultanément les notions du surplus et de la duration sera également présentée.

II-1 Modèle de Kim et Santomero [1988]

Ce modèle donne plus d'importance aux conditions spécifiques à l'entreprise concernée. Celles-ci sont représentées par un ratio d'effet de levier : « *surplus/valeur de marché initiale de l'actif* ».

Il s'applique aux portefeuilles constitués de plusieurs classes d'actifs. Le but du modèle est de minimiser le risque de perte du surplus (mesuré par la variance de la rentabilité du surplus) sous contraintes de rentabilité et de poids des actifs.

Formulation théorique

La notion de rentabilité du surplus utilisée est la variation relative du surplus sur la période (année par exemple) :

$$R_s = \frac{S_1 - S_0}{S_0}$$

avec le surplus initial S_0 et S_1 le surplus final.

Nous considérons un portefeuille avec M classes d'actifs risqués et un passif évalué en valeur actuelle. Le vecteur des rentabilités attendues est : $R = [R_A, R_P]_{(M+1,1)}$

$(R_A)_{M \times 1}$ représente le vecteur rentabilités des actifs.

R_P la variation relative de la valeur actuelle du passif.

Le vecteur des poids est noté : $w = [w_A, w_P]$

Le poids de l'actif i est mesuré par la valeur de marché initiale de cet actif divisée par la valeur initiale du surplus.

$$w_A^i = \frac{A_0^i}{S_0} \text{ poids de l'actif } i.$$

$$w_P = -\frac{P_0}{S_0} \text{ poids du passif.}$$

Nous vérifions que la somme des poids est égale à un.

La rentabilité du surplus est donc :

$$R_s = w'_A R_A + w_P R_P$$

Ce modèle dépend d'un ratio $k_0 = \frac{S_0}{A_0}$ appelé ratio de solvabilité initial.

Nous vérifions que $k_0 = 1 - \frac{1}{rf_0}$ avec rf_0 le ratio de financement initial (valeur de marché des actifs divisée par la valeur actuelle du passif). Ce modèle permet d'obtenir une frontière efficiente en utilisant la même démarche que Markowitz [1952] : minimiser la variance de la rentabilité du surplus σ_s^2 pour un niveau de rentabilité du surplus donné avec une contrainte sur les poids (somme des poids égale à un).

Notations

Dans leurs travaux, Kim et Santomero [1988] utilisent les notations suivantes :

$$\begin{aligned} A &= e' V_A^{-1} R_A & B &= R_A' V_A^{-1} R_A > 0 \\ C &= e' V_A^{-1} e > 0 & D &= BC - A^2 > 0 \\ F_0 &= R_A' V_A^{-1} V_{AP} & F_1 &= e' V_A^{-1} V_{AP} \\ F_2 &= V_{AP}' V_A^{-1} V_{AP} > 0 \end{aligned}$$

Avec :

$$e' = (1, \dots, 1)$$

V_A la matrice covariance des actifs

V_{AP} le vecteur covariance entre les actifs et le passif.

Equation analytique de la frontière efficiente

Le couple solution du problème vérifie la relation suivante, cette relation est l'équation d'une hyperbole dans le repère écart-type du surplus/rentabilité du surplus (cf. graphique 31) :

$$R_S = R_{SM} \pm \frac{1}{C} [DC(\sigma_S^2 - \sigma_{SM}^2)]^{1/2}$$

R_{SM} et σ_{SM} sont les caractéristiques du portefeuille de variance minimale (Coordonnées du sommet de l'hyperbole).

$$\begin{aligned} R_{SM} &= \frac{1}{k_0} \frac{A}{C} + \left(\frac{1}{C} \right) \left(\frac{1}{k_0} - 1 \right) (CF_0 - AF_1) - \left(\frac{1}{k_0} - 1 \right) R_P \\ R_{SM} &= \left(1 - \frac{1}{k_0} \right)^2 (Var(R_p) - F_2) + \left(\frac{1}{C} \right) \left(\frac{1}{k_0} - \left(\frac{1}{k_0} - 1 \right) F_1 \right) \end{aligned}$$

La partie supérieure de l'hyperbole représente une frontière efficiente qui dépend du niveau du ratio k_0 (solvabilité de l'entreprise).

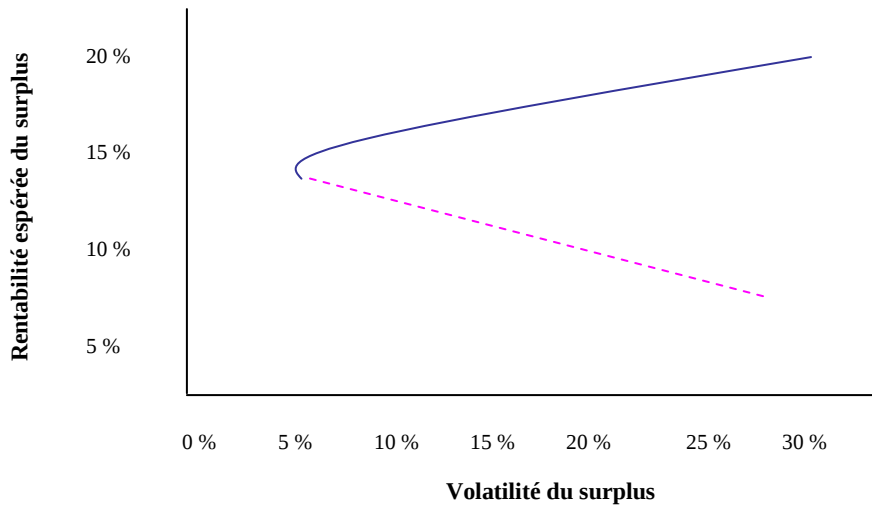


Fig. 31 : *Frontière efficiente (Kim&Santomero)*

Contrainte de déficit

Nous supposons que le vecteur des rentabilités $R = [R_A, R_P]_{(M+1,1)}$, est gaussien (la variable rentabilité du surplus suit donc une loi normale). Nous introduisons la contrainte de déficit suivante : « La probabilité pour que la rentabilité du surplus soit inférieure à un certain seuil ne doit pas dépasser une probabilité donnée ». Ce qui revient à limiter la probabilité de perte d'une partie du surplus initial.

Cette contrainte est représentée par une droite de déficit dans le repère (écart-type de la rentabilité du surplus, rentabilité du surplus espérée). L'intersection entre cette droite et la frontière efficiente donne le portefeuille efficient vérifiant la contrainte (cf. graphique 32).

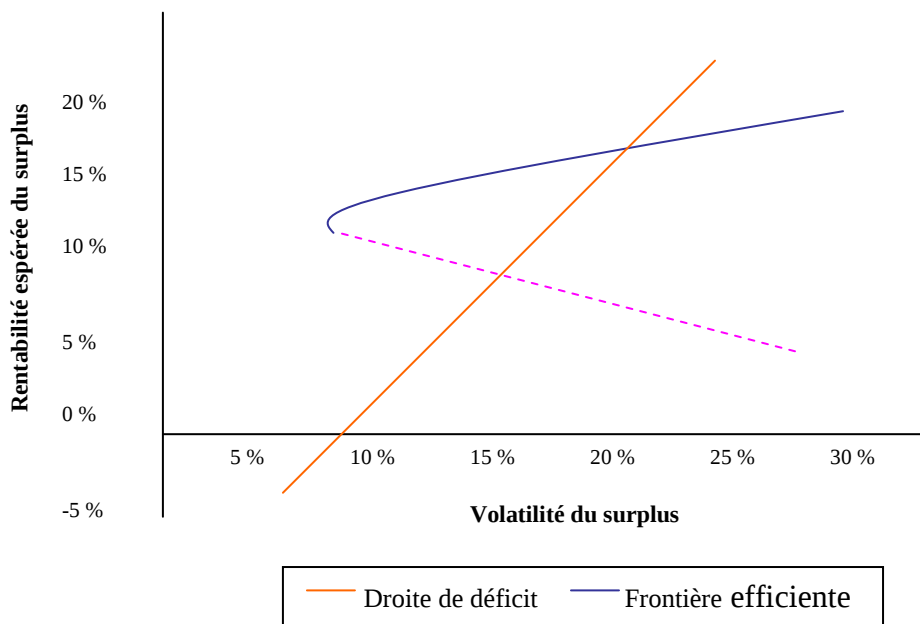


Fig. 32 : *Frontière efficiente et contrainte de déficit (Kim&Santomero)*

Nous pouvons aussi, en utilisant les techniques de Roy [1952], chercher le portefeuille efficient qui minimise la probabilité de perdre un pourcentage du surplus initial. Cette approche a deux avantages :

- La prise en compte de la corrélation entre la valeur de marché de l'actif et la valeur actuelle du passif
- L'introduction d'une relation entre la solvabilité de l'entreprise et l'allocation optimale d'actifs.

II-2 Modèle de Sharpe et Tint [1990]

Ce modèle s'applique aussi à un portefeuille composé de plusieurs classes d'actifs. Le passif est évalué en valeur actuelle et l'actif en valeur de marché. Il est basé sur la notion de surplus S défini comme étant la différence entre la valeur de marché des actifs et la valeur actuelle du passif. Le but du modèle est de minimiser le risque de perte du surplus (mesuré par la variance de la « rentabilité du surplus ») pour un niveau de « rentabilité du surplus » donné et sous un ensemble de contraintes (même démarche que Markowitz [1952]).

La notion de rentabilité du surplus utilisée est la variation du surplus divisée par la valeur initiale des actifs :

$$R_s = \frac{S_1 - S_0}{A_0}$$

L'ensemble des portefeuilles vérifiant ce programme d'optimisation pour différents niveaux de rentabilité espérée, décrit *une hyperbole* dans le repère (risque mesuré par l'écart type de la rentabilité du surplus, rentabilité espérée du surplus). La partie supérieure de cette hyperbole représente *la frontière efficiente*.

Formulation théorique

Le portefeuille est composé de n actifs risqués de rentabilité $R_i (i = 1, \dots, n)$

La stratégie d'investissement est donnée par le choix du portefeuille $x = (x_i)_{i=1, \dots, n}$ avec

$$\sum_{i=1}^n x_i = 1$$

$$R_A(x) = \sum_{i=1}^n x_i R_i : \text{la rentabilité du portefeuille des actifs.}$$

R_p : la variation relative de la valeur actuelle du passif.

rf_0 : ratio de financement initial

$$R_s(x) = R_A(x) - \frac{1}{rf_0} R_p : \text{la rentabilité du surplus.}$$

$V = [cov(R_i, R_j)]$ la matrice de variance covariance supposée régulière,

$\delta = (cov(R_i, R_p))$ le vecteur des covariances entre les actifs et le passif,

$\mu = (E(R_i))$ le vecteur des rentabilités espérées des actifs,

$e = (1, \dots, 1)$: le vecteur unité transposé

x^* : Portefeuille optimal,

x^{min} : Portefeuille à variance minimale,

Les portefeuilles efficients

Le portefeuille x^* solution du problème d'optimisation est donné par la relation suivante :

$$x^* = x^{min} + \lambda z^*$$

λ dépend du niveau de rentabilité du surplus.

$x^{min} = \frac{V^{-1}e}{e'V^{-1}e} + \frac{1}{rf_0} \left(V^{-1}\delta - \frac{e'V^{-1}\delta}{e'V^{-1}e} V^{-1}e \right)$ représente le portefeuille a variance minimale.

$$z^* = V^{-1}\mu - \frac{e'V^{-1}\mu}{e'V^{-1}e} V^{-1}e$$

La frontière efficiente

En remplaçant le portefeuille efficient par son expression, nous obtenons les deux équations suivantes :

$$\begin{aligned} Var(R_s(x^*)) &= \lambda^2 Var(R_A(z^*)) + Var(R_s(x^{min})) \\ E(R_s(x^*)) &= \lambda E(R_A(z^*)) + E(R_s(x^{min})) \end{aligned}$$

En éliminant le λ dans les deux équations, nous obtenons l'équation d'une hyperbole. La partie supérieure de cette courbe représente la frontière efficiente dans le repère (volatilité du surplus, rentabilité attendue du surplus).

Contrainte de déficit

« La probabilité pour que la rentabilité du surplus soit inférieure à un certain seuil (u) ne doit pas dépasser une certaine probabilité (p) ».

Si nous supposons que le vecteur des rentabilités est gaussien, la variable aléatoire rentabilité du surplus suit une loi normale. La contrainte de déficit se traduit par l'équation de droite suivante dans le plan (volatilité du surplus, rentabilité espérée du surplus) :

$$E(R_s) = q_p^{N(0,1)} * \sigma_s + u$$

$q_p^{N(0,1)}$: représente le quantile de la loi normale centrée réduite associé à la probabilité $1-p$.

L'intersection entre la droite de déficit et la frontière efficiente donne le **portefeuille efficient de rentabilité maximale et vérifiant cette contrainte**.

Remarque: Une contrainte de déficit sur le ratio de financement

Le ratio de financement à l'instant final (grandeur inconnue) est donné par la relation suivante :

$$rf_1(x) = rf_0 \frac{1 + R_A(x)}{1 + R_p}$$

La contrainte de déficit est la suivante: «la probabilité pour que le ratio de financement final soit inférieur à 100 % ne doit pas dépasser une certaine probabilité p ».

Cette contrainte se traduit par l'équation de droite suivante :

$$E(R_s) = q_p^{N(0,1)} * \sigma_s + \left(\frac{1}{rf_0} - 1 \right)$$

Les courbes obtenues dans ce cadre (cf. graphique 33) auront ainsi la même allure obtenue dans le graphique précédent (avec le modèle de Kim&Santomero).

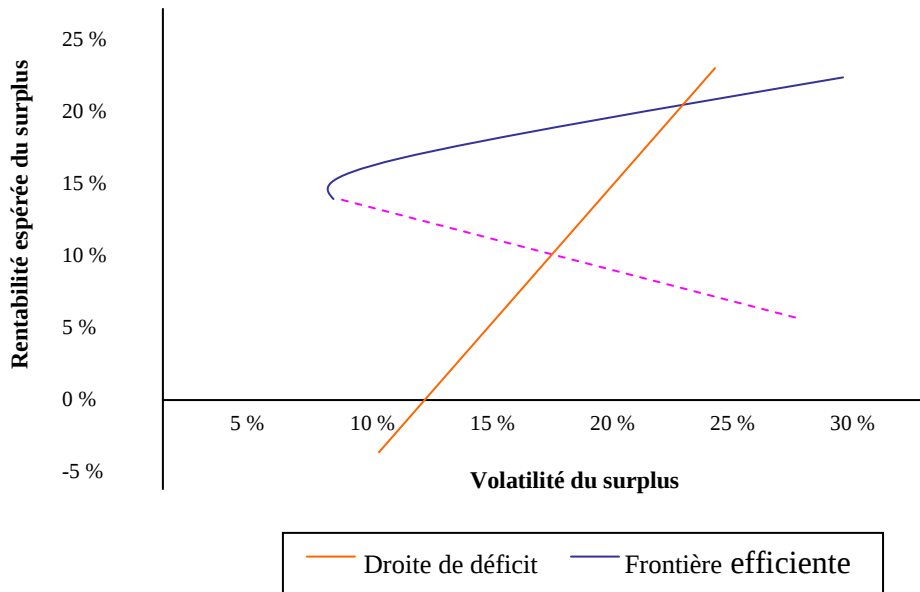


Fig. 33 : *Frontière efficiente et contrainte de déficit (Sharpe & Tint)*

II-3 Modèle de Leibowitz [1992]

Ce modèle s'applique à un portefeuille composé de deux actifs: actions et obligations. Il consiste à déterminer d'une part le pourcentage d'actions et d'autre part la « durée » de la composante obligataire qui permettent de maximiser la « rentabilité du surplus » et donc la rentabilité du portefeuille tout en respectant des contraintes sur l'actif et le passif. Ces contraintes sont modélisées en probabilité « d'insuffisance » et concernent des grandeurs et des rentabilités exprimées en valeur de marché ou valeur actuelle.

a - Hypothèses et notations :

R_A : rentabilité des actions (suit une loi normale).

$\bar{R}_A$: rentabilité attendue des actions.

σ_A : volatilité des actions, écart-type de la rentabilité R_A .

R_O : rentabilité des obligations (suit une loi normale).

$\bar{R}_O$: rentabilité attendue des obligations.

σ_O : volatilité des obligations, variable endogène au modèle.

ρ : corrélation (R_A, R_O)

R_{pf} : rentabilité du portefeuille (suit une loi normale de moyenne $\bar{R}_{pf}$).

$R_{pf} = \alpha R_A + (1 - \alpha) R_O$

α : pourcentage d'actions.

σ_{pf} : volatilité du portefeuille. $\sigma_{pf}^2 = \alpha^2 \sigma_a^2 + (1 - \alpha)^2 \sigma_0^2 + 2\alpha(1 - \alpha)\sigma_A\sigma_O\rho$

R_p : Variation relative des valeurs actuelles du passif.

rf_0 : Ratio de financement initial (valeur de marche initiale de l'actif divisée par la valeur actuelle initiale du passif).

D_0 : duration des obligations.

La volatilité de la composante obligataire est supposée proportionnelle à sa duration.

$$\sigma_0 = D_0 \sigma_{taux1an}$$

$\sigma_{taux1an}$: écart-type du taux d'intérêt pour les emprunts 1 an.

Les obligations de toutes les maturités fournissent le même rendement.

La rentabilité du surplus : mesure de risque actif / passif

La rentabilité du surplus est représentée dans ce modèle par le rapport entre la variation du surplus et la valeur actuelle initiale du passif.

$$R_s = \frac{S_1 - S_0}{P_0}$$

S_0 le surplus initial, S_1 le surplus final et P_0 la valeur actuelle initiale du passif. En décomposant l'expression du surplus, la rentabilité du surplus devient : $R_s = rf_0 * R_{pf} - R_p$

b - La contrainte sur l'actif

« La probabilité pour que la rentabilité du portefeuille d'actifs soit inférieure à un certain seuil ne doit pas dépasser une probabilité donnée »

$$P(R_{pf} < u') < p'$$

La probabilité d'insuffisance p' et le seuil critique u' sont fixés arbitrairement selon la tolérance au risque de l'investisseur. Cette contrainte dans le plan risque/rentabilité espérée est caractérisée par une droite de « déficit ». (Tous les portefeuilles au dessus de cette droite vérifient cette contrainte).

$$\bar{R}_{pf} = u' + \sigma_{pf} * q_{p'}^{N(0,1)}$$

c - La contrainte sur le surplus

« La probabilité pour que la rentabilité du surplus soit inférieure a un certain seuil (u) ne doit pas dépasser une probabilité donnée (p) ».

$$P(R_s < u) < p$$

Nous supposons que R_{pf} et R_p sont normalement distribuées, et que toute combinaison linéaire de ces deux variables aléatoires est aussi normalement distribuée ($R_{pf} \approx N(\bar{R}_{pf}, \sigma_{pf})$ et $R_p \approx N(\bar{R}_p, \sigma_p)$).

Nous avons donc :

$$R_s \approx N(\bar{R}_s, \sigma_s)$$

De la même façon que la contrainte de « déficit » sur l'actif, nous déduisons l'inéquation suivante :

$$u + \sigma_s q_p^{N(0,1)} < \bar{R}_s$$

Démarche théorique :

La frontière de la courbe qui décrit cette contrainte est caractérisée par :

$$u + \sigma_s q_p^{N(0,1)} = \bar{R}_s$$

$$\begin{aligned} \bar{R}_s &= rf * \bar{R}_{pf} - \bar{R}_p \\ \sigma_s^2 &= rf^2 * \sigma_{pf}^2 + \sigma_p^2 - 2 * rf * \sigma_{pf} * \sigma_p * \text{Corr}(R_{pf}, R_p) \end{aligned}$$

L'équation de la frontière devient après substitution des expressions ci-dessus :

$$\begin{aligned} &[rf^2(1-\alpha)^2] \sigma_0^2 + [2rf^2\alpha(1-\alpha)\sigma_A\rho - 2rf(1-\alpha)\sigma_p] \sigma_0 \\ &- \left[\left(\frac{rf(\alpha\bar{R}_A + (1-\alpha)\bar{R}_O) - \bar{R}_p - u}{q_p^{N(0,1)}} \right)^2 - \sigma_p^2 - rf^2\alpha^2\sigma_A^2 + 2rf\alpha\sigma_A\sigma_p\rho \right] = 0 \end{aligned}$$

Les seules inconnues restantes sont le pourcentage d'action α et la volatilité de l'obligation à détenir dans le portefeuille σ_0 .

Pour un α donné, cela revient à résoudre une équation du second degré en σ_0 . La condition de positivité sur le discriminant de cette équation permet d'obtenir un ensemble de valeurs de α à partir desquelles nous calculons les σ_0 correspondants. Grâce aux valeurs de ces deux variables, nous reconstituons un ensemble de portefeuilles de coordonnées $(\sigma_{pf}, \bar{R}_{pf})$ dans le repère risque/rentabilité dont les surplus associés vérifient la contrainte de déficit. La contrainte sur le surplus est représentée par une courbe convexe en forme d'œuf.

d - Portefeuille optimal vérifiant les deux contraintes

La partie supérieure à la droite de déficit et à l'intérieur de « l'œuf » représente tous les portefeuilles qui vérifient les deux contraintes. Le portefeuille intersection de la droite et de « l'œuf » de surplus, de caractéristiques $\bar{R}_{pf}^*$ et σ_{pf}^* , correspond au portefeuille de rentabilité espérée la plus grande qui remplit les deux demandes de « déficit ».

La part optimale d'actions à détenir en portefeuille :

$$\alpha^* = \frac{\bar{R}_{pf}^* - \bar{R}_O}{\bar{R}_A - \bar{R}_O}$$

La volatilité optimale des obligations à détenir en portefeuille :

$$\sigma_O^* = \frac{\sqrt{(\sigma_{pf}^*)^2 + (\alpha^*)^2 \sigma_A^2 (\rho^2 - 1)} - \alpha^* \sigma_A \rho}{1 - \alpha^*}$$

e- L'impact des changements des paramètres

Variation du seuil de rentabilité de l'actif :

Les changements du seuil de rentabilité de l'actif entraînent des changements parallèles de la droite de déficit comme le montre le graphique suivant :

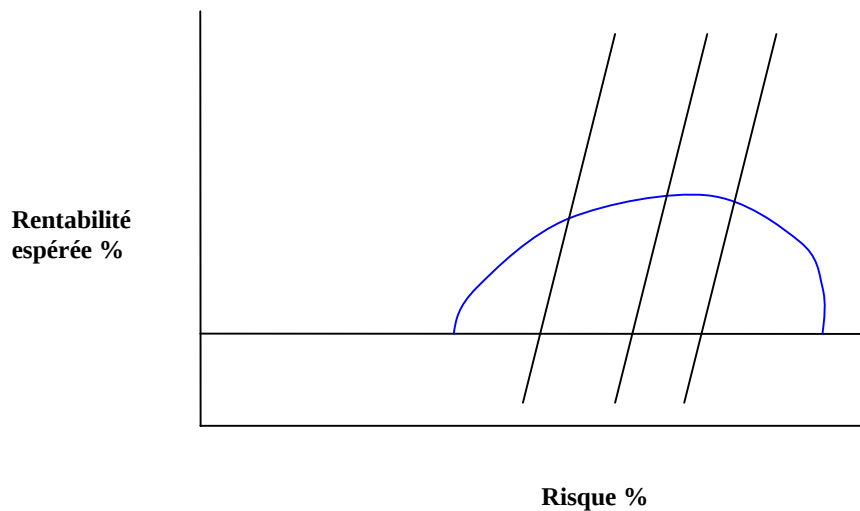


Fig. 34 : Exemple de l'impact de la variation du seuil de rentabilité de l'actif

Variation de la probabilité de déficit de l'actif :

La pente de la droite de «déficit» croît lorsque la probabilité p' diminue, alors que le point d'intersection de cette droite avec l'axe des ordonnées reste fixe (cf. graphique 35).

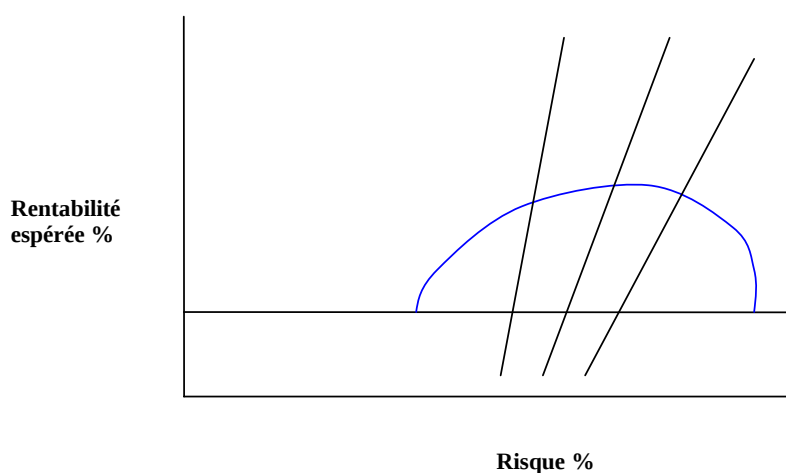


Fig. 35 : Exemple de l'impact de la variation de la probabilité de déficit de l'actif

Variation du seuil de rentabilité du surplus :

Lorsque le seuil de rentabilité du surplus augmente, la taille de l'œuf diminue et la rentabilité du portefeuille optimal aussi (pour une contrainte fixe sur l'actif). Cela est illustré dans le graphique suivant :

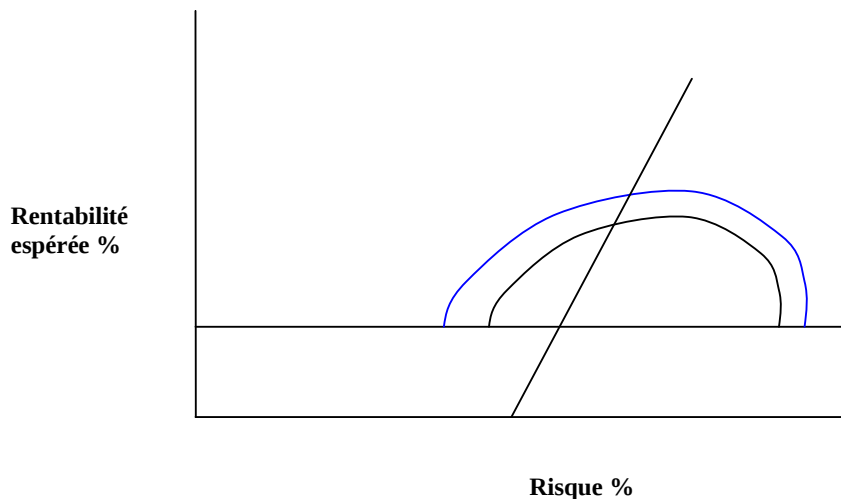


Fig. 36 : Exemple de l'impact de la variation du seuil de rentabilité du surplus

Variation de la probabilité de déficit du surplus :

L'œuf de déficit varie de la même façon que précédemment.

Variation du ratio de financement initial :

Plus le ratio de financement initial diminue, plus l'œuf de déficit rétrécit et se déplace vers la droite (cf. graphique 37).

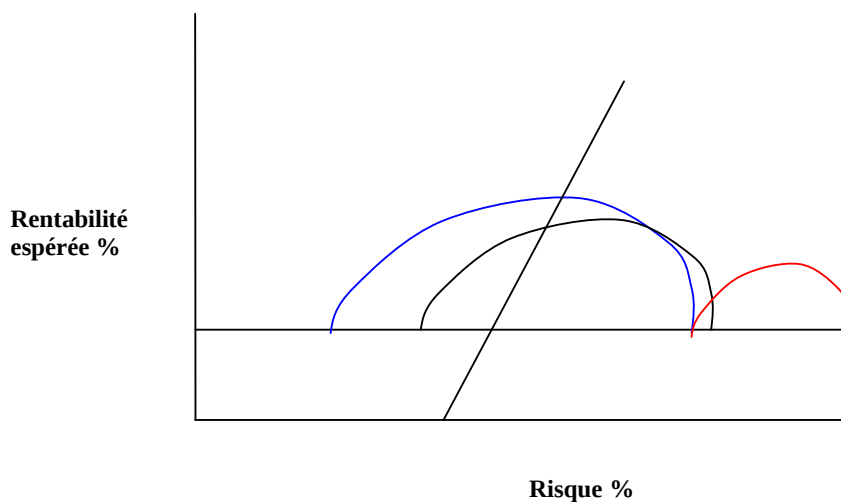


Fig. 37 : Exemple de l'impact de la variation du ratio de financement initial

f- Contrainte sur la rentabilité relative par rapport à un benchmark

« La probabilité pour que la rentabilité du portefeuille d'actifs soit inférieure d'un certain seuil à la rentabilité du benchmark ne doit pas dépasser une certaine probabilité ».

$$P(R_{rel} < u'') = p''$$

R_{rel} : La rentabilité relative, c'est l'écart entre la rentabilité du portefeuille d'actifs et la rentabilité du benchmark. Le benchmark considéré est composé de deux actifs : actions et obligations. Le pourcentage en actions et la durée de la composante obligataire du benchmark sont connus.

Les rentabilités du portefeuille d'actifs et du benchmark sont données par :

$$R_{ptf} = aR_A + (1-a)R_{O_1}$$

$$R_{ben} = bR_A + (1-b)R_{O_2}$$

R_A : La rentabilité de l'action qui est la même pour le portefeuille d'actifs et le benchmark,

R_{O_1} : La rentabilité de la composante obligataire du portefeuille,

R_{O_2} : La rentabilité de la composante obligataire du benchmark,

a : le pourcentage d'actions dans le portefeuille, à déterminer.

b : le pourcentage d'actions dans le benchmark, connu.

La rentabilité relative a pour expression :

$$R_{rel} = (a-b)R_A + (1-a)R_{O_1} + (1-b)R_{O_2}$$

La rentabilité relative espérée est donc :

$$\bar{R}_{rel} = (a-b)\bar{R}_A + (1-a)\bar{R}_{O_1} + (1-b)\bar{R}_{O_2}$$

Nous supposons que les deux obligations ont la même rentabilité espérée ($\bar{R}_{O_1} = \bar{R}_{O_2}$) :

$$\bar{R}_{rel} = (a-b)(\bar{R}_A - \bar{R}_{O_1})$$

C'est donc le produit de « l'excès » d'actions ($a-b$) et de la prime de risque.

La variance de la rentabilité relative est donnée par :

$$\sigma_{rel}^2 = E[(R_{rel} - \bar{R}_{rel})^2]$$

Avec l'hypothèse de normalité, et en posant :

$$\text{corr}(A, O_1) = \text{corr}(A, O_2) = \rho$$

nous obtenons :

$$\sigma_{rel}^2 = (a-b)^2 \sigma_A^2 (1-\rho^2) + [(a-b)\sigma_A \rho + (1-a)\sigma_{O_1} - (1-b)\sigma_{O_2}]^2$$

Cette relation donne une formule usuelle pour la variance de la rentabilité relative et permet de trouver des portefeuilles qui rencontrent la contrainte de déficit.

$$P(R_{rel} < u'') = p''$$

Cette contrainte se traduit par la relation suivante (hypothèse de normalité) :

$$\bar{R}_{rel} = u'' + q_{p''}^{N(0,1)} \sigma_{rel}$$

Les seules inconnues sont le pourcentage d'actions a et la volatilité de l'obligation à détenir dans le portefeuille σ_{O_1} . Pour un a donné, cela revient à résoudre une équation du second

degré en σ_{O_1} . La condition de positivité sur le discriminant du trinôme permet d'obtenir un ensemble de valeurs de « a » à partir duquel nous calculons les σ_{O_1} correspondants. Les valeurs de ces deux variables permettent de reconstituer un ensemble de portefeuilles de coordonnées $(\sigma_{pf}, \bar{R}_{pf})$ dans le repère risque/rentabilité espérée vérifiant la contrainte de déficit. L'ensemble des portefeuilles vérifiant la contrainte sur rentabilité relative décrit une courbe convexe en forme d'œuf comme le montre le graphique suivant :

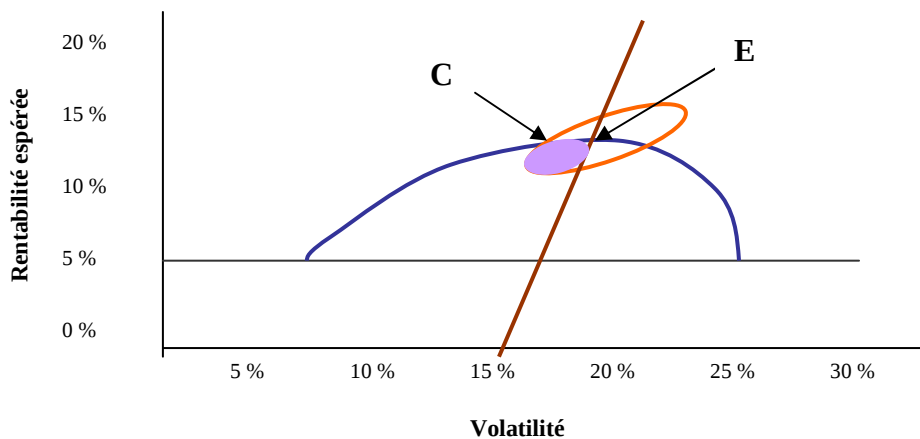


Fig. 38 : Représentation graphique des portefeuilles vérifiant la contrainte sur la rentabilité relative

L'ensemble des portefeuilles optimaux respectant les trois contraintes de Leibowitz se situe dans la zone grise. Les portefeuilles efficients se situent alors sur l'arc [CE]. Le portefeuille E correspond au portefeuille efficient répondant à la recherche d'un rendement optimal. Le point C indique un portefeuille efficient qui a un risque inférieur. Puisque les prévisions de l'investisseur ne sont jamais sûres à 100 %, il est raisonnable de préférer un point compris entre E et C à la place d'un point extrême.

g- Limite du modèle et test de robustesse

Les limites de ce modèle résident dans le fait qu'il est basé sur des hypothèses très fortes :

- la volatilité de la composante obligataire est proportionnelle à sa « durée ».
- les obligations de toutes les maturités fournissent le même rendement.
- les rentabilités sont normalement distribuées.

Ce modèle économique ne respecte pas les contraintes réglementaires, techniques, et comptables. Il surestime donc la capacité de prise de risque. Pour tester la robustesse du modèle, nous faisons varier les hypothèses sur les entrées (corrélations, volatilité et rentabilités attendues) et nous mesurons l'impact de ses variations sur le pourcentage d'actions et la durée de la composante obligataire.

h- Une variante du modèle de Leibowitz

Pour mesurer le risque actif/passif, nous remplaçons la rentabilité du surplus par une mesure plus concrète : la « rentabilité du ratio de financement ». C'est la variation relative du ratio de financement sur la période considérée. Si nous notons rf_0 et rf_1 les ratios de financement initial et final, la « rentabilité du ratio de financement » notée RRF est donnée par :

$$RRF = \frac{rf_1 - rf_0}{rf_0}$$

Cette mesure dépend uniquement de la rentabilité de l'actif et de la variation relative de la valeur actuelle du passif et ne dépend pas du ratio de financement initial.

Nous notons :

A_0 et A_1 les valeurs de marché de l'actif en début et fin de période,

P_0 et P_1 les valeurs actuelles du passif en début et fin de période,

$$R_{pf} = \frac{A_1 - A_0}{A_0} \quad \text{et} \quad R_p = \frac{P_1 - P_0}{P_0}$$

Nous obtenons :

$$RRF = \frac{R_{pf} - R_p}{1 + R_p}$$

Nous pouvons intégrer la rentabilité du ratio de financement dans l'analyse d'allocation d'actif sous forme de contrainte de déficit : « la probabilité pour que la rentabilité du ratio de financement soit inférieure à un certain seuil ne doit pas dépasser une probabilité donnée ». Le seuil de déficit dépend contrairement à la RRF du ratio de financement initial. Nous modifions le modèle de Leibowitz en remplaçant la contrainte sur la rentabilité du surplus par la contrainte sur la RRF .

$$P(RRF < u) < p$$

Cette contrainte est représentée par une courbe convexe.

II-4 Duration du Surplus

Le surplus est la différence entre la valeur de marché de l'actif et la valeur actuelle du passif.

Nous introduisons les notations suivantes :

- A : valeur de marché de l'actif
- P : valeur actuelle du passif
- $S = A - P$: surplus
- $rf = \frac{A}{P}$: le ratio de financement
- r : le taux d'intérêt
- $D_A = -\frac{1}{A} \frac{\Delta A}{\Delta r} (1+r)$ la duration de l'actif
- $D_P = -\frac{1}{P} \frac{\Delta P}{\Delta r} (1+r)$ la duration du passif

Nous définissons la « duration » du surplus de la même façon :

$$D_S = -\frac{1}{S} \frac{\Delta S}{\Delta r} (1+r)$$

En remplaçant S par son expression, nous obtenons la formule suivante :

$$SD_S = AD_A - PD_P$$

En divisant cette expression par le surplus S , nous obtenons :

$$D_S = D_P + \frac{rf}{rf - 1} (D_A - D_P)$$

Certaines remarques peuvent être avancées :

- Un adossement par la duration n'élimine pas le risque de taux sur le surplus. En effet, la formule précédente donne une duration du surplus égale à la duration du passif lorsque le ratio de financement est différent de 100 % : $D_S = D_A = D_P$.
- La duration de l'actif qui permet d'éliminer le risque de taux sur le surplus (duration du surplus nulle) est donnée par : $D_A^* = \frac{1}{rf} D_P$

Pour un ratio de financement supérieur à 100 %, l'actif obtenu est plus court que le passif.

- Lorsque le surplus est nul, la duration du surplus n'est pas définie.
- Nous pouvons représenter graphiquement la duration du surplus en fonction de la différence entre la duration de l'actif et la duration du passif.

La courbe obtenue est une droite qui traduit tous les niveaux de risque de taux sur le surplus en fonction de cette différence.

Pour un ratio de financement supérieur à 100 % (pente de la droite positive), nous obtenons la courbe suivante :

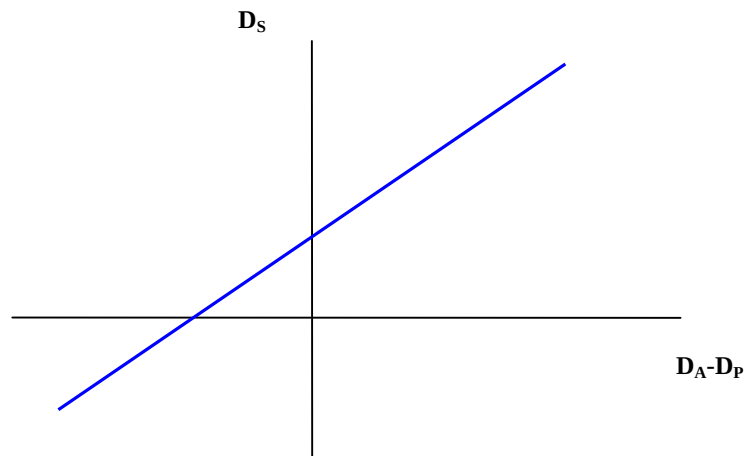


Fig. 39 : Exemple de la variation de la duration du surplus en fonction de l'écart entre la duration de l'actif et celle du passif

Notons que nous pouvons définir de la même façon la «convexité» du surplus notée C_S :

$$C_S = \frac{1}{2S} \frac{\Delta^2 S}{\Delta r^2}$$

Si nous notons respectivement C_A et C_P les convexités de l'actif et du passif, nous obtenons la formule suivante :

$$C_S = C_P + \frac{rf}{rf - 1} (C_A - C_P)$$

Conclusion du chapitre 1

L'aspect mono-périodique qui caractérise les différents modèles classiques de gestion actif-passif (que ce soit celles basées sur les techniques d'immunisation ou celles basées sur la notion du surplus) limite leur capacité à refléter de façon fiable les perspectives d'évolution des différentes variables financières, en particulier sur le long terme. Cela devient d'autant plus compliqué que les variables à considérer sont plus nombreuses et que la nécessité de prendre en compte les corrélations entre les différentes variables est plus importante.

Nous notons comme même que les modèles classiques de la gestion actif-passif, servant en particulier pour l'allocation stratégique d'actif, étaient de loin l'outil de référence sur le marché et ce depuis les années 1950. Les limites que présentent ces modèles ainsi que le développement des supports informatiques au cours des années 1980 ont permis récemment de passer vers d'autres types de modèles : les modèles stochastiques d'ALM. Ces derniers seront l'objet du chapitre suivant.

Chapitre 2 : Allocation stratégique d'actifs et stratégie *Fixed-Mix*

Les modèles stochastiques d'ALM utilisent les techniques de simulation stochastique (*Monte Carlo*) pour modéliser l'évolution des différents éléments, que ce soit au niveau des actifs financiers et des engagements, ou au niveau des variables de marché et des variables démographiques. Nous pourrions ainsi estimer les lois de probabilité associées aux résultats du fonds de retraite sur le long terme. Un passage vers le cadre multi-périodique est effectué. Dans un cadre d'une allocation stratégique stochastique, la revue de certains récents travaux peut être utile.

Frauendorfer [2007] montre comment un critère moyenne-variance peut être appliqué dans un cadre multi-périodique afin d'obtenir les portefeuilles efficients d'une gestion actif-passif. Le modèle d'optimisation tient compte des coûts de transaction et des volatilités stochastiques aussi bien des actifs que du passif. De plus, un outil général permettant la projection des engagements du fonds de pension ainsi que la projection des rendements des actifs est présenté. Dans une étape suivante, la dynamique de la structure par maturité des engagements (*liability maturity structure*) est modélisée comme un indice personnalisé (*customized index*) dont la volatilité et la corrélation avec les rendements de l'actif deviennent une composante intégrale de l'approche de changement de régime appliquée. Les résultats numériques illustrent la diversification des actifs et la relation entre la variabilité du rendement et la dynamique des engagements.

Martellini [2006] considère le problème de sélection de portefeuille dans un cadre intertemporel, en présence de contraintes sur le passif. En utilisant la valeur des engagements comme un numéraire naturel (*natural numeraire*), il trouve que la solution à ce problème induit un théorème de séparation en trois fonds. Cela constitue une justification à certaines pratiques appelées *Liability-Driven-Investment* (LDI), offertes par différentes banques d'investissement et des sociétés de gestion de portefeuille. La LDI se base sur l'investissement en deux types de fonds (en plus de l'actif sans risque) : le portefeuille de performance et le portefeuille de couverture du passif. Autrement, cette stratégie consiste à séparer le portefeuille en deux parties, dont l'une recourt à l'immunisation contingente et dont l'autre cherche le rendement absolu. Pour cette seconde partie le contrôle de risque passe impérativement par les méthodes de tests de stress capables d'évaluer les risques extrêmes.

Waring [2004] fait la revue, tout en les mettant à jour, des différentes techniques de détermination des frontières efficientes de surplus ainsi que l'allocation d'actif du surplus (*surplus asset allocation*). L'actualisation de ces techniques se fait à travers la prise en compte des caractéristiques systématiques (béta) et non systématiques (alpha). Il développe ainsi une vision économique des engagements, en termes d'alpha et de béta. Cela nous donne une mesure des engagements qui est plus intéressante pour résoudre le problème d'allocation d'actif et ce par rapport à ce qui est fourni dans les approches standards. Avec ces outils nous pouvons améliorer le contrôle des risques pour les fonds de pension, à travers le contrôle de la couverture du passif par les actifs. En plus, cette vision favorise, de façon appropriée, le recours à l'alpha et plus généralement à la mesure du risque dynamique provenant de la gestion active.

Munk et al. [2004] présentent la stratégie optimale d'allocation d'actifs pour un investisseur qui peut investir en *cash* (monétaire), en obligations nominales et en actions (ou indices d'actions). Le modèle suppose le retour à la moyenne des rendements des actions et des taux de rendement réels aléatoires. Le modèle de marché est calibré pour le cas des données

historiques des Etats-Unis (actions, obligations et inflation). En plus, afin d'illustrer les allocations d'actifs optimales, Munk et al. [2004] présentent une méthode de calibrage où les paramètres d'aversion au risque et les horizons d'investissement sont estimés de façon à avoir les recommandations optimales pour différents groupes d'investisseurs : « agressif », « modéré » et « conservatif » à différentes horizons d'investissement.

A ce niveau, nous nous proposons de distinguer deux sous-groupes de modèles d'ALM basés sur les techniques stochastiques. L'élément clé de distinction sera si oui ou non les poids des différents actifs reviennent périodiquement à ceux de l'allocation stratégique définie initialement (si oui, les modèles seront appelés modèles à poids constants ou stratégie *Fixed-Mix*).

- Pour le premier sous-groupe de modèles à poids constants et malgré les avancées réalisées avec ces techniques (surtout au niveau de l'implémentation informatique), l'aspect dynamique de l'allocation stratégique reste encore marginalisé. En fait, ces modèles permettent de comparer des allocations constantes dans le temps (statiques) indépendamment des opportunités liées aux évolutions inter-temporelles des marchés (cf. Merton [1990], Kouwenberg [2001], Infanger [2002], Dempster et al. [2003]). L'étude de l'allocation stratégique d'actifs dans le cadre des modèles à poids constants est l'objet de ce chapitre.
- Le deuxième sous-groupe de modèles, et le plus récent, est principalement inspiré de la théorie du choix de la consommation et de portefeuille développée par Merton [1971]. Il s'agit des modèles d'allocation dynamique. Par exemple, à partir de la définition de la fonction objectif pour l'investisseur, ces modèles permettent la détermination d'une trajectoire des poids des différents actifs jusqu'à la date d'échéance (l'ajustement des poids est fonction des évolutions projetées du marché et de la règle de gestion prédéfinie). L'allocation stratégique retenue sera l'allocation optimale d'actifs à la date initiale t_0 . Le cadre d'utilisation de ces modèles récents se heurte au problème d'implémentation vu la complexité des outils mathématiques employés (cf. Hainaut et al. [2005], Hainaut et al. [2007], Rudolf et al. [2004] et Yen et al. [2003]).

Dans ce chapitre, nous considérons une approche d'allocation stratégique d'actifs basée sur la stratégie dite « à poids constants » ou *Fixed-Mix*. Nous proposons une modélisation du régime-type de retraite et étudions certains critères d'allocation stratégique d'actifs en fonction du type du régime (provisionné, partiellement provisionné, etc.) sont étudiés. Nous passons ensuite à l'illustration de la stratégie *Fixed-Mix* avec une application sur les réserves d'un régime de retraite partiellement provisionné. Les résultats obtenus sont discutés et différents tests de sensibilité sont mises en place : ces tests sont liés principalement à l'impact des hypothèses de rendement ou de corrélation retenues (le graphique 40 illustre les étapes générales nécessaires pour l'aboutissement d'un processus d'ALM stochastique).

Les difficultés rencontrées lors de la mise en place de la stratégie *Fixed-Mix*, dues essentiellement à la multiplicité du nombre de classes d'actifs considérées et au nombre de scénarios économiques simulés, nous ont mené à nous pencher sur les aspects d'optimisation numérique. Le point de départ est notre constat d'un temps de calcul significatif dû

simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Dans ce cadre, une présentation détaillée des travaux menés lors de la rédaction de l'article de Rullière et al. [2010] est effectuée. Un algorithme d'optimisation globale d'une fonction non convexe et bruitée est présenté. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement et compare les résultats obtenus avec l'algorithme de Kiefer-Wolfowitz- Blum (cf. Blum [1954], Kiefer et Wolfowitz [1952]).

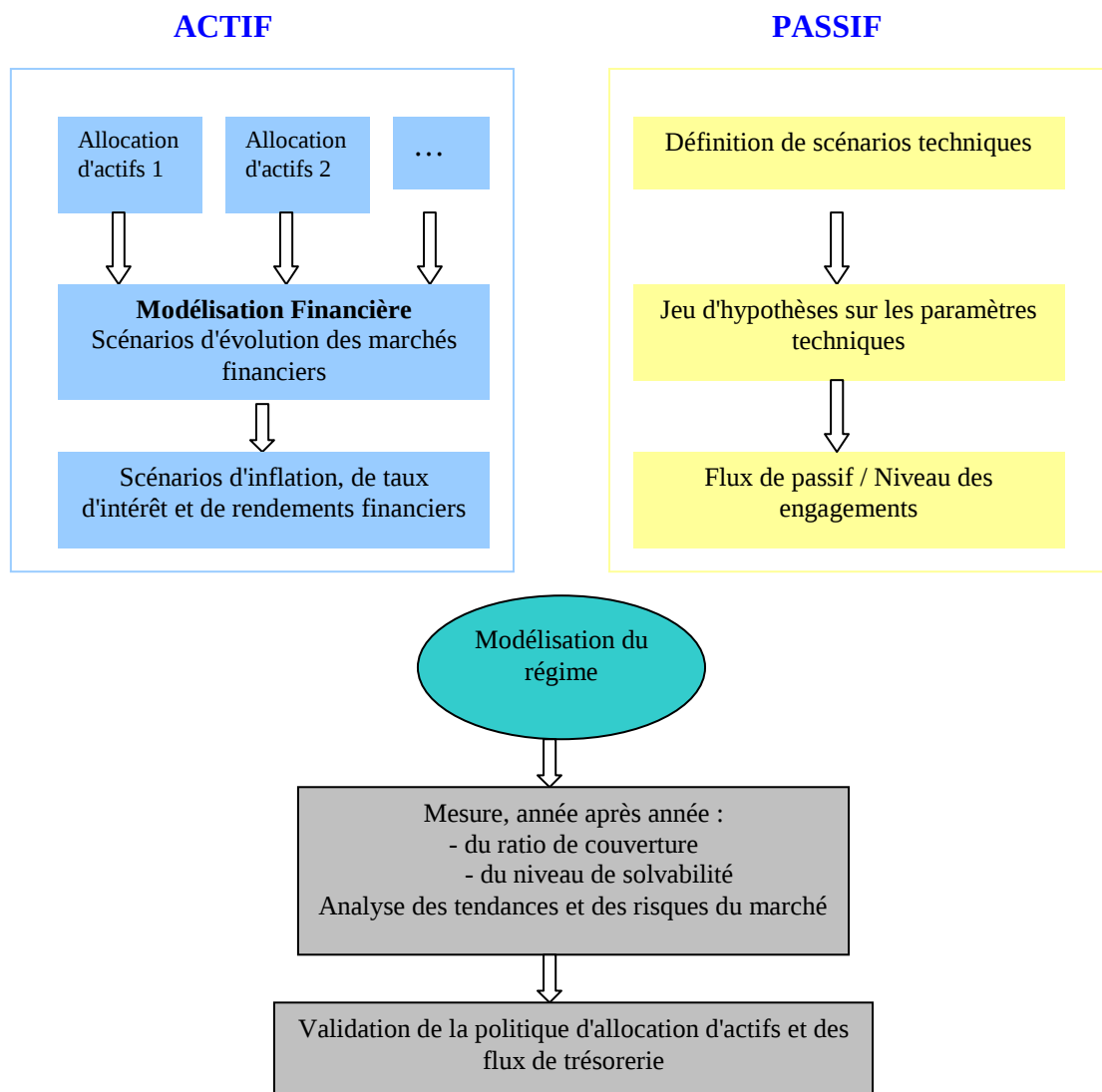


Fig. 40 : Illustration du modèle de gestion actif-passif dans le cadre stochastique

I- Etude de l'allocation d'actifs dans le cadre de la stratégie *Fixed-Mix*

I-1 Présentation

I-2-6 Résultats obtenus

Dans le cadre de l'application de la stratégie *Fixed-Mix*, nous optons pour la majoration des poids des différentes classes d'actifs dans l'allocation optimale de la manière suivante : pondération maximale de 40 % pour la classe des obligations à taux fixe, la classe des actions, l'immobilier et les obligations indexées et pondération maximale de 10 % pour le monétaire. L'intérêt de ce plafonnement des poids des classes d'actifs est d'ordre pratique : il permet de réduire le nombre d'allocations possibles, et donc de minorer les temps de calculs du moteur ALM.

L'objectif étant d'illustrer le caractère opérationnel du moteur ALM, la matrice d'allocation est générée en fixant un pas de maillage (de 5 %) et des bornes pour les poids de chaque actif. Les bornes ont été choisies en prenant en compte deux facteurs : le premier consiste à limiter la taille de la matrice des allocations possibles et le second consiste à couvrir un ensemble d'allocations pertinentes et vraisemblables. Les bornes retenues pour l'illustration des résultats sont donc les suivantes :

	Borne min	Borne sup
Monétaire	0,00 %	30,00 %
ZC 5 ans (nominal)	10,00 %	30,00 %
ZC 10 ans (nominal)	10,00 %	30,00 %
ZC 15 ans (nominal)	0,00 %	0,00 %
ZC 8 ans (réel)	10,00 %	30,00 %
ZC 15 ans (réel)	0,00 %	0,00 %
Obligation crédit	0,00 %	30,00 %
Actions	25,00 %	40,00 %
Immobilier	0,00 %	10,00 %

Tab. 29 : Bornes retenues pour les poids de différentes classes d'actifs

D'un autre côté, deux scénarios économiques ont été considérés : un scénario de stress et un scénario alternatif. En pratique, le but de ce type de scénario est de tester la robustesse des allocations optimales obtenues avec le scénario central (présenté dans la section III du chapitre 2 de la partie I). Ci-après les paramètres relatifs à un scénario de stress (avec forte baisse des rendements des actions : chute de 30 % sur un an) et au scénario alternatif (baisse continue de l'inflation moyenne de 100 pb). Il est rappelé que les paramètres utilisés pour les différents scénarios sont purement illustratifs.

Paramètre	Scénario central	Scénario stress (forte baisse des rendements des actions -30%)	Scénario alternatif (baisse de l'inflation moyenne de 100 bp)
Inflation à long terme			
Vitesse de retour à la moyenne	0,01%	0,01%	0,01%
Volatilité	5,00%	5,00%	5,00%
Moyenne	2,00%	2,00%	1,00%
Valeur initiale	2,00%	2,00%	2,00%
Inflation à court terme			
Vitesse de retour à la moyenne	10,00%	10,00%	10,00%
Volatilité	1,00%	1,00%	1,00%
Valeur initiale	1,00%	1,00%	1,00%
Borne inférieur	-1,00%	-1,00%	-1,00%
Borne supérieur	3,00%	3,00%	3,00%
Taux d'intérêt réel long			
Vitesse de retour à la moyenne	35,00%	35,00%	35,00%
Volatilité	2,50%	2,50%	2,50%
Moyenne	3,00%	3,00%	3,00%
Valeur initiale	2,50%	1,50%	2,50%
Taux d'intérêt réel court			
Vitesse de retour à la moyenne	40,00%	40,00%	40,00%
Volatilité	2,50%	2,50%	2,50%
Valeur initiale	1,00%	1,00%	1,00%
Borne inférieur pour les taux	-1,00%	-1,00%	-1,00%
Borne supérieur pour les taux	4,00%	4,00%	4,00%
Action			
Valeur initiale	100,00%	100,00%	100,00%
Moyenne regime 1 (faible volatilité)	7,50%	chute 30% sur un an	6,50%
Volatilité du régime 1 (faible volatilité)	16,00%	16,00%	16,00%
Moyenne régime 2 (forte volatilité)	-10,00%	chute 30% sur un an	-11,00%
Volatilité du régime 2 (forte volatilité)	27,00%	27,00%	27,00%
Immobilier			
Vitesse de retour à la moyenne	15,00%	15,00%	15,00%
Volatilité	10,00%	11,00%	10,00%
Moyenne	6,00%	6,00%	5,00%
Valeur initiale	5,00%	5,00%	5,00%
Crédit			
<i>Coefficient d'battement*</i>			
<i>Maturité 5 ans</i>	97,18%	97,18%	97,18%
<i>Maturité 4 ans et 9 mois</i>	97,17%	97,17%	97,17%
<i>Valeur initiale</i>	-	-	-
<i>Moyenne regime 1 (faible volatilité)</i>	-	-	-
<i>Volatilité du régime 1 (faible volatilité)</i>	-	-	-
<i>Moyenne régime 2 (forte volatilité)</i>	-	-	-
<i>Volatilité du régime 2 (forte volatilité)</i>	-	-	-
Monétaire (lognormal)			
Moyenne	3,50%	3,50%	2,50%
Volatilité	1,50%	1,50%	1,50%

Tab.30 : Paramètres retenus pour les différents scénarios

Il est à préciser que le TRI et le médian de la réserve sont calculés sur un horizon égal à l'horizon d'allocation (9 ans). Par ailleurs, la VaR du ratio de viabilité et la VaR du ratio de solvabilité sont établis pour un horizon d'allocation de 9 ans (sachant par ailleurs que l'horizon de la contrainte de solvabilité est de 20 ans et que l'horizon de la contrainte de viabilité est de 30 ans). La performance et la volatilité sont des valeurs moyennes calculées sur l'ensemble de la projection (50 ans).

Au niveau des résultats sur l'allocation optimale, nous observons à première vue que le poids du monétaire est nul. Aussi, nous observons que les poids cumulé des obligations nominales et l'obligation crédit est de 40 %. Cette valeur est de 10 % pour les obligations indexées à l'inflation et de 50 % pour les actifs risqués (actions et immobilier). En outre l'allocation optimale présentée ci-dessous ne comporte pas d'obligation nominale et d'obligation indexée à l'inflation de maturité 15 ans, étant donné que ces poids sont nuls dans la matrice des allocations testées.

Le tableau ci-après présente les dix meilleures allocations parmi les 2 985 allocations testées avec le scénario économique central, la meilleure allocation est marquée en bleu :

TRI moyen (9 ans)	10,631%	10,630%	10,629%	10,630%	10,629%	10,631%	10,631%	10,632%	10,630%	10,630%
Médian de la réserve (9 ans)	10133	10137	10124	10125	10136	10126	10138	10138	10125	10133
VaR Ratio de pérennité (9 ans)	1,06	1,06	1,06	1,06	1,06	1,06	1,06	1,06	1,06	1,06
VaR Ratio de financement (9 ans)	2,786	2,788	2,792	2,792	2,787	2,792	2,788	2,792	2,792	2,786
Performance moyenne à long terme (50 ans)	5,551%	5,553%	5,545%	5,546%	5,552%	5,548%	5,553%	5,549%	5,547%	5,550%
Volatilité à long terme (50 ans)	7,396%	7,414%	7,385%	7,370%	7,382%	7,373%	7,381%	7,375%	7,369%	7,370%
Monétaire	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
ZC 5 ans (nominal)	10%	15%	30%	25%	20%	15%	10%	10%	20%	15%
ZC 10 ans (nominal)	15%	20%	10%	10%	20%	10%	20%	10%	10%	15%
ZC 15 ans (nominal)	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
ZC 8 ans (réel)	10%	10%	10%	10%	10%	10%	10%	10%	10%	10%
ZC 15 ans (réel)	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
Obligation crédit	15%	5%	0%	5%	0%	15%	10%	20%	10%	10%
Actions	40%	40%	40%	40%	40%	40%	40%	40%	40%	40%
Immobilier	10%	10%	10%	10%	10%	10%	10%	10%	10%	10%

Tab.31 : Liste des dix meilleures allocations obtenues dans l'exemple illustratif

Nous notons que le poids du monétaire est nul dans les 10 allocations. Nous observons également sur l'ensemble des allocations une répartition équilibrée entre les actifs risqués (actions et immobilier), qui ont un poids total de 50 %, et les actifs moins risqués (obligations), qui ont également un poids total de 50 %.

En complément, un test est mené sur les 10 meilleures allocations obtenues. Le test sur ces dix allocations est réalisé pour s'assurer de leur robustesse en cas de scénarios de stress (forte baisse des actions sur un an) ou alternatif (diminution du niveau d'inflation de long terme) : il s'agit ainsi de s'assurer que ces 10 meilleures allocations vérifient les contraintes dans les deux scénarios complémentaires retenus (scénario de stress et scénario alternatif). En pratique, cela consiste à lancer la partie du code relatif aux contraintes considérées en retenant en *input* la matrice des 10 meilleures allocations.

Dans le cadre de ce test, il s'avère que les 10 allocations présentées comme les meilleures dans le scénario central vérifient bien les contraintes dans les deux scénarios complémentaires. Aussi, les allocations optimales proposées semblent résister à des scénarios

de projections des actifs différents de ceux retenus dans le scénario central. La conclusion de ce test confirme le caractère robuste des dix allocations testées.

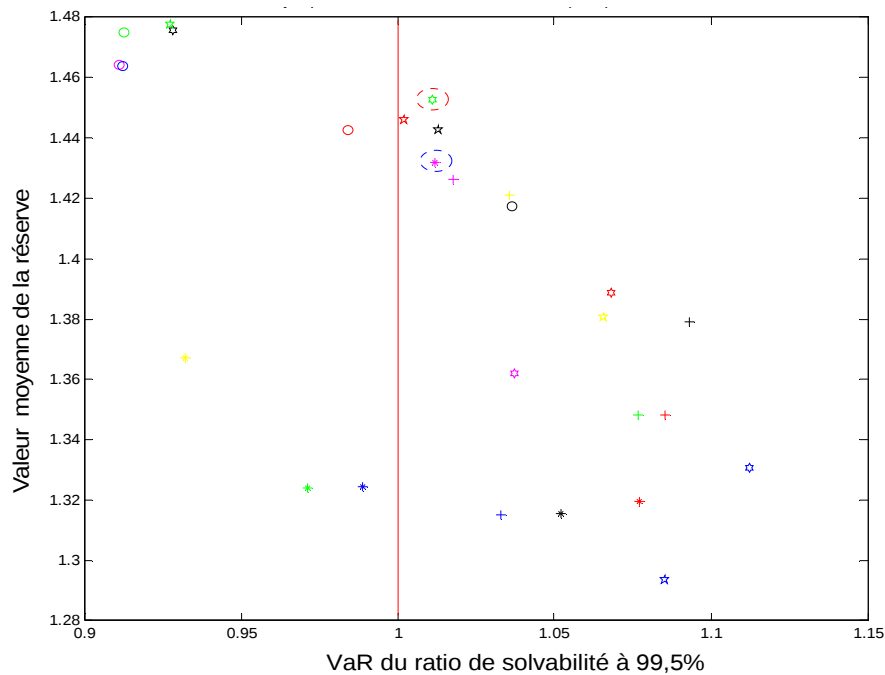


Fig. 44 : *Ratio de couverture des flux jusqu'en 2029 sur un horizon de 9 ans (2018) et un niveau de confiance de 97,5 %*

Les résultats suivants sont indiqués à titre d'illustration : Valeur à Risque (VaR) du ratio de solvabilité à 99,5 % et valeur moyenne de la réserve dans le graphique 44, ceci pour un échantillon donné d'allocation. On obtient alors une courbe d'efficacité au sens de la théorie de Markowitz [1952] : un certaines allocations permettent d'avoir plus de rendement (valeur moyenne de la réserve) pour un niveau de risque plus réduit (VaR du ratio à 99,5 %). Il s'agit du cas, par exemple, des deux allocations ayant le couple de rendement/risque encerclé dans le graphique 44 et qui garantissent un niveau de VaR supérieur à 1.

Le graphique 45 illustre, quant à lui, un test effectué sur la sensibilité de l'allocation optimale par rapport aux changements des niveaux des rendements espérés des actions.

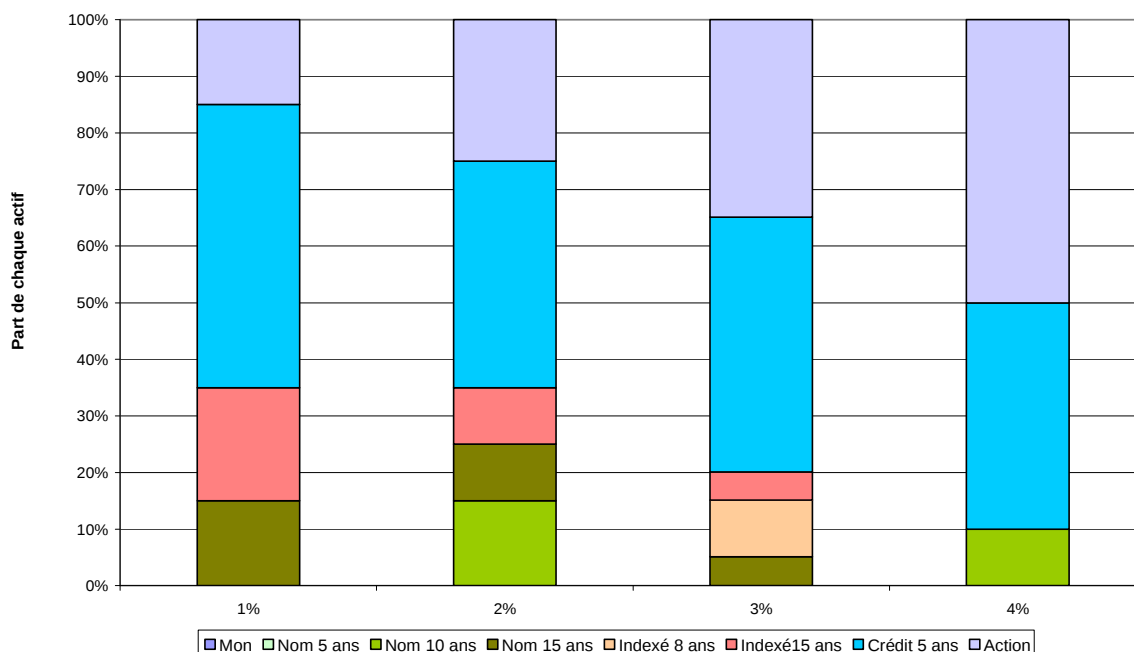


Fig. 45 : Sensibilité de l'allocation optimale par rapport à l'espérance de rendement des actions

Deux constats peuvent être avancés à ce titre :

- une sensibilité significative de l'allocation stratégique d'actifs par rapport aux hypothèses de départ sur les rendements des actions.
- une sensibilité croissante : plus les hypothèses de rendement sont élevées plus la variation des poids de l'allocation stratégique optimale est significative dans notre cas.

Dans ce qui suit nous nous intéressons à l'impact du recours à deux matrices de corrélation différentes (cf. section I, chapitre 2 de la partie I) entre les actions et les autres classes d'actifs : une première matrice reflétant la structure dépendance dans un contexte « normal » et une deuxième matrice reflétant une situation de « crise » (appelée matrice de « corrélation à risque »).

Pour cela, un premier tableau (cf. tableau 32) est présenté indiquant l'allocation optimale obtenue dans le cas d'un seul régime (1 matrice de corrélation) et dans le cas de deux régimes (2 matrices de corrélation), toutes choses étant égales par ailleurs. Nous constatons que, dans ce dernier cas, nous sommes moins exposés aux classes d'actifs risqués (exposition totale aux actions et à l'immobilier à la hauteur de 65 % contre 70 % dans le cas d'un seul régime. Cela constitue une réponse à une aversion au risque supplémentaire prise en compte via la matrice de « corrélation à risque »).

	Alloc optim 1 matrice de corr	Alloc optim 2 matrice de corr
Nom 5 ans	0%	0%
Nom 10 ans	10%	15%
Nom 15 ans	5%	10%
Indexées 8 ans	0%	0%
Indexées 15 ans	15%	10%
Crédit 5 ans	40%	40%
Actions	30%	25%

Tab. 32 : *Allocations optimales en fonction du nombre de régimes de corrélations*

Le même travail a été effectué au niveau des VaR (Valeur à risque) correspondant au montant de la réserve totale observée à une date donnée (horizon de l'allocation). Ainsi, pour deux niveaux de confiance possibles (97,5 % et 99,5 %), nous avons calculé ces valeurs dans le cas où nous avons un seul régime de corrélation et dans le cas où nous avons deux régimes de corrélation (cf. tableau 33). Les résultats confirment une moindre exposition dans le cas de deux régimes et montrent une sensibilité significative du niveau de la VaR (en montant) par rapport à l'approche retenue (dans le cas du seuil de 97,5 %, nous remarquons bien un doublement du montant extrême en passant de l'approche à un seul régime à l'approche à deux régimes de corrélation).

	VaR 97,5 %	VaR 99,5 %
Alloc optim 1 matrice de corr	203	-3740
Alloc optim 2 matrice de corr	435	-3255

Tab.33 : *Valeur à Risque de la valeur de la réserve en fonction du nombre de régimes de corrélations*

L'histogramme ci-dessous confirme également ce constat via des queues extrêmes relativement plus épaisses (et plus étendues vers la gauche) dans le cas d'un seul régime par rapport à ceux observés dans le cas de deux régimes.

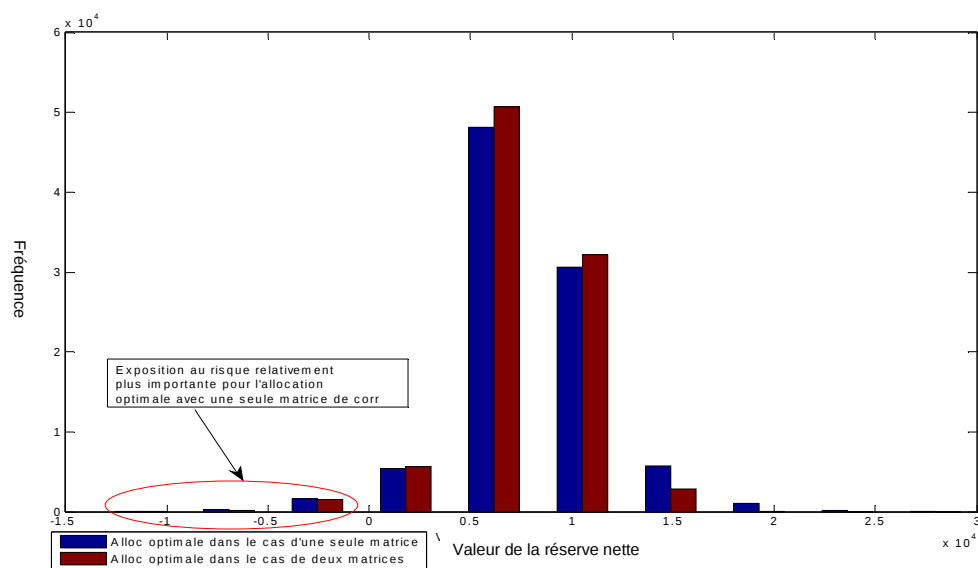


Fig. 46 : Distribution de la valeur de la réserve totale (nette des flux de trésorerie) dans le cas de projection de scénarios avec deux matrices de corrélation

II- Proposition de méthodes numériques

La résolution du problème posé peut être présentée en trois étapes. La première consiste à générer des trajectoires pour chaque classe d'actifs. La deuxième consiste à projeter chaque allocation en fonction de la règle de gestion du dispositif et à en retenir une, qui conduit à une trajectoire de rendement de l'actif synthétique. La troisième étape vise à déterminer la valeur de la fonction objective (le passif étant déterministe). Les allocations susceptibles d'être retenues seront alors celles qui maximisent la fonction objective et respectent les contraintes du dispositif.

Ce paragraphe présente les méthodes de mise en œuvre opérationnelle de l'allocation stratégique.

Le nombre d'allocations possibles avec les neuf classes (notamment lorsque le pas de l'allocation est égal à 1 %) risqueraient d'engendrer des temps de calcul importants, ce qui limiteraient les performances opérationnelles du modèle. Dans ce contexte, la réduction du nombre de classes ou l'augmentation du pas dans le choix des allocations possibles améliorerait les performances opérationnelles du modèle. Néanmoins, la réduction du nombre de classes d'actifs ne peut aller au-delà d'un minimum de 2 classes d'actifs sur les taux nominaux (hors monétaire) et les taux réels, à savoir une maturité moyen terme et une maturité long terme. Cela se justifie par plusieurs raisons :

- pouvoir tester dans l'allocation, au moins l'échéance de réinvestissement et de « refixation » du rendement en regard de la couverture à l'inflation ;
- les engagements des régimes pouvant être très longs, il convient de tenir compte un minimum du positionnement sur la courbe des taux.

Quelle que soit la méthode retenue, le nombre d'allocations possibles nécessite de recourir à des algorithmes d'optimisation efficaces, l'approche consistant à tester de manière exhaustive l'ensemble des allocations possibles n'étant pas envisageable.

Cette partie traite donc un problème d'optimisation dans le cas où la fonction objectif est estimée à l'aide de simulations. Il présente un algorithme d'optimisation globale d'une fonction non convexe et bruitée. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés. Une application numérique illustre la conformité du comportement de cet algorithme à celui prévu théoriquement. Les performances de l'algorithme sont analysées au moyen de différents critères de convergence. Des zones de confiance sont également proposées. Enfin, une comparaison avec un algorithme classique d'optimisation stochastique est menée.

II-1 Introduction

II-1-1 Le problème d'optimisation

Nous allons chercher à déterminer l'optimum global d'une fonction réelle définie sur un ensemble Θ soumise à un bruit, ainsi que les paramètres de Θ conduisant à cet optimum. Ainsi, dans le contexte de l'allocation stratégique optimale d'actifs, une allocation

particulière, $\theta \in \mathfrak{R}^d$, $d \in \mathbb{N}^*$ peut nécessiter de simuler de nombreuses trajectoires de plusieurs actifs, sur une succession de périodes, afin de déterminer un unique indicateur de risque ou de gain, bruité, $F(\theta)$: l'indicateur en question peut être un ratio de financement, un indicateur de gain pénalisé en fonction du risque, ou un indicateur synthétique de compromis risque/gain. Le calcul de chaque $F(\theta)$ est coûteux en terme de temps de calcul, et les allocations conduisant à un optimum de $f(\theta) = E[F(\theta)]$ sont recherchées pour l'allocation optimale.

Nous considérons ici une fonction objectif réelle $f(\theta)$ d'un paramètre $\theta \in \Theta$, $\Theta \subset \mathbb{R}^d$, $d \in \mathbb{N}^*$:

$$f : \Theta \rightarrow \mathbb{R}, \quad \Theta \subset \mathbb{R}^d$$

Nous supposons que f est continue et bornée, non nécessairement dérivable, et que Θ est une union finie de d -simplexes, typiquement un ensemble de pourcentages d'allocation possibles, inclus dans $[0,1]^d$. Enfin, f n'est pas nécessairement une fonction convexe de θ , de sorte que la fonction peut posséder plusieurs optima locaux.

En outre, nous supposons que f n'est pas directement connue, mais qu'elle est estimée à l'aide de simulations. En tout point $\theta \in \mathfrak{R}^d$, du fait des erreurs d'estimation, l'observateur ne peut accéder qu'à des réalisations d'une variable aléatoire $F(\theta) = f(\theta) + \varepsilon(\theta)$, où $\varepsilon(\theta)$ représente un bruit d'espérance nulle, dont nous postulons l'existence d'une variance finie, non nécessairement homogène en θ . Nous postulons que les $\{\varepsilon(\theta)\}_{\theta \in \Theta}$ sont mutuellement indépendants. La fonction bruitée F est donc telle que :

$$\begin{aligned} \forall \theta \in \Theta : \\ \begin{cases} E[F(\theta)] = f(\theta) \\ V[F(\theta)] < \infty \end{cases} \end{aligned}$$

Nous nous placerons dans le cadre d'une minimisation. Sous ces hypothèses, chercher à optimiser la fonction bornée $f(\theta) = E[F(\theta)]$ revient alors à rechercher, à partir de réalisations ponctuelles indépendantes de la fonction bruitée F :

1. L'unique valeur minimale de la fonction objectif f ,

$$m^* = \inf_{\theta \in \Theta} E[F(\theta)]$$

2. L'ensemble des paramètres conduisant à une valeur proche de m^* :

$$S_x = \{\theta \in \Theta, E[F(\theta)] \leq x\}$$

pour tout réel x donné dans un voisinage de m^*

Le premier point est intéressant, mais s'agissant d'allocation d'actifs, le résultat le plus utile est naturellement le second. En un mot, nous cherchons tous les paramètres θ conduisant à un $f(\theta)$ proche de l'optimum.

En présence d'une incertitude sur l'estimateur de f , la recherche d'un paramètre θ^* conduisant à un optimum global supposé de f nécessite l'exploration de tous les paramètres susceptibles de conduire à un optimum global inférieur. Il sera ainsi nécessaire de connaître l'ensemble S_x des paramètres conduisant à une valeur estimée de f proche de m^* .

En outre, en présence de plusieurs paramètres solutions, seul l'utilisateur de l'algorithme peut décider lequel privilégier. C'est la raison pour laquelle nous rechercherons l'ensemble des paramètres solution, et non un seul point de cet ensemble.

Lorsque la fonction F est déterministe, dans le cas où $f(\theta) = F(\theta)$, différentes méthodes d'optimisation peuvent être proposées, comme les méthodes de descente de gradient, de Newton-Raphson, de Hooke et Jeeves, la méthode de Nelder, Mead [1965], ou des méthodes spécifiquement adaptées à certaines formes de f , comme lorsque f est convexe. Ces méthodes ne garantissent toutefois pas que l'optimum obtenu est un optimum global. S'agissant d'optimisation globale déterministe, les premières recherches datent d'environ trente ans et sont attribuées à Hansen [1979]. L'optimisation Lipschitzienne, l'algorithme de Schubert (cf. Schubert [1972]) ou l'algorithme DIRECT (cf. Jones et al. [1993]) sont des méthodes largement utilisées dans un cadre déterministe. D'une façon plus générale, ces algorithmes peuvent s'intégrer dans le cadre d'algorithmes de type *Branch and Bound* (cf. Lawler et al. [1966]), où la zone à explorer (ici Θ) est partitionnée en plusieurs zones (*branching*), dont certaines sont exclues de l'analyse selon certains critères (*bounding*).

Les critères d'exclusions de ce type d'algorithme reposent notamment sur des propriétés de la fonction objectif et sur une arithmétique d'intervalle (cf. Wolfe [1996]), nous parlons alors d'optimisation globale déterministe, permettant de garantir avec certitude l'absence d'optimum sur les zones exclues. En considérant que f est le résultat d'une expérience, généralement coûteuse en terme de temps de calcul, les techniques d'optimisation entrent dans le cadre de *Computer Experiments* et la représentation de la fonction en dehors des points d'observation peut s'appuyer sur différentes techniques de régression, de Krigage, ou de champs gaussien (cf. Krige [1951], Jones et al. [1998], Santner et al. [2003]). Une revue des différentes techniques utilisées dans le champ de l'optimisation globale pourra être trouvée dans Horst, Pardalos [1995], ainsi que dans plusieurs thèses récentes (cf. Emmerich [2005], Ginsbourger [2009] et Villemonteix [2009]).

Lorsque la fonction F est aléatoire, lors de la recherche d'un unique optimum non nécessairement global, une très large littérature existe sur les algorithmes stochastiques pour la recherche d'optimum. Lorsque $d = 1$ des algorithmes tels que celui de Kiefer-Wolfowitz (voir Kiefer, Wolfowitz [1952]) peuvent être utilisés. Dans le cas $d > 1$, une extension due à Blum peut être exploitée (voir Blum [1954]). Strugarek [2006] présente un ensemble d'algorithmes plus récents et plus détaillés portant sur l'optimisation stochastique. S'agissant de la recherche de racines, que nous pourrions également traiter avec l'algorithme ici présenté, des techniques classiques sont celles dérivées de l'algorithme de Robbins-Monro (Robbins, Monro [1951]). De même, l'algorithme de Cohen-Culioli permet de résoudre des problèmes très proches avec une contrainte de contrôle de la probabilité de ruine (voir Cohen et al. [1994]). Nous ne détaillerons pas ici l'ensemble des techniques utilisées dans les champs de l'optimisation stochastique (recherche d'optimum), ou de l'approximation stochastique (recherche de racines).

Enfin, lorsque la fonction F est aléatoire, et lorsque nous cherchons un minimum global d'une fonction non nécessairement convexe, le bruit complique encore le problème difficile de l'optimisation globale (voir Bulger et al. [2005] pour une discussion sur ce sujet). Certaines méthodes stochastiques ont été adaptées à un environnement bruité : nous pouvons citer les méthodes génétiques (voir Alliot [1996], Mathias et al. [1996] en présence de bruit), ainsi que les méthodes de recuit simulé (voir Aarts et al. [1985], Branke et al. [2008] en présence de bruit). Encore peu connues dans le domaine de l'actuariat, quelques extensions

d'algorithmes de type *Branch and Bound* sont également proposées dans un cadre stochastique (cf. par exemple Norkin et al. [1996]).

Le risque de se tromper d'optimum Lorsque F est aléatoire, nous pouvons envisager l'usage d'algorithmes stochastiques d'optimisation locale. Toutefois, tout comme les méthodes déterministes de descente de gradient, ces algorithmes requièrent de choisir un point de départ suffisamment proche d'un optimum global, et d'éviter les plus évidents optima locaux. Sauf pour quelques fonctions f particulières, il s'agit alors d'appréhender dans un premier temps la forme de la fonction f . Une exploration globale de la fonction s'avère donc souvent nécessaire, le risque étant moins de mal estimer la valeur d'un optimum global, que de se tromper d'optimum en choisissant indûment un mauvais optimum local.

Une seconde piste serait d'établir un grand nombre de tirages de F , afin de se ramener, à une marge d'erreur près, au cas déterministe. Cela autoriserait l'usage d'algorithmes d'optimisation locale déterministe, généralement rapides. Toutefois, la recherche d'un optimum local se ferait au prix d'une exploration préalable de la fonction. De surcroît, l'algorithme utilisé serait lourdement pénalisé par le grand nombre de tirages requis pour F .

De la même façon, cette pénalité frapperait les algorithmes d'optimisation globale déterministe. Enfin, les méthodes génétiques ou le recuit simulé ne visent pas directement à proposer des zones de confiance pour les optimiseurs de la fonction objectif. Il nous a semblé délicat de quantifier, avec ces méthodes, le risque de proposer un optimum local et d'ignorer un optimum global meilleur.

Explorer ou connaître A titre illustratif, plaçons nous un instant dans un cadre simplifié, en présence d'un unique minimiseur $\theta^* = \arg \min_{\theta \in \Theta} f(\theta)$. Supposons que l'observateur puisse réaliser n tirages de F en chaque point d'un ensemble $\Theta_m = \{\theta_1, \dots, \theta_m\}$. En supposant donnée la valeur n , qui détermine la précision de la connaissance de f , et Θ_m qui détermine l'ampleur de l'exploration de f , un estimateur envisageable $\hat{\theta}^*$ de l'unique minimiseur θ^* est le suivant :

$$\hat{f}_n(\theta) = \frac{1}{n} \sum_{i=1}^n F_i(\theta) \quad \theta \in \Theta_m$$

$$\hat{\theta}_{n,m}^* = \arg \min_{\theta \in \Theta_m} \hat{f}_n(\theta)$$

La question se pose ici de l'arbitrage entre le choix d'un n élevé ou d'un cardinal de Θ_m élevé. Les algorithmes usuels d'optimisation supposent généralement que l'ensemble des points $F(\theta)$ sont connus, et ignorent, à notre connaissance, la question pratique de l'optimisation du couple de paramètre (n, m) . En effet, la détermination d'une réalisation de la variable aléatoire F peut prendre un temps de calcul important. En conséquence, n réalisations de F en chaque point de Θ_m conduisent à un total de $(n.m)$ tirages de F . Or, les tirages étant coûteux en termes de temps de calcul, le nombre de tirages est nécessairement limité. Pour un budget de tirages fixé, le choix d'un nombre de simulations n élevé réduit la variance de $\hat{f}_n$, et donc améliore ponctuellement la connaissance de f , mais limite également l'exploration de la fonction en d'autres points.

Cette illustration simplifiée introduit la problématique de l'arbitrage entre l'exploration de la fonction en différents points et la connaissance ponctuelle de celle-ci. Trouver les minimiseurs de f va requérir d'une part d'explorer la fonction en différents points (nous parlerons d'exploration), et d'autre part d'opérer également plusieurs tirages de F en chaque point exploré pour obtenir un estimateur non biaisé et aussi peu dispersé que possible de l'espérance de la fonction (nous parlerons de connaissance ponctuelle). Nous proposerons dans cette partie un algorithme qui vise à répondre à cet arbitrage.

Objectif poursuivi L'objectif que nous poursuivrons ici est l'exploration des zones susceptibles de contenir un minimum global de f : nous chercherons donc à explorer sélectivement la fonction f , de façon à avoir une vision globale de celle-ci, tout en privilégiant certaines zones d'intérêts, comme les minima globaux, dans un souci d'économie du nombre de tirages de la fonction F . L'algorithme vise au final à estimer m^* ainsi que S_x , pour x donné dans un voisinage de m^* .

II-1-2 Approche retenue

L'approche proposée ici s'apparente à une approche de type *Branch and Bound*, mais aucune zone n'est jamais définitivement exclue de l'analyse, l'idée étant ici d'ordonner les zones en fonction de la probabilité qu'une zone contienne un optimum, selon un modèle que nous détaillerons.

Nous nous placerons dans le cadre de la minimisation d'une fonction, nous chercherons ici à quantifier la probabilité qu'une zone contienne un minimum plus petit qu'un minimum observé, pour au final construire l'ensemble des zones susceptibles de contenir un minimum global.

Grille à pas fixe Une solution simple et très commune est l'utilisation d'une grille de simulation, à pas fixe. Selon cette solution, la fonction f est explorée sur un ensemble Θ_m . Θ_m est une grille de pas δ , toutes les composantes de $\theta \in \Theta_m$ parcourent l'ensemble des valeurs multiples de δ dans Θ : $\Theta_m = \left\{ \theta \in \Theta, \frac{1}{\delta} \theta \in \mathbb{N}^d \right\}$. Le nombre n de tirages de $F(\theta)$ est identique pour chaque $\theta \in \Theta_m$. Cette solution est néanmoins très onéreuse en terme de temps de calcul, dans la mesure où n simulations seront conduites sur chacun des points, y compris ceux très éloignés d'un optimum, pour lesquels la meilleure connaissance de f n'apporte quasiment rien.

Grille à pas variable Nous chercherons donc à mettre en place une grille à pas variable, où les retirages de F se feront principalement dans les zones susceptibles d'accueillir le minimum. Il va s'agir d'une part d'estimer f en de nouveaux points θ (sommets), et d'autre part de répartir de nouvelles simulations entre les anciens sommets et les nouveaux, en fonction de l'intérêt que peut avoir, sur la connaissance du minimum global, l'ajout de simulations en chacun de ces points. Les positions des sommets envisagés ne seront plus régulièrement réparties, et le nombre de simulations conduites en chaque sommet va différer selon les sommets. Une première difficulté est le choix d'une forme convenable pour les zones de recherche, pour les cellules de la grille. Une autre difficulté est que nous avons besoin de deviner où conduire les futures simulations. Pour cela, il faut avoir une idée de comment va

évoluer la fonction entre les points explorés : il faut en un sens deviner quel pourra être l'impact de futures simulations avant même de les réaliser.

Algorithmes présentés et structure du document Nous allons aborder dans cette étude deux algorithmes réalisant une grille à pas variable. Dans la sous-section II-2, nous envisagerons le cas où le nombre de tirages en chaque point est fixe, l'algorithme cherchant alors uniquement à déterminer les prochains points où évaluer la fonction F . Dans la sous-section II-3, nous introduirons la possibilité d'opérer des retirages en des points déjà explorés. Enfin, dans une dernière sous-section II-4, nous présenterons des illustrations simples et visuelles du comportement de l'algorithme proposé sur une fonction bimodale élémentaire. Une comparaison avec un algorithme classique d'optimisation stochastique sera également menée.

II-2 Une grille à pas variable par subdivision systématique

II-2-1 Forme des zones de recherche

Zone initiale de recherche En pratique, par exemple lors de l'optimisation de choix d'investissements de $d + 1$ actifs, les proportions investies dans chacun des actifs se somment à un. Il suffit alors de rechercher d pourcentages d'allocation, l'allocation numéro $d + 1$ se déduisant des d allocations précédentes. Il s'agit donc d'optimiser une fonction de R^d dans R .

Nous nommerons $Z_0 \subset R^d$ la zone initiale de recherche, simplexe standard orthogonal constitué des sommets $(1, 0, \dots, 0)$, $(0, 1, 0, \dots, 0)$, $\dots$, $(0, \dots, 0, 1)$, ainsi que du sommet $(0, \dots, 0)$.

Cela correspondra bien à la situation où la seule contrainte est d'avoir une somme des composantes inférieure à un, par exemple d pourcentages d'allocation d'actifs dont la somme est inférieure à 100 %, la différence avec 1 formant l'allocation de l'actif numéro $d + 1$:

$$Z_0 = \{(x_1, \dots, x_d) \in R^d, x_1 + \dots + x_d \leq 1, x_1 \geq 0, \dots, x_d \geq 0\}$$

Dans les situations où l'optimum est à rechercher sur une zone plus complexe soumise à de nombreuses contraintes, l'algorithme proposé sera applicable si cette zone peut être représentée dès l'initialisation par une union finie de simplexes. Dans tous les cas, nous supposerons que l'ensemble des sommets de la ou des zones initiales ont été explorés, par la réalisation de tirages de $F(\theta)$ en chacun des sommets de l'enveloppe de ces zones.

Mécanisme de scission L'idée d'une grille à pas variable est de séparer la zone de recherche de l'optimum en plusieurs zones de différentes tailles. A chaque étape, la scission d'une zone peut se faire en opérant des tirages de F en un ou plusieurs points non encore explorés de Θ .

Certains algorithmes d'optimisation globale fonctionnent, dans un contexte particulier déterministe, par subdivision de la zone à explorer. Nous pouvons notamment citer l'algorithme DIRECT (pour *D*IViding *R*ECtangles, cf. Jones et al. [1993]), qui subdivise un pavé de R^d en plusieurs sous-pavés. Si la subdivision d'une zone nécessite l'exploration de n^+ nouveaux points, il paraît préférable de choisir $n^+ = 1$. Dans ce seul cas, le choix d'un nouveau point de subdivision se fait alors en connaissance de tous les précédents tirages de la fonction. Nous avons ici choisi un mode de division qui d'une part nous semblait plus adapté

aux problèmes définis sur une union de simplexes, et d'autre part ne nécessitait l'exploration que d'un unique point à chaque subdivision.

Nous supposons que les zones de recherche sont des ensembles convexes, de façon à faciliter d'éventuelles interpolations au coeur de chaque zone. Imaginons un pavé de $\mathbb{R}^d$ dont on connaît les sommets, et à l'intérieur duquel nous ajouterions un point θ_c . Comment découper rapidement ce pavé en une partition d'ensembles convexes dont l'enveloppe convexe contiendrait θ_c ? La réponse n'étant pas si évidente, nous opterons pour le choix suivant :

- Les zones considérées seront délimitées par $d + 1$ sommets.
- Chaque nouveau point sera ajouté sur un segment de l'enveloppe convexe de la zone.
- A chaque étape, chaque zone sera éventuellement scindée en deux ensembles convexes.

Tout point à l'intérieur d'une zone appartiendra à une zone convexe délimitée par $d + 1$ points, et le choix d'un nombre de sommets égal à $d + 1$ facilitera par la suite la séparation d'une zone en plusieurs zones. D'autre part, le choix d'un nombre fixe de sommets délimitant chaque zone sera de nature à faciliter l'implémentation de l'algorithme.

Nous appellerons zone un d -simplexe, c'est-à-dire un ensemble convexe inclus dans Θ , avec $\Theta \subset \mathbb{R}^d$, dont l'enveloppe convexe est déterminée par $d + 1$ sommets distincts (les sommets formant un repère affine de $\mathbb{R}^d$). Par la suite, nous noterons $S(Z)$ l'ensemble des $d + 1$ sommets délimitant une zone convexe Z . La zone initiale de recherche de l'optimum est notée Z_0 . A l'issue de l'étape numéro k , la zone de recherche est subdivisée en un ensemble de zones recouvrant Z_0 , dont l'intersection est de mesure nulle. Cet ensemble de zones est noté Z_k .

Considérons un ensemble de sommets E délimitant une zone Z et deux points distincts θ_1 et θ_2 de cet ensemble. Supposons que θ_c soit le barycentre (par exemple équipondéré) entre ces deux points. Lorsque la décision est prise de scinder cette zone en deux zones Z_1 et Z_2 , les ensembles de sommets délimitant ces deux zones seront :

$$\begin{aligned} E_1 &= (E \setminus \{\theta_1\}) \cup \{\theta_c\} \\ E_2 &= (E \setminus \{\theta_2\}) \cup \{\theta_c\} \end{aligned}$$

Une preuve de ce lemme est donnée dans Rullière et al. [2010]. Par souci de simplicité, nous nommerons par la suite sommets d'une zone Z l'ensemble des $d + 1$ sommets délimitant l'enveloppe convexe de la zone Z .

Le graphique 47 illustre en dimension 2 la façon dont une zone peut être progressivement subdivisée. Nous pouvons remarquer que le choix d'une décomposition de l'ensemble Θ en un ensemble de zones contenant des points de tirage $\{\theta_i\}_{i=1,2,\dots,n}$ est un problème de triangulation, classique en géométrie algorithmique. Une technique de triangulation très connue est la triangulation de Delaunay (cf. De Berg et al. [2008]). Le problème étant ici de choisir un point de tirage au sein d'une zone et non pas de construire des zones pour séparer des points de tirage existants, les zones seront scindées par la simple bisection présentée.

Un des avantages de la bisection présentée est de permettre l'exploration des frontières de la zone initiale. D'autres choix de subdivision de zones sont naturellement possibles, par exemple autour de l'isobarycentre d'une zone. Comme nous pouvons l'observer sur le graphique 47 ainsi que sur les différentes illustrations de la sous-section II-4, l'union des enveloppes convexes des zones scindées forme un ensemble beaucoup plus vaste que l'enveloppe convexe du simplexe initial, de sorte que les points de tirage ne sont pas condamnés à rester dans l'enveloppe de la zone initiale.

Dans le cas de barycentres équipondérés, nous pouvons d'ailleurs montrer par récurrence qu'il est possible d'atteindre en un nombre fini d'étapes tout point de Θ de coordonnées :

$$\left(\frac{i_1}{2^{k_1}}, \dots, \frac{i_d}{2^{k_d}} \right) \text{ avec pour tout } j \in \{1, \dots, d\}, k_j \in \mathbb{N}, i_j \in \mathbb{N}, i_j \leq 2^{k_j}$$

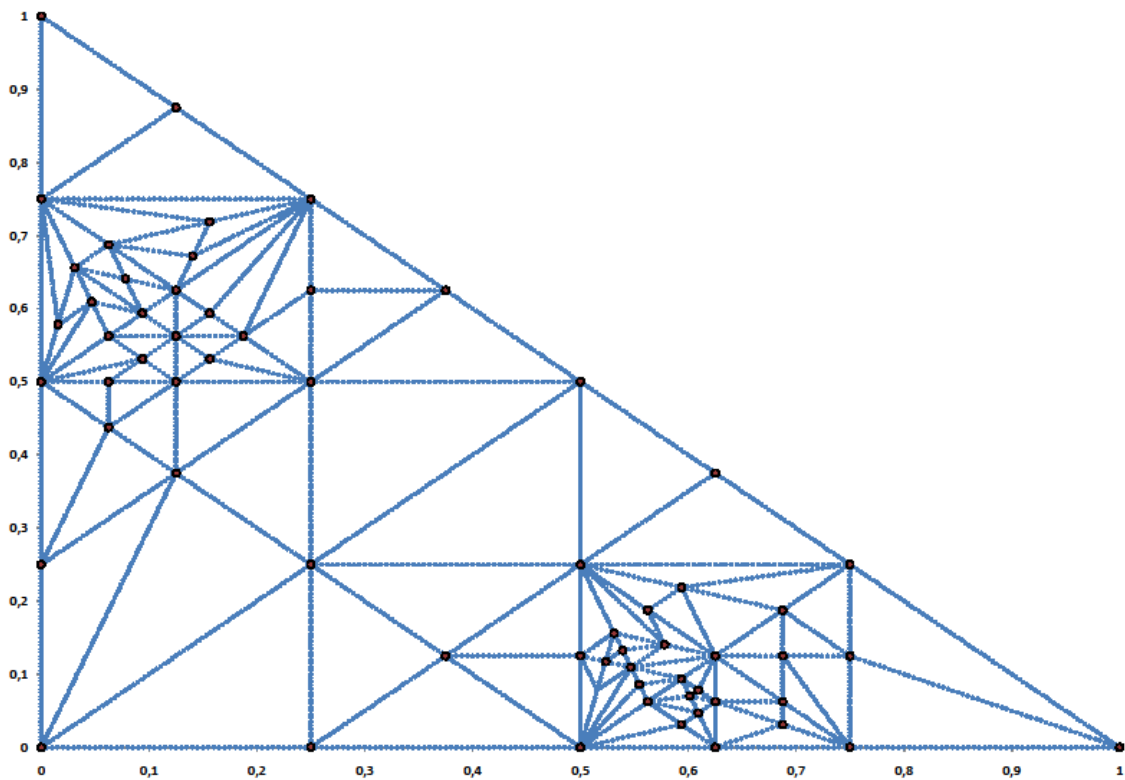


Fig. 47 : Un exemple de partition d'une zone dans $\mathbb{R}^2$ en 60 zones. ($d = 2$)

II-2-2 Choix de la zone à explorer ou segmenter

Idée du potentiel d'une zone L'étude de la possibilité pour f d'atteindre un minimum global sur une zone non explorée nécessite de fixer des hypothèses : si f est supposée extrêmement erratique, f pourra franchir un seuil inférieur sur à peu près n'importe quelle zone, et les tirages opérés de F n'apporteront que très peu d'information. Une solution classiquement retenue, dans le domaine de l'optimisation globale déterministe, est le choix d'une forme Lipschitzienne pour f (cf Jones et al. [1993]) : cela revient à dire, dans un cadre déterministe,

que f pourra atteindre un minimum global sur une zone si la pente nécessaire pour cette atteinte est inférieure à un certain seuil, ou que f ne pourra pas atteindre le minimum global.

Cette logique binaire permet d'indiquer de façon certaine si f appartient ou non à un intervalle donné, et correspond à une logique d'arithmétique d'intervalle (cf. Wolfe [1996]). En l'absence d'information très précise sur f , nous avons préféré une logique probabiliste, malgré une part de subjectivité qu'elle peut engendrer : plutôt qu'une indicatrice de franchissement possible à valeur dans $\{0,1\}$, nous allons rechercher une mesure sur $[0,1]$ quantifiant la probabilité, selon la représentation de f par l'observateur, que f franchisse le seuil sur une zone. Si cette solution introduit une nécessaire subjectivité sur le calcul de cette probabilité, elle offre néanmoins l'avantage de maintenir une hiérarchie entre différentes zones, ce qui est notamment utile lorsque beaucoup de zones sont susceptibles de contenir l'optimum, et permet de ne jamais exclure *a priori* de zone.

Comme le temps de calcul de F est supposé beaucoup plus important que le temps d'exécution de l'algorithme de choix des zones à explorer, nous choisirons à chaque étape de l'algorithme une unique zone à scinder en deux ou à explorer davantage. L'objectif étant de limiter les risques qu'une zone non explorée contienne un optimum, nous piocherons à chaque étape une zone à explorer, avec une probabilité proportionnelle à un coefficient β spécifique à la zone.

Pour chaque zone Z , le coefficient $\beta(Z)$ déterminera s'il est plausible (dans un sens que nous préciserons) que la zone contienne un minimum plus petit que le minimum estimé m^* . Afin d'approcher ce coefficient sur toute la zone, nous aurons besoin de le déterminer en un unique point. Selon un modèle probabiliste que nous préciserons, nous nommerons "potentiel" d'un point une mesure d'autant plus grande que la fonction objectif est susceptible, en ce point, d'être plus petite que le minimum global observé. Une représentation probabiliste du potentiel d'un point passe par la modélisation de la connaissance incertaine de la fonction f au vu des tirages opérés en différents points, par exemple au moyen de champs aléatoires conditionnels.

Nous présentons ici une approche plus simple partant d'une modélisation de l'incertitude en dimension 1, puis agrégeant différentes incertitudes compte tenu des sommets explorés.

Potentiel d'un point en direction d'un sommet Considérons un sommet θ' déjà exploré d'une zone Z , et un point θ de cette zone, distinct de θ' . Par souci de simplicité, nous étudierons ici l'évolution de la fonction f entre ces deux points θ et θ' en connaissance de la seule exploration en θ' .

Même en l'absence d'erreur d'échantillonnage au sommet θ' , l'évolution de la fonction f est naturellement inconnue sur le segment $]\theta'; \theta]$: nous parlerons d'*erreur de grille*. Nous supposons que, à partir de la seule connaissance de f au point θ' , l'observateur représente $f(\theta)$ par une variable aléatoire $\tilde{f}_{\theta'}(\theta)$, choisie de loi normale, d'espérance $f(\theta')$ et dont l'écart-type est une fonction croissante de la distance $d(\theta, \theta')$. Nous choisirons :

$$\begin{aligned}\tilde{f}_{\theta'}(\theta) - f(\theta') &\approx N(0, \sigma_g^2) \\ \sigma_g(\theta, \theta') &= \sigma_K d(\theta, \theta')^\alpha \quad \sigma_K > 0, \alpha > 0\end{aligned}$$

Cela revient à décrire l'incertitude maximale pesant sur $\tilde{f}$, lorsque σ_g est soumis à une condition de type Hölder (ou de type Lipschitz lorsque $\alpha = 1$).

A titre d'illustration, pour $\theta_0 \in [\theta'; \theta]$, la représentation $\tilde{f}_{\theta'}(\theta_0) - f(\theta')$ par un mouvement brownien correspond à $\alpha = \frac{1}{2}$. La représentation de la pente $(\tilde{f}_{\theta'}(\theta_0) - f(\theta'))/d(\theta_0, \theta')$ par un mouvement brownien correspond à $\alpha = \frac{3}{2}$. Nous évoquerons plus en détail le choix de α et σ_K au paragraphe II-2-3, ainsi que dans la sous-section d'application numérique II-4.

En présence d'erreur d'échantillonnage, nous disposons d'un estimateur $\hat{f}(\theta')$ de $f(\theta')$. Si $\hat{f}(\theta') - f(\theta')$ est elle-même une variable aléatoire de loi normale centrée et de variance $\sigma_e^2(\theta')$, indépendante de l'erreur de grille $\tilde{f}_{\theta'}(\theta) - f(\theta')$, alors :

$$\tilde{f}_{\theta'}(\theta) - \hat{f}(\theta') \approx N(0, \sigma_e^2(\theta') + \sigma_g^2(\theta, \theta'))$$

En connaissance d'une réalisation de $\hat{f}(\theta')$, nous noterons :

$$\frac{1}{2} L(\theta, \theta') = P[\tilde{f}_{\theta'}(\theta) \leq m^*]$$

où m^* représente la valeur du minimum global de $f(\theta)$.

La quantité $L(\theta, \theta')$ sera nommée potentiel directionnel du point θ en direction de θ' . Le facteur de normalisation $1/2$ n'a aucune incidence sur les comparaisons des potentiels entre eux (comme tout facteur strictement positif), mais permettra par la suite à la quantité L de pouvoir atteindre toutes les valeurs de $[0, 1]$, et non pas seulement $[0, 1/2]$. Le potentiel directionnel s'interprète alors comme la probabilité que $\tilde{f}_{\theta'}$ franchisse un minimum global connu sachant que $\tilde{f}_{\theta'}$ décroît depuis le point d'accroche θ' .

Il est possible de tenir compte de l'erreur d'estimation du minimum m^* : en supposant que l'estimateur $\hat{m}^*$ est une variable aléatoire indépendante de $\tilde{f}_{\theta'}(\theta) - \hat{f}(\theta')$, de loi normale, d'espérance m^* et de variance $\sigma_{m^*}^2$,

$$\frac{1}{2} L(\theta, \theta') = 1 - \Phi\left(\frac{\hat{f}(\theta') - \hat{m}^*}{\sigma_T(\theta, \theta')}\right), \quad \sigma_T(\theta, \theta') > 0$$

avec $\sigma_T^2(\theta, \theta') = \sigma_{m^*}^2 + \sigma_e^2(\theta') + \sigma_K^2 d(\theta', \theta)^{2\alpha}$

La fonction Φ désigne la fonction de répartition d'une loi normale centrée réduite. Nous prendrons la convention $L(\theta, \theta') = 1_{\hat{f}(\theta') = \hat{m}^*}$ dans le cas où $\sigma_T(\theta, \theta') = 0$, en l'absence de bruit et lorsque $\sigma_K = 0$ ou $\theta = \theta'$. Nous n'évoquerons pas en détail la détermination très classique, en un point exploré θ , de l'estimateur $\hat{f}(\theta)$ de $f(\theta)$, ni de l'estimateur de la variance de $\hat{f}(\theta)$.

L'usage de la moyenne et de la variance empirique non biaisée donnerait, à partir de n_0 observations de $F(\theta)$ notées $F_1(\theta), \dots, F_{n_0}(\theta)$, $n_0 \geq 2$:

$$\hat{f}(\theta) = \frac{1}{n_0} \sum_{i=1}^{n_0} F_i(\theta)$$

$$\sigma_e^2(\theta) = \frac{1}{n_0} \hat{\sigma}_F^2(\theta), \text{ avec } \hat{\sigma}_F^2(\theta) = \frac{1}{n_0 - 1} \sum_{i=1}^{n_0} (F_i(\theta) - \hat{f}(\theta))^2$$

$\sigma_F^2(\theta)$ désigne ici la variance de $F(\theta)$, et $\hat{\sigma}_F^2(\theta)$ un estimateur de cette variance. Des raffinements peuvent être envisagés afin de tenir compte de l'erreur d'estimation de σ_F . Notons surtout la nécessité d'un paramètre $n_0 \geq 2$, nombre de tirages requis pour l'estimation de la variance empirique $\sigma_e^2(\theta)$. Dans le cas déterministe, nous pouvons fixer $n_0 = 1$, $\hat{f}(\theta) = f(\theta)$ et $\sigma_e^2(\theta) = 0$. Mais dans le cas général, ce paramètre n_0 restera une entrée de l'algorithme. Eventuellement, l'usage d'une hypothèse de répartition gaussienne pour $F(\theta)$ peut permettre de ne pas mémoriser l'intégralité des tirages de $F(\theta)$, mais de simplement mémoriser le nombre de tirages précédents, leur somme et la somme de leurs carrés.

Pour l'estimation de m^* et de σ_m^{*2} , nous retiendrons en première approche les valeurs obtenues pour $\hat{f}(\theta^*)$ et $\sigma_e^2(\theta^*)$ au dernier point optimal θ^* rencontré parmi les tirages réalisés itérativement par l'algorithme, bien que là encore des perfectionnements puissent être suggérés.

Remarquons que le potentiel peut être adapté à d'autres recherches que la recherche d'optimum. Ainsi, si l'on recherchait les paramètres θ tels que $f(\theta)$ appartienne à un ensemble $A \subset \mathbb{R}$, on pourrait utiliser un potentiel proportionnel à $P[\tilde{f}_\theta(\theta) \in A]$. L'algorithme peut ainsi être très facilement adapté à la recherche de racines.

Potentiel d'un point Considérons un point θ à l'intérieur d'une zone Z , et cherchons à définir une mesure d'autant plus grande qu'un minimum global de f peut être observé au point θ . Ce point θ est muni de $d + 1$ potentiels directionnels en direction des $d + 1$ sommets délimitant la zone Z , sommets notés ici θ_i , $i \in \{1, \dots, d + 1\}$. Ces potentiels marquent chacun la probabilité, selon la modélisation de l'observateur, que la fonction f franchisse le seuil m^* sur l'un des segments $[\theta, \theta_i]$. Les potentiels de points de différentes zones sont destinés à être comparés entre eux. La principale exigence est ici de respecter la logique booléenne selon laquelle si le minimum ne peut pas être atteint en direction d'un sommet, alors il est exclu qu'il soit atteint sur la zone. Une mesure commode respectant cette exigence est naturellement le produit, qui correspond bien au ET booléen lorsque les quantités $L(\theta, \theta_i)$ appartiennent à $\{0, 1\}$.

Nous appellerons potentiel du point θ dans la zone Z la quantité :

$$\beta_z(\theta) = \left(\prod_{i=1}^{d+1} L(\theta, \theta_i) \right)^\gamma$$

Nous avons choisi ici d'introduire une fonction lien $g(x) = x^\gamma$, pour $\gamma > 0$, qui préserve les propriétés requises de logique booléenne du produit, et permet de modifier le comportement de l'algorithme. Le choix de γ , régissant la convexité de g , permet de distordre les potentiels et de moduler ainsi l'importance relative accordée à la zone de plus grand potentiel. Au niveau de l'impact de la distorsion, pour un coefficient γ très grand, le tirage d'une zone reviendrait au tirage systématique de la zone dont le potentiel est le plus grand. Un coefficient γ faible conduirait à davantage explorer des zones de potentiel médian. Le facteur γ s'interprète comme un facteur de priorité attribuée aux zones de grand potentiel.

Nous avons constaté que l'effet de γ était modeste pour des valeurs raisonnables de paramètres (cf. sous-section II-4 de ce chapitre). Considérant par ailleurs une certaine redondance de ce paramètre avec les paramètres (σ_K, α) du potentiel, nous opterons sauf mention contraire pour :

$$\gamma = 1/(d+1)$$

qui a le mérite de clarifier l'interprétation du potentiel : $\beta_z(\theta)$ correspondra simplement à la moyenne géométrique des potentiels directionnels. Le potentiel d'un point s'interprétera alors comme un potentiel directionnel moyen.

Le choix d'une mesure pour le potentiel d'un point justifierait de nombreuses études. S'agissant d'agrégation de potentiels directionnels et de fonctions liens, d'autres distorsions de probabilités facilement utilisables pourront être trouvées dans Bienvenue et Rullière [2010]. Une autre piste pour cette agrégation est l'usage de copules, ou encore celle de la logique floue (cf. Zadeh [1965]). Nous étudions actuellement d'autres mesures de ce potentiel (cf. Rullière et al. [2010]), basés sur l'agrégation des variances directionnelles $\sigma_T(\theta, \theta')$, ainsi que la modélisation de $\tilde{f}$ par des champs aléatoires, en lien avec la théorie du Krigeage (cf. par exemple Krige [1951], Jones et al. [1998]). Ces mesures présentent l'avantage d'éliminer l'agrégation de potentiels directionnels, au prix d'un modèle d'une complexité parfois accrue.

Il faut néanmoins tempérer l'impact du choix d'une mesure de potentiel : les potentiels des différents points serviront à établir une hiérarchie des zones les plus susceptibles de contenir un optimum global, afin de choisir laquelle explorer en priorité. Le choix de la fonction lien g croissante modifiera les priorités d'exploration, mais ne modifiera pas la hiérarchie elle-même, et celle-ci dépendra en très grande partie de l'éloignement de f à la valeur de l'optimum.

Si les potentiels directionnels appartenaient à $\{0,1\}$, un potentiel permettrait simplement de dire si il est possible ou non que f possède un minimum global sur une zone, compte tenu des sommets adjacents et de la valeur estimée du minimum. Cela correspond à l'idée développée dans les algorithmes mettant en oeuvre une arithmétique d'intervalle (cf. Wolfe [1996]). Le potentiel proposé ici doit être vu comme une simple mesure, à valeur sur $[0, 1]$, permettant

d'étendre une logique d'arithmétique d'intervalle qui conduirait à des potentiels définis sur $\{0,1\}$.

Potentiel d'une zone En pratique, une zone sera d'autant plus susceptible de contenir un minimum plus petit que la valeur estimée de m^* si sa surface est grande et si ses points ont un potentiel élevé. Une mesure logique de la "surface probable" d'une zone Z est donnée par :

$$\bar{\beta}(Z) = \int_{\theta \in Z} \beta_z(\theta) d\theta$$

Bien que le calcul de cette intégrale puisse être approché par des techniques de simulation, nous avons préféré retenir comme mesure du potentiel d'une zone la mesure suivante :

$$\beta(Z) = V(Z) \cdot \beta_{\theta_{\beta_z}}(\theta_{\beta_z})$$

θ_{β_z} isobarycentre des sommets de la zone Z ,

où $V(Z)$ est le volume de la zone Z (hypervolume dans le cas $d > 3$).

Cette solution a le mérite de la simplicité, offre l'avantage d'être très rapide et de ne pas être aléatoire. Un calcul d'un $\beta(Z)$ précédemment mené n'aura donc pas à être réitéré si la zone Z n'est pas modifiée. En outre, l'objectif est de comparer les coefficients β entre eux : le calcul fin de l'intégrale a peu de raisons explicites de beaucoup perturber la hiérarchie entre les différentes zones. Enfin, l'isobarycentre nous a paru bien rendre compte de l'erreur de grille au sein de la zone, et facilitera l'interprétation future du potentiel d'une zone.

S'agissant du calcul du volume $V(Z)$, dans le cas où $d = 2$, les zones sont des triangles et un volume $V(Z)$ est donné par la formule de Héron. Dans le cas général, le déterminant de Cayley-Menger donne le volume exact de la zone (cf. Sommerville [1958]). Plus simplement, dans le cas d'une séparation d'une zone en deux volumes égaux, ce qui sera ici le cas, la simple mémorisation du volume de la zone à subdiviser permet de déduire immédiatement le demi-volume de chaque zone fille.

Choix d'une zone en fonction des potentiels Comme nous l'avons évoqué, le choix d'une zone Z^+ parmi un ensemble z de zones se fera de la façon suivante : la probabilité de piocher une zone sera proportionnelle au potentiel de chaque zone.

Notons $z = \{Z_1, \dots, Z_n\}$ l'ensemble des zones dans lequel doit être piochée Z^+ . Si $\{U_v\}_{v=1,2,\dots}$ désigne une suite de variables aléatoires de loi uniforme sur $[0,1]$, mutuellement indépendantes, piocher une zone parmi n avec une probabilité fixée au prorata de son potentiel revient à choisir, à une étape v de l'algorithme :

$$Z^+ = Z_{k^*(U_v)}$$

avec :

$$k^*(u) = \min \{k \in \{1, \dots, n\}, B_k \geq u \cdot B_n\}$$

et :

$$B_k = \sum_{i=1}^k \beta(Z_i), \quad k \in \{1, \dots, n\}$$

Ce choix découle de plusieurs idées. D'une part, l'idée d'explorer de façon uniforme la zone de recherche lorsque les potentiels sont égaux, d'autre part l'idée de préserver, à la façon d'un

recuit simulé, la possibilité d'exploration de zones *a priori* peu prometteuses. Nous imaginons, en présence d'un très grand nombre de très petites zones, que cette solution limite le risque de confiner les points de tirage dans le voisinage d'un unique minimiseur, et favorise ainsi une certaine prudence dans l'exploration de la fonction. D'autres choix possibles sont évoqués dans la section d'applications numériques (sous-section II-4 de ce chapitre).

II-2-3 Choix des paramètres (σ_K, α) du potentiel

Au moyen du potentiel que nous avons défini, nous avons transféré une part de la subjectivité du choix du prochain point de tirage sur le choix de quelques paramètres, au premier rang desquels se trouvent les coefficients de type Hölder α et σ_K . Le choix ou l'estimation de ces paramètres est un problème délicat qui nécessiterait à lui seul une étude poussée, et que nous présentons ici de façon simplifiée.

Le choix des paramètres α et σ_K dépend de la connaissance de la fonction f considérée, ainsi que, lorsque celle-ci s'avère insuffisante, de la prudence de l'observateur. Le choix du paramètre α est un choix de modèle : envisageons-nous que f puisse varier brutalement sur un très petit intervalle comme le ferait une trajectoire de mouvement brownien, ou de façon plus régulière, comme une fonction lipschitzienne ?

Considérons deux points distincts θ et θ' , en ignorant dans un premier temps l'erreur d'échantillonnage aux sommets explorés. Si seul θ' est exploré, l'observateur suppose *a priori* que $Y_{\theta, \theta'} = \tilde{f}_{\theta'}(\theta) - f(\theta')$ est distribué selon une loi normale centrée, d'écart-type $\sigma_K d^\alpha$, où $d = d(\theta, \theta')$, $d > 0$. Après exploration des deux points, l'observateur dispose d'une observation (d, y) de ce couple (d, Y) . En supposant que sont collectées n réalisations (d_i, y_i) , en les supposant de surcroît mutuellement indépendantes, une estimation maximum de vraisemblance de σ_K en connaissance de α conduirait à :

$$\sigma_K^2 = \frac{1}{n} \sum_{i=1}^n \frac{y_i^2}{d_i^{2\alpha}}$$

et une estimation maximum de vraisemblance de α en connaissance de σ_K conduirait à α tel que :

$$\sum_{i=1}^n \frac{y_i^2 \ln(d_i)}{\sigma_K^2 d_i^{2\alpha}} = \sum_{i=1}^n \ln(d_i)$$

Si l'on considère que, en tout point exploré θ , $\hat{f}(\theta)$ subit une erreur d'estimation, on peut observer $Y'_{\theta, \theta'} = (\tilde{f}_{\theta'}(\theta) + N_\theta) - (f(\theta') + N_{\theta'})$, où N_θ et $N_{\theta'}$ sont des variables aléatoires indépendantes, de lois normales, de variances respectives σ_θ^2 et $\sigma_{\theta'}^2$. Alors $Y'_{\theta, \theta'}$ est supposée distribuée selon une loi normale centrée, de variance $\sigma_e^2 + \sigma_K^2 d^{2\alpha}$, où $\sigma_e^2 = \sigma_\theta^2 + \sigma_{\theta'}^2$ représente une erreur d'estimation, connue après exploration des deux points puisque mesurée aux points d'observation.

L'estimation maximum de vraisemblance de σ_K (sachant α) et de α (sachant σ_K), conduit respectivement à σ_K et α tels que :

$$\sum_{i=1}^n \frac{d_i^{2\alpha}}{\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha}} = \sum_{i=1}^n \frac{y_i^2 d_i^{2\alpha}}{(\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha})^2}$$

$$\sum_{i=1}^n \frac{d_i^{2\alpha} \ln(d_i) y_i^2}{(\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha})^2} = \sum_{i=1}^n \frac{d_i^{2\alpha} \ln(d_i)}{\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha}}$$

Dans la pratique, la répétition des étapes d'estimation de σ_K sachant α puis de α sachant σ_K a rapidement convergé dans tous les cas que nous avons testé (cf sous-section II-4 de ce chapitre).

Cette estimation nous a conduits à des résultats très proches de ceux obtenus en maximisant directement la log-vraisemblance de l'échantillon :

$$\ln V(\sigma_K, \alpha) = -\frac{n}{2} \ln(2\pi) - \sum_{i=1}^n \frac{1}{2} \ln(\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha}) - \sum_{i=1}^n \frac{y_i^2}{2(\sigma_{e_i}^2 + \sigma_K^2 d_i^{2\alpha})}$$

Au fur et à mesure des nouveaux tirages, chaque nouveau point θ_0 exploré conduit, en direction des $d + 1$ sommets de la zone, à $d + 1$ nouvelles réalisations de variables aléatoires de loi normale centrée, d'écart-type fonction de la distance, et il est alors possible de corriger à chaque étape une valeur *a priori* de α et de σ_K .

Il faut ici noter que rien n'interdit de faire varier les coefficients α et σ_K en fonction de la zone considérée, et l'estimation des coefficients pourrait se faire sur chaque zone Z en affectant chaque réalisation de la variable aléatoire $Y_{\theta, \theta'}$ (sur d'autres zones) de poids d'autant plus élevés que la θ et θ' sont proches de la zone considérée Z . L'estimation de coefficients de Hölder à partir d'observations de tirages de F est un vaste sujet (cf. Blanke [2002] pour un article traitant de ce type d'estimation sur des processus). Les choix numériques concrets des paramètres de l'algorithme seront détaillés dans la sous-section d'application numérique II-4.

Enfin, le raisonnement tenu jusqu'à présent se basait sur le fait que pour deux points d'une zone, l'observateur représentait la variation des pentes entre ces deux points par une variable aléatoire dont l'écart-type était une fonction croissante de la distance. Même si cela n'est pas requis pour l'implémentation de l'algorithme, un élément susceptible de fournir a posteriori des informations sur l'hypothèse prise est le suivant : pour une zone Z , nous définissons une variable aléatoire IR_α :

$$IR_\alpha(Z) = \frac{f(\theta_1) f(\theta_2)}{d(\theta_1, \theta_2)^\alpha}$$

θ_1, θ_2 vecteurs aléatoires distincts issus de tirages indépendants, uniformes sur Z .

Précisons ici les modalités des tirages. D'une part, le tirage uniforme d'un vecteur θ sur une zone Z peut s'opérer facilement en choisissant des coordonnées barycentriques de façon uniforme dans un simplexe unité : ainsi, si ω_j est la $j^{\text{ème}}$ coordonnée barycentrique de θ dans

la zone Z , $j \in \{1, \dots, d+1\}$, nous pourrions prendre $\omega_j = e_j / e_{tot}$, avec $e_{tot} = \sum_{j=1}^{d+1} e_j$, pour un ensemble $\{e_j\}_{j \in \{1, \dots, d+1\}}$ de variables aléatoires mutuellement indépendantes, de loi exponentielle de paramètre 1 (ce qui revient à prendre les intervalles successifs de statistiques d'ordre d'un vecteur uniforme). Le lecteur intéressé par ce type de tirages pourra consulter notamment Smith et Tromble [2004]. D'autre part, nous conviendrons que (θ_1, θ_2) est le premier couple distinct issu de tirages uniformes (du fait de la précision arithmétique finie des ordinateurs, la probabilité que les deux points soient confondus peut être non nulle en pratique, bien qu'extrêmement faible).

Selon le modèle présenté, à défaut d'autres informations, l'observateur suppose qu'il existe une quantité $\alpha \geq 0$ telle que la variation de f entre θ_1 et θ_2 ne s'explique que par la distance $d(\theta_1, \theta_2)^\alpha$, et pour lequel l'écart-type de $IR_\alpha(Z)$ serait fini (ce qui n'est pas acquis pour tout α bien sûr). Sous l'hypothèse prise, cet écart-type fournira une information sur les variations de pentes auxquelles l'utilisateur pourra s'attendre. En particulier, si $IR_\alpha(Z)$ est supposé distribué de façon gaussienne, alors son écart-type empirique correspondra à un estimateur de σ_K en connaissance de α . Des distributions de IR_α seront illustrées dans la sous-section II-4.

II-2-4 Choix du sommet de scission

Une fois une zone Z^+ choisie parmi un ensemble de zones z , il va s'agir de scinder la zone en ajoutant un point à l'intérieur de celle-ci. De par nos choix précédents, ce point se situera sur l'un des segments de l'enveloppe convexe de la zone.

Pour la zone Z^+ , nous noterons $C = \{(\theta_{j_1}, \theta_{j_2})\}_{j_1, j_2}$ l'ensemble des couples de sommets distincts de Z^+ . Le sommet de séparation retenu peut être choisi de nombreuses façons.

La solution que nous retiendrons consiste ici à sélectionner, parmi l'ensemble des couples possibles, l'un de ceux qui maximisent la longueur du segment. Si $\theta_1, \dots, \theta_{d+1}$ désignent les sommets de la zone Z^+ , nous choisirons comme sommet de séparation de la zone le sommet θ^+ , isobarycentre du segment $(\theta_{j_1}, \theta_{j_2})$ tel que :

$$(\theta_{j_1}, \theta_{j_2}) = \arg \max_{(\theta_i, \theta_j) \in C} d(\theta_i, \theta_j)$$

$$C = \{(\theta_i, \theta_j)\}_{i \in \{1, \dots, d\}, j \in \{i+1, \dots, d+1\}}$$

Dans le cas de plusieurs segments de longueur maximale, nous conviendrons que $\arg \max$ désignera un segment choisi de façon uniforme parmi l'ensemble (fini) des segments de longueur maximale.

Choisir le point de segmentation sur un segment de longueur maximale a le mérite de limiter l'apparition de simplexes très déséquilibrés dans leurs longueurs de segment : en effet, la connaissance de l'appartenance de l'optimum à une zone est moins utile si cette zone est très allongée dans certaines directions. Le choix d'un sommet conduisant également à la partition des zones adjacentes, l'apparition de simplexes très déséquilibrés, si elle est ainsi limitée,

n'est toutefois pas absolument exclue. Une idée de la forme des zones ainsi obtenues en dimension 2 est illustrée par la figure 47, qui a été construite par scission du segment de plus grande longueur de la zone choisie.

Lorsque des simulations sont opérées au point θ^+ du segment choisi de séparation, la scission des zones se fait ainsi :

- L'ensemble des zones z^+ contenant ce segment est déterminé (celles qui contiennent les deux sommets du segment à la fois).
- Chacune des zones Z de l'ensemble z^+ est scindée en deux zones Z_1 et Z_2 comme indiqué dans le paragraphe II-2-1 de ce chapitre.

II-2-5 Schéma de l'algorithme de scission systématique

Les données en entrée de l'algorithme sont les paramètres de régularité α et σ_K , le nombre de simulations n_0 à opérer en chaque point (nombre de tirages requis pour l'estimation de la variance empirique $\sigma_e^2(\theta)$), le nombre d'étapes n de l'algorithme. Nous supposons également que la zone de recherche Z_0 est connue et que les premiers tirages ont été réalisés aux sommets de cette zone.

L'algorithme 1 récapitule le procédé général de scission des zones. L'algorithme permet de construire à chaque étape j l'ensemble des zones z_j formant une partition de la zone de recherche initiale Z_0 . Par construction, les nouveaux sommets de scission sont toujours explorés, de sorte que l'on a exploré, en fin d'algorithme, l'ensemble des sommets de la partition finale z_n . Nous rappelons ici que l'ensemble des sommets d'une zone Z est noté $S(Z)$. L'algorithme est ici présenté de façon synthétique. Chacune des étapes de l'algorithme est détaillée dans les sections précédentes. L'exploitation des nombreuses données disponibles en sortie de l'algorithme est détaillée dans la sous-section II-2-6 ci-après.

Algorithme 1 : algorithme à scission systématique

Entrée: σ_K, α, n_0, n

Entrée: $z_0 = \{Z_0\}$

pour $j = 0$ à $n - 1$

choix d'une zone Z^+ de z_j (cf. § II-2-2)

calculer $\hat{m}^*$ et σ_{m^*}

$\forall Z_i \in z_j$, calculer $\beta(Z_i)$

piocher une zone Z^+ en fonction des $\{\beta(Z_i)\}_{Z_i \in z_j}$

choix d'un sommet de scission θ^+ de Z^+ (cf. § II-2-4)

piocher un sommet de scission θ^+

scinder les zones contenant θ^+

calculer z_{j+1} , ensemble des nouvelles zones

tirages au sommet $\theta^+ \in Z^+$

calculer n_0 tirages de $F(\theta^+)$

mise à jour facultative du couple (σ_K, α) , éventuellement par zone (cf. § II-2-3)

fin pour

Sortie: $\hat{m}^*, \sigma_{m^*}$

Sortie: $\forall Z \in z_n, V(Z), \beta(Z)$

Sortie: $\forall Z \in z_n, \forall \theta \in S(Z), \hat{f}(\theta), \sigma_e(\theta), \beta_z(\theta)$

II-2-6 Résultat final et critère de convergence

Zones de confiance Nous supposons ici que l'algorithme de scission systématique est terminé, par atteinte du critère d'arrêt proposé (nombre suffisant de tirages ici) : plus aucun tirage de F ne sera donc opéré.

Un estimateur de la valeur m^* de l'unique minimum global de f est fourni par l'algorithme. Pour autant, nous recherchons essentiellement à agir sur les paramètres, c'est-à-dire à obtenir l'ensemble des paramètres susceptibles de conduire à ce minimum, ce qui nous donnera également une indication sur la fiabilité du résultat obtenu et sur les investigations futures à opérer, par exemple pour départager deux candidats potentiels. Nous cherchons donc un ensemble discret (car il doit être traité numériquement) s'approchant (selon une mesure qui sera définie ultérieurement) de :

$$S_x = \{\theta \in \Theta, E[F(\theta)] \leq x\}$$

pour tout x appartenant à un voisinage de m^* .

A l'issue de l'algorithme, le domaine de recherche initial Z_0 est scindé en un ensemble de zones z . La définition des zones de confiance sera facilitée si nous définissons le potentiel d'un point y compris sur les frontières qui peuvent appartenir à plusieurs zones à la fois (sur les faces des simplexes). Les points appartenant à plus d'une zone formeront un ensemble de volume nul. Toutefois, afin de lever toute ambiguïté pour ces points particuliers, nous définirons pour chaque point une unique valeur de potentiel :

$$\beta_z(\theta) = \max_{Z \in z, \theta \in Z} \beta_Z(\theta)$$

L'ensemble des paramètres admissibles, que nous nommerons également zone de confiance, sera défini pour un seuil $s \in [0,1]$ comme l'ensemble des points candidats θ pour lesquels $f(\theta)$ est potentiellement inférieur à m^* :

$$\hat{S}_{m^*,s} = \{\theta \in \Theta_0, \beta_z(\theta) \geq s\}$$

L'ensemble Θ_0 des points candidats pourra être l'ensemble des points pour lesquels des tirages ont été opérés, ou bien, lorsque cela facilite l'usage futur de $\hat{S}_{m^*,s}$, un ensemble de points régulièrement espacés, pour un pas donné strictement positif δ :

$$\Theta_0 = \{\theta \in \Theta, \frac{1}{\delta} \theta \in \mathbb{N}^d\}$$

Lorsque s est égal à 0, tous les points de Θ_0 sont retenus. Lorsque s augmente, l'ensemble $\hat{S}_{m^*,s}$ se restreint à l'ensemble des paramètres pour lesquels $\hat{f}(\theta)$ est très proche de m^* .

En résumé, nous obtenons finalement :

- la valeur m^* de l'optimum de f ,
- l'ensemble des paramètres $\hat{S}_{m^*,s}$ susceptibles de conduire à cet optimum.

Critère de convergence en connaissance du résultat Nous allons dans un premier temps chercher un critère de convergence de l'algorithme lorsque nous connaissons l'ensemble de solutions S_{m^*} . Un tel critère facilitera la compréhension de l'algorithme sur des fonctions de test, et constituera un outil de comparaison de différents algorithmes.

Nous cherchons à produire un ensemble $\hat{S}_{m^*}$ qui donne une représentation fidèle de S_{m^*} . Par fidèle, nous imaginons d'une part que tout point du véritable ensemble solution S_{m^*} doit être proche d'un point de l'ensemble solution proposé $\hat{S}_{m^*}$: $\hat{S}_{m^*}$ doit être suffisamment grand (condition n° 1). D'autre part, l'ensemble proposé ne doit pas non plus contenir de points trop éloignés des véritables solutions, et tout point de l'ensemble proposé $\hat{S}_{m^*}$ doit être proche d'un point de S_{m^*} : $\hat{S}_{m^*}$ ne doit pas être trop grand (condition n° 2). Cela nous conduira à proposer, comme distance entre les deux ensembles S_{m^*} et $\hat{S}_{m^*}$ la distance de Hausdorff, pour $X \subset \Theta$, $Y \subset \Theta$:

$$d^H(X, Y) = \max \left\{ \sup_{y \in Y} \inf_{x \in X} d(x, y), \sup_{x \in X} \inf_{y \in Y} d(x, y) \right\}$$

Si $X = \hat{S}_{m^*,s}$ est l'ensemble proposé et $Y = S_{m^*}$ est l'ensemble cible, alors majorer le premier terme du max indique que tout point de la cible Y est proche d'un point de X (condition n° 1). Majorer le second terme du max indique que tout point de X est proche d'un point de Y (condition n° 2). Finalement, un point que nous proposons comme solution ne doit pas être trop éloigné d'une solution réelle, et une solution réelle ne doit pas être trop éloignée d'un point proposé comme solution. La distance de Hausdorff est donc parfaitement adaptée à l'objectif recherché. Cette distance représente, pour le pire point de l'un des ensembles, la distance de ce point à l'autre des deux ensembles : elle fournit directement une idée de l'incertitude sur l'ensemble des paramètres conduisant à l'optimum. Nous obtenons donc un critère de convergence en connaissance du résultat recherché S_{m^*} :

$$\rho_1(s) = d^H(\hat{S}_{m^*,s}, S_{m^*})$$

Le critère précédant dépendant de la mesure choisie pour le potentiel, nous proposons également l'usage d'une distance de Hausdorff partielle :

$$\rho_2 = \sup_{y \in S_{m^*}} \inf_{x \in \Theta_e} d(x, y)$$

Où Θ_e représente l'ensemble des points explorés. Ce critère fournit la pire distance d'un point solution au plus proche point exploré. Il indique donc si tous les points solutions ont bien été explorés, et serait naturellement très bon si Θ_e recouvrait Θ . Ce critère n'a de sens que dans la mesure où le nombre de tirages total de F est limité : il pourra notamment servir à comparer différents algorithmes pour un même budget de tirages. Son avantage est de ne requérir que l'ensemble des points de tirage successifs de F . Une limite de ce critère est qu'un algorithme peut converger très vite vers un optimum global sans avoir exploré les zones candidates, et donc en ayant pris un risque important : trouver rapidement le vrai optimum

global d'une fonction n'indique pas si la méthode est prudente ou non, et seule la considération de la variabilité de la fonction entre les points de tirage permet de trancher cette question (en un sens l'usage d'un potentiel). Une autre limite de ce critère est qu'il ne tient pas compte de la précision de l'estimation de f aux points explorés. Nous avons néanmoins tenu à le présenter dans la mesure où il s'agit d'un des critères permettant de comparer l'algorithme proposé à d'autres algorithmes ne permettant pas le calcul de potentiels.

Il faut noter que du fait du caractère aléatoire de la fonction observée et de l'algorithme proposé, ces critères devraient varier selon les exécutions de l'algorithme. En toute rigueur, pour un critère ρ , la comparaison de plusieurs algorithmes sur une fonction test, pour un même budget de tirages autorisé, devrait requérir l'obtention des distributions de ρ pour chaque algorithme, puis l'usage d'un indicateur de risque, comme un quantile de ρ . Une piste pour le choix opérationnel des paramètres de l'algorithme est l'étude approfondie du comportement du critère choisi en fonction de ces paramètres.

Critère de convergence hors connaissance du résultat Dans la pratique, nous ne connaissons pas l'ensemble cible S_{m^*} , puisque nous cherchons précisément à en cerner les contours.

Si nous cherchons à assurer que la fonction est suffisamment bien connue sur chaque zone, un indicateur de convergence est simplement la quantité :

$$\rho_3 = \max_{Z \in \mathcal{Z}} \beta(Z)$$

avec le choix ici opéré $\beta(Z) = V(Z) \cdot \beta_Z(\theta_{B_Z})$, θ_{B_Z} isobarycentre de la zone Z (cf. section II-2 de ce chapitre).

Majorer cet indicateur garantira en effet que pour toute zone $Z \in \mathcal{Z}$:

- Soit que la présence d'un minimum global au point central est très peu probable, $\beta_Z(\theta_{B_Z})$ étant suffisamment faible.
- Soit que la zone est suffisamment petite et que le minimum potentiel a bien été exploré, $V(Z)$ étant suffisamment faible.

Enfin, si nous cherchons à assurer que chaque point ne conduira pas à un minimum plus petit que celui connu, pour un écart η donné, un critère pourra être :

$$\rho_4 = \max_{\theta \in \Theta_0} \beta_z^{(m^*-\eta)}(\theta)$$

Nous notons ici $\beta_z^{(m^*)}(\theta)$ le potentiel au point θ obtenu pour un minimum global estimé m^* (l'indice supérieur (m^*) était resté implicite jusqu'à présent pour ne pas alourdir les notations). D'autres critères peuvent naturellement être envisagés. Un élément incontournable dans la définition d'un critère est la modélisation de $f(\theta)$ ou de l'incertitude de $\hat{f}(\theta)$ en dehors des points θ d'observation : si f est ou est supposée extrêmement erratique, toute zone sera susceptible de contenir un point conduisant à un optimum global, et le critère devra être très différent de celui obtenu en supposant f très régulière. C'est un avantage des critères proposés ici, à l'exception de ρ_2 , mais c'est également une limitation puisque ces critères ne s'appliqueront qu'aux algorithmes proposant une partition de l'espace en zones affectées d'un potentiel.

II-3 Une grille à pas variable avec retraitage possible

Jusqu'à présent, chaque zone était systématiquement subdivisée en deux. Or, il peut être plus avantageux, plutôt que de subdiviser une zone, d'ajouter des simulations en des points déjà explorés. Ce sera l'objet de cette sous-section.

II-3-1 Critères de scission

Supposons que nous estimons une espérance et un intervalle de confiance de la fonction $f(\theta)$, aux points explorés θ_1 et θ_2 . Supposons que θ_c soit le barycentre équipondéré (isobarycentre) entre ces deux points.

La question que l'on se pose est la suivante : au vu des intervalles de confiance de $f(\theta)$, connus aux points θ_1 et θ_2 , et connaissant la variabilité de l'évolution de $f(\theta)$ sur ce segment, mesurée par les paramètres $(\alpha; \sigma_K)$, nous cherchons un critère permettant de déterminer la décision à prendre entre les suivantes :

- diminuer l'incertitude sur $\hat{f}(\theta_i)$ par de nouveaux tirages de F en θ_i , $i \in \{1, 2\}$.
- scinder le segment en deux parties, en opérant des tirages de F au nouveau point θ_c .

Plusieurs critères peuvent être retenus pour opérer ce choix. Le choix du critère retenu au final dépend essentiellement de l'objectif poursuivi : s'assurer que le minimum ne peut être présent dans une zone sous l'hypothèse que la fonction objectif ne varie pas trop brusquement, minimiser le maximum des potentiels d'une zone, etc. Ce choix pouvant dépendre du problème considéré, nous détaillerons ci-dessous deux critères dont nous étudierons la performance dans la sous-section II-4 de ce chapitre.

Comparaison des erreurs de grille et d'estimation Considérons une zone Z . En un point θ de Z , le potentiel est la mesure retenue pour quantifier la probabilité que $\tilde{f}(\theta)$ soit plus petite que l'optimum global $\hat{m}^*$ estimé. Cette mesure dépend à la fois de l'erreur d'estimation aux sommets explorés adjacents, et à la fois de l'erreur de grille, liée à l'éloignement entre θ et les sommets adjacents (et dépendant de la régularité supposée de $\tilde{f}$, mesurée par α et σ_K). Par souci de simplicité, l'incertitude liée à l'estimation du minimum global m^* ne sera pas ici scindée en erreur d'estimation et erreur de grille.

Considérons l'*erreur d'échantillonnage*. Cette erreur est liée à la mauvaise connaissance de la fonction f aux points explorés, du fait du nombre réduit de tirages et du bruit frappant f . Elle dépend essentiellement de la probabilité de présence du minimum aux sommets θ_i de Z du fait de trop larges intervalles de confiance aux points explorés. En l'absence de doute lié à l'évolution de la fonction entre ces points, cette probabilité peut être mesurée par le potentiel $\beta(\theta)$ calculé en l'absence d'erreur de grille :

$$\beta_e(\theta) = \beta(\theta) /_{\sigma_K=0}$$

Considérons l'*erreur de grille*. Cette erreur dépend essentiellement de la possibilité que la fonction évolue brusquement entre un sommet de Z et le point θ considéré. En l'absence de doute lié à la mauvaise connaissance de la fonction aux extrémités du segment, compte tenu de la seule mauvaise exploration de la fonction, cette erreur peut être estimée au point θ par le

potentiel en θ , calculé en l'absence d'erreur d'estimation aux $d + 1$ sommets θ_i de la zone Z :

$$\beta_g(\theta) = \beta(\theta) / \sigma_{e_i=0, i=1, \dots, d+1}$$

Envisageons de scinder un segment autour d'un point barycentre non encore exploré θ_c . Si le potentiel en ce point est grand essentiellement à cause de l'erreur de grille, il apparaîtra raisonnable de scinder le segment et d'explorer θ_c . Si le potentiel est grand essentiellement à cause de l'erreur d'estimation, il paraîtra plus raisonnable au contraire de réduire cette dernière en explorant davantage les sommets du segment. Le critère de scission que nous retiendrons sera donc le suivant :

$$\text{Critère n° 1 : Scission si } \beta_g(\theta_c) \geq \beta_e(\theta_c)$$

En dimension $d = 1$, si θ_c est l'isobarycentre de θ_1 et θ_2 , alors nous pouvons noter $d_0 = d(\theta_c, \theta_1) = d(\theta_c, \theta_2)$. Si nous supposons $\sigma_e = \sigma_e(\theta_1) = \sigma_e(\theta_2)$, compte tenu de la définition retenue ici pour le potentiel d'une zone, nous montrons facilement que l'équivalence suivante est obtenue :

$$(\beta_g(\theta_c) \geq \beta_e(\theta_c)) \Leftrightarrow (\sigma_e \leq \sigma_K d_0^\alpha)$$

Ce dernier critère s'interprète alors très aisément : la scission est opérée si l'erreur de grille ($\sigma_g = \sigma_K d_0^\alpha$) est supérieure à l'erreur d'estimation (σ_e). En dimension $d > 1$, le critère $\beta_g(\theta) \geq \beta_e(\theta)$ permet de tenir compte en un point θ des différences éventuelles de distances aux sommets adjacents, et des différences d'erreurs d'estimation aux différents sommets de la zone Z : Ce critère revient à opérer une moyenne particulière entre les erreurs d'estimation σ_e et les erreurs de grille σ_g en direction des différents sommets.

En toute logique, si F est aléatoire et si $\sigma_K = 0$, alors la supposition de la connaissance parfaite de f entre les points de tirage conduit à seulement mieux estimer $\hat{f}$ aux sommets de la zone Z . Par ailleurs, en l'absence d'erreur d'estimation (par exemple si F est déterministe), le critère conduit bien à une scission systématique, il devient naturellement inutile d'effectuer un retraitage.

Meilleur potentiel après ajout de n_0 simulations Nous envisageons ici le cas où la répartition se fait en fonction de l'amélioration du potentiel maximal de la zone, scindée ou non. Nous supposons que n_0 simulations seront ajoutés sur un sommet de la zone choisie Z^+ . Les erreurs d'estimation correspondent à des écarts-type de variables aléatoires supposées de loi normales. Il est donc simple d'estimer quelle sera la réduction de la variance empirique en cas d'ajout de n_0 simulations, avant même la réalisation de ces simulations. Ainsi, en un point déjà exploré θ , si n_θ tirages de F ont déjà eu lieu et si $\sigma_F(\theta)$ représente la variance empirique de F en ce point, nous pourrions proposer $\sigma_e(\theta) = \sigma_F(\theta) / \sqrt{n_\theta + n_0}$.

- Si nous envisageons d'ajouter n_0 simulations en un sommet existant θ de la zone choisie, il est donc aisé d'estimer la nouvelle erreur d'échantillonnage $\sigma_e(\theta)$ en cas d'ajout.
- Si nous envisageons d'ajouter n_0 points sur le barycentre θ_c du segment choisi, nous pouvons alors estimer $\sigma_F(\theta_c)$ par interpolation linéaire, avant l'exploration du point θ_c , et en déduire une estimation de l'erreur d'échantillonnage $\sigma_e(\theta_c) = \sigma_F(\theta) / \sqrt{n_0}$ en cas d'exploration.

Quel que soit le sommet choisi pour l'ajout parmi les $d + 1$ sommets existants, nous pouvons donc donner un nouvel estimateur $\hat{\beta}(Z^+)$ du potentiel de la zone choisie après ajout de n_0 simulations. Si un sommet existant doit être privilégié pour l'ajout, il semble logique de convenir d'un ajout sur celui de ces $d + 1$ points qui conduit à minimiser cette valeur estimée, ce minimum étant noté $\hat{\beta}_{min}$:

$$\hat{\beta}_{min} = \min_{\theta_i \in S(Z^+)} \left\{ \hat{\beta}(Z^+) \right\}_{n_0 \text{ ajouts en } \theta_i}$$

Si au contraire une scission est envisagée, deux zones Z_1 et Z_2 seront formées. La scission pourra être privilégiée, par exemple, si les potentiels estimés sur ces deux zones sont tous deux inférieurs à $\hat{\beta}_{min}$:

$$\text{Critère n° 2 : Scission si } \max(\hat{\beta}(Z_1), \hat{\beta}(Z_2)) < \hat{\beta}_{min}$$

L'esprit étant alors qu'une scission ne doit pas conduire à augmenter le maximum des $\beta(Z_i)$ sur l'ensemble des zones $\{Z_i\}$ à l'étape considérée : l'ajout se porte ainsi sur le sommet dont nous estimons qu'il minimise le maximum des potentiels sur l'ensemble des zones.

Cette répartition est également pleinement cohérente avec la définition choisie des coefficients β , et elle conduit, par construction, à une diminution du potentiel maximum estimé sur l'ensemble des zones, ce qui est de nature à faciliter les futures démonstrations de la convergence de l'algorithme, notamment si le critère de convergence ρ_3 est utilisé.

II-3-2 Schéma de l'algorithme à tirage possible

Algorithme 2 : algorithme à tirage possible

Entrée: σ_K, α, n_0, n

Entrée: $z_0 = \{Z_0\}$

pour $j = 0$ à $n - 1$

choix d'une zone Z^+ de z_j (cf. § II-2-2)

calculer $\hat{m}^*$ et σ_{m^*}

$\forall Z_i \in z_j$, calculer $\beta(Z_i)$

piocher une zone Z^+ en fonction des $\{\beta(Z_i)\}_{Z_i \in z_j}$

si critère scission (Z^+) vrai (cf. § II-3-1) *alors*

choix d'un sommet de scission θ^+ de Z^+ (cf. § II-2-4)

piocher un sommet de scission θ^+
scinder les zones contenant θ^+
calculer z_{j+1} , ensemble des nouvelles zones
sinon
choix d'un sommet de retraitage θ^+ de Z^+ (cf. § II-3-1)
piocher un sommet déjà exploré $\theta^+ \in S(Z^+)$
 $z_{j+1} = z_j$, les zones sont inchangées
fin si
tirages au sommet $\theta^+ \in Z^+$
calculer n_0 tirages de $F(\theta^+)$
mise à jour facultative du couple (σ_K, α) , éventuellement par zone (cf. § II-2-3)
fin pour
Sortie: $\hat{m}^*$ et σ_m^*
Sortie: $\forall Z \in z_n, V(Z), \beta(Z)$
Sortie: $\forall Z \in z_n, \forall \theta \in S(Z), \hat{f}(\theta), \sigma_e(\theta), \beta_z(\theta)$

L'algorithme 2 récapitule le procédé général de scission des zones. L'algorithme permet de construire à chaque étape j l'ensemble des zones z_j formant une partition de la zone de recherche initiale Z_0 . Cet algorithme diffère de l'algorithme 1 par sa faculté d'opérer un nouveau tirage en un sommet déjà exploré de la zone à explorer Z^+ , plutôt que de systématiquement scinder Z^+ .

II-4 Applications numériques

II-4-1 Fonction test utilisée

Nous présenterons ici une application partant d'une fonction test connue, afin notamment de vérifier le bon comportement de l'algorithme. Par souci de lisibilité, cette application sera tout d'abord présentée en dimension $d = 2$ (mais l'algorithme présenté dans ce travail est proposé pour toute dimension $d \in \mathbb{N}^*$, et le problème de la montée en dimension est évoqué avec plus de détail dans Rulière et al. [2010]).

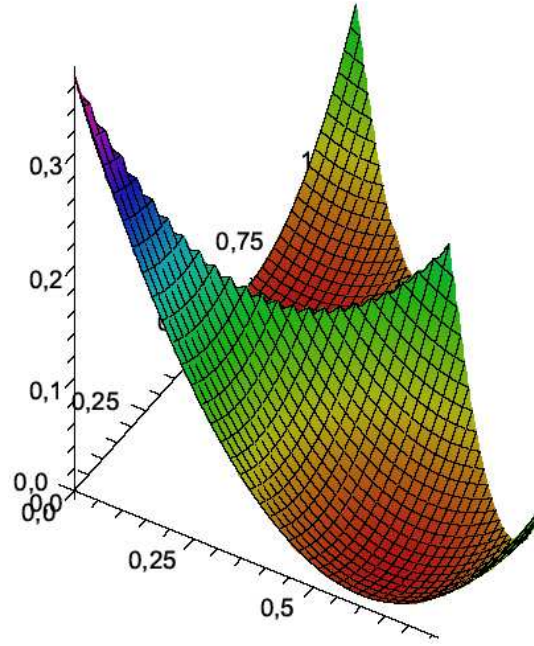


Fig. 48 : Allure générale de la fonction test $f(\theta)$ pour $\theta \in Z_0$

Nous considérerons ici la fonction suivante dans le cas $d = 2$ (correspondant par exemple à 3 poids d'allocation d'actifs se sommant à 1).

$$f(\theta) = (\min(x, y) - 0.1)^2 + (\max(x, y) - 0.6)^2$$

$$F(\theta) = f(\theta) + \sigma_B(U - 0.5)$$

$$\theta = (x, y)$$

avec U une variable aléatoire uniforme sur $[0, 1]$. L'espérance f de la fonction F admet deux minima, l'un en $\theta_1^* = (0.1, 0.6)$, l'autre en $\theta_2^* = (0.6, 0.1)$, la valeur de f étant alors 0. A titre indicatif, afin d'imaginer les variations possibles de cette fonction, la fonction f atteint son maximum sur Z_0 en $\theta = (0, 0)$ et nous avons alors $f(\theta) = 0.37$ (mais nous cherchons ici le minimum, non le maximum). Un bruit pouvant conduire à des variations d'amplitude de 0.1 entre deux points proches est donc assez élevé au regard du domaine de variation $[0, 0.37]$ de la fonction sur la zone initiale de recherche. Un aperçu de la fonction f est donné dans le graphique 48.

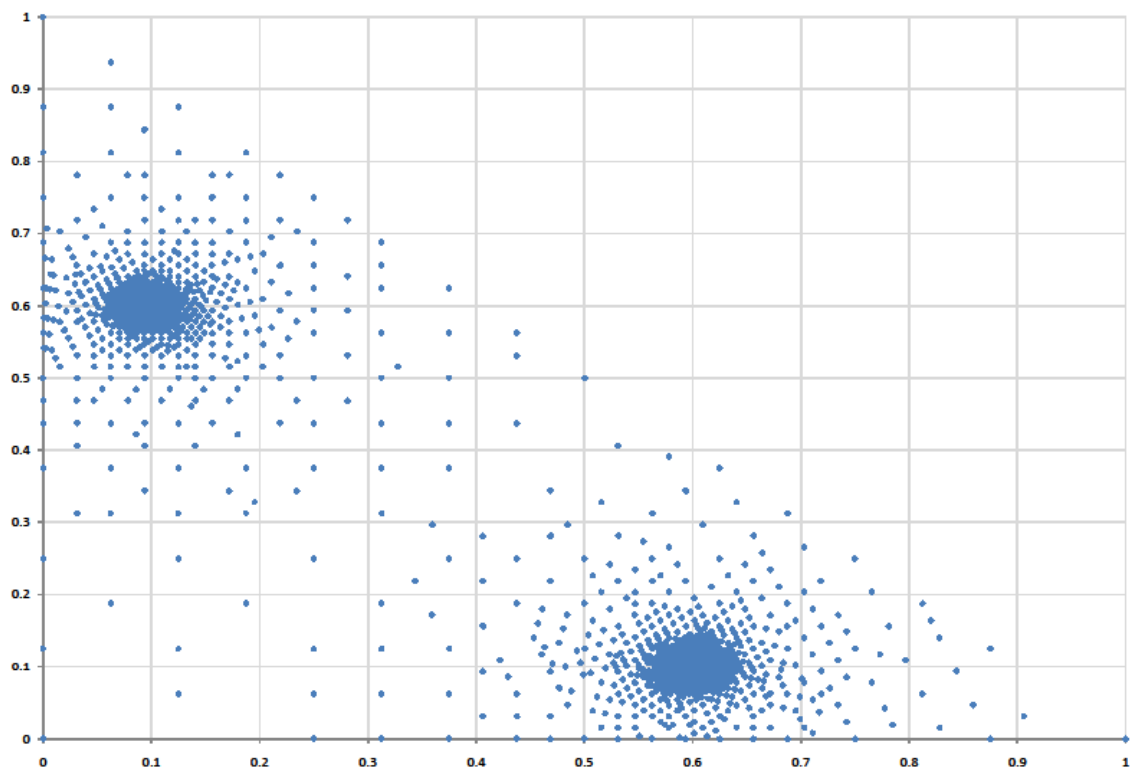


Fig. 49 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0$ (scission systématique)

Dans le cadre de l'algorithme à scission systématique, nous utiliserons ici toujours les paramètres suivants :

- le nombre n de tirages réalisés, sera $n = 2000$,
- le nombre de tirages en chaque point sera $n_0 = 10$,
- le coefficient de priorité sera $\gamma = 1/(d + 1) = 1/3$, conduisant à une moyenne géométrique des potentiels directionnels pour le potentiel d'un point,
- le seuil η du critère de convergence ρ_4 sera $\eta = 0.01$.

Par ailleurs, pour chaque illustration, nous préciserons

- le niveau du bruit σ_B frappant la fonction f ,
- les coefficients de variabilité σ_K et α ,
- le seuil s du critère de convergence ρ_1 .

A titre de remarque, le graphique 47 présentée précédemment a été obtenue avec la fonction F évoquée ci-dessus, à partir des paramètres $(\sigma_K, \alpha) = (1, 1.5)$ et $\sigma_B = 0.1$ (et $\gamma = 1$ dans ce seul cas).

II-4-2 Comportement de l'algorithme à scission systématique

Dans cette sous-section, nous allons observer le comportement de l'algorithme n° 1 à scission systématique, lorsque le niveau de variabilité σ_K varie, pour différents niveaux de bruit σ_B .

Dans ce paragraphe, le paramètre α est ici fixé, égal à 1.5, et correspond à un modèle où les pentes sont supposées suivre un mouvement brownien depuis un point d'accroche (cf. définition du potentiel directionnel dans les paragraphes II-2-2 et II-2-3 de ce chapitre).

A titre indicatif, nous mentionnerons dans la légende de certains graphiques la valeur obtenue pour le vecteur critère de convergence, $\vec{\rho} = (\rho_1(s), \rho_2, \rho_3, \rho_4(\eta))$, avec s égal à 10 % du potentiel maximal observé, et $\eta = 0.01$. Le détail du calcul de ces critères ainsi qu'une analyse plus détaillée des valeurs numériques obtenues pour ceux-ci seront abordés dans le paragraphe II-4-5 de ce chapitre.

Pour un niveau de variabilité supposée de la fonction $\sigma_K = 2$, illustré dans les graphiques 49, 50, et 51, lorsque le niveau de bruit passe respectivement par $\sigma_B = 0$, $\sigma_B = 0.1$ et $\sigma_B = 0.3$, nous observons bien l'exploration préférentielle des zones à proximité des minimiseurs de f . Bien entendu, cette exploration est plus étendue lorsque le niveau de bruit augmente, puisque du fait de ce bruit, l'assurance de l'absence d'autres optima locaux nécessite l'exploration d'une zone plus vaste. Il faut noter que dans la pratique, ce niveau de bruit de F est généralement une donnée exogène sur laquelle il n'est pas possible d'agir.

Pour un niveau de bruit fixé à $\sigma_B = 0.1$, illustré dans les graphiques 50, 52 et 53, lorsque la variabilité supposée de la fonction passe successivement par $\sigma_K = 2$, $\sigma_K = 20$ et $\sigma_K = 100$, nous observons un résultat parfaitement logique : si le comportement attendu de la fonction entre les points de tirage est supposé peu variable, l'algorithme se concentre sur les zones proches des minimiseurs supposés. Dans le cas inverse, une variabilité extrême $\sigma_K = 100$ conduit à une répartition quasi uniforme des points d'exploration : avec une telle variabilité, l'optimum global de la fonction peut en effet se trouver dans n'importe quelle zone. Il faut toutefois remarquer que la fixation d'un seuil de variabilité σ_K trop faible conduirait par définition à supposer que le comportement de la fonction est globalement connu entre les points explorés, ce qui conduit à négliger des zones pourtant susceptibles de contenir un optimum. L'arbitrage entre la rapidité de convergence vers les optima globaux et le risque d'avoir négligé une zone dépendra donc largement de ce paramètre σ_K .

Pour α fixé, la fixation du paramètre σ_K est donc un problème délicat, que nous aborderons dans le paragraphe II-4-4. Elle pourra se baser en particulier sur les explorations précédentes de la fonction, sur d'autres connaissances de celles-ci, ainsi que sur les contraintes d'implémentation et de niveau de risque accepté. Il faut noter à ce sujet que rien dans le modèle n'oblige à fixer cette valeur constante sur l'ensemble de la zone de recherche, et qu'il est également possible de tenir compte d'une variabilité particulière de la fonction sur certaines zones.

En résumé, pour l'algorithme n° 1 à scission systématique, l'algorithme vise essentiellement à répartir des points d'exploration, en délaissant temporairement certaines zones non susceptibles de contenir le minimum. Selon les paramètres choisis et la fonction optimisée, le comportement attendu est le suivant, conforme à ce que nous avons pu observer :

- Lorsque le bruit est important devant les variations de la fonction, ou lorsque l'incertitude sur la variabilité de la fonction σ_K est élevée, l'algorithme explore assez

uniformément la fonction, de façon similaire aux traditionnelles grilles de recherche à pas fixe évoquées précédemment.

– Lorsque le bruit n'est pas trop élevé et lorsque la variabilité de la fonction σ_K est faible, les points de tirage se concentrent sur les zones contenant les points solutions supposés, le comportement de la fonction étant supposé sans surprise entre les points déjà explorés.

Le choix du coefficient σ_K de variabilité de la fonction f dépend donc du but poursuivi : trop élevé, l'exploration de la fonction sera très poussée, et le temps de calcul pourra être élevé si nous visons une bonne connaissance d'un optimum de la fonction. Pour un σ_K trop faible, l'algorithme se focalisera très vite sur un optimum local, donc permettra un gain en terme de temps de calcul, à précision égale, mais le risque de ne pas explorer une zone susceptible de contenir un autre optimum, éventuellement meilleur, sera plus élevé.

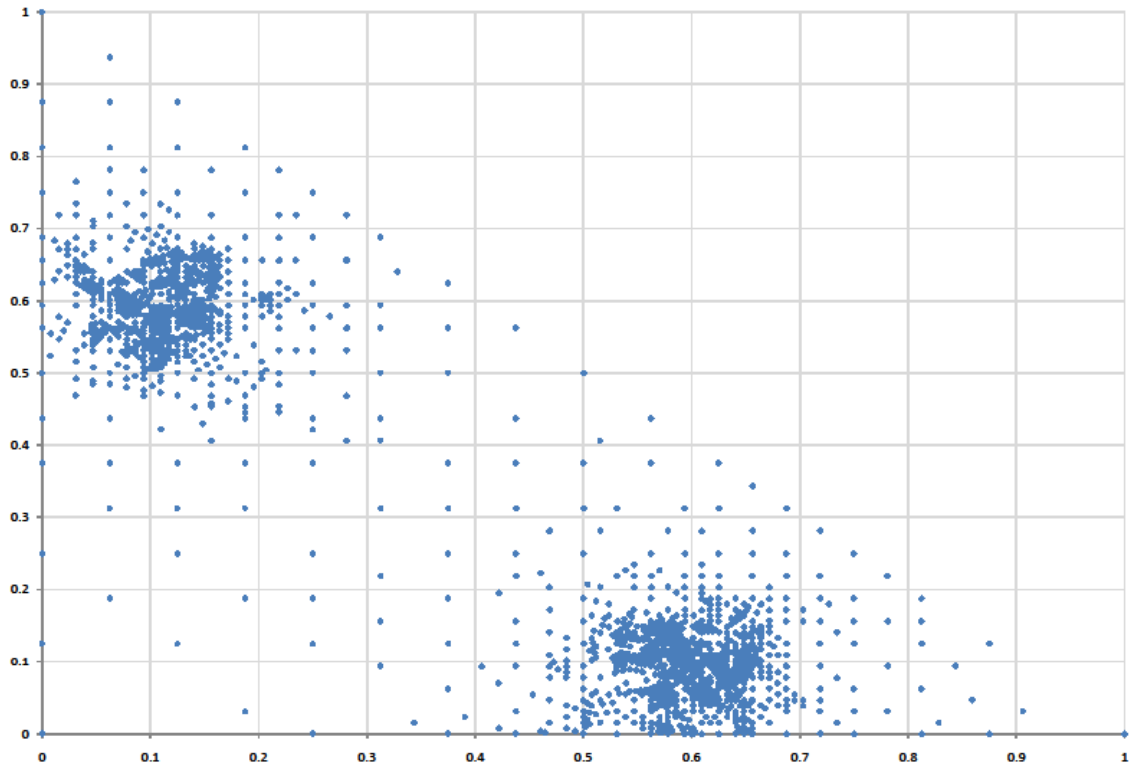


Fig. 50 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\vec{p} = (0.10, 3.5E-03, 6.8E-07, 1.9E-2)$

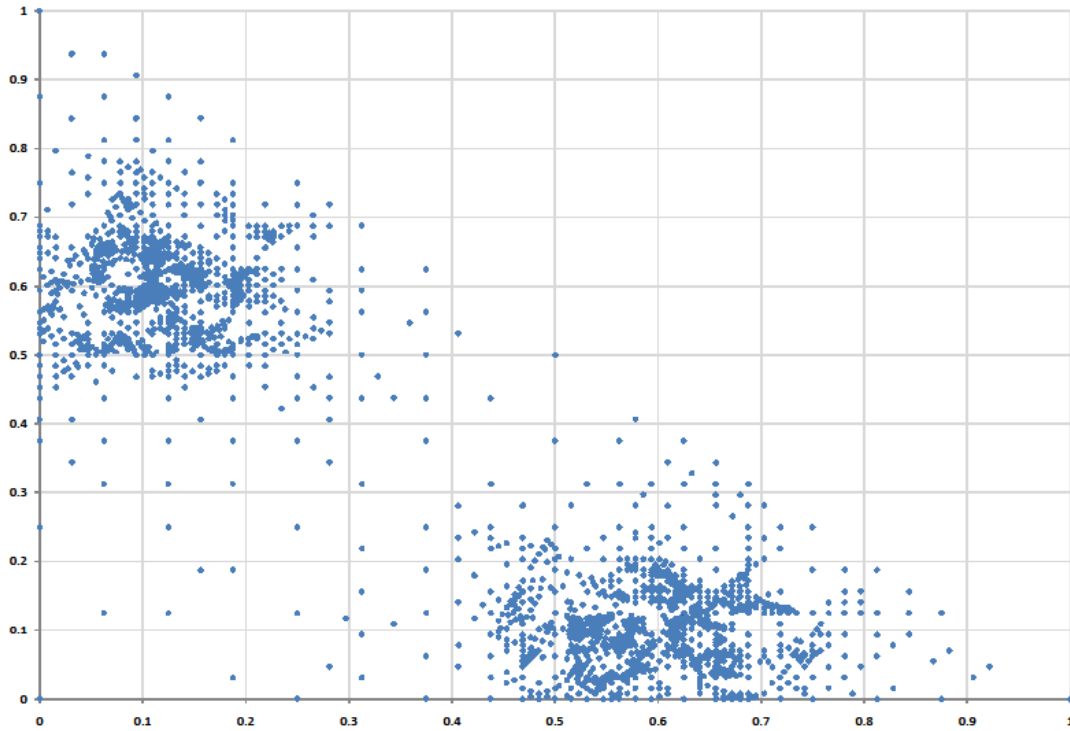


Fig. 51 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0.3$ (scission systématique), $\vec{\rho} = (0.13, 5.9E-03, 1.6E-06, 5.8E-2)$

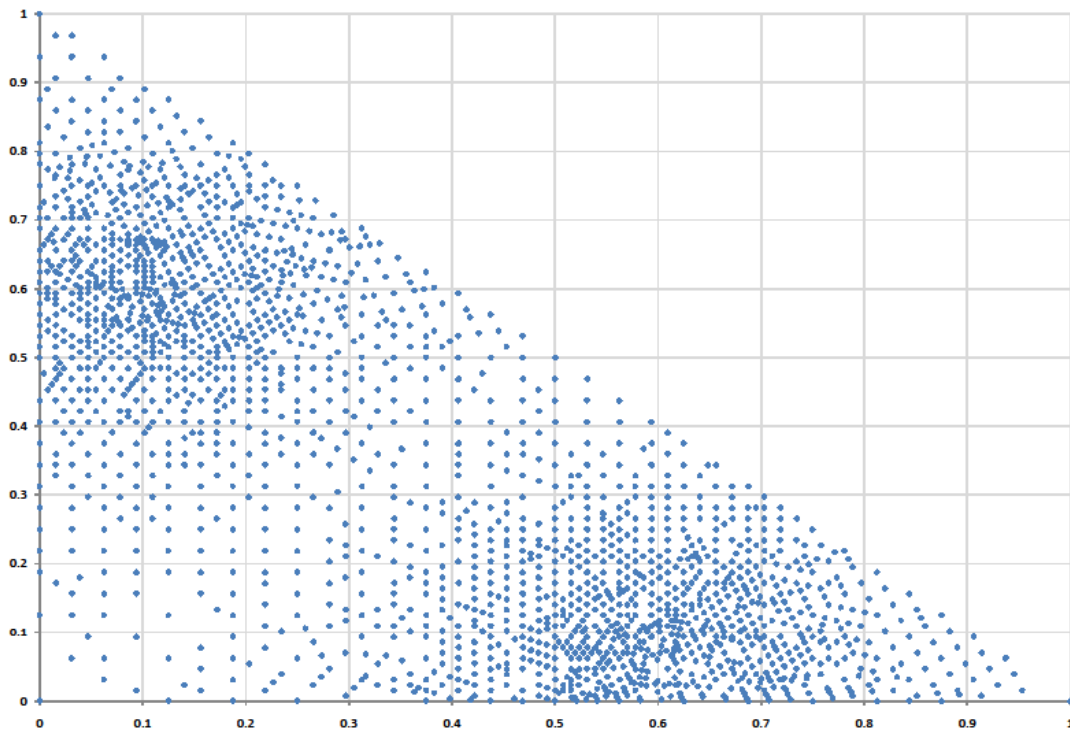


Fig. 52 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 20$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\vec{\rho} = (0.45, 2.2E-03, 7.3E-05, 0.72)$

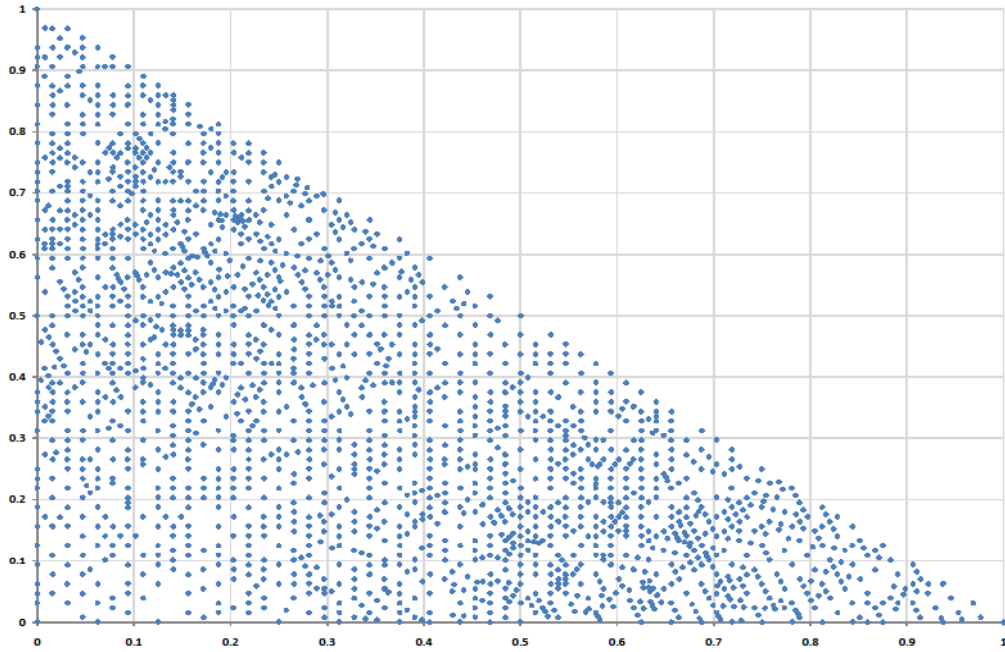


Fig. 53 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 100$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\bar{\rho} = (0.60, 8.8E-03, 5.9E-04, 0.97)$. La recherche de l'optimum d'une fonction constante bruitée conduit à un nuage de points d'allure proche de celui-ci.

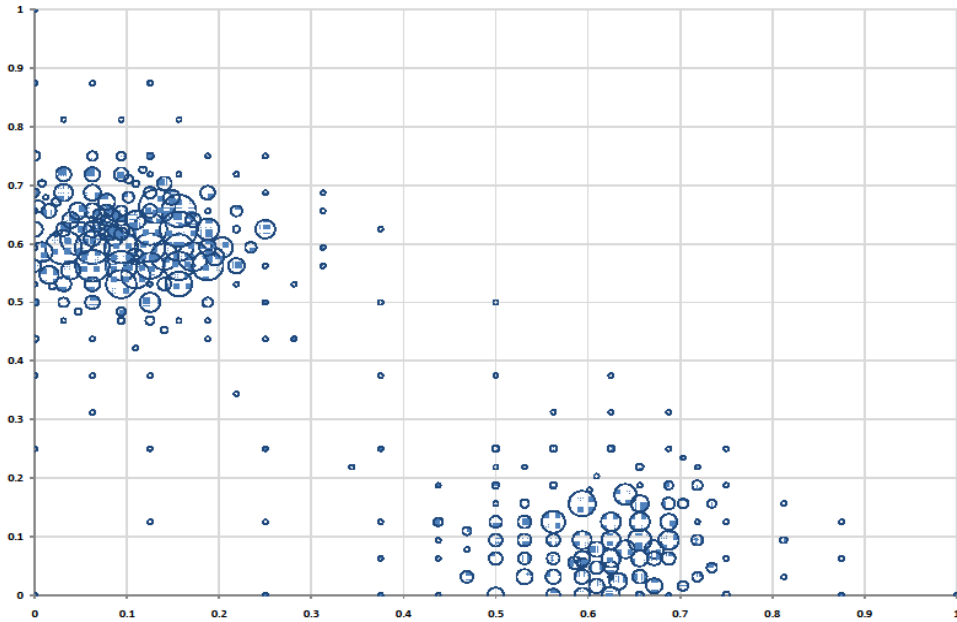


Fig. 54 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.1$, critère de scission n° 2, pour une variabilité $(\sigma_K; \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\bar{\rho} = (0.085, 8.8E-03, 3.1E-05, 8.2E-04)$

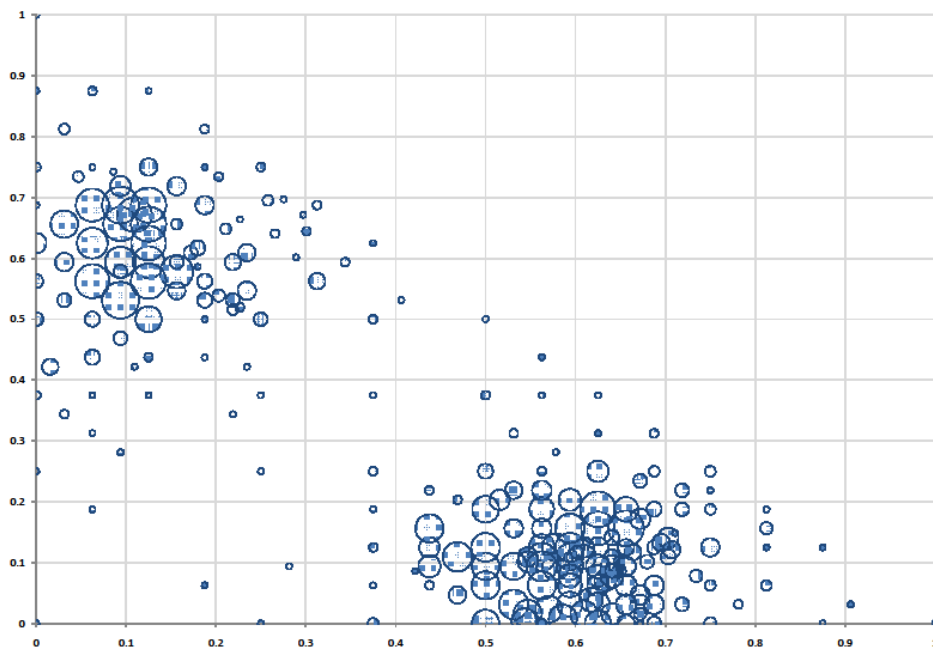


Fig. 55 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.2$, critère de scission n° 2, pour une variabilité $(\sigma_K, \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\vec{p} = (0.10, 8.8E-03, 6.7E-05, 6.2E-03)$

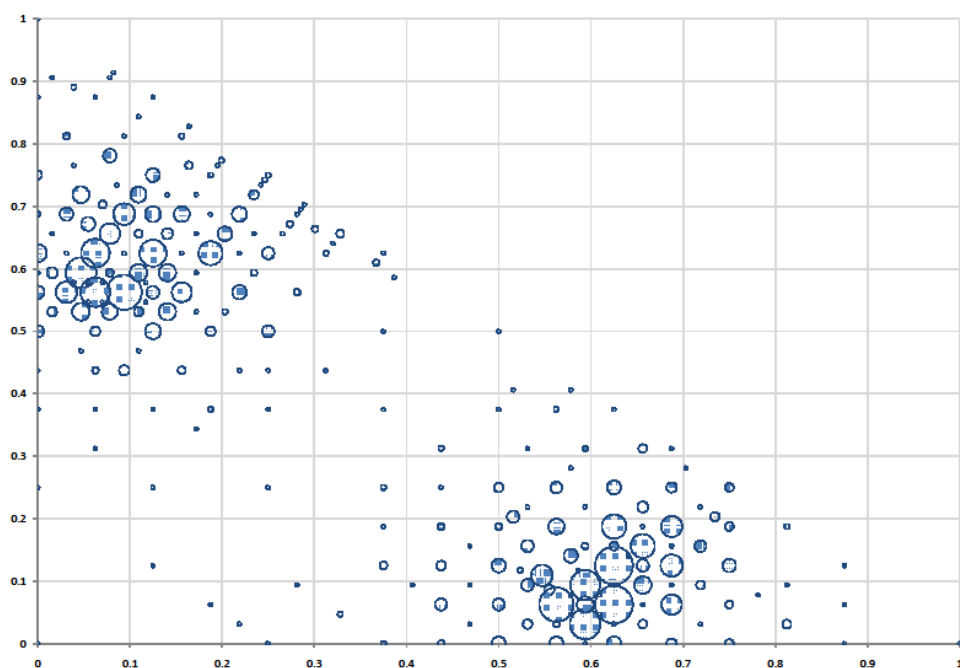


Fig. 56 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.2$, critère de scission no 1, pour une variabilité $(\sigma_K, \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\vec{p} = (0.10, 0.011, 1.5E-05, 1.6E-02)$

II-4-3 Comportement de l'algorithme avec retirages possibles

Les deux graphiques 54 et 55 illustrent les retirages qui peuvent s'opérer sur des sommets déjà explorés, afin d'améliorer la connaissance de la fonction f en ces points. Pour ces figures, nous avons utilisé le critère de scission n° 2 envisageant l'amélioration du potentiel après ajout de n_0 simulations (cf. § II-3-1 de ce chapitre). Lorsque le niveau de bruit est $\sigma_B = 0.1$, le graphique 54 montre que les retirages ont essentiellement lieu autour des points solutions supposés, lorsqu'il n'est plus seulement nécessaire d'avoir une idée de la zone où se situe un minimiseur, mais qu'il faut également estimer avec précision la valeur m^* de l'optimum atteint. Lorsque le bruit augmente, cela tend naturellement à disperser un peu les points d'exploration, et favorise le retraitage, dans la mesure où l'incertitude en certains points, du fait du bruit, devient plus forte que l'incertitude liée à la variation de la fonction entre les points de la grille.

Le graphique 56 illustre un choix de scission sur critère n° 1 comparant directement erreur de grille et erreur d'estimation (cf. § II-3-1 de ce chapitre) : elle donne une idée de l'impact du choix du critère de scission, qui conduit à un comportement proche pour les deux critères. Les retirages sur des sommets existants se font lorsque la présence du minimum est suffisamment vraisemblable et que l'incertitude pesant sur les sommets de la zone doit être diminuée.

Pour les deux critères abordés ici, il faut enfin noter que pour un niveau de bruit nul, aucun retraitage n'intervient (les graphiques illustrant ce résultat, qui conduisent au graphique 49, ont toutefois été omis). Cela est logique dans la mesure où ces retirages n'apporteraient rien à la connaissance de la fonction f .

Pour l'algorithme avec retirages possibles, le comportement attendu est le suivant, conforme à ce qui a été observé :

- Lorsque le bruit est faible, ou à plus forte raison lorsque F est déterministe, la seule incertitude réside dans le comportement de la fonction entre les points, et l'algorithme crée systématiquement de nouveaux points de tirage.
- En présence d'un bruit, les points de retraitage se concentrent dans les zones petites, généralement proches d'un point solution. Pour ces points, la réduction de l'erreur de grille, faible sur de courtes distances, est moins prioritaire que la réduction de l'erreur d'estimation.

	$\sigma_B = 0$ scission syst.		$\sigma_B = 0.1, \sigma_e = 0$ scission syst.		$\sigma_B = 0.1, \sigma_e$ estimé scission syst.		$\sigma_B = 0.1, \sigma_e$ estimé avec retirages.	
	σ_K	α	σ_K	α	σ_K	α	σ_K	α
n=10	0.635	1.505	0.475	1.212	0.474	1.214	0.449	1.277
n=100	0.946	1.592	0.392	1.063	0.645	1.374	0.947	1.585
n=1000	1.074	1.566	0.112	0.515	0.554	1.221	0.599	1.203

Tab. 34 : Coefficients σ_K et α estimées par maximum de vraisemblance, en présence ou non d'un bruit σ_B , à partir de n points explorés

II-4-4 Estimation des paramètres (σ_K, α)

Le tableau 34 représente les estimations des coefficients σ_K et α , opérées à partir d'un nombre n de points explorés, pour des données soumises à un bruit σ_B . L'estimation est réalisée par maximum de vraisemblance, comme indiqué au paragraphe II-2-3 de ce chapitre, à l'aide d'un point fixe qui a convergé dans toutes les situations testées. Ici, seuls les segments délimitant les zones ont été retenus pour l'estimation, et non pas les segments joignant deux points explorés de zones distinctes. Les zones en question ont été obtenues par l'algorithme avec des paramètres initiaux $(\sigma_K, \alpha) = (0.3, 0.9)$.

Dans la première colonne de ce tableau 34, lorsque $\sigma_B = 0$ (et en conséquence $\sigma_e = 0$), nous pouvons constater une relative stabilité des paramètres estimés, même dans le cas où l'estimation se base sur un nombre restreint de segments (lorsque $n = 10$ et en l'absence de considération de segments inter-zones). Dans la deuxième colonne, en italique, lorsque $\sigma_B = 0.1$ et pour tout point $\sigma_e(\theta) = 0$, l'estimation est menée comme si l'estimateur $\hat{f}$ correspondait à la fonction f aux points explorés. Cette colonne 2 est donnée à titre indicatif, dans la mesure où il est parfaitement possible de tenir compte des erreurs d'estimation aux points explorés (cf. colonnes 3 et 4).

La non prise en compte des erreurs d'estimation explique le coefficient α inférieur dans le cas bruité, qui correspond bien à une incertitude accrue sur les petites distances, du fait du bruit supporté par $\hat{f}$. Ce phénomène est d'autant plus marqué que les zones sont petites (ici lorsque n est grand). Dans le cas bruité, le coefficient α inférieur est toutefois compensé par une pente σ_K inférieure. Enfin, dans les troisième et quatrième colonnes, l'estimation est menée sur des données bruitées, en tenant compte cette fois du bruit σ_e aux points explorés. La troisième colonne présente une estimation issue d'un découpage par scission systématique des zones, tandis que la quatrième présente une estimation issue de l'algorithme avec retirages possibles. Les paramètres observés sur ces colonnes 3 et 4 peuvent varier un peu, notamment lorsque le nombre de segments utilisés pour l'estimation est réduit.

Toutefois, dans les observations que nous avons pu mener, une valeur moins élevée de α (volatilité supérieure sur de courtes distances) est systématiquement compensée par une valeur moins élevée de σ_K (volatilité globale inférieure) : nous aborderons ce point par l'observation des écarts-types de la variable aléatoire IR_α dans le graphique 58.

Le comportement de l'algorithme initialisé avec les différents paramètres optimaux obtenus est resté assez stable dans les exemples que nous avons testés. C'est ce qu'indique le graphique 57, qui présente les points d'exploration obtenus, lorsque $\sigma_B = 0.1$, pour l'algorithme initialisé avec les paramètres estimés hors bruit pour $n = 1000$, $(\sigma_K, \alpha) = (1.0743, 1.5665)$ (à gauche) ou estimés en présence de bruit $(\sigma_K, \alpha) = (0.5545; 1.2209)$ (à droite). Un coefficient α supérieur permet de focaliser les recherches un peu plus rapidement sur de petites zones, mais la différence de valeur entre les coefficients α testés est ici trop ténue pour que l'effet global sur le comportement de l'algorithme en soit beaucoup affecté.

Enfin, nous avons déterminé, indépendamment de l'algorithme proposé, l'écart-type empirique de la variable aléatoire IR_α :

$$IR_\alpha(Z) = \frac{f(\theta_1) - f(\theta_2)}{d(\theta_1, \theta_2)^\alpha}$$

θ_1, θ_2 vecteurs aléatoires distincts issus de tirages indépendants, uniformes sur Z .

Le détail de la cette variable aléatoire et des tirages uniformes sur une zone Z est donné dans la sous-section II-2-3 de ce chapitre. L'écart-type de cette variable aléatoire fournit des informations sur la variation des pentes observées sur la zone en fonction de la distance considérée. Comme nous pouvons le voir dans le graphique 58, ces écarts-type sont légèrement supérieurs dans le cas de la demi-zone Z_1 correspondant à la partie de Θ située au dessus de la droite d'équation $y = x$.

Cela est logique dans la mesure où f a la forme d'une cuvette sur cette zone et non plus de deux cuvettes, les pentes sont ici en moyenne un peu plus abruptes que lorsque nous relient deux points d'une même cuvette, plutôt que deux points de cuvettes différentes. Nous pouvons remarquer que les ordres de grandeurs trouvés par maximum de vraisemblance sont tout à fait conformes à l'écart-type empirique de IR_α , notamment de celui de $IR_\alpha(Z_1)$, dans la mesure où dès la première itération de l'algorithme, le simplexe initial est scindé en une zone Z_1 et sa zone complémentaire.

La variable IR_α n'est pas utilisée par l'algorithme, mais fournit une indication sur l'erreur faite *a priori* sur la régularité de la fonction objectif selon le modèle proposé. Même il ne s'agit pas ici de fournir une étude exhaustive de la distribution de IR_α , qui dépend naturellement de la fonction objectif utilisée, il nous a paru intéressant de montrer quelle pouvait être la nature de cette erreur de modèle sur les données ici testées. Le résultat de cette analyse apparaît dans le graphique 59.

Sur notre fonction objectif, la distribution a posteriori de IR_α s'est révélée finalement assez proche de la distribution *a priori* supposée gaussienne, notamment pour $\alpha = 1.2$. Pour des α plus élevés, la présence de nombreuses valeurs très élevées se traduit par une légère sous-évaluation de la fréquence des pentes très élevées, et une sur-évaluation de la fréquence des pentes moyennes : la distribution observée a posteriori semble leptokurtique. Il est également possible que la distribution de IR_α évolue au fur et à mesure des découpages.

Ces résultats indiquent qu'il appartient à l'observateur d'intégrer une éventuelle prudence, en surestimant par exemple le coefficient σ_K ou en sous-estimant le coefficient α , et qu'il est aussi envisageable de modifier a posteriori la distribution supposée des pentes, ou le *kurtosis* de la distribution, l'hypothèse gaussienne pouvant servir de distribution a priori lors d'une inférence bayésienne.

Un écart important peut également indiquer une erreur dans le choix du coefficient α , le modèle devant conduire à un IR_α proche d'une loi normale pour un α à déterminer, non pour tous. Remarquons toutefois que, convolués avec l'erreur d'estimation en chaque point, les écarts observés peuvent en partie s'estomper. D'autre part, si la hiérarchie entre les zones les

plus susceptibles de contenir un optimum global peut se trouver affectée par le décalage observé, elle n'est pas non plus radicalement remise en cause, le coefficient d'aplatissement de la distribution utilisé ayant peu d'incidence pour des valeurs de f proches de la valeur estimée de l'optimum global (pour une zone dont les sommets conduiraient à $\hat{f}(\theta_i) = \hat{m}^*$, le potentiel en tout point de la zone serait égal à 1, quel que soit le coefficient d'aplatissement de la distribution utilisée) : le potentiel est utilisé pour opérer un arbitrage entre différents points d'exploration possibles, non pour prévoir précisément le comportement de la fonction en ces points.

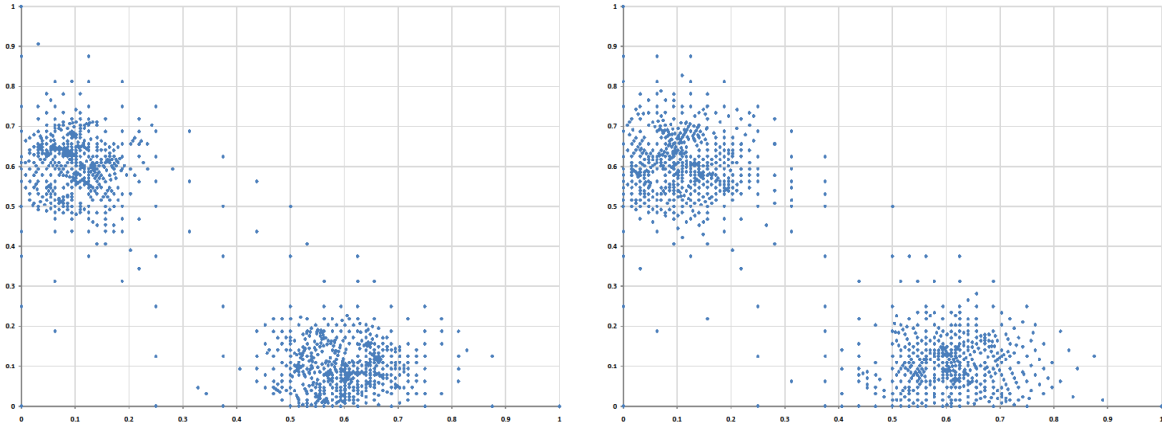


Fig. 57 : Points d'explorations obtenus, lorsque $\sigma_B = 0.1$, avec les paramètres estimés hors bruit, $(\sigma_K, \alpha) = (1.0743, 1.5665)$ (à gauche) ou estimés en présence de bruit $(\sigma_K, \alpha) = (0.5545; 1.2209)$ (à droite)

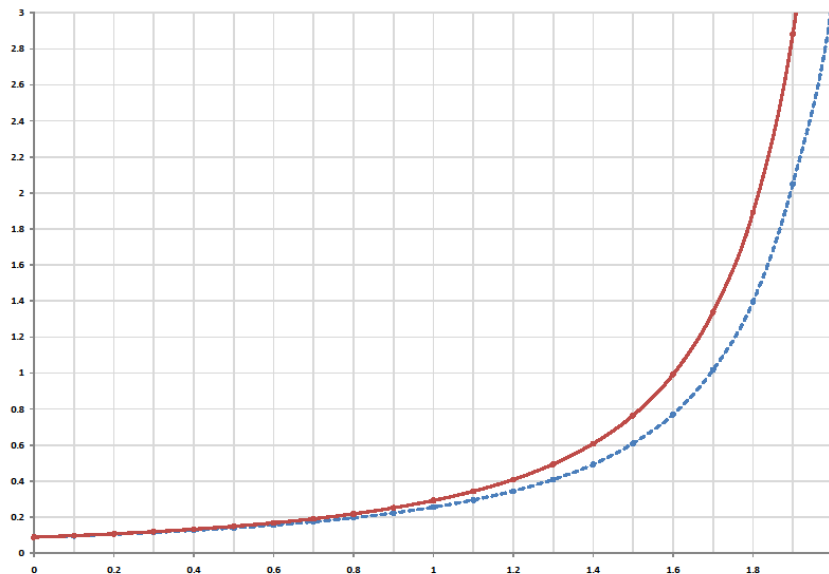


Fig. 58 : Ecart-type empirique de $IR_\alpha(Z)$ obtenu pour la zone Z_0 correspondant au simplexe orthogonal standard initial (en pointillé) et pour la demi-zone Z_1 correspondant à la partie de Z_0 située au dessus de la première bissectrice (trait plein)

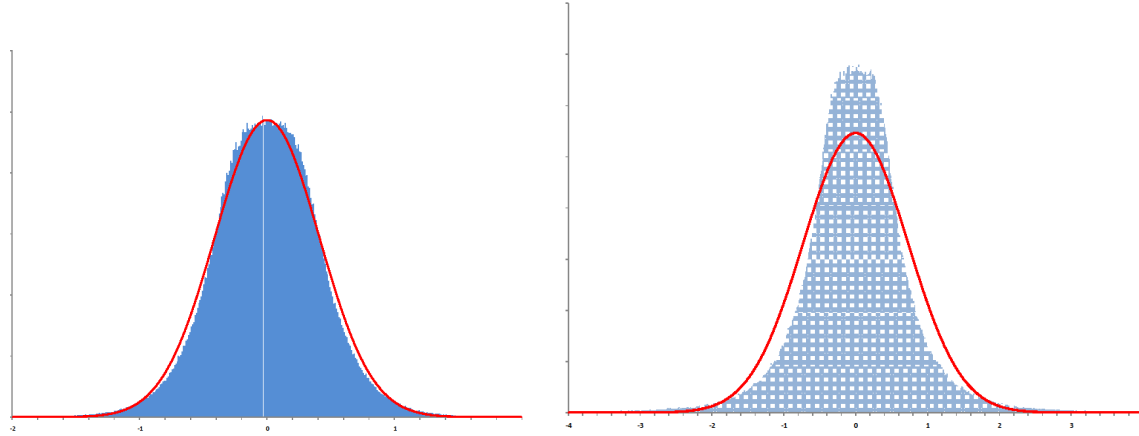


Fig. 59 : Distribution a priori de IR_α (courbe gaussienne continue au premier plan) et histogramme a posteriori, pour la demi-zone Z_1 obtenue après la première itération de l'algorithme. Cas $\alpha = 1.2$ (à gauche) et $\alpha = 1.5$ (à droite). Les écarts-type empiriques obtenus à partir de l'histogramme sont respectivement $\sigma = 0.41$ et $\sigma = 0.73$

II-4-5 Critères de convergence

critères de convergence calculés Le critère $\rho_1(s)$ est défini dans la sous-section II-2-6 de ce chapitre. Il s'appuie sur une zone de confiance proposée pour l'ensemble des solutions, zone comprenant un ensemble de points de potentiel supérieur à un seuil s . Le critère $\rho_1(s)$ indique la plus grande distance possible entre un point de cette zone et le point solution le plus proche : il peut donc s'interpréter comme un rayon maximal de la zone de confiance depuis chaque point solution. Le critère ρ_4 , dépend quant-à-lui d'un seuil η qui sera, dans les applications numériques présentées, toujours fixé ainsi $\eta = 0.01$. Ce critère indique la plus grande probabilité (selon le modèle choisi), qu'un minimum plus petit que $(m^* - \eta)$ soit observé en un point θ (cf. sous-section II-2-6 pour une définition précise).

Dans les applications numériques, pour le calcul des critères ρ_1 et ρ_4 , l'ensemble des points candidats Θ_0 a été construit sans opérer de nouveaux tirages de F , en piochant a posteriori 50 points dans chaque zone (de façon uniforme, cf. Smith et Tromble [2004]), et en calculant leur potentiel en fonction des points explorés aux sommets de chaque zone. Pour le critère ρ_1 , l'ensemble des sommets proposés $\hat{S}_{m^*,s}$ a été extrait de Θ_0 en ne retenant que les points dont le potentiel était supérieur au seuil s fixé. Il est enfin rappelé que le critère ρ_2 est construit sur l'ensemble Θ_e des sommets explorés, et le critère ρ_3 sur l'ensemble des zones construites (cf. § II-2-6) :

$$\hat{S}_{m^*,s} = \{\theta \in \Theta_0, \beta_z(\theta) \geq s\}$$

$$\rho_1(s) = d^H(\hat{S}_{m^*,s}, S_{m^*})$$

$$\rho_2 = \sup_{y \in S_{m^*}} \inf_{x \in \Theta_e} d(x, y)$$

$$\rho_3 = \max_{Z \in Z_n} \beta(Z)$$

$$\rho_4 = \max_{\theta \in \Theta_0} \beta_z^{(m^*-\eta)}(\theta)$$

Fonction objectif non bruitée Bien que cette situation idéale ne soit pas celle nous désirons traiter dans la pratique, nous avons tout d'abord vérifié numériquement la convergence de l'algorithme sur la fonction test non bruitée, dans le cas $\sigma_B = 0$.

Les résultats numériques obtenus pour les critères de convergence apparaissent dans le tableau 35. Du fait de l'absence de bruit, le minimum estimé $\hat{m}^*$ était supposé ne pas souffrir d'erreur d'estimation. Nous avons à titre indicatif ajouté dans le tableau 35 une ligne indiquant les résultats obtenus tenant compte d'une erreur de grille σ_{m^*} pour ce minimum, fonction du diamètre d_* de la zone sur laquelle il était présent, $\sigma_{m^*} = \sigma d_*^\alpha$. La localisation du minimum est alors plus dispersée, conduisant alors à une exploration accrue dans une zone plus étalée autour des minima trouvés, au détriment de la vitesse de convergence.

Nous pouvons notamment constater dans le tableau 35 le très bon comportement de l'algorithme en l'absence de bruit, avec des points explorés à une distance d'ordre 10^{-5} de chacune des solutions, et des zones de confiance proposées dans un rayon d'ordre $\rho_1(5\%) \cong 4.10^{-3}$ ou autour de $\rho_1(90\%) \cong 3.10^{-4}$ ces solutions. La zone de confiance proposée est naturellement plus réduite dans le cas où s est élevé.

	$\rho_1(s)$	ρ_2	ρ_3	ρ_4
$s=90\%$, $\sigma_{m^*}=0$	2.53 E-04	3.45 E-05	2.43 E-08	0
$s=5\%$, $\sigma_{m^*}=0$	4.32 E-03	3.45 E-05	2.43 E-08	0
$s=5\%$, $\sigma_{m^*}>0$ (erreur de grille)	6.46 E-03	1.38 E-04	2.28 E-07	0

Tab. 35 : Critères de convergence obtenus par l'algorithme lorsque $\sigma_B = 0$, pour $\alpha = 1$ et $\alpha = 1.6$, avec ou sans prise en compte d'erreur de grille pour m^*

	$\rho_1(s)$	ρ_2	ρ_3	ρ_4
cas a, $\sigma_B=0.1$, sans retraitage	7.83 E-02	2.21 E-03	1.32 E-06	3.24 E-02
cas b, $\sigma_B=0.1$, retirages	0.11	8.84 E-03	1.82 E-05	3.99 E-02
cas c, $\sigma_B=0.1$, retirages	5.92 E-02	2.21 E-03	6.14 E-06	4.58 E-02

Tab. 36 : Critères de convergence obtenus pour les algorithmes avec ou sans retraitage

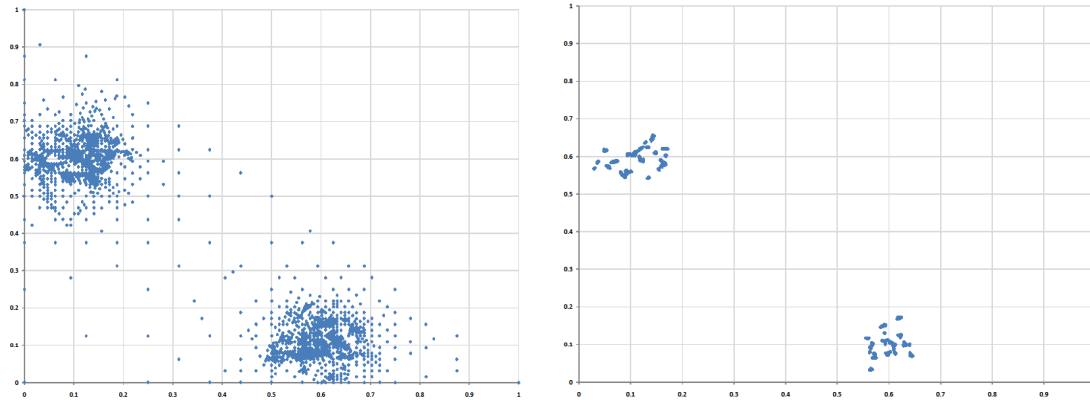


Fig. 60 : Points d'explorations obtenus pour l'algorithme à scission systématique, cas a , lorsque $\sigma_B = 0.1$, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*, 5\%}$ avec un seuil de 5 % (à droite)

Fonction objectif bruitée Nous avons comparé les résultats obtenus par l'algorithme à scission systématique, ainsi que par l'algorithme à tirage pour les critères de scission n° 1 et n° 2. Les graphiques 60, 61 et 62 illustrent cette comparaison. Ces figures ont été obtenues pour un niveau de bruit $\sigma_B = 0.1$, des paramètres de régularités $(\sigma_K, \alpha) = (0.6, 1.2)$, pour $n = 2000$ points de tirage ou tirage, avec un seuil minimal de $n_0 = 10$ tirages par points. Dans ces figures apparaissent les nuages de points de tirage, ainsi que les zones de confiance proposées $S_{m^*, 5\%}$. Du fait de l'absence de prise en compte d'erreur de grille pour m^* , et du fait du seuil s utilisé, il est possible qu'un point solution soit très proche mais en dehors d'une zone de confiance.

Les trois algorithmes se comportent correctement, et conduisent à la proposition de zones de confiance dans le voisinage direct des points solutions. Du fait du bruit frappant la fonction f , nous observons que des zones très proches peuvent être tantôt exclues, tantôt englobées dans $S_{m^*, s}$. Les zones de confiance ainsi bâties sont donc très fractionnées, mais bien localisées. Le critère n° 2 basé sur l'estimation des potentiels après ajout de n_0 tirages, a conduit à une zone de confiance légèrement moins dispersée autour des solutions connues.

Les résultats sont récapitulés pour ces mêmes tirages dans le tableau 36. Au regard du critère $\rho_1(5\%)$, l'algorithme à tirage fournit sur ces données les meilleurs résultats pour le critère n° 2.

Ces résultats peuvent néanmoins varier d'une exécution à l'autre du fait du caractère aléatoire du bruit pesant sur f et du caractère stochastique de l'algorithme. L'étude de la distribution de $\rho_1(s)$ constitue à cet égard un champ d'investigation intéressant.

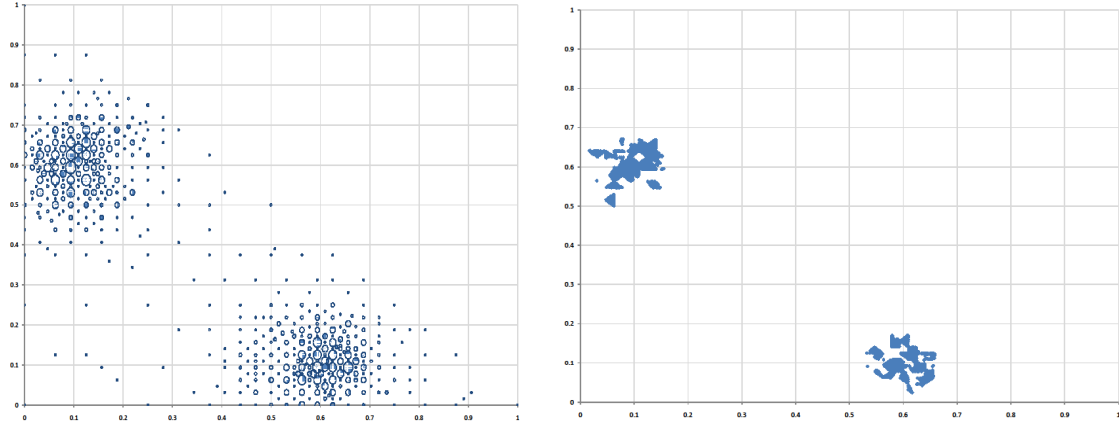


Fig. 61 : Points d'explorations obtenus pour l'algorithme avec tirage, cas b, critère de scission n° 1 de comparaison des erreurs de grille et d'estimation, lorsque $\sigma_B = 0.1$, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*, 5\%}$ avec un seuil de 5 % (à droite)

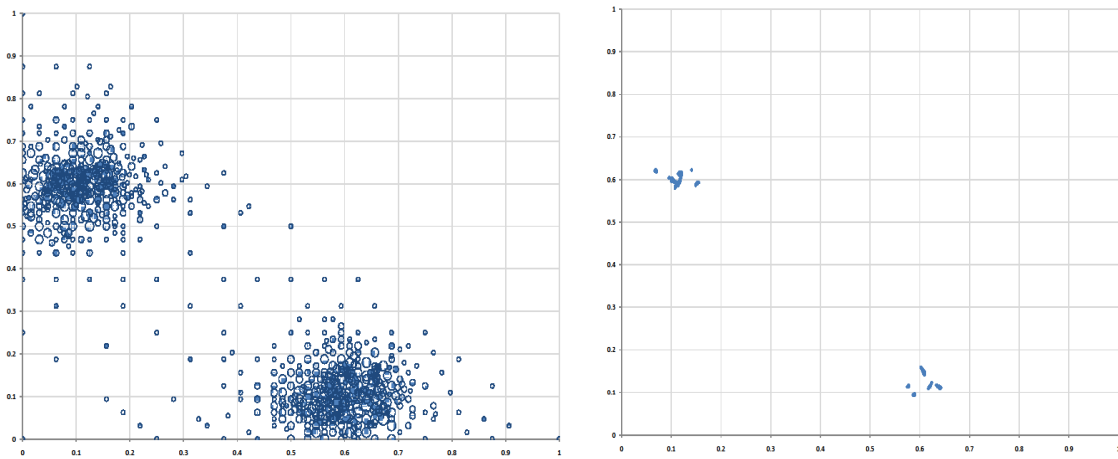


Fig. 62 : Points d'explorations obtenus pour l'algorithme avec tirage, cas c, lorsque $\sigma_B = 0.1$, critère de scission n° 2 considérant l'ajout de n_0 observations, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*, 5\%}$ avec un seuil de 5 % (à droite).

Priorité	bruit σ_B	(σ_K, α)	$\rho_1(s)$	ρ_2	ρ_3	$\rho_4(\eta)$
absolue	0.1	(0.6, 1.2)	9.23 E-02	2.21 E-03	6.46 E-08	1.08 E-02
relative	0.1	(0.6, 1.2)	8.87 E-02	2.76 E-03	3.78 E-07	9.57 E-03
absolue	0	(0.6, 1.2)	3.61 E-02	5.52 E-04	1.17 E-06	0
relative	0	(0.6, 1.2)	6.18 E-02	5.52 E-04	6.77 E-06	8.05 E-05
absolue	0.1	(1, 1.6)	9.79 E-02	2.21 E-03	2.04 E-07	4.72 E-02
relative	0.1	(1, 1.6)	8.69 E-02	1.61 E-03	1.01 E-06	1.79 E-02
absolue	0	(1, 1.6)	1.26 E-03	3.45 E-05	3.17 E-09	0
relative	0	(1, 1.6)	2.57 E-03	3.45 E-05	2.57 E-08	0

Tab. 37 : Comparaison des critères de convergence pour les choix priorité relative (exploration d'une zone avec une probabilité proportionnelle à son potentiel) ou priorité absolue (exploration de la zone de meilleur potentiel)

Impact de la priorité Nous avons omis ici les graphiques qui permettaient de discuter du choix de γ , paramètre présenté dans la sous-section II-2-2. Nos observations nous conduisaient à un impact limité de ce paramètre. Dans les lignes suivantes, nous allons néanmoins évoquer plus précisément le choix de la priorité accordée à la zone de meilleur potentiel.

Jusqu'à présent, la zone à subdiviser ou explorer était choisie avec une probabilité proportionnelle au potentiel de la zone (choix qualifié ici "*priorité relative*"). Il est également envisageable de ne plus choisir la prochaine zone de façon stochastique, mais de choisir de façon déterministe zone de meilleur potentiel (choix qualifié ici "*priorité absolue*"). Cela revient à traiter une question déjà abordée : choisir entre explorer davantage la fonction, ou privilégier la vitesse en explorant en priorité une zone prometteuse, au risque d'ignorer un optimum global.

Les paramètres (σ_K, α) permettent justement d'opérer cet arbitrage. Si ces paramètres sont en partie estimés, ce dernier choix de priorité *relative* versus *absolue* conduit à un arbitrage naturellement différent. Dans les applications numériques opérées, ce choix n'a pas eu beaucoup d'incidence, notamment en présence de bruit.

Le tableau 37 récapitule les résultats obtenus avec le choix de priorité relative ou le choix absolue. Nous nous sommes placé pour cette application dans un cadre de scission systématique (pas de retirages), avec un nombre de tirages $n = 2000$. Pour le critère ρ_1 , nous avons bâti l'ensemble solution proposé à l'aide d'un seuil s égal à 10 % du meilleur potentiel observé, en conservant ainsi une part notable des solutions possibles.

Rappelons que le choix d'un seuil supérieur conduit à rétrécir les zones de confiance autour des optima trouvés, et peut améliorer de façon importante le critère de convergence ρ_1 , comme nous pouvons le constater en comparant ces résultats avec ceux du tableau 36. Ce critère ρ_1 n'est donc pas comparable à d'autres critères qui seraient obtenus avec d'autres seuils. Par ailleurs l'augmentation du seuil s augmenterait le risque d'ignorer un optimum global, et le critère ρ_1 pourrait alors être brutalement dégradé.

Remarquons tout d'abord dans le tableau 37 la bonne exploration des voisinages de chacun des points solutions : dans tous les cas, après 2000 itérations, des points ont été tirés à une

distance ρ_2 inférieure à 3.10^{-03} de chacune des solutions, avec des zones de confiance d'un rayon ρ_1 raisonnable autour des solutions (qu'il est possible de réduire en augmentant le seuil s). En l'absence de bruit, des tirages sont obtenus, avec $(\sigma_K, \alpha) = (1, 1.6)$, à une distance très réduite d'environ 3.10^{-05} de chaque point solution.

Au vu des mesures effectuées dans ce tableau 37, le choix *priorité relative* ou *priorité absolue* a surtout un impact dans les situations non bruitées où les paramètres $(\sigma_K, \alpha) = (1, 1.6)$ privilégient une moindre exploration. Dans ce cas, la convergence est assez rapide, et le choix *priorité absolue* conduit à une zone de confiance de diamètre faible $\rho_1 = (1.2).10^{-03}$, plus faible que dans le cas *priorité relative*. Dans les autres cas, il apparaît que les mesures ne sont pas radicalement perturbées par ce choix, notamment dans le cas d'un bruit σ_B non nul.

II-4-6 Comparaison avec l'algorithme de Kiefer-Wolfowitz-Blum

Notre algorithme vise à opérer une optimisation à partir d'une fonction F bruitée, et il n'est pas en pratique possible d'éliminer ce bruit. Pour cette raison, nous ne comparerons pas notre algorithme avec les algorithmes d'optimisation classiques en l'absence de bruit. Nous avons donc choisi de comparer les résultats obtenus avec ceux que donnerait un autre algorithme d'optimisation de fonction bruitée.

L'algorithme utilisé pour la comparaison est un des algorithmes d'optimisation stochastique parmi les plus classiques : il s'agit de l'extension multidimensionnelle, proposée par Blum [1954] de l'algorithme de Kiefer et Wolfowitz [1952]. Rappelons tout d'abord que l'algorithme de Kiefer- Wolfowitz vise l'obtention d'un unique optimum, non la garantie d'absence d'autres optima. Tel que décrite dans Broadie et al. [2009], la version de l'algorithme de Blum consiste à déterminer une suite de points $\theta^{(1)}, \theta^{(2)}, \dots$ convergeant vers la solution θ^* , supposée unique, minimisant $f(\theta) = E[F(\theta)]$. En dimension $d = 2$, si l'on note $\theta^{(k)} = (\theta_x^{(k)}, \theta_y^{(k)})$, $k \in \mathbb{N}$ la suite de points est telle que :

$$\begin{aligned}\theta_x^{(n+1)} &= \theta_x^{(n)} - a_n \frac{F(\theta^{(n)} + c_n e_x) - F(\theta^{(n)})}{c_n} \\ \theta_y^{(n+1)} &= \theta_y^{(n)} - a_n \frac{F(\theta^{(n)} + c_n e_y) - F(\theta^{(n)})}{c_n}\end{aligned}$$

où $e_x = (1, 0)$ et $e_y = (0, 1)$, et où $\{a_n\}_{n \in \mathbb{N}}$ et $\{c_n\}_{n \in \mathbb{N}}$ sont des suites réelles décroissantes en n . Les contraintes auxquelles sont soumises ces constantes, ainsi que les conditions générales d'application et de convergence de l'algorithme, sont détaillées dans Broadie et al. [2009].

Remarquons qu'à chaque étape de l'algorithme, le budget de tirage de F est amputé de $d + 1 = 3$ tirages : d tirages pour estimer le gradient, ici en $\theta^{(n)} + c_n e_x$ et $\theta^{(n)} + c_n e_y$, et un tirage pour le nouveau point $\theta^{(n+1)}$. Pour un budget de n points de tirage de la fonction F supposée coûteuse, l'algorithme ne pourra pas faire appel à un nombre d'étapes supérieur à $(n/3)$. A cet égard, cette version de l'algorithme, faisant appel à des différences finies sur un seul côté, dépense moins de tirages que d'autres versions pour l'estimation du seul gradient.

Par ailleurs, les suites $\{a_n\}$ et $\{c_n\}$ semblent délicates à choisir, et les écrits de Broadie et al. [2009] mentionnent une grande sensibilité du comportement de l'algorithme en fonction de ces choix. En l'absence supposée d'autres informations sur f , nous avons opté pour un choix par défaut classique de $a_n = 1/n$ et $c_n = 1/n^{1/4}$, choix proposé en première page dans l'article originel Kiefer, Wolfowitz [1952]. Les paramètres $a_n = a_0/n$, $c_n = c_0/n^{1/3}$, $n \geq 1$, présentés dans Broadie et al. [2009] conduisaient ici à des résultats décevants pour $a_0 = c_0 = 1$ et les quelques choix alternatifs testés, mais nous n'avons pas cherché spécifiquement à trouver a posteriori les meilleurs paramètres, pour une fonction f supposée inconnue.

En conséquence, les résultats numériques obtenus ici sont représentatifs d'un type de trajectoire classiquement obtenue lors d'une descente stochastique de gradient, mais certainement améliorables en terme de vitesse de convergence, les suites $\{a_n\}$ et $\{c_n\}$ pouvant être modifiées à cette fin.

Dans la plupart des illustrations précédentes, pour l'algorithme à scission systématique, $n_0 = 10$ tirages étaient opérés en chacun des points explorés, pour un budget de 2000 points d'exploration. Le budget global de tirage de F était donc de 2000 n_0 tirages. Le graphique 63 rend compte des points de tirage obtenus par l'algorithme de Kiefer-Wolfowitz- Blum, pour 2000 points de tirage d'une fonction F soumise à un bruit $\sigma_B / \sqrt{n_0}$ (figure de gauche), ou pour 2000 n_0 tirages d'une fonction F soumise à un bruit σ_B (figure de droite). Le nombre global de tirage correspond ainsi à celui des illustrations précédentes. Le point initial a été tiré aléatoirement, de façon uniforme, sur le simplexe initial (cf. Smith et Tromble [2004]).

Sur le graphique 63, nous constatons tout d'abord la présence de trois séries de points, correspondant à l'ensemble des points suggérés, ainsi qu'aux points décalés verticalement et horizontalement pour l'estimation du gradient (points décalés respectivement vers le haut et vers la droite). Les points sont initialement assez espacés (a_n grand) puis sont de plus en plus proches au fil des tirages (a_n petits, lorsque nous sommes en principe proche d'une solution). Sur les données moins bruitées (figure de gauche), les gradients estimés sont moins erratiques, les directions choisies d'un point à l'autre étant plus stables.

Ce graphique 63 est essentiellement intéressant dans la mesure où il marque bien les différences d'approches et de finalité entre les algorithmes de descente stochastique de gradient et l'algorithme proposé dans le présent travail. Concernant l'algorithme de Kiefer-Wolfowitz- Blum, nous pouvons notamment faire les constats suivants :

- En dehors de l'espace occupé par les trois trajectoires qui se dessinent, la fonction F est très peu explorée sur le reste de la zone de recherche, et est donc susceptible d'ignorer des optima globaux. Cela est logique dans la mesure où nous n'exigeons normalement pas d'un algorithme d'optimisation locale qu'il fournisse un optimum global.
- Par ailleurs, même en présence supposée d'un unique optimum, la construction de zones de confiance pour l'optimum, à partir des trois trajectoires obtenues, semble ici délicate.

– Des points à l’extérieur du simplexe initial Z_0 ont pu être tirés, dans la mesure où $f(\theta)$ pouvait prendre ici des valeurs en dehors du simplexe initial, mais l’exploitation de cet algorithme supposerait l’introduction préalable de contraintes.

– L’algorithme se comporte ici de façon cohérente dans la mesure où la fonction f est suffisamment régulière sur une grande partie du simplexe initial (en dehors de la première diagonale). La situation serait bien différente en cas de non-dérivabilité de f sur l’essentiel du domaine Z_0 .

S’agissant des critères de convergence, l’algorithme de Kiefer-Wolfowitz-Blum ne permet malheureusement pas *a priori* de partitionner la zone de recherche initiale Z_0 en plusieurs zones.

En conséquence, nous ne calculerons pas ici de potentiel. En supposant que Θ_e désigne l’ensemble des points de tirage obtenus, nous réutiliserons donc le critère ρ_2 :

$$\rho_2 = \sup_{y \in S_{m^*}} \inf_{x \in \Theta_e} d(x, y)$$

Ce critère permet de savoir s’il existe des points de tirage proches de chacune des solutions.

L’algorithme de Kiefer-Wolfowitz-Blum recherchant une unique solution est très fortement pénalisé par ce critère, car il peut converger vers une solution tout en ayant des points d’exploration très éloignés des autres solutions. Un second critère retenu spécifiquement pour cet algorithme sera :

$$\tilde{\rho}_2 = \inf_{y \in S_{m^*}} \inf_{x \in \Theta_e} d(x, y)$$

Ce dernier critère indique juste s’il existe des points de tirage proches de l’une des solutions. Les résultats obtenus sont indiqués dans le tableau 38, qui correspond aux courbes du graphique 63.

D’une part, l’algorithme de Kiefer-Wolfowitz-Blum ne recherche qu’un optimum local. Nous devrions donc obtenir un critère ρ_2 de convergence vers un optimum global très dégradé par rapport à l’algorithme de scission systématique, ce qui est bien le cas.

D’autre part, la recherche d’un optimum local doit intuitivement être moins coûteuse que la recherche d’un optimum global. Les points explorés près d’une unique solution devraient être plus proches de la solution pour un algorithme d’optimisation locale et nous devrions obtenir un critère ρ_2 meilleur. Ce n’est pas le cas ici : les mesures de ρ_2 obtenues avec l’algorithme à scission systématique ont conduit (cf. tableau 37) à des points toujours tirés à une distance ρ_2 inférieure à 3.10^{-3} de chaque solution. Or, la distance $\tilde{\rho}_2$ à la solution la mieux explorée est encore inférieure à ρ_2 . L’algorithme de Kiefer-Wolfowitz Blum a conduit quant-à-lui à des résultats de l’ordre de 0.2 pour la distance ρ_2 , et 3.10^{-2} pour la distance $\tilde{\rho}_2$. Notre algorithme a donc été ici localement et globalement plus performant. Ce dernier constat est sans doute lié à l’absence d’optimisation spécifique des séquences $\{a_n\}$ et $\{c_n\}$, qui reste toutefois délicate à opérer *a priori* sur une fonction supposée inconnue.

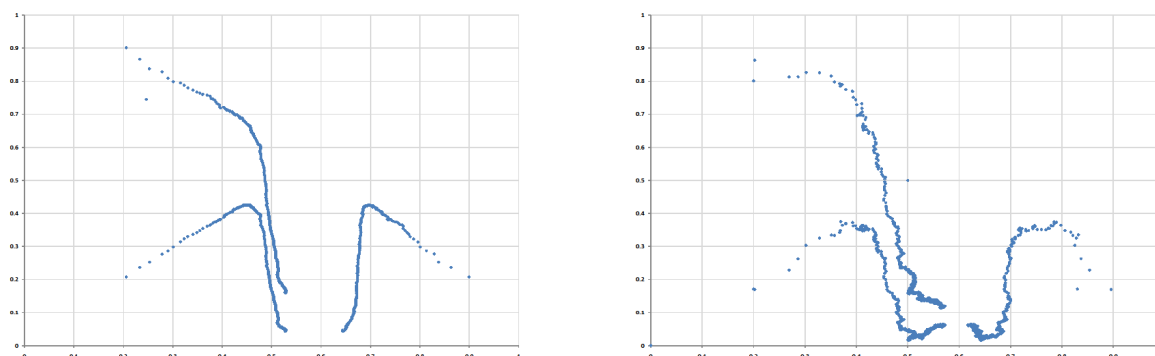


Fig. 63 : Points d'explorations obtenus pour l'algorithme Kiefer-Wolfowitz-Blum, pour 2000 tirages avec $\sigma_B = 0.1/\sqrt{10}$ (à gauche), ou 20000 tirages avec $\sigma_B = 0.1$ (à droite)

	ρ_2	$\tilde{\rho}_2$
Cas $n=2000$, $\sigma_B=0.1/\sqrt{10}$	0.2058	6.35 E-02
Cas $n=20000$, $\sigma_B=0.1$	0.2248	3.34 E-02

Tab. 38 : Critères de convergence obtenus avec l'algorithme de Kiefer-Wolfowitz (pour des séquences $\{a_n\}$ et $\{c_n\}$ non spécifiquement optimisées)

II-5 Conclusion

Nous avons présenté un algorithme permettant de rechercher des paramètres de $\mathfrak{X}^d$ conduisant à minimiser, globalement, une fonction réelle bruitée. L'algorithme permet également de construire des zones de confiance autour des paramètres solutions supposés. L'approche retenue s'appuie sur la définition du potentiel d'une fonction en un point, et sur la mesure de l'incertitude frappant la fonction objectif entre les points déjà explorés. Cette incertitude découle d'une part des *erreurs d'estimation* aux points explorés, d'autre part des *erreurs de grille* liées à la distance entre un point considéré et les points explorés alentours.

L'algorithme permet de moduler facilement, au moyen des deux paramètres σ_K et α , la réitération de tirages dans un voisinage des points solutions découverts, ou à l'inverse l'exploration de la fonction dans des zones encore peu explorées. En outre, une procédure d'estimation de la régularité fonction objectif f permet d'ajuster les valeurs de ces deux paramètres.

Pour des dimensions pas trop élevées, une étude a également été effectuée (cf. Rulière et al. [2010]). Les résultats obtenus montrent que le comportement de l'algorithme est assez convaincant sur de nombreux points, comme sa capacité d'exploration globale du domaine, sa capacité à privilégier les zones proches des optima, sa faible exigence sur la régularité de la fonction objectif, son adaptation à un bruit nul comme à un bruit élevé. Toutefois, lorsque la dimension du problème devient grande, l'algorithme conduit à des résultats proches de ceux obtenus par une exploration uniforme du domaine initial. L'analyse détaillée de la vitesse de convergence de l'algorithme et l'amélioration éventuelle de cette vitesse en grande dimension constituent une suite logique de cette étude.

Conclusion du chapitre 2

Le paysage des régimes de retraite en France est diversifié. Une panoplie large de contraintes et d'objectifs peut être rencontrée en pratique. Dans ce cadre, nous avons proposé une modélisation du régime-type de retraite tout en étudiant certains critères d'allocation stratégique d'actifs (en fonction du type du régime : provisionné, partiellement provisionné, etc.). La stratégie *Fixed-Mix*, une des référence en matière d'allocation d'actifs, a été ensuite explorée et mise en œuvre dans le cas d'un régime de retraite partiellement provisionné. Le souci de réduction des temps de calcul lors de cette mise en œuvre nous a conduits à proposer un algorithme d'optimisation numérique « par exploration sélective ».

Il s'agit d'un algorithme qui vise à montrer l'intérêt et la faisabilité technique d'une grille à pas variable, susceptible de répondre au problème du choix entre exploration et connaissance d'une fonction aléatoire. L'approche proposée ici se voulait simple, avec notamment la définition d'un potentiel directionnel pour une unique dimension, puis l'agrégation de potentiels directionnels sur l'ensemble des dimensions. Il est possible que de nombreux autres choix puissent être faits pour la mesure du potentiel d'un point et d'une zone. Les techniques d'agrégation et de tirages peuvent également être amendées de façon à améliorer l'efficacité de l'algorithme sur des architectures parallèles.

Dans le chapitre suivant, nous essayons d'étudier, sur un autre plan, le problème d'allocation stratégique d'actifs. L'objectif sera, cette fois-ci, de modifier la stratégie d'allocation et d'étudier l'effet sur le résultat final. En particulier, nous explorons différentes techniques permettant de nous positionner dans un contexte d'allocation dynamique, différent de celui considéré jusque là avec l'allocation à poids constants (*Fixed-Mix*).

Chapitre 3 : Allocation stratégique d'actifs et modèles d'ALM dynamique

CONCLUSION GÉNÉRALE

Synthèse et analyse des résultats

L'étude de l'utilité des réserves pour un système par répartition et *a fortiori* de leur gestion reste un sujet peu exploré. Pour apprécier la solidité d'un régime de retraite, les experts se sont longtemps limités à une approche binaire, selon laquelle il convenait de tout provisionner dans le cas d'un système par capitalisation et rien dans le cas d'un système par répartition. C'est dans ce contexte que le présent travail a été focalisé sur les modèles d'allocation stratégiques d'actifs et sur leurs applications pour la gestion des réserves financières des régimes de retraite par répartition.

Au terme de cette étude, nous pouvons tenter de tirer un certain nombre de conclusions et de points clés des travaux présentés :

- Concernant les générateurs de scénarios économiques (GSE) :
 - Le choix des composantes d'un GSE est lié à sa vocation finale, que ce soit en tant qu'outil d'évaluation des produits financiers (*pricing*) ou en tant qu'outil de projection et de gestion des risques.
 - Dans le cas où le GSE est utilisé comme un outil de gestion des risques, il ne faut pas perdre de vue qu'il ne s'agit finalement que d'un outil d'aide à la prise de décision. De ce point de vue, la performance du GSE ne peut être mesurée que sur la base de sa contribution à la prise de la « bonne » décision et non sur la base du degré de correspondance des scénarios projetés par rapport à la distribution réelle des variables du GSE, une distribution qui reste peu accessible quel que soit les outils d'estimation et d'approximation employés. A ce niveau, certains indicateurs de mesure de performance ont été étudiés, à savoir : l'indicateur de stabilité et l'indicateur d'absence de biais.
 - La distinction entre les aspects théoriques et la mise en œuvre lors de l'élaboration d'un GSE est nécessaire. Par aspects théoriques, nous notons en particulier la conception et le choix de la structure de dépendance. Par mise en œuvre nous notons en particulier la détermination de la structure de projection des scénarios, le choix de la méthodologie de leur génération et le calibrage du générateur.
 - Au-delà de la problématique des queues de distribution des variables modélisées, l'instabilité dans le temps de la corrélation entre les différents facteurs de risque a été considéré comme un élément important lors du choix du niveau de finesse d'un GSE. Cette importance provient de l'impact significatif des hypothèses de corrélation sur les résultats finaux obtenus et ainsi sur la décision finale de l'investisseur (cas de l'allocation d'actifs). Par exemple, l'omission des régimes qui caractérisent les processus générateurs de la rentabilité des actions dans l'analyse de la relation « inflation-rentabilité peut introduire des biais.

Nous avons mis en évidence la nécessité de ne pas sous estimer l'impact de l'instabilité dans le temps des corrélations et ainsi l'importance de ne pas se contenter des modèles log-normaux usuels et/ou de corrélations linéaires. A

ce niveau, nous avons essayé, en construisant un GSE, d'apporter des améliorations par rapport à ce qui se pratique usuellement en proposant un modèle de changement de régime plus global par rapport au modèle initial de Hardy [2001].

En particulier, la notion de changement de régime a été élargie pour concerner non seulement les rendements et les volatilités des actions, comme dans le cas du modèle initial de Hardy [2001], mais aussi la matrice des corrélations entre les différentes variables. Cette matrice devient fonction du régime de marché observé à une date donnée : cela nous place dans un cadre de dépendance non linéaire. L'apport de cette approche au niveau de l'allocation stratégique d'actifs a été mis en évidence via une application numérique. Les résultats obtenus confirment une aversion au risque supplémentaire par l'investisseur lors de l'utilisation de telles approches.

- Dans le cadre de notre construction d'un GSE, il apparaît délicat de prétendre obtenir des hypothèses pérennes sur les paramètres de celui-ci et de retenir une approche de calibrage plutôt qu'une autre de manière définitive. Le processus de calibrage dépend fortement de différents facteurs (nature de données, historique retenu, etc.) et les résultats obtenus par le GSE sont dépendent, bien évidemment, des hypothèses initiales sur les différents paramètres. Cela ne peut être remédié *a priori* que par un recours à des *stress test* sur les hypothèses initiales. Dans le cadre du calibrage du GSE construit, un ensemble d'outils a été présenté de manière à permettre l'estimation de ses différents paramètres.
- Concernant les modèles d'élaboration de l'allocation d'actif
 - L'aspect mono-périodique qui caractérise les différents modèles classiques déterministes de gestion actif-passif (que ce soit celles basées sur les techniques d'immunisation ou celles basées sur la notion du surplus) limite leur capacité à refléter de façon fiable les perspectives d'évolution des différentes variables financières, en particulier sur le long terme. Cela devient d'autant plus compliqué que les variables à considérer sont plus nombreuses et que la nécessité de prendre en compte les dépendances entre les différentes variables est plus importante.
 - Lors du choix des critères de l'allocation stratégique, il est nécessaire de tenir compte du type de régime de retraite considéré (régime provisionné, régime partiellement provisionné, régime non provisionné). En outre, quel que soit le régime, dans le développement de l'allocation stratégique, une attention particulière peut également être portée à la distinction entre la phase de constitution et la phase de restitution du régime.
 - En guise d'illustration, la sensibilité des résultats de l'allocation d'actifs par rapport aux hypothèses retenues initialement sur les paramètres du GSE a été mise en évidence dans le cadre de la stratégie *Fixed-Mix*.
 - L'approche dynamique pour l'allocation stratégique d'actifs présente l'avantage théorique de la robustesse face aux changements de régime des

marchés. L'autorisation du changement des poids des différentes classes d'actifs, sur la base d'une règle de gestion bien définie, constitue un élément intéressant. Au-delà du fait que cela permet l'ajustement des expositions aux différentes classes d'actifs (suite à l'évolution des conditions de marché), cette approche peut constituer dans certaine mesure une version simplifiée et réaliste de l'approche *Fixed-Mix* (à poids constants). Certains de ces modèles dynamiques ont été mis en œuvre et détaillées (notamment, les techniques d'assurance de portefeuille)

- Parmi les modèles dynamiques d'allocation d'actifs, les techniques de « programmation stochastique » (cf. Dantzig et al. [1990]) sont peu explorées, en particulier au niveau des régimes de retraite par répartition partiellement provisionnés. Nous sommes ainsi partis de la description du principe de base de ces techniques pour finalement mettre en œuvre des techniques novatrices d'ALM.

Parallèlement, une nouvelle méthodologie pour la génération de l'arbre des scénarios a été adoptée. La méthodologie a le mérite de tenir compte de la corrélation entre les distributions des différentes variables projetées et permet de lier l'aspect linéaire de génération des scénarios (tel que décrit dans la partie I de ce travail) avec les techniques de PS (sujet profondément étudié dans la partie II).

Une étude comparative du modèle d'ALM ainsi développé avec celui basé sur la stratégie *Fixed-Mix* a été effectuée. L'écart constaté en termes de résultats peut *a priori* être expliqué par les paramètres d'entrée de chacun des deux modèles, qui ne permettent pas forcément de converger vers un même résultat.

Différents tests de sensibilité ont été par ailleurs mis en place pour mesurer l'impact du changement de certaines variables clés d'entrée sur les résultats produits par notre modèle d'ALM. Les résultats obtenus, dans le cadre de ces tests, indiquent une sensibilité significative à certains paramètres ce qui signifie qu'une attention particulière doit être donnée aux hypothèses de départ.

Concernant la règle de gestion qu'adopte, implicitement, notre modèle d'ALM basé sur la programmation stochastique, nous avons constaté qu'au fur et à mesure que nous rencontrons des scénarios favorables, l'allocation optimale préconisée par le modèle commence à retenir de plus en plus d'actifs risqués (ex. les actions selon nos hypothèses), dans le cas inverse le portefeuille se charge en actifs non risqués.

De façon générale, la programmation stochastique est un domaine en pleine effervescence, permettant d'étendre l'optimisation classique au niveau des problèmes stochastiques et dynamiques tout en gardant la possibilité de résoudre ce type de problème. L'approche de programmation stochastique développée dans ce travail est claire, théoriquement efficiente et relativement facile à mettre en œuvre (implémentation).

L'application de cette approche peut améliorer les techniques actuelles utilisées par les fonds de retraite en France en matière d'allocation stratégique d'actifs, souvent basées sur la simple simulation et/ou sur l'expérience. Le modèle d'ALM développé (en se basant sur les techniques de PS) pourra également servir comme une approche de référence, dont les résultats sont à comparer avec ceux de la stratégie *Fixed-Mix*. Cette stratégie peut aussi constituer un outil pour tester la fiabilité et l'optimalité de l'allocation de long terme obtenue avec la stratégie *Fixed-Mix*.

Cette thèse a également accordé une attention particulière aux techniques numériques de recherche de l'optimum, qui demeurent des questions essentielles pour la mise en place d'un modèle d'allocation. Le point de départ était le constat d'un temps de calcul significatif dû simultanément à un nombre élevé de scénarios économiques générés et à un nombre d'allocations d'actifs testées également élevé.

Cela nous a conduits à une réflexion sur un algorithme d'optimisation globale d'une fonction non convexe et bruitée. L'algorithme est construit après une étude de critères de compromis entre, d'une part, l'exploration de la fonction objectif en de nouveaux points (correspondant à des tests sur de nouvelles allocations d'actifs) et d'autre part l'amélioration de la connaissance de celle-ci, par l'augmentation du nombre de tirages en des points déjà explorés (correspondant à la génération de scénarios économiques supplémentaires pour les allocations d'actifs déjà testées). Une application numérique a illustré la conformité du comportement de cet algorithme à celui prévu théoriquement.

L'algorithme permet de moduler facilement, au moyen de deux paramètres, la réitération de tirages dans un voisinage des points solutions découverts, ou à l'inverse l'exploration de la fonction dans des zones encore peu explorées. En outre, une procédure d'estimation de la régularité fonction objectif f permet d'ajuster les valeurs de ces deux paramètres. Il s'agit d'un algorithme qui vise à montrer l'intérêt et la faisabilité technique d'une grille à pas variable, susceptible de répondre au problème du choix entre exploration et connaissance d'une fonction aléatoire.

Perspectives

Les réserves contribuent à part entière à la solidité du régime. L'ensemble des éléments présentés fournit des composants essentiels pour élaborer l'allocation stratégique d'un régime de retraite partiellement provisionné. Une première perspective de prolongement de ce travail pourra être d'examiner l'apport de chacune des classes d'actifs dites « non classiques » (tels que l'immobilier, les obligations indexées sur l'inflation, etc.) en tant qu'alternative stratégique d'investissement. Nous rappelons que dans le cas de l'allocation stratégique d'actifs d'un fonds de retraite la priorité est souvent donnée aux actions, aux obligations et au monétaire, dites classes « classiques » (cf. Campbell et al. [2001]).

En matière de projection de grandeurs réelles sur le long terme, les actifs non classiques sont souvent traités avec prudence et ne suscitent pas la grande part de l'intérêt des décideurs. Ces classes d'actifs se heurtent souvent aux problèmes d'un historique peu profond, d'une liquidité insuffisante et de données confidentielles (cas des fonds de couverture). Cela ne remet pas en cause le potentiel qu'elles présentent en tant que source de performance et/ou de couverture supplémentaire pour les réserves gérées : dans ce sens leurs apports à la gestion des réserves pourront faire l'objet d'études spécifiques sur le court terme. La littérature

présente à ce niveau différentes techniques pouvant servir de base à cette étude en particulier les techniques d'Analyse par Composante Principale (cf. Roncalli [1998]).

Par ailleurs, nous notons qu'un régime de retraite par répartition provisionné (ou partiellement provisionné) doit faire face au moins à deux risques essentiels : l'incertitude sur l'évolution future de ses flux techniques et le risque de placement. Le présent travail a été focalisé sur la gestion et le contrôle du deuxième risque. Il apparaît ainsi intéressant la réalisation d'analyses de sensibilité à la tendance des flux techniques futurs (risque de longévité, etc.).

Cela nous place dans un cadre d'étude plus global permettant de définir les termes d'un pilotage technique du régime. Par pilotage technique on entend les règles d'actualisation, au fil du temps, des paramètres du régime notamment les taux de cotisation et de prestation en fonction de l'évolution de la solidité financière du régime. L'étude des possibilités d'application des techniques de programmation stochastique avec recours pour ce type de problématique constitue *a priori* un sujet prometteur.

Enfin, nous travaillons à enrichir l'algorithme d'optimisation numérique par exploration sélective. Le cadre de son application a été jusque là réduit à un cas simplifié pour des raisons pratiques de mise en œuvre. L'un des prochains défis sera donc de l'exploiter dans un univers plus réaliste et plus global.

Index des graphiques

Fig. 1 : Pourcentages de la performance globale expliqués par certaines composantes de l'allocation d'actifs selon l'étude de Brinson et al. [1991].....	8
Fig. 2 : Les modules de la formule standard du SCR.....	13
Fig. 3 : Structure par cascade dans le modèle de Wilkie [1986].....	26
Fig. 4 : Structure du modèle d'Ahlgrim et al. [2005].....	28
Fig. 5 : Comparaison entre deux structures schématiques de projection des scénarios.....	31
Fig. 6 : Exemple détaillé de la structure d'arbre des scénarios.....	31
Fig. 7 : Arbre de scénarios avec différents nombres de nœuds-enfants pour les nœuds d'un même niveau.....	32
Fig. 8 : Projection du rendement des actions selon une structure schématique linéaire.....	47
Fig. 9 : <i>Box plot</i> de la distribution des valeurs de la fonction objectif pour différentes tailles de scénarios.....	48
Fig. 10 : Les allocations optimales obtenues pour les différents ensembles de scénarios (30 ensembles) ayant chacun la taille de 10 000 scénarios.....	49
Fig. 11 : Evolution de la corrélation (glissante) sur 5 ans entre l'inflation et le rendement des actions aux Etats-Unis depuis 1970.....	56
Fig. 12 : La corrélation entre les taux d'inflation et les taux nominaux courts au Royaume-Uni.....	59
Fig. 13 : Illustration du phénomène de retour à la moyenne des taux réels dans le cadre du modèle de Hull et White à deux facteurs.....	60
Fig. 14 : Illustration de la structure du GSE construit.....	72
Fig. 15 : Illustration de la structure schématique linéaire dans le cas de la projection des rendements des actions.....	93
Fig. 16 : Illustration du passage de browniens indépendants à des browniens corrélés négativement.....	95
Fig. 17 : Courbe des taux des obligations indexées sur l'inflation d'une maturité de quinze ans.....	100
Fig. 18 : Structure par terme des taux nominaux, des taux réels et des taux d'inflation en cas de taux réels longs élevés par rapport à ceux de l'équilibre : Mise en évidence de la courbure de la STTI.....	101
Fig. 19 : Evolution dans le temps de la STTI : Mise en évidence du phénomène de retour à la moyenne.....	102
Fig. 20 : Evolution des pentes moyennes de taux (écarts absolus moyens entre les taux de deux maturités différentes) tout au long de la période de projection.....	103
Fig. 21 : Histogramme des taux nominaux de court terme (3 mois) à la fin de la première période.....	104
Fig. 22 : Histogramme des taux nominaux de long terme (10 ans) à la fin de la première période.....	104
Fig. 23 : Histogramme des taux réels sur le court terme (3 mois) à la fin de la première période.....	104
Fig. 24 : Histogramme des taux réels de long terme (10 ans) à la fin de la première période.....	104
Fig. 25 : Histogramme des taux d'inflation anticipés de court terme (3 mois) à la fin de la première période.....	104
Fig. 26 : Histogramme des taux d'inflation anticipés de long terme (10 ans) à la fin de la première période.....	104
Fig. 27 : Histogramme de la distribution des obligations nominales gouvernementales de duration 5 ans sur la première période.....	105

Fig. 28 : Histogramme de la distribution des obligations indexées inflation de duration 8 ans sur la première période.....	105
Fig. 29 : Histogramme de la distribution des rendements des obligations nominales non gouvernementales de duration 5 ans sur la première période.....	105
Fig. 30 : Evolution du marché des actions entre le régime 1 (actions à faible volatilité) et le régime 2 (actions à forte volatilité) sur une trajectoire simulée.....	108
Fig. 31 : Frontière efficiente (Kim&Santomero).....	127
Fig. 32 : Frontière efficiente et contrainte de déficit (Kim&Santomero).....	127
Fig. 33 : Frontière efficiente et contrainte de déficit (Sharpe&Tint).....	130
Fig. 34 : Exemple de l'impact de la variation du seuil de rentabilité de l'actif.....	133
Fig. 35 : Exemple de l'impact de la variation de la probabilité de déficit de l'actif.....	133
Fig. 36 : Exemple de l'impact de la variation du seuil de rentabilité du surplus.....	134
Fig. 37 : Exemple de l'impact de la variation du ratio de financement initial.....	134
Fig. 38 : Représentation graphique des portefeuilles vérifiant la contrainte sur la rentabilité relative.....	136
Fig. 39 : Exemple de la variation de la duration du surplus en fonction de l'écart entre la duration de l'actif et celle du passif.....	138
Fig. 40 : Illustration du modèle de gestion actif-passif dans le cadre stochastique.....	142
Fig. 41 : Exemple d'évolution des flux nets de trésorerie (flux désinflatés).....	158
Fig. 42 : Illustration de la structure schématique linéaire dans le cas de la projection des rendements des actions.....	159
Fig. 43 : Illustration du modèle de gestion actif-passif stochastique objet de l'application dans le cadre de la stratégie <i>Fixed-Mix</i>	160
Fig. 44 : Ratio de couverture des flux jusqu'en 2029 sur un horizon de 9 ans (2018) et un niveau de confiance de 97,5 %.....	164
Fig. 45 : Sensibilité de l'allocation optimale par rapport à l'espérance de rendement des actions.....	165
Fig. 46 : Distribution de la valeur de la réserve totale (nette des flux de trésorerie) dans le cas de projection de scénarios avec deux matrices de corrélation.....	167
Fig. 47 : Un exemple de partition d'une zone dans $\mathbb{R}^2$ en 60 zones. ($d = 2$).....	175
Fig. 48 : Allure générale de la fonction test $f(\theta)$ pour $\theta \in Z_0$	192
Fig. 49 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0$ (scission systématique).....	193
Fig. 50 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\vec{\rho} = (0.10, 3.5E-03, 6.8E-07, 1.9E-2)$	195
Fig. 51 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 2$ et un niveau de bruit $\sigma_B = 0.3$ (scission systématique), $\vec{\rho} = (0.13, 5.9E-03, 1.6E-06, 5.8E-2)$	196
Fig. 52 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 20$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\vec{\rho} = (0.45, 2.2E-03, 7.3E-05, 0.72)$	196
Fig. 53 : Ensemble des points de simulations de F pour une variabilité $\sigma_K = 100$ et un niveau de bruit $\sigma_B = 0.1$ (scission systématique), $\vec{\rho} = (0.60, 8.8E-03, 5.9E-04, 0.97)$. La recherche de l'optimum d'une fonction constante bruitée conduit à un nuage de points d'allure proche de celui-ci.....	197
Fig. 54 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.1$, critère de scission n° 2, pour une variabilité $(\sigma_K ; \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\vec{\rho} = (0.085, 8.8E-03, 3.1E-05, 8.2E-04)$	197

Fig. 55 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.2$, critère de scission n° 2, pour une variabilité $(\sigma_K, \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\vec{\rho} = (0.10, 8.8E-03, 6.7E-05, 6.2E-03)$	198
Fig. 56 : Ensemble des points de simulation de F pour un niveau de bruit $\sigma_B = 0.2$, critère de scission no 1, pour une variabilité $(\sigma_K, \alpha) = (1, 1.5)$, et un nombre minimal de tirage $n_0 = 10$. La surface des bulles est proportionnelle au nombre de tirages de F en chaque point. $\vec{\rho} = (0.10, 0.011, 1.5E-05, 1.6E-02)$	198
Fig. 57 : Points d'explorations obtenus, lorsque $\sigma_B = 0.1$, avec les paramètres estimés hors bruit, $(\sigma_K, \alpha) = (1.0743, 1.5665)$ (à gauche) ou estimés en présence de bruit $(\sigma_K, \alpha) = (0.5545; 1.2209)$ (à droite).....	202
Fig. 58 : Ecart-type empirique de $IR_\alpha(Z)$ obtenu pour la zone Z_0 correspondant au simplexe orthogonal standard initial (en pointillé) et pour la demi-zone Z_1 correspondant à la partie de Z_0 située au dessus de la première bissectrice (trait plein).....	202
Fig. 59 : Distribution <i>a priori</i> de IR_α (courbe gaussienne continue au premier plan) et histogramme <i>a posteriori</i> , pour la demi-zone Z_1 obtenue après la première itération de l'algorithme. Cas $\alpha = 1.2$ (à gauche) et $\alpha = 1.5$ (à droite). Les écarts-type empiriques obtenus à partir de l'histogramme sont respectivement $\sigma = 0.41$ et $\sigma = 0.73$	203
Fig. 60 : Points d'explorations obtenus pour l'algorithme à scission systématique, cas a, lorsque $\sigma_B = 0.1$, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*,5\%}$ avec un seuil de 5 % (à droite).....	205
Fig. 61 : Points d'explorations obtenus pour l'algorithme avec retraitage, cas b, critère de scission n° 1 de comparaison des erreurs de grille et d'estimation, lorsque $\sigma_B = 0.1$, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*,5\%}$ avec un seuil de 5 % (à droite).....	206
Fig. 62 : Points d'explorations obtenus pour l'algorithme avec retraitage, cas c, lorsque $\sigma_B = 0.1$, critère de scission n° 2 considérant l'ajout de n_0 observations, avec les paramètres $(\sigma_K, \alpha) = (0.6, 1.2)$ (à gauche). Zone finalement retenue $S_{m^*,5\%}$ avec un seuil de 5 % (à droite).....	206
Fig. 63 : Points d'explorations obtenus pour l'algorithme Kiefer-Wolfowitz-Blum, pour 2000 tirages avec $\sigma_B = 0.1/\sqrt{10}$ (à gauche), ou 20000 tirages avec $\sigma_B = 0.1$ (à droite).....	211
Fig. 64 : Ratio de Sharpe du surplus final en fonction de m	220
Fig. 65 : VaR à 95 % du surplus final en fonction de m	221
Fig. 66 : Surplus final moyen en fonction de m	221
Fig. 67 : Probabilité de surplus négatif en cours de vie du fonds (en fonction de m).....	221
Fig. 68 : Probabilité de surplus final négatif en fonction de m	222
Fig. 69 : Trajectoire moyenne du portefeuille (actif et garantie) pour $m=2.5$	222
Fig. 70 : Allure des 1000 trajectoires simulées du surplus ($m=2.5$).....	222
Fig. 71 : Illustration du problème de planification à deux étapes (dans le cas d'une seule période).....	231
Fig. 72 : Illustration du problème de planification à trois étapes (dans le cas de deux périodes).....	233
Fig. 73 : Exemple de la structure schématique d'arborescence pour les trajectoires simulées.....	243

Fig.74 : Illustration du lien entre le niveau de probabilité cumulée α_i d'un quantile q_i et le poids associé à ce dernier dans l'arbre des scénarios (cas où $n=3$).....	246
Fig.75 : Illustration du lien entre le niveau de probabilité cumulée α_i d'un quantile q_i et le poids associé à ce dernier dans l'arbre des scénarios (cas où $n=3$ et α_i est régulier) : $p_1 = p_2 = p_3$	247
Fig. 76 : Flux de passif (déinflatés) sur 20 ans.....	251
Fig. 77 : Logigramme des traitements effectués dans le cadre de l'étude du modèle d'ALM développé.....	252
Fig. 78 : Présentation de l'arbre des allocations (« chemin ») obtenu par le modèle d'optimisation avec PS (cas des paramètres standards, en particulier la structure des nœuds [1 5 4 4] et la structure des périodes {1 4 15}).....	257
Fig. 79 : Présentation de l'arbre des allocations (« chemin ») obtenu par le modèle d'optimisation avec PS (cas de paramètres alternatifs, en particulier une structure des nœuds alternative [1 5 3 3] et la structure des périodes {1 4 15}).....	258

Bibilographie

- Aarts E.H.L., Laarhoven V. [1985] « Statistical cooling : a general approach to combinatorial optimization problems », *Philips Journal of Research*, 40 [4], 193-226.
- Adam A. [2007] *Handbook of Asset and Liability Management: From Models to Optimal Return Strategies*, Wiley, Octobre 2007, 576 pages.
- Ahlgrim K.C., D'Arcy S.P., Gorvett R.W. [2005] « Modeling Financial Scenarios: A Framework for the Actuarial Profession », *Proceedings of the Casualty Actuarial Society*, 177-238. <http://www.casact.org/pubs/proceed/proceed05/05187.pdf>
- Ahlgrim K.C., D'Arcy S.P., Gorvett R.W. [2008] « A Comparison of Actuarial Financial Scenario Generators », *Variance*, 2:1, 2008, pp. 111-134.
- Albeanu G., Ghica M., Popentiu-Vladicescu F. [2008] « On using bootstrap scenario generation for multi-period stochastic programming applications », *Proceedings of ICCCC 2008*, 15-17 May 2008 , pp. 156-161.
- Alexander C. [2001] *Market Models: A Guide to Financial Data Analysis*, Chichester, England: John Wiley & Sons.
- Alexander C., Leigh C. [1997] « On the covariance matrices used in Value-at-Risk models », *Journal of Derivatives*, 4, 50-62.
- Alliot J.M. [1996] *Techniques d'optimisation stochastique appliquées aux problèmes du contrôle aérien*, INPT, Habilitation à Diriger des Recherches.
- Arbulu P. [1998] *Le marché parisien des actions au XIXe siècle: performance et efficience d'un marché émergent*, Thèse de doctorat ès sciences de gestion, Université d'Orléans.
- Armel K., Planchet F., Kamega A. [2010] « Quelle structure de dépendance pour un générateur de scénarios économiques en assurance ? », *Les cahiers de recherche de l'ISFA*, WP2134.
- Artus P., [1998] « Faut-il introduire les prix d'actifs dans la fonction de réaction des banques centrales ? », *Document de travail*, Caisse des Dépôts et Consignations, n°1998-26, juin.
- Bassetto M. [2004] « Negative Nominal Interest Rates », *American Economic Review (Papers and Proceedings)* 94.2 (2004): 104-108.
- Beasley J.E. [1990] « OR-Library: distributing test problems by electronic mail », *Journal of the Operational Research Society*, 41(11) (1990) pp1069-1072. <http://people.brunel.ac.uk/~mastjjb/jeb/info.html>
- Beder T. [1995] « VaR: Seductive but Dangerous », *Financial Analysts Journal*, 51, 12-24.
- Bellman R.E. [1957] *Dynamic Programming*, Princeton University Press, Princeton, NJ.

- Bhansali V., Wise M. [2001] « Forecasting portfolio risk in normal and stressed markets », *Journal of Risk*, 4, 91-106.
- Bienvenüe A., Rullière D. [2010] « Iterative adjustment of survival functions by composed probability distortions », *soumis*.
- Birge J.R., Louveaux F. [1997] *Introduction to Stochastic Programming*, Springer Series in Operations Research and Financial Engineering Heidelberg, 1997.
- Black F., Litterman R. [1992] « Global Portfolio Optimization », *Financial Analysts Journal*, September/October, 28-43.
- Black F., Scholes M. [1973] « The pricing of options and corporate liabilities », *Journal of Political Economy*, 1973, pp. 637-654.
- Blake D., Lehmann B., Timmermann A. [1999] « Asset Allocation Dynamics and Pension Fund Performance », *Journal of Business*, 72, 429-461.
- Blanke D. [2002] « Estimation du coefficient de régularité locale d'une trajectoire de processus », *Comptes rendus de l'Académie des sciences*, Paris, Serie I, 334, 145-148.
- Blum J.R. [1954] « Multidimensional stochastic approximation methods », *Annals of Mathematical Statistics*, 25, 737-744.
- Bluhm C., Overbeck L. [2007] « Calibration of PD term structures: to be Markov or not to be », *Risk*, 20(11), pages 98-103.
- Bollerslev T. [1986] « Generalized autoregressive conditional heteroscedasticity », *Journal of Econometrics*, 31, 307-327.
- Bouroche J.M., Saporta G., [1980] *L'Analyse des Données*, collection Que sais-je, PUF.
- Brandimarte P. [2006] *Numerical Methods in Finance and Economics: A MATLAB-Based Introduction*, Wiley-Interscience, deuxième édition.
- Branke J., Meisel S., Schmidt C. [2008] « Simulated annealing in the presence of noise », *Journal of Heuristics*, vol. 14, no 6, pp.627-654.
- Brennan M., Schwartz E. [1982] « An Equilibrium model of Bond Pricing and a Test of Market Efficiency », *Journal of Financial and Quantitative Analysis*, 17, 3, 301-329.
- Brinson G. P., Singer B.D., Beebower G.L. [1991] « Determinants of portfolio performance II : An update », *Financial Analysts Journal*, May-June 1991.
- Broadie M., Cicek D.M., Zeevi A. [2009] « An adaptative multidimensional version of the kiefer-wolfowitz stochastic approximation algorithm », *Proceeding of the 2009 Winter Simulation Conference*.
- Bulger D.W., Romeijn H.E. [2005] « Optimising noisy objective functions », *Journal of Global Optimization*, 31 : 599-600.

Byström H.N. [2002] « Using simulated currency rainbow options to evaluate covariance matrix forecasts », *Journal of International Financial Markets, Institutions and Money*, 12, 216-230.

Campbell J.Y., Viceira L.M. [2001] *Strategic Asset Allocation: Portfolio Choice for Long-Term Investors*, Oxford University Press.

Campbell J.Y., Viceira L.M. [2005] « The term structure of the risk-return trade off », *Financial Analysts Journal*, 61, 34-44.

Cariño D.R., Kent T., Myers D.H., Stacy C., Sylvanus M., Turner A., Watanabe K., Ziemba W. T. [1994] « The Russell-Yasuda Kasai Model : An Asset Liability Model for a Japanese Insurance Company using Multi-stage Stochastic Programming », *Interfaces*, 24, Jan-Feb 1994, 29-49.

Cariño, D.R., Turner A. [1998] « Multiperiod Asset Allocation with Derivative Assets », dans *Worldwide Asset and Liability Modelling*, Ziemba & Mulvey, Cambridge University Press, Cambridge , UK, pp.182-204.

Carter J. [1991] « The derivation and application of an Australian stochastic investment model », *Transactions of the Institute of Actuaries of Australia*. 1991, I 315-428.

Castro J. [2009], « A stochastic programming approach to cash management in banking », *European Journal of Operational Research*, 2009, vol. 192, issue 3, pages 963-974.

Chan L.K., Karceski J., Lakonishok J. [1999] « On portfolio optimization: forecasting covariances and choosing the risk model », *Review of Financial Studies*, 12, 937-974.

Chan K.C., Karolyi G.A., Longstaff F.A., Sanders A.B. [1992] « An empirical comparison of alternative models of short-term interest rate », *Journal of Finance*, 47:1209-1227.

Chourdakis K. [2004] « Characteristic functions and Regime Switching Models », *CCFEA Workshop on Computational Finance 2004*, <http://thePonyTail.net>

Cohen G., Culioli J.-C. [1994] *Optimisation stochastique sous contraintes en espérance*, Rapport interne Centre Automatique et Systèmes, Ecole des Mines de Paris, no. A-288.

Consigli G., Dempster M.A.H. [1998] *The CALM stochastic programming model for dynamic asset-liability management*, dans *Worldwide Asset and Liability Modelling*, Ziemba & Mulvey, Cambridge University Press, Cambridge , UK, pp 464-500.

COSP - Stochastic Programming Community : <http://stoprog.org/> [17, 101]

Cox J.C., Huang C.F. [1989] « Optimum Consumption and Portfolio Policies When Asset Prices Follow a Diffusion Process », *Journal of Economic Theory*, 49, 33-83.

Cox J.C., Ingersoll J.E., Ross S.A. [1985] « A Theory of the Term Structure of Interest Rates », *Econometrica*, 53:385-407.

Dantzig G.B. [1955] « Linear programming under uncertainty », *Management Science* 1, (1955) pp. 197-206.

Dantzig G.B., Glynn P.W. [1990] « Parallel processors for planning under uncertainty », *Annals of Operations Research*, 22, 1-22.

Date P., Wang C. [2009] « Linear Gaussian Affine Term Structure Models with Unobservable Factors: Calibration and Yield Forecasting », *European Journal of Operations Research*, 195 : 156-166 BURA.

Daykin C.D., Hey G.B. [1990] « Managing uncertainty in a general insurance company », *Journal of the Institute of Actuaries*, 117, 173-277

De Berg M., Cheong O., van Kreveld M., Overmars M. [2008] *Computational Geometry : Algorithms and Applications*, Springer-Verlag.

Décamps J.P. [1993] « Une Formule Variationnelle pour les Obligations du Secteur Privé », *Finance*, 14, 2, 61-77.

Delarue A. [2001] « Evaluation des réserves prudentielles en répartition », Lettre de l'Observatoire des retraites, n°12 mars 2001.

Deler B., Nelson B.L. [2001] « Modeling and generating multivariate time series with arbitrary marginals and autocorrelation structures », *Proceedings of the 2001 Winter Simulation Conference*, Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, 275-282.

Dempster M., Evstigneev I., Schenk-Hoppé K. [2003] « Exponential Growth of Fixed Mix Assets in Stationary Markets », *Finance and Stochastics*, 7, 263-276.

Dert C. [1995] *A dynamic model for asset/ liability management for defined benefit pension funds*, dans *Worldwide Asset and Liability Modelling*, Ziemba & Mulvey, Cambridge University Press, Cambridge , UK, pp 501-536.

Dogan K., Goetschalkx M. [1999] « A primal decomposition method for the integrated design of multi-period production–distribution systems », *IIE Transactions*, Volume 31, Number 11, November 1999 , pp. 1027-1036(10).

Dupacova J., Consigli G., Wallace S. [2000] « Generating scenarios for multistage stochastic programs », *Annals of Operations Research*, 100 (2000) 25-53 [17]

Dupacova J., Polivka J. [2009] « Asset-liability management for Czech pension funds using stochastic programming », *Annals of Operations Research*. Volume 165, Number 1 / janvier 2009, pages 5-28.

Durbin J., Watson G. S. [1951] « Testing for Serial Correlation in Least Squares Regression, II » *Biometrika* 38, 159–179.

Dyson A.C.L., Exley C.G. [1995] « Pension fund asset valuation and investment », *British Actuarial Journal*, 1: 471-557.

Elton E. J., Gruber M.J. [1973] « Estimating the dependence structure of share prices – implications for portfolio selection », *Journal of Finance*, 28, 1203-1232.

Emmerich M.T.M. [2005] *Single and Multi-objective evolutionary design optimization assisted by Gaussian Random Field Metamodels*, Thèse de doctorat, Université de Dortmund, Dortmund.

Engle R.F. [1982] « Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflation », *Econometrica*, 50, 987-1007.

Escudero L.F., Kamesam P.V., King A., Wets R.J.B. [1993] « Production planning via scenario modeling », *Annals of Operations Research*, 43, 311-335.

Faleh A. [2011] « Un modèle de programmation stochastique pour l'allocation stratégique d'actifs d'un régime de retraite partiellement provisionné », *soumis*.

Faleh A., Planchet F., Rullière D. [2010] « Les générateurs de scénarios économiques : de la conception à la mesure de la qualité », *Assurances et gestion des risques*, n° double avril/juillet, Vol. 78 (1/2).

Fama E.F., [1981] « Stock Returns, Real Activity, Inflation, and Money », *American Economic Review*, September, pp. 545-565.

Fisher I. [1930] *The Theory of Interest*, New York: Macmillan.

Fisher R.A. [1954] *Statistical Methods for Research Workers*, 12^{ème} édition, Edinburgh, Oliver & Boyd, 339 p.

Föllmer H., Kabanov Y. [1998] « Optional Decomposition and Lagrange Multipliers », *Finance and Stochastics*, Vol 2 No 1, 1998.

Fourer R., Gay D.M., Kernighan B.W. [2002] *AMPL : A Modelling Language for Mathematical Programming*, 2nd édition, Duxbury Press.

Frauendorfer K., Jacoby U., Schwendener A. [2007] « Regime switching based portfolio selection for pension funds », *Journal of Banking & Finance* 31 (2007) 2265–2280.

Friggit J. [2001] *Prix des logements, produits financiers immobiliers et gestion des risques*, Paris : Economica.

Friggit J. [2007] « Long Term (1800-2005) Investment in Gold, Bonds, Stocks and Housing in France – with Insights into the USA and the UK : a Few Regularities », *CGPC Working paper*. <http://www.adeq.org/statistiques/index.htm>

Gibson M.S., Boyer B.H. [1998] « Evaluating forecasts of correlation using option pricing », *Journal of Derivatives*, 6, 13-38.

Ginsbourger D. [2009] *Multiple métamodèles pour l'approximation et l'optimisation de fonctions numériques multivariées*, Thèse de doctorat de mathématiques appliquées, Ecole nationale supérieure des mines de Saint-Etienne, no 519MA.

- Godfrey L. [1979] « Testing the Adequacy of a Time Series Model », *Biometrika*, 64, 67-72.
- Gordon M., Shapiro E. [1956] « Capital Equipment Analysis : The Required Rate of Profit », *Management Science*, 3, p. 102-110.
- Hainaut D., Devolder P. [2005] « Management of a pension fund under a VaR constraint », *working paper*.
- Hainaut D., Devolder P. [2007] « Management of a pension fund under mortality and financial risks », *Insurance: Mathematics and Economics*, 41 (2007) 134–155.
- Hamilton J. D. [1989] « A New Approach to the Economic Analysis of Nonstationary Time Series and the Business Cycle », *Econometrica*, 57, 357-384
- Hamilton J. D., Susmel R. [1994] « Autoregressive Conditional Heteroskedasticity and Changes in Regime », *Journal of Econometrics*, 64, 307-333.
- Hansen E.R. [1979] « Global optimization using interval analysis : the one dimensional case », *Journal of Optimization Theory and Applications*, 29 :331-344.
- Hardy M.R., [2001] « A Regime Switching Model of Long-Term Stock Returns », *North American Actuarial Journal*, 5 (2), 41-53. http://www.soa.org/library/journals/north-american-actuarial-journal/2001/april/naaj0104_4.pdf
- Heath D., Jarrow R.A., Merton A. [1992] « Bond Pricing and the Term Structure of Interest Rates: A New Methodology for Contingent Claims Valuation », *Econometrica*, 60:77-105.
- Hibbert J., Mowbray P., Turnbull C. [2001] *A Stochastic Asset Model & Calibration for Long-Term Financial Planning Purposes*, Rapport Barrie & Hibbert Limited. http://www.actuaries.org.uk/__data/assets/pdf_file/0014/26312/hibbert.pdf
- Hicks J.R. [1946] *Value and capital*, Second Edition, Oxford : Clarendon Press. Edition française publiée en 1956, "Valeur et Capital", Dunod.
- Hilli P., Koivu M., Pennanen T., Ranne A. [2007] « A stochastic programming model for asset and liability management of a Finnish pension company », *Annals of Operations Research*, 152(2007), pp. 115-139
- Ho T.S.Y., Lee S.B. [1986] « Term Structure Movements and Pricing Interest Rate Contingent Claims », *Journal of Finance*, 41:1011-1029.
- Hochreiter R., Pflug G. Ch. [2002] *Scenario tree generation as a multidimensional facility location problem*, Rapport technique AURORA, TR2002-33, December 2002.
- Holmer M. [1995] *Integrated asset/ liability management: an implementation case study*, dans *Worldwide Asset and Liability Modelling*, Ziemba & Mulvey, Cambridge University Press, Cambridge , UK.
- Horst R., Pardalos P.M. [1995] « Handbook of Global Optimization », Kluwer Academic Publishers, Dordrecht Boston London.

Horta P., Mendes C., Vieira I. [2008] « Contagion effects of the US Subprime Crisis on Developed Countries », CEFAGE-UE, *Working Paper* 2008/08.

Hoyland K., Kaut M., Wallace S.W. [2003] « Heuristic for moment-matching scenario generation », *Computational Optimization and Applications*, 24(2-3):169–185, 2003.

Huber P. [1995] « A review of Wilkie's stochastic asset model », *British Actuarial Journal*, 1, 181-211.

Hull J. [2000] *Options, Futures, and Other Derivative Securities*, New Jersey: Prentice-Hall, 5th. Edition.

Hull J.C., White A. [1990] « Pricing Interest Rate Derivative Securities », *Review of Financial Studies*, 3:573-592.

Hull J.C., White A. [1994] « Numerical Procedures for Implementing Term Structure Models II: Two-Factor Models », *Journal of Derivatives* (Winter), 37-48.

Jarque C.M., Bera A.K. [1980] « Efficient tests for normality, homoscedasticity and serial independence of regression residuals », *Economics Letters* 6 (3): 255–259.

Johnson N. L. [1949] « Systems of frequency curves generated by methods of translation », *Biometrika* 36 (1): 297-304.

Jones D.R., Pertunen C.D., Stuckman B.E. [1993] « Lipschitzian optimization without the Lipschitz constant », *Journal of Optimization Theory and Applications*, 79[1], 157-181.

Jones D.R., Schonlau M., Welch W.J. [1998] « Efficient global optimization of expensive blackbox functions », *Journal of Global Optimization*, 13, 455-492.

Kallberg J.G., White R.W., Ziemba W.T. [1982] « Short Term Financial Planning under Uncertainty », *Management Science*, Vol. 28, No. 6, June 1982, pp. 670-682.

Kaut M., Wallace S.W., [2003] « Evaluation of scenario-generation methods for stochastic programming », SPEPS, *Working Paper*, 14 (<http://edoc.hu-berlin.de/series/speps/2003-14/PDF/14.pdf>), 2003.

Kharoubi-Rakotomalala C. [2008] *Les fonctions copules en finance*, Publications De La Sorbonne, collection Sorbonensia Oeconomica.

Kiefer J., Wolfowitz J. [1952] « Stochastic estimation of the maximum of a regression function », *Annals of Mathematical Statistics*, 23, 462-466.

Kim D., Santomero A. [1988] « Risk in banking and capital regulation », *Journal of Finance*, Vol. 43 (5), p.1219-1233.

Kim J.R. [2003] « The stock return-inflation puzzle and asymmetric causality in stock returns, inflation and real activity », *Discussion paper*, Economic Research Centre of the Deutsche Bundesbank, 03/03, January.

- Klassen P. [1997] « Solving stochastic programming models for asset/ liability management using iterative disaggregation », *Research Memorandum, 0011*, VU University Amsterdam, Faculty of Economics, Business Administration and Econometrics.
- Kouwenberg R. [2001] « Scenario Generation and Stochastic Programming Models for Asset Liability Management », *European Journal of Operational Research*, vol. 134, 51-64.
- Krige D.G. [1951] *A statistical approach to some mine valuations and allied problems at the Witwatersrand*, Thèse de master de l'Université de Witwatersrand.
- Kusy M.I., Ziemba W.T. [1986] « A Bank Asset and Liability Management Model », *Operations research*, Vol. 34, No. 3, May-June 1986, pp. 356-376.
- Lawler E.L., Wood D.E. [1966] « Branch and Bound methods : a survey », *Operations Research*, Vol. 14, no 4, pp 699-719.
- Leibowitz M.L., Kogelman S., Bader L.N. [1992] « Asset Performance and Surplus Control: A Dual-Shortfall Approach », *The Journal of Portfolio Management*, WINTER.
- Le Vallois F., Palsky P., Paris B., Tosetti A. [2003] *Gestion actif passif en assurance vie : Réglementation - Outils – Méthodes*, Economica.
- Macaulay F.R. [1938] « The Movements of Interest Rates. Bond Yields and Stock Prices in the United States since 1856 », New York: *National Bureau of Economic Research*.
- Mandelbrot B. [2005] *Une approche fractale des marchés : risquer, perdre et gagner*, Editions Odile Jacob.
- Malliari A. G., Brock W.A. [1982] *Stochastic Methods in Economics and Finance*, New York: Elsevier Science Publishers.
- Markowitz H. M. [1952] « Portfolio Selection », *The Journal of Finance*, Vol. 7, No.1, March.
- Martellini L. [2006] *Managing Pension Assets: from Surplus Optimization to Liability-Driven Investment*, EDHEC Risk and Asset Management Research Centre.
- Martellini L., Priaulet P., Priaulet S. [2003] *Fixed-Income Securities: Valuation, Risk Management and Portfolio Strategies*, Wiley, 2003.
- Mathias K., Whitley D., Kusuma A., Stork C. [1996] « An empirical evaluation of genetic algorithms on noisy objective functions », *Genetic Algorithms for Pattern Recognition* S.K. Pal, ed. pp: 65-86. CRC Press, 1996.
- Mazzola J.B., Schantz R.H. [1996] « Multiple-facility loading under capacity-based economies of scope », *Naval Research Logistics (NRL)*, Volume 44 Issue 3, Pages 229 – 256.
- Merton R. [1971] « Optimum Consumption and Portfolio Rules in a Continuous Time Model », *Journal of Economic Theory*, 3, 373-413.

Merton R. [1976] « Option pricing when underlying stock returns are discontinuous », *Journal of Financial Economics*, 3, pp.125-144.

Merton R. [1990] *Continuous-Time Finance*, Oxford, UK.

Meucci A. [2005] *Risk and asset allocation*, New York : Springer.

Mitra S. [2006] « A White Paper on Scenario Generation for Stochastic Programming », *OptiRisk Systems:White Paper Series*, juillet.

Mulvey J.M., Thorlacius A.E. [1998] *The Towers Perrin Global Capital Market Scenario Generation System: CAP:Link*, dans *Worldwide Asset and Liability Modelling*, Ziemba & Mulvey, Cambridge University Press, Cambridge , UK.

Munk C., Sorensen C., Vinther T.W. [2004] « Dynamic asset allocation under mean-reverting returns, stochastic interest rates and inflation uncertainty », *International Review of Economics & Finance*, Volume 13, Février 2004,141-166.

Nelder J., Mead R. [1965] « A simplex method for function minimization », *Computer Journal*, vol. 7, no 4, p.308-313.

Norkin V., Pflug G.Ch., Ruszczyński A. [1996] « A branch and bound method for stochastic global optimization », *Mathematical Programming*, vol 83, no 1-3, pp 452-450.

Pennanen T., Koivu M. [2002] « Integration quadratures in discretization of stochastic programs », *Stochastic Programming*, E-Print Series, <http://www.speps.info> , May 2002.

Perold A. F, Sharpe W. F. [1988] « Dynamic Strategies for Asset Allocation », *Financial Analysts Journal*, 44, no. 1.

Pierre S.N. [2009] *Passif social, construction du portefeuille d'investissement et couverture du risque de taux* , Mémoire d'actuaire, CNAM.

Pirkul H., Jayaraman V. [1996] « Production, Transportation, and Distribution Planning in a Multi-Commodity Tri-Echelon System », *Transportation Science*, 30(4), 291-302.

Planchet F. [2009] « Quel modèle d'actifs en assurance ? », *la Tribune de l'Assurance* (rubrique « le mot de l'actuaire »), n°136 du 01/05/2009.

Planchet F., Thérond P.E. [2005] *Modèles financiers en assurance – Analyses de risque dynamiques*, Economica.

Planchet F., Thérond P.E. [2007] *Pilotage technique d'un régime de rentes viagères*, Economica.

Planchet F., Thérond P.E., Kamega A. [2009] *Scénarios économiques en assurance - Modélisation et simulation*, Paris : Economica.

Portrait R., Poncet P. [2008] *Finance de marché – Instruments de base, produits dérivés, portefeuilles et risques*, Dalloz.

Rambaruth G. [2003] *A comparison of Wilkie-type stochastic investment models*, Mémoire de master, City University, UK.

Ranne A. [1998] « The Finnish Stochastic Investment Model », *Transactions of 26th ICA*, 213-238.

Rasmussen K.M., Clausen J. [2004] « Mortgage Loan Portfolio Optimization Using Multi Stage Stochastic Programming », *Journal of Economic Dynamics and Control*, 31, 742-766.

Redington F.M. [1952] « Review of the principles of life office valuations », *Journal of the Institute of Actuaries*, 78: 1-40.

Robbins H., Monro S. [1951] « A Stochastic approximation method », *Annals of Mathematical Statistics*, 22, 400-407.

Roncalli T. [1998] *La structure par terme des taux zéro : modélisation et implémentation numérique*, Thèse de doctorat, Université de Bordeaux IV.

Roy A.D. [1952] « Safety-first and the holding of asset », *Econometrica*, 20, 431-449.

Rudolf M., Ziemba W.T. [2004] « Intertemporal surplus management », *Journal of Economic Dynamics & Control*, 28 (2004) 975 – 990.

Rullière D., Faleh A., Planchet F. [2010] « Un algorithme d'optimisation par exploration sélective », *soumis*.

Rullière D., Faleh, A., Planchet F. [2010] « Fonction d'information et optimisation globale d'une fonction bruitée », *preprint*.

Santner T., Williams B., Notz W. [2003] *The design and analysis of computer experiments*, Springer Verlag, New York.

Santosa F., Alexandrov O. [2005] « A topology-preserving level set method for shape optimization », *Journal of Computational Physics*, Volume 204 , Issue 1 (March 2005), Pages: 121 – 130.

Schubert B. [1972] « A sequential method seeking the global maximum of a function », *SIAM Journal on Numerical Analysis*, 9 :379-388.

Sharpe W.F., Tint L.G. [1990] « Liabilities - A new approach », *Journal of Portfolio Management*, Winter, 1990, 5-10.

Shapiro A., Verweij B., Shabbir A., Kleywegt A.J., Nemhauser G. [2003] « The Sample Average Approximation Method Applied to Stochastic Routing Problems: A Computational Study », *Computational Optimization and Applications*, Volume 24, Numbers 2-3 / février 2003, pages 289-333.

Shiller R.J. [1981] « Do Stock Prices Move Too Much to be Justified by Subsequent Changes in Dividends? », *NBER Working Papers* 0456, National Bureau of Economic Research, Inc.

Skintzi V.D., Spyros X.S. [2007] « Evaluation of correlation forecasting models for risk management », *Journal of Forecasting*, 26:7, 497-526.

Smith N.A., Tromble R.W. [2004] *Sampling uniformly from the unit simplex*, Rapport technique, Johns Hopkins University.

Smith A.D. [1996] « How actuaries can use financial economics », *British Actuarial Journal*, 2: 1057-1193.

Sommerville D.M.Y. [1958] *An introduction to the Geometry of n dimensions*, New York, Dover edition.

Strugarek, C. [2006] *Approches variationnelles et autres contributions en optimisation stochastique*, ENPC, Thèse de doctorat.

Student (ou Gosset W.S.) [1908] « The Probable Error of a Mean », *Biometrika*, 6, 1–25.

Tanaka S., Inui K. [1995] « Modelling Japanese financial markets for pension ALM simulations », *5th AFIR colloquium*, 563-584

Thomson R.J. [1994] « A stochastic investment model for actuarial use in South Africa », *Transactions of the Actuarial Society of South Africa* 1.

Tuckman B., [2002] *Fixed Income Securities: Tools for Today's Markets*, John Wiley & Sons, Inc, 10-01-2002, 528 pages.

Vasicek O. [1977] « An equilibrium Characterization of the Term structure », *Journal of Financial Economics*, 5:177-188.

Villemonteix J. [2009] *Optimisation de fonctions coûteuses*, Thèse de doctorat de physique, Université Paris Sud 11, Faculté des sciences d'Orsay, no 9278.

Wainscott C. B. [1990] « The stock-bond correlation and its implications for asset allocation », *Financial Analysts Journal*, 46, 55-60.

Waring B. [2004] « Liability-Relative Investing II: Surplus Optimization with Beta, Alpha, and an Economic View of the Liability », *working paper*, September 2004.

Wilkie D. [1984] « Steps towards a comprehensive stochastic model », *Occasional Actuarial Research Discussion Paper*, The Institute of Actuaries, London, 36:1-231.

Wilkie D. [1986] « A Stochastic Investment Model for Actuarial Use », *Transactions of the Faculty of Actuaries*, 39:341-403.

Wilkie D. [1995] « More on a Stochastic Model for Actuarial Use », *British Actuarial Journal*, pp. 777-964. <http://www.ressources-actuarielles.net/jwa/documentation/1226.nsf>.

Wolfe M.A. [1996] « Interval Methods for Global Optimization », *Applied Mathematics and Computation*, Vol. 75, Issues 2-3, pp. 179-206.

Yen S.H., Hsu Ku Y.H. [2003] « Dynamic Asset Allocation Strategy for Intertemporal, Pension Fund Management with Time-Varying Volatility », *Academia Economic Papers*, 31: 3 (September 2003), 229–261.

Zadeh L.A. [1965] « Fuzzy Set », *Information and Control*, 8 [3] : 338-353.

Zenios S.A. [2007] *Practical Financial Optimization : Decision Making For Financial Engineers*, Blackwell Publishers, Oxford, 2007.

Ziemba W.T. [2003] *The Stochastic Programming Approach to Asset-Liability and Wealth Management*, AIMR-Blackwell.

Ziemba W. T., Mulvey J.M. [1998] *Worldwide Asset and Liability Modelling*, Cambridge University Press, UK.

Table des matières

Sommaire.....	1
Introduction générale.....	4
Partie I : Les générateurs de scénarios économiques (GSE).....	17
Introduction	18
Chapitre 1 : Etude des GSE.....	19
I- Présentation théorique	19
I-1 Problèmes théoriques de modélisation liés aux GSE.....	19
I-1-1 Modèles des taux d'intérêt	19
I-1-1-1 <i>Les modèles d'absence d'opportunité d'arbitrage (AOA)</i>	19
I-1-1-2 <i>Les modèles d'équilibre général</i>	21
I-1-2 Modèles de rendement des actions	23
I-1-3 Autres classes d'actifs	24
I-2 Littérature sur les structures des GSE.....	24
I-2-1 Modèles à structure par cascade.....	25
I-2-2 Modèles basés sur les corrélations.....	27
II- Mise en œuvre d'un GSE.....	30
II-1 Structure schématique de projection de scénarios pour un GSE.....	30
II-2 Méthodologies pour la génération de scénarios économiques.....	34
II-2-1 Les approches basées sur l'échantillonnage.....	34
II-2-2 Les approches basées sur le <i>matching</i> des propriétés statistiques.....	35
II-2-3 Les approches basées sur les techniques de <i>Bootstrapping</i>	36
II-2-4 L'utilisation de l'Analyse en Composantes Principales.....	36
II-3 Probabilité risque-neutre Vs probabilité réelle.....	36
II-4 L'exclusion de valeurs négatives pour les modèles de taux.....	37

III- Mesure de qualité.....	38
III-1 Mesure qualitative de la qualité d'un GSE.....	38
III-2 Mesure quantitative de la qualité d'un GSE.....	40
III-2-1 Présentation du problème.....	40
III-2-2 Le test de la qualité des décisions obtenues par le GSE.....	42
III-2-2-1 <i>La condition de stabilité</i>	42
III-2-2-2 <i>La condition d'absence de biais</i>	44
III-3 Application numérique.....	45
Conclusion du chapitre 1.....	49
Chapitre 2 : Construction d'un générateur de scénarios économiques.....	50
I- Conception et composantes	50
I-1 Conception théorique	50
I-2 Composantes	53
I-2-1 Structure de corrélation	53
I-2-2 Choix des modèles de diffusions.....	58
I-2-2-1 <i>Le modèle de structure par terme des taux d'intérêt</i>	58
I-2-2-2 <i>Le modèle des actions</i>	65
I-2-2-3 <i>Le modèle des obligations "crédit"</i>	67
I-2-2-4 <i>Le modèle de l'immobilier</i>	71
I-2-2-4 <i>Le monétaire</i>	71
II- Calibrage.....	73
II-1 Présentation du contexte de l'étude	73
II-1-1 Description des données retenues.....	73
II-1-2 Approche et méthode de calibrage.....	75
II-2 Calibrage du modèle de l'inflation.....	76
II-2-1 Méthode.....	76
II-2-2 Données.....	77
II-2-3 Estimation des paramètres.....	77
II-2-4 Tests d'adéquation du modèle.....	78
II-3 Calibrage du modèle des taux réels.....	79
II-3-1 Méthode.....	79
II-3-2 Données.....	80

II-3-3 Estimation des paramètres.....	80
II-3-4 Tests d'adéquation du modèle.....	81
II-4 Calibrage du modèle des actions.....	83
II-4-1 Méthode.....	83
II-4-2 Données.....	88
II-4-3 Estimation des paramètres.....	88
II-5 Calibrage du modèle de l'immobilier.....	89
II-5-1 Méthode.....	89
II-5-2 Données.....	89
II-5-3 Estimation des paramètres.....	89
II-5-4 Tests d'adéquation du modèle.....	90
II-6 Calibrage du modèle de crédit.....	91
III- Projection des scénarios	93
III-1 Éléments sur la mise en œuvre.....	93
III-2 Tests du modèle : paramètres et matrice de corrélation.....	96
III-2-1 Tests sur les paramètres et sur les développements informatiques.....	97
III-2-2 Tests sur les matrices de corrélation.....	98
III-3 Résultats de projection.....	99
Conclusion du chapitre 2.....	110
Conclusion de la partie I.....	110
Partie II : L'allocation stratégique d'actifs dans le cadre de la gestion actif-passif (ALM).....	112
Introduction.....	113
Chapitre 1 : Les modèles d'ALM classiques (déterministes).....	118
I- Immunisation du portefeuille.....	118
I-1 Adossement des flux de trésorerie.....	119
I-2 Adossement par la duration.....	120

II- Modèles basés sur le surplus.....	125
II-1 Modèle de Kim et Santomero [1988].....	125
II-2 Modèle de Sharpe et Tint [1990]	128
II-3 Modèle de Leibowitz [1992].....	130
II-4 Duration du Surplus.....	137
Conclusion du chapitre 1.....	139
Chapitre 2 : Allocation stratégique d'actifs et stratégie <i>Fixed-Mix</i>.....	140
I- Etude de l'allocation d'actifs dans le cadre de la stratégie <i>Fixed-Mix</i>.....	143
I-1 Présentation.....	143
I-2 Application.....	143
I-2-1 Modélisation du régime-type.....	144
I-2-1-1 <i>Notations</i>	144
I-2-1-1 <i>Modélisation</i>	144
I-2-2 Etude des critères d'allocation.....	146
I-2-2-1 <i>Analyse générale de la solvabilité du régime</i>	147
I-2-2-2 <i>Critères d'allocation en fonction du type du régime</i>	153
I-2-2-3 <i>Critères d'allocation en fonction de la phase dans laquelle évolue le régime</i>	153
I-2-3 Etude de la question du taux d'actualisation des engagements.....	153
I-2-4 Hypothèses du modèle.....	155
I-2-5 Méthodologie retenue.....	158
I-2-6 Résultats obtenus	161
II- Proposition de méthodes numériques	168
II-1 Introduction.....	168
II-1-1 Le problème d'optimisation.....	168
II-1-2 Approche retenue.....	172
II-2 Une grille à pas variable par subdivision systématique.....	173
II-2-1 Forme des zones de recherche.....	173
II-2-2 Choix de la zone à explorer ou segmenter.....	175
II-2-3 Choix des paramètres (σ_K, α) du potentiel.....	181
II-2-4 Choix du sommet de scission.....	183
II-2-5 Schéma de l'algorithme de scission systématique.....	184
II-2-6 Résultat final et critère de convergence.....	185
II-3 Une grille à pas variable avec retraitage possible.....	188

II-3-1 Critères de scission.....	188
II-3-2 Schéma de l'algorithme à tirage possible.....	190
II-4 Applications numériques.....	191
II-4-1 Fonction test utilisée.....	191
II-4-2 Comportement de l'algorithme à scission systématique.....	193
II-4-3 Comportement de l'algorithme avec tirages possibles.....	199
II-4-4 Estimation des paramètres (σ_K, α).....	200
II-4-5 Critères de convergence.....	203
II-4-6 Comparaison avec l'algorithme de Kiefer-Wolfowitz-Blum.....	208
II-5 Conclusion.....	211
Conclusion du chapitre 2.....	212
Chapitre 3 : Allocation stratégique d'actifs et modèles d'ALM dynamique.....	213
I - L'allocation stratégique d'actifs dans le cadre des modèles classiques d'ALM dynamique.....	214
I-1 Présentation générale.....	214
I-2 Techniques d'assurance de portefeuille	215
I-2-1 Principe.....	215
I-2-2 Illustration.....	216
I-2-2-1 Contexte.....	216
I-2-2-2 Modélisation de l'actif.....	217
I-2-2-3 Modélisation du passif.....	217
I-2-2-4 Calibrage des paramètres de marché.....	218
I-2-2-5 Calibrage des paramètres de gestion.....	218
I-2-2-6 Résultats.....	219
I-3 Programmation dynamique.....	223
II- L'allocation stratégique d'actifs dans le cadre de la programmation stochastique	227
II-1 Introduction	227
II-2 Présentation de la programmation stochastique.....	228
II-3 Illustration de la programmation stochastique avec recours dans le cas de la planification de la production.....	230
II-3-1 Cas mono-périodique (à deux étapes)	230
II-3-1 Cas multi-périodique (multi- étapes)	232
II-4 Formulation mathématique de la programmation stochastique avec recours.....	236
II-5 Résolution des programmes stochastiques.....	240

III- Nouvelle approche d'ALM par la discrétisation des scénarios économiques.....	242
III-1 Méthode des quantiles de référence pour le GSE.....	242
III-2 Modèle d'optimisation basé sur la programmation stochastique.....	247
III-3 Discussion des résultats.....	253
Conclusion du chapitre 3.....	261
Conclusion de la partie II	261
Conclusion générale.....	262
Index des graphiques.....	268
Bibliographie.....	272
Table des matières.....	284