



Mesure et gestion des risques d'assurance : analyse critique des futurs référentiels prudentiel et d'information financière

Pierre-Emmanuel Thérond

► To cite this version:

Pierre-Emmanuel Thérond. Mesure et gestion des risques d'assurance : analyse critique des futurs référentiels prudentiel et d'information financière. Gestion des risques [q-fin.RM]. Université Claude Bernard - Lyon I, 2007. Français. NNT : . tel-00655896

HAL Id: tel-00655896

<https://theses.hal.science/tel-00655896>

Submitted on 3 Jan 2012

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.

N° d'ordre : 76-2007

Année 2007

THESE
présentée
devant l'Université Claude Bernard – Lyon 1
pour l'obtention
du DIPLOME DE DOCTORAT
(arrêté du 7 août 2006)

présentée
par
Pierre-Emmanuel THÉRON

**Mesure et gestion des risques d'assurance : analyse critique
des futurs référentiels prudentiel et d'information financière**

Directeur de thèse : Professeur Jean-Claude AUGROS

Soutenue publiquement le 25 juin 2007, devant le jury composé de :

Jean-Claude AUGROS (Professeur, Université Claude Bernard – Lyon 1)
François EWALD (Professeur, Conservatoire National des Arts et Métiers)
Hans GERBER (Professeur, Université de Lausanne-HEC), Rapporteur
Jean-Paul LAURENT (Professeur, Université Claude Bernard – Lyon 1)
Etienne MARCEAU (Professeur, Université Laval), Rapporteur
Frédéric PLANCHET (Professeur associé, Université Claude Bernard – Lyon 1)

Sommaire

Sommaire.....	1
INTRODUCTION GÉNÉRALE.....	3
PARTIE I NOUVELLES APPROCHES COMPTABLE, PRUDENTIELLE ET FINANCIÈRE DES RISQUES EN ASSURANCE.....	9
Chapitre 1 Traitement spécifique du risque : aspects théoriques	13
1. Les outils mathématiques de l'analyse des risques	13
2. Un traitement différencié du risque	29
3. Contrats en unités de compte : les garanties plancher.....	36
4. Conclusion	42
Bibliographie	44
Chapitre 2 Traitement spécifique du risque : aspects pratiques	47
1. De nouveaux référentiels distincts	47
2. Des incidences opérationnelles	52
3. Cas pratique : portefeuille d'assurance vie	64
Bibliographie	70
Chapitre 3 Incidence sur la gestion technique d'un assureur.....	71
1. Modélisation de la société d'assurance.....	72
2. Critère de maximisation des fonds propres économiques.....	73
3. Recherche de l'allocation optimale.....	77
4. Conclusion	88
Annexe A : Démonstration des résultats mathématiques.....	91
Annexe B : Simulation des réalisations de la charge de sinistres.....	93
Bibliographie	94
PARTIE II MODÉLISATIONS AVANCÉES EN ASSURANCE	95
Chapitre 4 Limites opérationnelles : la prise en compte des extrêmes	97
1. Calcul de VaR en assurance.....	99
2. Notations.....	101
3. Estimation de quantiles extrêmes.....	102
4. Application du bootstrap.....	105
5. Robustesse du SCR.....	111
6. Conclusion	117

Annexe A : Loi de Pareto généralisée (GPD)	118
Annexe B : Résultats probabilistes	121
Annexe C : Estimation du paramètre de queue	125
Bibliographie	129
Chapitre 5 Prise en compte de la dépendance.....	133
1. Analyse mathématique de la dépendance	134
2. Rappels sur la dépendance linéaire	143
3. La théorie des copules.....	145
4. Rachat de contrats d'épargne : un modèle ad hoc.....	161
5. Capital de solvabilité : les méthodes d'agrégation.....	165
Bibliographie	168
Chapitre 6 Techniques de simulation.....	171
1. Introduction.....	171
2. Discrétisation de processus continus.....	172
3. Estimation des paramètres	180
4. Génération des trajectoires.....	185
5. Simulation de la mortalité d'un portefeuille d'assurés.....	193
6. Conclusion	197
Annexe : Tests d'adéquation à une loi	198
Bibliographie	200
CONCLUSION GÉNÉRALE	201
Annexe Solvabilité 2 - les modèles proposés par QIS 3	205
1. Modèles d'évaluation : l'approche standard	206
2. Provisions techniques en assurance vie.....	209
3. Provisions techniques en assurance non-vie	210
4. Capital éligible.....	210
5. Capital de solvabilité (SCR) : formule standard	211
BIBLIOGRAPHIE GÉNÉRALE.....	217
Notations utilisées	225
Table des illustrations	227
Table des matières	229

INTRODUCTION GÉNÉRALE

Les prémices de l'assurance moderne

Les premiers dispositifs d'assurance modernes remontent au XIV^e siècle et concernent l'assurance maritime. Le commerce maritime s'était jusqu'alors développé grâce aux *prêts à la grosse aventure* qui consistaient, pour l'armateur, à emprunter une somme d'argent gagée sur la valeur des marchandises qui devaient être expédiées par delà les mers. En cas d'arrivée à bon port, l'emprunteur remboursait cette somme majorée d'un intérêt très élevé venant en contrepartie du risque des voyages maritimes puisque, si la cargaison était perdue, ni l'intérêt, ni le capital n'étaient remboursés.

Cette opération n'était pas à proprement parler une opération d'assurance mais plutôt une opération de crédit dans laquelle le prêteur porte le risque de la perte du capital. Le fait que l'intérêt stipulé soit fixé arbitrairement sans réelle considération au risque effectivement encouru a conduit au développement d'une activité spéculative sur ce type de contrats. Aussi, le recours au prêt à la grosse aventure a pratiquement disparu suite à l'interdiction de la pratique de l'usure par l'Église catholique en 1234. Il a alors fallu trouver un autre moyen de financer ces expéditions qui participaient au développement du commerce et de l'économie. C'est ainsi que naît la convention d'assurance sous sa forme moderne : des banquiers acceptent de garantir la valeur du bateau et de sa cargaison (le capital sous risque) en cas de naufrage du navire (le risque) contre le paiement d'une somme fixe (la prime). Le développement de ce type de contrat connaît un réel succès compte tenu, notamment, des conséquences de la survenance des risques de la mer sur la solvabilité des armateurs.

Le deuxième type d'assurance à connaître un essor considérable est l'assurance incendie. Le début de son développement est généralement associé au gigantesque incendie qui frappa Londres au XVII^e siècle. En effet, le 2 septembre 1666, un feu se déclare dans Londres et ravage les quatre cinquièmes de la ville en une semaine. Le bilan¹ est de 13 200 maisons et 87 églises détruites dont la cathédrale Saint-Paul. Le coût de la reconstruction est considérable, aussi l'État britannique encourage la création de sociétés d'assurance pour couvrir ce risque dans le futur.

Les facteurs du développement de l'assurance

Ces deux exemples historiques illustrent parfaitement les deux axes qui ont conduit ces derniers siècles au développement de l'assurance.

Le premier est celui du développement économique : des projets de grande envergure comportent un grand nombre de risques différents dont la survenance d'un seul de ces risques peut mettre en péril l'intégralité de l'entreprise. Les opérations d'assurance, en transférant le risque de l'assuré vers l'assureur, permettent d'envisager de telles initiatives qui seraient impossibles sinon. C'est ainsi que des projets qui n'étaient concevables auparavant que par les États du

¹. Cf. Ewald et Lorenzi (1998) pour les données chiffrées sur l'incendie de 1666.

fait de leur capacité à lever l'impôt et de la solvabilité que cela leur procure, ont pu, au fur et à mesure du développement de l'activité d'assurance, être envisagés par des sociétés privées qui, par définition, disposent de ressources limitées.

Par exemple, la construction du viaduc de Millau (400 M€) a été intégralement financée par la Compagnie Eiffage du viaduc de Millau et réalisée par des sociétés du groupe Eiffage en échange d'une concession de l'exploitation pour une durée de 75 ans. C'est donc cette société de droit privé qui supporte les risques inhérents à la construction. Une société privée ne peut envisager de supporter le coût d'un tel édifice sans la mise en place de dispositifs d'assurance ad hoc.

On peut trouver d'autres exemples dans des domaines tels que l'envoi de satellites spatiaux, initialement réservé à des usages militaires et scientifiques, qui est maintenant largement ouvert aux sociétés de communication.

Le second axe de développement est celui de la sécurisation financière de la société en général et de l'individu en particulier. Lorsque les autorités britanniques favorisent le développement des sociétés d'assurance après l'incendie de Londres de 1666, cela permet aux londoniens d'avoir la possibilité de se prémunir de ce risque qui pouvait emporter la totalité de leur richesse. De ce point de vue, le contrat d'assurance ouvre une option à l'individu ou la compagnie : celle de ne pas subir le risque comme une fatalité. L'offre d'assurance met donc l'agent économique, quel qu'il soit, devant sa responsabilité en lui offrant une issue lui permettant d'échapper, non pas à la survenance du risque, mais à ses conséquences financières.

Dans certains cas, le processus de sécurisation de la société va plus loin, puisque l'État a été amené à imposer l'obligation d'assurance dans un certain nombre de cas : responsabilité civile professionnelle, responsabilité civile du conducteur, assurance maladie ou retraite pour les salariés, etc. Le caractère obligatoire de ces assurances vise à protéger les assurés (pour les assurances de personnes obligatoires) et les individus à qui ils seraient susceptibles de causer des dommages (assurances de responsabilité).

Un rôle croissant dans l'économie

Sous l'impulsion des deux facteurs précédemment évoqués, le marché de l'assurance connaît un développement continu depuis ses prémices.

En France, le développement de l'assurance s'avère initialement plus lent qu'en Angleterre. Pour preuve, en 1681, l'ordonnance de Colbert sur les activités liées à la mer est muette en matière d'assurance. De plus, la même année, les opérations d'assurance vie sont prohibées car jugées immorales. Dans le même temps, la première table de mortalité construite sur de réelles bases scientifiques voit le jour en 1690, en Angleterre. Cette date marque le début de l'assurance vie. Cependant, en France, ce n'est qu'en 1787, qu'un arrêt du Conseil d'État autorise la création d'une Compagnie royale des assurances sur la vie. Autorisation de courte durée, puisque six ans plus tard, les entreprises d'assurances sont supprimées par la révolution car considérées comme spéculatives. Il faudra attendre la restauration pour voir la renaissance des sociétés d'assurance.

La deuxième partie du XIXe et le XXe siècles voient l'expansion de la diversité de l'offre en assurance : de plus en plus de risques sont concernés. Ainsi, le montant des capitaux assurés par les sociétés françaises d'assurance vie progresse, en moyenne, de 35 % par an entre 1907 et 1913 (cf. le rapport du Sénat de Lambert (1998)).

Si le montant des capitaux assurés progresse, c'est également le cas des provisions techniques que doivent constituer les assureurs du fait de l'inversion de leur cycle de production. En effet, les assureurs percevant les primes avant de régler les prestations, ils doivent constituer des réserves de manière à être capable de payer les sinistres en leur temps. Ces provisions sont investies sur les marchés financiers de manière à dégager des produits financiers.

Le niveau de ces encours est un indicateur du développement de l'activité d'assurance. Ainsi entre 1976 et 1995, ils ont augmenté de 17,8 % en moyenne par an (cf. Ewald et Lorenzi (1998)). À la fin de l'année 2005, la valeur² des actifs gérés par les assureurs s'élève à 1 285 milliards d'euros. Cette somme est investie, pour moitié, dans des obligations ou des actions d'entreprises industrielles et commerciales. Ainsi l'assurance est *de facto* un acteur majeur du développement économique en cela qu'il participe fortement au financement des projets industriels et commerciaux. Le secteur de l'assurance est également un créancier des états puisque 36 % de leur actif est constitué d'obligations d'état. Le reste étant investi en immobilier et en monétaire.

Un contrôle spécifique

En vertu de leur rôle essentiel dans la sécurisation financière des individus comme des sociétés, les compagnies d'assurance font l'objet d'une attention particulière des pouvoirs publics avec pour objectif de préserver l'intérêt des assurés, des souscripteurs et des bénéficiaires de contrats. Le contrôle des assurances se développe à partir des années 30 dans un souci de protection des assurés en vue de parer à des situations telles que celle de la caisse Lafarge mise en place à la fin du XVIIIe siècle et qui s'est terminée en 1888. Cette affaire symbolise tout ce qu'il ne faut pas faire en assurance : mauvaise information des assurés qui, au surplus, ne comprennent pas les mécanismes en jeu, utilisation d'hypothèses actuarielles (mortalité) extravagantes, opportunités d'arbitrages des assurés, malversations suspectées des créateurs de la tontine et affaires de la révolution...

La veille prudentielle de l'État sur les opérations d'assurance s'organise en deux niveaux.

Le premier niveau d'intervention du législateur est celui des contrats. En effet, le Code des assurances oblige les assureurs à inscrire dans les contrats des clauses de protection des assurés (telles que les clauses de rachat de contrat, de renonciation, etc.) et impose un formalisme censé procurer à l'assuré une bonne visibilité des droits contractuels dont il bénéficie.

Le deuxième niveau relève du contrôle de la solidité financière des assureurs et donc de leur capacité à honorer leurs engagements. Le droit

². Valeur estimée par la Fédération Française des Sociétés d'Assurances (FFSA) dans son rapport annuel.

national et la transposition des directives européennes a conduit à la situation actuelle dans laquelle la solvabilité des entreprises d'assurance est appréciée annuellement par l'autorité de contrôle³ au moyen de ses comptes sociaux, d'états réglementaires spécifiques et de divers rapports (rapport de réassurance, rapport de solvabilité, etc.).

Il est d'ailleurs intéressant de noter que les actuelles règles de solvabilité se sont enrichies dans le temps avec l'évolution de la gestion des risques et notamment la prise en compte récente de certains risques⁴ pesant sur les assureurs. Aussi le système de solvabilité s'est construit en « empilant » des nouvelles contraintes au fur et à mesure de l'identification de nouveaux risques (cf. Thérond (2005)). C'est ainsi que sont apparus récemment de nouveaux états réglementaires tels que :

- les états trimestriels T3 de suivi actif-passif ;
- le test d'exigibilité (état C6bis) qui permet d'apprécier la solidité de l'entreprise à des scénarios adverses d'actifs (chute des marchés boursiers et immobilier, hausse des taux d'intérêt) et de passifs (triplement des rachats, sur-sinistralité, etc.) ;
- les stress-tests sur les dispositifs de réassurance (états C8 et C9).

De nouveaux enjeux

Aujourd'hui, le secteur de l'assurance est confronté à une triple mutation :

- prudentielle avec l'avènement du futur cadre prudentiel européen qui résultera du projet Solvabilité 2 ;
- du reporting financier avec le recours de plus en plus massif aux méthodes d'*European Embedded Value* de valorisation de compagnie d'assurance ;
- comptable avec la préparation de la phase II de la norme internationale IFRS consacrée aux contrats d'assurance.

Dans ce contexte, les assureurs sont invités, pour chacun de ces trois aspects, à mieux identifier, mesurer et gérer les risques auxquels ils sont soumis.

Ces trois référentiels s'inscrivent dans une même logique d'uniformisation internationale (à tout le moins communautaire) et de transparence. Pour cela, la référence « au marché » est omniprésente : dès que c'est possible, c'est en référence au marché que les engagements d'assurance doivent être valorisés. En particulier, les risques financiers doivent être traités de la même manière que des instruments financiers qui seraient cotés sur un marché financier liquide. Ce principe n'est pas sans poser des problèmes conceptuels et opérationnels.

Pour s'en convaincre, il suffit de considérer un « simple » contrat d'assurance vie de type épargne. Ces contrats présentent le plus souvent des taux garantis (taux minimum garanti, taux technique ou taux annuel garanti) et

³. L'Autorité de Contrôle des Assurances et des Mutuelles (ACAM).

⁴. La plupart de ces risques ne sont pas à proprement parler nouveaux : c'est généralement l'augmentation de la vraisemblance de leur survenance et l'évolution des outils techniques qui a entraîné leur prise en compte.

une clause de participation aux bénéfices techniques et financiers (clause réglementaire du Code des assurances ou contractuellement plus favorable pour l'assuré). Ainsi, chaque fin d'année, le montant de l'épargne accumulée est revalorisé des intérêts techniques et de la participation aux bénéfices. Dans les faits, la participation aux bénéfices accordée aux assurés s'avère souvent supérieure à la clause réglementaire ou contractuelle. Par ailleurs, la plupart des nouveaux contrats ne disposent plus de taux technique ou de taux minimum garanti sur la durée du contrat mais plutôt des taux garantis annuellement en fonction des performances de l'année écoulée. Or pour ces contrats, l'engagement de taux est relativement faible et la PB discrétionnaire constitue une grande partie de la revalorisation. La revalorisation provient donc, en grande partie de la décision de l'assureur. Cette décision intervient de deux manières :

- dans le pilotage du rendement financier de la compagnie,
- dans le choix du niveau de revalorisation des contrats.

La latitude mentionnée au premier point réside dans le fait que les clauses de participation aux bénéfices font référence au rendement comptable des actifs en représentation des engagements⁵. Ainsi un contrat pourra prévoir que les assurés bénéficieront d'au moins 95 % des bénéfices financiers. Si cette règle semble constituer un minimum objectif, l'assureur dispose d'une marge de manœuvre dans la constitution du rendement qui sert d'assiette à ce calcul. En effet, du fait du mode de comptabilisation des actifs financiers (la règle générale étant le coût historique dans la réglementation française), un assureur qui dispose de titres en plus ou moins-values latentes peut décider, par exemple, de vendre des titres en plus-values latentes, pour doper son rendement ou, au contraire, de les conserver comme matelas de sécurité pour des temps plus difficiles.

Concernant le deuxième point, de manière discrétionnaire⁶, l'assureur peut revaloriser davantage l'épargne pour certains produits que pour d'autres, en favorisant, par exemple, les plus rentables de manière à conserver ces contrats. De plus, les assurés bénéficient du droit de racheter leur contrat à n'importe quel moment moyennant une pénalité de rachat qui ne peut pas être supérieure à 5 % de la valeur de l'épargne.

On voit donc que, pour un contrat aussi simple que celui présenté et dont on aurait pu dire, *a priori*, qu'il s'agissait uniquement d'un produit dérivé sur l'actif de l'assureur, le dénouement du contrat et l'évolution de l'épargne va non seulement dépendre de celle des rendements financiers mais également du comportement de l'assureur (revalorisation notamment) et de l'assuré (rachat). Par ailleurs, on se rend rapidement compte que comportements des assurés et de l'assureur ne sont pas indépendants : si l'assureur ne revalorise pas suffisamment l'épargne, il risque de voir les assurés racheter leurs contrats et ne pas en souscrire d'autres. De plus, en pratique, on observe des comportements que l'on pourrait juger irrationnels : rachats de contrats avec des taux garantis plus élevés que les taux actuels de marché (en vue de financer un emprunt par

⁵. Le plus souvent l'actif général de la compagnie et dans certains cas des actifs cantonnés.

⁶. L'assureur est tenu de revaloriser l'épargne selon un minimum contractuel mais cela ne constitue qu'un minimum, il peut leur accorder une revalorisation plus favorable.

exemple), par exemple. Aussi l'assimilation, à des fins de valorisation, d'une option de rachat sur un contrat d'épargne en euros à une option financière à caractère européen n'est pas évidente puisque, *de facto*, les assurés ne se comportent pas comme des investisseurs sur un marché financier liquide.

Les comportements observés peuvent être rapprochés, pour leur aspect irrationnel, de la *prospect theory* ou théorie des perspectives (1979) de Kahneman⁷ et Tversky qui repose sur l'observation de comportements asymétriques des individus selon qu'ils viennent de subir un gain ou une perte. En effet, après des gains, les personnes ont tendance à consolider leurs gains en vendant leurs actifs financiers par exemple de manière à matérialiser le gain et se prémunir contre une baisse. Alors qu'après une perte, surtout si elle est conséquente, la tendance est à surinvestir dans des placements risqués, puisqu'il n'y a « plus rien à perdre » de manière à conserver une probabilité de « se refaire ».

Ainsi les méthodes d'inspiration financière, prônées par les nouveaux référentiels, se heurtent parfois à la réalité des observations. En effet, ces méthodes ont été élaborées dans des contextes où les instruments que l'on cherche à valoriser sont des dérivés de titres cotés dont l'évolution du cours est indépendante du comportement des parties prenantes. En assurance, la situation est différente. Ces méthodes de valorisation nécessitent donc la modélisation du comportement des assurés et de l'assureur. Or ceux-ci ont des comportements parfois différents de ceux d'investisseurs sur un marché financier.

Organisation de la thèse

L'objectif de cette thèse est de présenter les principes communs sur lesquels reposent les trois nouveaux référentiels, d'illustrer la limite de leur application en assurance et de proposer des modèles pour leur mise en œuvre effective.

Pour cela, la première partie est consacrée à la présentation des nouveaux référentiels prudentiel, comptable et de communication financière et insiste plus particulièrement sur la manière dont les risques sont valorisés et l'incidence de ces principes d'évaluation en termes de gestion d'une compagnie d'assurance.

La seconde partie aborde ces nouvelles normes sous un angle plus opérationnel en identifiant un certain nombre de problèmes pratiques auxquels leur mise en œuvre confronte l'actuaire et en proposant des modèles permettant de surmonter ces difficultés.

⁷. Kahneman reçut le prix Nobel d'économie en 2002 pour ses travaux en finance comportementale.

PARTIE I

NOUVELLES APPROCHES COMPTABLE, PRUDENTIELLE ET FINANCIÈRE DES RISQUES EN ASSURANCE

La première partie de cette thèse est dédiée aux risques portés par les sociétés d'assurance, leurs caractéristiques et leur traitement. En effet, l'activité d'assurance est née du besoin de se prémunir contre le risque (les agents économiques sont généralement averses aux risques qui peuvent réduire leur patrimoine), ce que permet l'opération d'assurance en transférant les risques de l'assuré vers l'assureur qui, en vertu de la loi des grands nombres, bénéficie de les effets de la mutualisation et est donc relativement moins exposé au risque que l'assuré.

Les évolutions récentes ou à venir amènent les assureurs à reconsidérer, au moins pour partie, leur vision des risques qu'ils assurent. Ainsi, qu'il s'agisse des nouvelles dispositions réglementaires (Solvabilité 2), de communication financière (EEV/MCEV) ou comptables (IFRS), l'objectif est similaire : identifier les risques et les analyser le plus finement possible.

Le passage d'un système où les hypothèses sont exogènes et prudentes, car contraintes par la réglementation, à un système où les hypothèses les plus réalistes doivent être privilégiées conduit à prendre en considération de « nouveaux risques ». Ces risques ne sont généralement pas à proprement parler « nouveaux » : la plupart du temps, ils existaient déjà mais n'avaient pas été soit étudiés plus avant du fait de leur caractère secondaire par rapport aux risques principaux, soit identifiés. Par exemple, dans le cas du risque de mortalité, un assureur qui veut étudier ce risque va, dans un premier temps, considérer son portefeuille et l'historique des données correspondant de manière à établir des statistiques descriptives de suivi du risque. Sur des portefeuilles d'assureurs, compte-tenu de la taille des échantillons, de telles études mettront en évidence le phénomène de fluctuation d'échantillonnage autour de la tendance centrale qui est le risque principal, mais certainement pas les risques systématiques de mortalité (mortalité stochastique et risque de longévité) qui s'avèrent relativement plus petits (cf. Planchet et Thérond (2007a) pour une étude du risque de mortalité sur un portefeuille de rentiers). Ces deux risques ne pourront être identifiés que par des études plus poussées, en étudiant par exemple, en parallèle les statistiques nationales de l'évolution au cours du temps de la mortalité.

Cette partie s'organise en trois chapitres.

Le premier chapitre s'intéresse aux aspects théoriques du traitement du risque. Après une première partie consacrée à l'analyse mathématique des risques (leur mesure et leur comparaison), nous verrons, dans un deuxième temps, que différents modèles de valorisation co-existent en assurance et que le recours aux modèles économiques issus de la théorie financière est de plus en plus fréquent. Il convient néanmoins de remarquer qu'associer une valeur à un risque et le gérer de manière effective relèvent de deux démarches distinctes. Ce point est illustré dans le cas d'une garantie plancher en cas de décès de l'assuré sur un contrat d'épargne en unités de compte.

Le deuxième chapitre s'attache à identifier les divergences entre les différentiels précédemment évoqués de manière à en tirer les conclusions adéquates en termes opérationnels. En effet, même s'ils reposent sur un socle de principes communs, la diversité des finalités des référentiels conduit à des options différentes dans la modélisation des produits d'assurance.

En particulier, un des principes fondamentaux commun aux trois approches est l'utilisation d'hypothèses *best estimate*, i.e. le recours aux hypothèses les plus réalistes compte-tenu de l'information dont dispose l'assureur. Ce point est fondamental car il diffère du contexte traditionnel de l'assurance qui repose sur des hypothèses prudentes. Par exemple, le taux d'actualisation d'un régime de rentiers ne doit pas, selon la réglementation française, être supérieur à 60 % du taux moyen des emprunts de l'État français (TME) quand bien même une société d'assurance investirait intégralement en OAT disposerait d'un rendement (certain) supérieur à ce taux d'actualisation. À titre illustratif, une attention particulière est portée sur l'évolution récente des tables de mortalité pour les risques viagers. Cet exemple montre que sur une période de temps relativement réduite, l'estimation de l'évolution de tel ou tel phénomène (l'espérance résiduelle de vie à 60 ans pour un assuré né en 1950 par exemple) peut être révisée en profondeur et avoir un impact important sur les niveaux de provisions techniques. De plus dans certains cas, un même phénomène sera modélisé sur des bases différentes selon que l'on cherche à valoriser un portefeuille de contrats ou à assurer sa solvabilité.

Par ailleurs, la valorisation des portefeuilles d'assurance nécessite fréquemment la modélisation du comportement de l'assureur et des assurés, particulièrement en assurance vie. Aussi les modèles implémentés ont de réels impacts sur les valorisations obtenues. Un exemple dans le cas de la gestion d'un portefeuille financier vient illustrer cela.

Enfin ce chapitre se conclut sur la modélisation et la valorisation d'un portefeuille d'assurance vie.

Les exigences quantitatives prévues dans le Pilier I de Solvabilité 2 prévoient notamment des exigences de fonds propres en référence au risque global supporté par l'assureur. Cette démarche impose des contraintes fortes en termes de gestion technique. Le troisième chapitre met ainsi en évidence les conséquences du changement de référentiel prudentiel sur la gestion des actifs de la société. Le projet Solvabilité 2 fixant les exigences quantitatives de fonds propres en fonction du risque global supporté par la compagnie, n'importe quel acte de gestion modifiant la structure ou la forme de ce risque a pour conséquence automatique et immédiate de modifier l'exigence minimale de

capitaux propres. Nous étudierons cela dans le cas du choix d'une allocation stratégique d'actifs et observons notamment la manière dont le processus de fixation de l'allocation évolue entre la réglementation prudentielle actuelle et Solvabilité 2.

Chapitre 1

Traitement spécifique du risque : aspects théoriques

L'activité d'assurance repose sur le concept de transfert de risque : moyennant une prime, l'assuré se protège d'un aléa financier. Mesurer le risque assuré s'avère donc inévitable puisque cette information est nécessaire dans le cadre de la tarification pour déterminer les chargements de sécurité à ajouter à la prime pure et dans une approche de solvabilité pour déterminer le niveau des réserves et des fonds propres dont doit disposer l'assureur pour être solvable.

En effet, bien que bénéficiant de l'effet de mutualisation, l'assureur ne peut se contenter de demander la prime pure des risques qu'il assure. Ce pour une raison évidente : la mutualisation ne saurait être parfaite et dès lors ne demander que la prime pure reviendrait à ce que, en moyenne, la société d'assurance soit en perte près d'un exercice sur deux¹.

Le niveau de fonds propres vient ensuite comme un matelas de sécurité destiné à amortir une sinistralité excessive mais aussi des placements risqués.

Le but de ce chapitre est, dans un premier temps, de présenter les outils permettant de comparer les risques et d'apprécier leur dangerosité, puis d'analyser le traitement du risque qui est effectué selon que l'on suit une démarche financière de valorisation ou une démarche assurantielle de contrôle et de gestion du risque.

Ces concepts sont illustrés dans le cas d'une garantie plancher en cas de décès adossée à un contrat d'assurance vie en unités de compte.

1. Les outils mathématiques de l'analyse des risques

Ce premier paragraphe a pour objectif de résumer en quelques pages les outils mathématiques usuels de mesure et de comparaisons des risques. Une attention particulière est portée aux caractéristiques particulières que l'on peut attendre des mesures de risque pour pouvoir être utilisées à des fins de solvabilité et de mesure de capitaux économiques.

1.1. Les mesures de risque

Après avoir défini ce qu'est une mesure de risque, nous rappelons les principales propriétés qu'elles doivent respecter pour être jugées satisfaisantes puis faisons un rapide tour d'horizon des mesures de risque les plus utilisées.

¹. D'après le théorème de la limite centrale le débours moyen de l'assureur converge avec le nombre de polices vers une variable gaussienne et donc symétrique.

1.1.1. Définition et propriétés

Nous reprenons ici la définition d'une mesure de risque telle qu'elle est formalisée dans Denuit et Charpentier (2004).

Définition 1. Mesure de risque

On appelle mesure de risque toute application ρ associant un risque X à un réel $\rho(X) \in \mathbb{R}_+ \cup \{+\infty\}$.

En particulier, cette définition nous permet d'établir que lorsqu'ils existent, l'espérance, la variance ou l'écart-type sont des mesures de risque.

Si un grand nombre d'applications répondent à la définition de mesure de risque, pour être jugée « satisfaisante » il est souvent exigé d'une mesure de risque d'avoir certaines propriétés dont les plus fréquentes sont rappelées *infra*.

1.1.1.1 Chargement de sécurité

La notion de chargement de sécurité est étroitement liée à celle de tarification : un principe de prime contient un chargement de sécurité s'il conduit à exiger une prime supérieure à celle qui est exigée si la mutualisation des risques est parfaite (cf. Partrat et Besson (2005)).

Définition 2. Chargement de sécurité

Une mesure de risque ρ contient un chargement de sécurité si pour tout risque X , on a $\rho(X) \geq E[X]$.

Nous verrons dans la suite qu'une Tail-Value-at-Risk (TVaR), lorsqu'elle existe, contient un chargement de sécurité ce qui n'est pas le cas d'une Value-at-Risk (VaR).

1.1.1.2 Mesure de risque cohérente

La définition d'une mesure de risque est très générale puisque toute fonctionnelle réelle positive d'une variable aléatoire peut être considérée comme étant une mesure de risque. Aussi, en pratique, on exige de telles mesures qu'elles disposent de propriétés mathématiques dont la transcription conceptuelle permette de les jauger. En pratique, on exige fréquemment qu'une mesure de risque ρ possède une partie des caractéristiques suivantes :

- Invariance par translation : $\rho(X + c) = \rho(X) + c$ pour toute constante c .
- Sous-additivité : $\rho(X + Y) \leq \rho(X) + \rho(Y)$ quels que soient les risques X et Y .
- Homogénéité : $\rho(cX) = c\rho(X)$ pour toute constante positive c .
- Monotonie : $\Pr[X < Y] = 1 \Rightarrow \rho(X) \leq \rho(Y)$ quels que soient les risques X et Y .

Ces caractéristiques trouvent une interprétation naturelle dans la situation où la mesure de risque doit permettre de définir un capital de solvabilité d'une société d'assurance.

Ainsi la sous-additivité représente l'effet de la diversification : une société qui couvre deux risques ne nécessite pas davantage de capitaux que la somme de ceux obtenus pour deux entités distinctes se partageant ces deux risques.

La monotonie traduit quant à elle le fait que si le montant résultat d'un risque est systématiquement (au sens presque sûr) inférieur à celui résultant d'un autre risque, le capital nécessaire à couvrir le premier risque ne saurait être supérieur à celui nécessaire pour couvrir le second.

L'association de ces quatre axiomes a donné naissance au concept de cohérence d'une mesure de risque dans Artzner et al. (1999).

Définition 3. Mesure de risque cohérente

Une mesure de risque invariante par translation, sous-additive, homogène et monotone est dite cohérente.

Cette notion de cohérence n'est toutefois pas ce que l'on attend a minima d'une mesure de risque. Ainsi certaines mesures de risque parmi les plus exploitées actuellement ne le sont pas. C'est notamment le cas de la Value-at-Risk (VaR) ou encore de la variance.

1.1.1.3 Mesure de risque comonotone additive

Rappelons qu'un vecteur aléatoire $(X_1; X_2)$, de fonctions de répartition marginales F_1, F_2 , est un vecteur comonotone s'il existe une variable aléatoire U de loi uniforme sur $[0;1]$ telle que $(X_1; X_2)$ a la même loi que $(F_1^{-1}(U); F_2^{-1}(U))$.

Définition 4. Mesure de risque comonotone

On appelle mesure de risque comonotone additive toute mesure de risque ρ telle que : $\rho(X_1 + X_2) = \rho(X_1) + \rho(X_2)$ pour tout vecteur comonotone $(X_1; X_2)$.

Une mesure de risque comonotone additive intègre donc le fait que deux risques comonotones ne se mutualisent pas. Pour une mesure utilisée pour déterminer un capital de solvabilité, cette propriété est souhaitable puisque, dès que le risque U se produit, les risques X_1 et X_2 se produisent également avec une ampleur croissante avec celle de U .

1.1.2. Mesures de risque usuelles

L'objet de ce paragraphe est de présenter les mesures de risque les plus usuelles. On s'attardera particulièrement sur la Value-at-Risk et la Tail-Value-at-Risk dont l'utilisation en assurance va être pérennisée par le futur système de solvabilité européen (Solvabilité 2) puisqu'elles seront à la base de la détermination du niveau prudentiel des provisions techniques et du besoin en fonds propres : le « capital cible » ou *Solvency Capital Requirement* (SCR).

1.1.2.1 L'écart-type et la variance

Ce sont les premières mesures de risque à avoir été utilisées ; on les retrouve notamment dans le critère de Markowitz (moyenne-variance) qui sert

de socle aux premières théories d'évaluation des actifs (MEDAF). Toutefois ce critère n'est pas bien adapté à l'activité d'assurance, notamment parce qu'il est symétrique et pénalise autant les « bonnes variations » que les « mauvaises ».

1.1.2.2 La Value-at-Risk (VaR)

La notion de *Value-at-Risk* ou valeur ajustée au risque s'est originellement développée dans les milieux financiers avant d'être largement reprise dans les problématiques assurantielles. Elle est notamment la mesure de risque sur laquelle repose le nouveau référentiel prudentiel Solvabilité 2. De plus, elle est évoquée dans les travaux de l'IAS Board comme une des mesure envisageables pour le calibrage des marges pour risque des provisions techniques.

Définition 5. Value-at-Risk (VaR)

La Value-at-Risk (VaR) de niveau α associée au risque X est donnée par :

$$VaR(X, \alpha) = \inf \{x \mid \Pr[X \leq x] \geq \alpha\}.$$

On notera que $VaR(X, \alpha) = F_X^{-1}(\alpha)$ où F_X^{-1} désigne la fonction quantile de la loi de X . Rappelons que, dans le cas général, la fonction quantile est la pseudo-inverse de la fonction de répartition, soit

$$F_X^{-1}(p) = \inf \{x \mid F(x) \geq p\}.$$

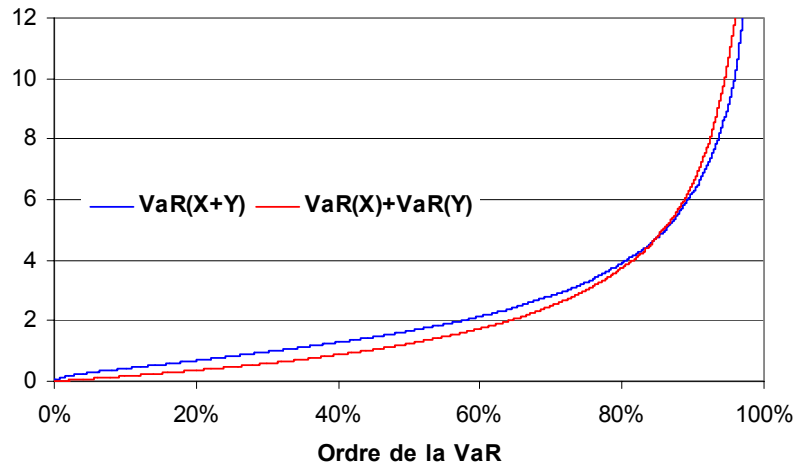
Cette mesure de risque a le mérite de reposer sur un concept simple et facilement explicable : $VaR(X, \alpha)$ est le montant qui permettra de couvrir le montant de sinistres engendré par le risque X avec une probabilité α . Ce concept est directement lié à celui de probabilité de ruine puisque si une société, disposant d'un montant de « ressources » égal à $VaR(X, \alpha)$, assure un unique risque X , sa probabilité de ruine est égale à $1 - \alpha$.

Comme évoqué précédemment, la VaR n'est pas cohérente car elle n'est pas sous-additive.

Ce résultat peut se démontrer à l'aide d'un contre-exemple. Soit X et Y deux variables aléatoires indépendantes de lois de Pareto de paramètres $(2 ; 1)$ et $(2 ; 2)$, alors

$$\exists \alpha \in]0;1[, VaR_\alpha(X + Y) > VaR_\alpha(X) + VaR_\alpha(X),$$

comme l'illustre la Figure 1.

Figure 1 - Value-at-Risk de la somme de deux v.a. de Pareto

Rappelons qu'une variable aléatoire X de loi de Pareto $Par(\alpha; \theta)$ a pour fonction de répartition $F_X(x) = 1 - \left(\frac{\theta}{\theta + x}\right)^\alpha$ si $x > 0$, et $F_X(x) = 0$ sinon. Sa fonction quantile se calcule facilement et vaut :

$$F_X^{-1}(p) = \theta \left[(1-p)^{-1/\alpha} - 1 \right].$$

Les Value-at-Risk ont un certain nombre de « bonnes » propriétés mathématiques parmi lesquelles le fait que pour toute fonction g croissante et continue à gauche, on a : $VaR_\alpha(g(X)) = g(VaR_\alpha(X))$.

Il découle de cette propriété en prenant $g = F_1^{-1} + F_2^{-1}$ et $X = U$, que les VaR sont comonotones additives puisque pour tout $\alpha \in]0; 1[$, on a $VaR_\alpha((F_1^{-1} + F_2^{-1})(U)) = (F_1^{-1} + F_2^{-1})(VaR_\alpha(U))$.

1.1.2.3 La Tail-Value-at-Risk (TVaR)

Initialement présentée dans les travaux de la Commission Européenne comme une des deux alternatives possibles (avec la Value-at-Risk) comme critère de fixation de l'exigence de capitaux propres dans Solvabilité 2, la Tail Value-at-Risk a rencontré beaucoup de partisans parmi les techniciens arguant du fait que l'information sur la probabilité de ruine fournie par la VaR n'est pas suffisante mais qu'il faut aussi en connaître son ampleur. Néanmoins elle n'a finalement pas eu les faveurs du CEIOPS (cf. l'annexe dédiée à QIS 3) notamment du fait de sa difficulté de mise en œuvre : espérance sur les valeurs extrêmes d'une distribution que l'on n'observe pas directement la plupart du temps (cf. le Chapitre 4).

Définition 6. *Tail Value-at-Risk (TVaR)*

La Tail Value-at-Risk de niveau α associée au risque X est donnée par :

$$TVaR(X, \alpha) = \frac{1}{1-\alpha} \int_{\alpha}^1 F_X^{-1}(p) dp .$$

On remarque que la TVaR peut s'exprimer en fonction de la VaR :

$$TVaR(X, \alpha) = VaR(X, \alpha) + \frac{1}{1-\alpha} E \left[(X - VaR(X, \alpha))^+ \right] .$$

Il vient de cette réécriture que pour tout $\alpha \in]0; 1[$, $TVaR_\alpha(X) < +\infty \Leftrightarrow E[X] < +\infty$. Par ailleurs, le deuxième terme du membre de droite représente la perte moyenne au-delà de la VaR, la TVaR est donc très sensible à la forme de la queue de distribution.

Propriétés : Une Tail Value-at-Risk :

- est cohérente ;
- inclut un chargement de sécurité ;
- est comonotone additive.

Les deux derniers points résultent :

- pour le chargement de sécurité, du fait que pour tout $\alpha \geq 0$, $TVaR(X; \alpha) \geq TVaR(X; 0) = E[X]$;
- pour la propriété d'additivité pour des risques comonotones, du fait que la TVaR est une somme de VaR qui sont elles-mêmes comonotones additives.

Définition 7. *Expected Shortfall (ES)*

L'expected shortfall de niveau de probabilité α est la perte moyenne au-delà de la VaR au niveau α , i. e.

$$ES_\alpha(X) = E \left[(X - VaR(X, \alpha))^+ \right] .$$

On peut remarquer que si X représente la charge brute de sinistres, $ES_\alpha(X)$ est le montant de la prime Stop-Loss dont la rétention pour l'assureur est la VaR au niveau α .

Définition 8. *Conditionnal Tail Expectation (CTE)*

La Conditionnal Tail Expectation de niveau α est le montant de la perte moyenne sachant que celle-ci dépasse la VaR au niveau α , i. e.

$$CTE(X, \alpha) = E \left[X \mid X > VaR(X, \alpha) \right] .$$

Cette définition est très proche de celle de la Tail-Value-at-Risk et en particulier, ces deux mesures coïncident lorsque la fonction de répartition F_X du risque X est continue.

Remarquons que si la TVaR est comonotone additive, ce n'est, en général, pas le cas de la CTE (cf. Dhaene, Vanduffel et al. (2004) pour un contre-exemple).

1.1.2.4 Mesures de risque de Wang

Les mesures de risque de Wang (2002) utilisent l'opérateur espérance sur des transformations de la distribution de la variable aléatoire d'intérêt. L'idée est en effet d'alourdir la queue de la distribution de la variable d'intérêt afin d'engendrer un chargement par rapport à la prime pure. Cette transformation de la fonction de répartition sera effectuée à l'aide d'une fonction de distorsion.

Rappelons qu'une fonction de distorsion est une fonction non décroissante $g : [0;1] \rightarrow [0;1]$ telle que $g(0) = 0$ et $g(1) = 1$.

Définition 9. Mesure de risque de Wang

On appelle mesure de risque de Wang issue de la fonction de distorsion g ,

$$\text{la mesure } \rho_g \text{ définie par : } \rho_g(X) = \int_0^\infty g(\Pr[X > x]) dx.$$

On remarque que toute mesure de Wang peut s'écrire comme somme de VaR, i. e.

$$\rho_g(X) = \int_0^1 \text{VaR}(X, 1-\alpha) dg(\alpha).$$

En effet, si l'on note $\bar{F}_X(x) = 1 - \Pr[X \leq x]$, on a $g(\bar{F}_X(x)) = \int_0^{\bar{F}_X(x)} dg(\alpha)$

puisque $g(0) = 0$ donc $\rho_g(X) = \int_0^\infty \int_0^1 \mathbf{1}_{\{\alpha \leq \bar{F}_X(x)\}} dg(\alpha) dx$. On déduit du théorème de Fubini que :

$$\rho_g(X) = \int_0^\infty \int_0^\infty \mathbf{1}_{\{\alpha \leq \bar{F}_X(x)\}} dx dg(\alpha) = \int_0^1 F_X^{-1}(1-\alpha) dg(\alpha).$$

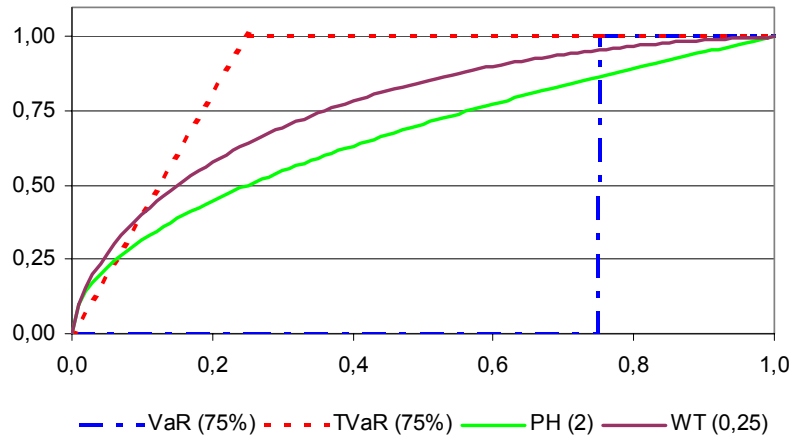
Définition 10. Wang-Transform (WT)

On appelle Wang-Transform (WT) la mesure de risque de Wang issue de la fonction de distorsion $g_\alpha(x) = \Phi[\Phi^{-1}(x) - \Phi^{-1}(\alpha)]$.

Certaines mesures de risque usuelles telles que la VaR ou la TVaR sont des mesures de risque de Wang, le tableau ci-dessous reprend les fonctions de distorsion correspondantes.

Tableau 1 - Fonctions de distorsion associées à quelques mesures de risque

Mesure de risque	Paramètre	Fonction de distorsion
Value-at-Risk	VaR_α	$g(x) = \mathbf{1}_{[\alpha; +\infty[}(x)$
Tail-Value-at-Risk	TVaR_α	$g(x) = \min(x/(1-\alpha); 1)$
Mesure de risque PH	PH_ξ	$g(x) = x^{1/\xi}$
Wang-transform	WT_α	$g(x) = \Phi[\Phi^{-1}(x) - \Phi^{-1}(\alpha)]$

Figure 2 - Quelques fonctions de distorsion

Les mesures de risque de Wang sont homogènes, invariantes par translation et monotones. Une mesure de Wang n'est en revanche sous-additive que si, et seulement si, la fonction de distorsion dont elle est issue est concave.

Ainsi une mesure de Wang est cohérente si, et seulement si, la fonction de distorsion correspondante est concave. En effet, si g est une fonction de distorsion concave, $g \circ \bar{F}_X$ est continue à droite et est donc la fonction de queue d'une variable aléatoire. Ainsi $\rho_g(X)$ est l'espérance mathématique de la variable aléatoire de fonction de queue $g \circ \bar{F}_X$.

Enfin comme une mesure de Wang peut s'écrire sous la forme d'une somme de VaR qui sont comonotones additives, elle est elle-même additive pour des risques comonotones.

La démarche des mesures de distorsion est à rapprocher de celle de la valorisation de dérivés financiers (cf. le paragraphe 2.2) : on détermine une espérance sur une déformation de la distribution de la variable d'intérêt. Dans l'approche financière, la transformation se déduit des prix observés sur le

marché : le choix de l'un opérateur de distorsion ne revient pas à l'actuaire, il lui est imposé par le marché.

1.1.2.5 Mesure de risque d'Esscher

La mesure de risque d'Esscher consiste à mesurer le risque comme étant la prime pure, *i. e.* l'espérance de la transformée d'Esscher du risque initial.

Définition 11. *Mesure de risque d'Esscher*

On appelle mesure d'Esscher de paramètre $h > 0$ du risque X , la mesure de risque donnée par :

$$Es(X; h) = \frac{E[X e^{hX}]}{E[e^{hX}]} = \frac{d}{dh} \ln E[e^{hX}].$$

La mesure de risque d'Esscher n'est pas cohérente car elle n'est ni homogène, ni monotone. En revanche, elle contient un chargement de sécurité puisque $Es(X; h)$ est une fonction croissante en h et $Es(X; 0) = E[X]$.

1.1.3. Choix d'une mesure de risque pour déterminer un capital économique

Les travaux en cours sur Solvabilité 2 évoquent l'utilisation de mesures de risque dans la détermination du capital cible nécessaire à une société d'assurance pour pérenniser son activité. Ce capital cible sera déterminé en référence à une mesure de risque appliquée au risque global supporté par la société. Ce risque global sera modélisé à partir d'une formule commune à toutes les compagnies d'assurance européenne ou à partir d'un modèle interne. Les mesures de risque les plus souvent citées sont la Value-at-Risk et la Tail-Value-at-Risk. Un des principaux arguments en faveur de la TVaR est notamment le fait qu'elle soit sous-additive, ce qui est assez intuitif s'agissant d'une mesure de risque destinée à calculer un capital de solvabilité.

Néanmoins les travaux récents de Dhaene, Laeven et al. (2004) montrent que la propriété de sous-additivité est parfois trop forte. En effet considérons deux risques X_1 et X_2 , et ρ la mesure de risque associée à la détermination du capital réglementaire. Si ρ est « trop » sous-additive, on peut se retrouver dans une situation où

$$E[(X_1 + X_2 - \rho(X_1 + X_2))^+] > E[(X_1 - \rho(X_1))^+] + E[(X_2 - \rho(X_2))^+].$$

Cette inégalité signifie que l'ampleur de la ruine moyenne d'une société pratiquant les risques X_1 et X_2 et disposant d'un capital de niveau $\rho(X_1 + X_2)$ est plus importante que la somme de l'ampleur de la ruine moyenne de deux sociétés de capitaux respectifs $\rho(X_1)$ et $\rho(X_2)$ couvrant respectivement les risques X_1 et X_2 . Dhaene, Laeven et al. (2004) illustrent cette inégalité lorsque la mesure de risque considérée est la TVaR, en prenant deux risques indépendants de loi de Bernoulli de même paramètre.

Dans le cadre du choix de la mesure de risque permettant de déterminer le capital cible (au sens de Solvabilité 2), ils proposent donc de ne retenir que les

mesures de risque qui respectent la *condition du régulateur*, à savoir les mesures de risques ρ telles que pour ε fixé dans $]0;1[$ et pour tous risques X_1 et X_2 , l'inégalité suivante soit vérifiée :

$$\mathbb{E}\left[(X_1 + X_2 - \rho(X_1 + X_2))^+\right] + \rho(X_1 + X_2)\varepsilon \leq \sum_{i=1}^2 \left\{ \mathbb{E}\left[(X_i - \rho(X_i))^+\right] + \rho(X_i)\varepsilon \right\}.$$

Notons que ε peut s'interpréter comme le coût de l'immobilisation du capital, puisque, le deuxième terme du premier membre de l'inégalité précédente peut s'interpréter comme le flux à destination de l'actionnaire de manière à le rémunérer du capital $\rho(X_1 + X_2)$ qu'il « prête » à la société. A partir de cette condition, ils démontrent les trois propriétés suivantes.

Soit $\varepsilon \in]0;1[$. La condition du régulateur en référence au niveau ε est satisfaite :

- pour les $TVaR_p$ telles que $p > 1 - \varepsilon$,
- pour la $VaR_{1-\varepsilon}$,
- pour toute mesure de risque sous-additive ρ telle que $\rho(X) \geq VaR(X, 1 - \varepsilon)$.

En conclusion, ils démontrent enfin que la VaR de niveau $1 - \varepsilon$ est la mesure de risque respectant la *condition du régulateur* qui conduit au plus petit niveau de capital.

1.1.4. Value-at-Risk

La Value at Risk (VaR) tend à devenir un indicateur de risque largement utilisé tant par les établissements financiers que par les compagnies d'assurance car elle permet d'appréhender le risque global dans une unité de mesure commune à tous les risques encourus, quelle que soit leur nature.

L'objectif de cette section est de fournir une présentation plus intuitive de la notion de VaR que la formalisation mathématique directe présentée *supra*. Pour simplifier la présentation, on se place dans le contexte d'un portefeuille financier.

Pour un horizon de gestion donné, la VaR correspond au montant de perte probable du portefeuille. Elle exprime la perte liée à des variations défavorables. Si l'on note α le seuil de confiance choisi, la VaR vérifie donc l'équation suivante :

$$\Pr[\text{perte} > VaR] = 1 - \alpha.$$

Afin de calculer la VaR, il est essentiel de spécifier la période sur laquelle la variation de valeur du portefeuille est mesurée et le seuil de confiance $1 - \alpha$.

Il existe en pratique trois méthodes de calcul de la VaR. Notons dès à présent que les deux premières méthodes utilisent les données du passé pour estimer les variations potentielles de la valeur du portefeuille. Cela suppose implicitement que le futur se comporte comme le passé : il faut donc faire l'hypothèse que la série temporelle des valeurs du portefeuille est stationnaire.

1.1.4.1 Les différentes approches

La VaR analytique

Dans ce modèle, la valeur algébrique d'un portefeuille est représentée par une combinaison linéaire de K facteurs gaussiens. Notons $P(t)$ la valeur du portefeuille en t , $F(t)$ le vecteur gaussien des facteurs et a le vecteur des sensibilités aux facteurs, de dimension K .

Supposons que $F(t) \sim \mathcal{N}(m, V)$. A la date t , la valeur du portefeuille est $P(t) = a'F(t)$.

En t , $P(t+1)$ est une variable aléatoire gaussienne de loi $\mathcal{N}(a'm, a'Va)$. La valeur de la VaR pour un seuil de confiance α correspond alors à :

$$\Pr(P(t+1) - P(t) \geq -VaR) = \alpha.$$

La différence $P(t+1) - P(t)$ représente la variation de valeur du portefeuille entre l'instant $t+1$ et l'instant t . On est donc en présence d'une perte si la valeur réalisée du portefeuille dans une période est inférieure à sa valeur d'aujourd'hui.

Comme $P(t+1)$ est une variable aléatoire gaussienne de moyenne $a'm$ et de variance $a'Va$, nous avons alors l'équation suivante :

$$\Pr\left(\frac{P(t+1) - a'm}{\sqrt{a'Va}} \geq \Phi^{-1}(1 - \alpha)\right) = \alpha,$$

où Φ^{-1} est l'inverse de la fonction de répartition d'une gaussienne centrée et réduite.

Comme $\Phi^{-1}(1 - \alpha) = -\Phi^{-1}(\alpha)$, nous obtenons alors que

$$VaR = P(t) - a'm + \sqrt{a'Va} \Phi^{-1}(\alpha).$$

Lorsque les facteurs F modélisent directement la variation du portefeuille, et comme nous supposons en général que $m = 0$, nous obtenons alors $VaR = \sqrt{a'Va} \Phi^{-1}(\alpha)$.

La VaR apparaît alors proportionnelle à l'écart type de la variation de valeur du portefeuille.

Cette méthode analytique repose sur trois hypothèses qui permettent de simplifier les calculs : l'indépendance temporelle des variations de la valeur du portefeuille, la normalité des facteurs de risque $F(t)$ et la relation linéaire entre les facteurs de risque et la valeur du portefeuille.

La principale difficulté de cette méthode est d'identifier les facteurs et d'estimer la matrice de covariance. Cette méthode analytique ne prend en compte que des relations linéaires entre les facteurs de risque et la valeur du portefeuille.

La VaR historique

Contrairement à la VaR analytique, la VaR historique est entièrement basée sur les variations historiques des facteurs de risque. Supposons que nous disposions d'un historique de taille N . En t_0 nous pouvons valoriser le portefeuille avec les facteurs de risque de l'historique en calculant pour chaque date $t = (t_0 - 1, \dots, t_0 - N)$ une valeur potentielle du portefeuille. On détermine ainsi N variations potentielles. Ainsi, à partir de l'historique, nous construisons ainsi une distribution empirique de laquelle il est possible d'extraire le quantile à α %. Pour cela, il faut ranger les N pertes potentielles par ordre croissant et prendre la valeur absolue de la $(1-\alpha)N$ -ème plus petite valeur. Lorsque $(1-\alpha)N$ n'est pas un nombre entier, on calcule la VaR par interpolation linéaire.

En général, la VaR historique n'impose pas d'hypothèses sur la loi des facteurs de risque à la différence de la méthode analytique. Mais il est tout de même nécessaire d'avoir un modèle sous-jacent pour estimer les facteurs de risque pour l'historique de longueur N . Cette méthode est très utilisée dans la pratique car elle est simple conceptuellement et est facile à implémenter. Cependant, elle présente quelques difficultés : en effet, l'estimation d'un quantile demande beaucoup d'observations, condition rarement réalisée en pratique en assurance. C'est particulièrement le cas lorsqu'il s'agit d'estimer des valeurs dans la queue d'une distribution. Le Chapitre 4 s'intéresse à ces problématiques, dans lesquelles, le recours à la théorie des extrêmes permet de disposer d'estimateurs performants comparés aux estimateurs empiriques.

La VaR aléatoire (ou VaR Monte-Carlo)

La VaR Monte Carlo est basée sur la simulation des facteurs de risque dont on se donne la distribution. Cette méthode consiste à valoriser le portefeuille en appliquant ces facteurs simulés. Il suffit alors de calculer le quantile correspondant tout comme pour la méthode de la VaR historique. La seule différence entre ces deux méthodes est que la VaR historique utilise les facteurs passés, alors que la VaR Monte Carlo utilise les facteurs simulés.

Il faut cependant noter que cette demande beaucoup de temps de calcul. De plus, elle demande un effort important de modélisation puisqu'elle détermine entièrement les trajectoires des facteurs de marché utilisés pour le calcul de la VaR.

1.1.4.2 Les difficultés d'adaptation de la VaR à l'assurance

La VaR est à l'origine un modèle bancaire, et son adaptation à l'assurance pose certaines difficultés. En effet, il existe, dans les contrats d'assurance, des options cachées qui font apparaître des problèmes de fluctuations des encours et donc de la valeur sur laquelle on calcule la VaR. De plus, les durées de vie des contrats d'assurance sont beaucoup plus longues que celles des contrats bancaires, ce qui amène à reconsidérer la question de l'horizon de calcul de la VaR.

La valeur sur laquelle on calcule la VaR

Alors qu'un portefeuille bancaire possède une valeur de marché facile à calculer, il est simple de calculer la valeur de l'actif d'un produit d'assurance (on peut prendre la valeur boursière des actifs détenus) mais il est souvent plus difficile de calculer la valeur de son passif. En effet, il n'existe pas de marché où cette valeur soit échangée (malgré l'apparition de certains produits dérivés et les tentatives de titrisation de certains risques). L'une des difficultés dans la valorisation du passif réside en particulier dans la valorisation des options cachées.

En effet, une caractéristique importante du risque en assurance-vie est la présence d'options incluses dans les contrats : option de rachat, de versements (pour les anciens contrats dont le taux est garanti à vie), et risque d'arbitrage pour les contrats multi-supports :

- *L'option de rachat* : le client qui a souscrit un contrat d'assurance vie a le droit de le résilier avant terme. Dans ce cas, si son retrait se fait après le seuil fiscal, il n'encourt en général aucune pénalité et il récupère le montant investi additionné des intérêts capitalisés. Comme c'est seulement après un certain nombre d'années courues que le contrat devient rentable pour l'assureur, les rachats anticipés sont très coûteux pour ce dernier, et ce d'autant plus si le rachat intervient dans un contexte de hausse des taux, obligeant l'assureur à vendre des actifs en situation de moins-value.
- *Le risque d'arbitrage pour les contrats multi-supports* : le client peut transférer une partie de sa provision mathématique vers des supports en unité de compte. Cela est préjudiciable à l'assureur car ce dernier ne sait pas quand ces transferts vont avoir lieu et donc quand ses engagements vont changer.
- *L'option de versement* : le détenteur d'un contrat d'assurance vie est autorisé à verser de l'argent quand il le désire sur son contrat. Si le contrat est ancien, il se peut qu'il garantisse un taux anormalement élevé par rapport au taux du marché. Dans ce cas, il est coûteux pour l'assureur de tenir ses engagements.

On pourrait résumer le problème en disant que la banque calcule un risque sur un montant fixe d'argent placé (système fermé) tandis que dans l'assurance, le montant placé n'est pas figé : il faut tenir compte des entrées et des sorties (système ouvert). Autrement dit, dans la banque il n'y a pas de problèmes de fluctuations des encours alors qu'en assurance, la bonne gestion de cet aspect est primordiale.

L'horizon de calcul de la VaR

Le risque en assurance n'est pas à court terme mais plutôt à un horizon de quelques années. En effet, en assurance-vie, le cycle de production est long. La durée avant rachat d'un contrat est en général au minimum égale à 8 ans puisque c'est seulement au-delà de ce délai que l'assuré peut reprendre ses intérêts sans être soumis à la fiscalité. De plus, en assurance-vie, les obligations mises en représentation des contrats du passif ont en général une échéance

lointaine. Si les taux baissent, les assureurs peuvent continuer à servir des taux qui ne sont plus accessibles sur le marché obligataire, parce qu'ils possèdent dans leur portefeuille des obligations anciennes de taux supérieurs aux taux du marché. Cependant, ils doivent renouveler leur portefeuille car les obligations arrivent à maturité, et pour cela ils ont accès à des taux moins élevés. Par conséquent, lorsqu'ils n'auront plus d'obligations anciennes dans leurs portefeuilles, ils proposeront des produits d'épargne moins attractifs. Il y a donc une inertie des taux servis par les assureurs qui fait que si le marché varie trop vite, ils ne pourront pas s'y adapter et les clients rachèteront leurs contrats. Mais le risque n'est pas à court terme, au contraire de la banque, où un portefeuille d'actifs est immédiatement (quotidiennement) sensible aux variations des cours.

1.2. Comparaison des risques

L'objet de ce paragraphe est de présenter des outils permettant de classer les risques selon leur « dangerosité ».

1.2.1. Relation associée à une mesure de risque

Dans le paragraphe 1, nous avons étudié un certain nombre de mesures de risque. Une idée naturelle pour comparer deux risques X et Y est de choisir une mesure de risque ρ et de comparer $\rho(X)$ et $\rho(Y)$, ce que l'on peut toujours faire puisque $\mathbf{R}$ est ordonné par la relation d'ordre totale $\leq$. Cette démarche nous permet d'introduire la relation $\prec_\rho$ définie comme suit.

$$X \prec_\rho Y \text{ si } \rho(X) \leq \rho(Y)$$

La relation $\prec_\rho$ issue de la mesure de risque ρ est réflexive et transitive. De plus il est toujours possible de comparer par $\prec_\rho$ deux lois de probabilité ou deux variables aléatoires.

N.B. Cette relation n'est pas une relation d'ordre car elle n'est pas antisymétrique. En effet, avoir simultanément $\rho(X) \leq \rho(Y)$ et $\rho(X) \geq \rho(Y)$ n'implique pas que $X =_d Y$ (et *a fortiori* que $X = Y$).

Le principal mérite de ce type de relation est qu'il est toujours possible de comparer deux risques. Il faut néanmoins rester prudent car l'on peut avoir simultanément $X \prec_\rho Y$ et $X \succ_{\rho'}, Y$ pour deux mesures de risque ρ et ρ' différentes. C'est pour cette raison que l'on préférera se tourner vers des ordres partiels qui permettent de disposer de davantage de propriétés.

1.2.2. Ordre stochastique

Définition 12. *Dominance stochastique*

On dit que X domine selon l'ordre stochastique Y ($Y \prec_{st} X$) si pour toute fonction de distorsion g , on a : $\rho_g(Y) \leq \rho_g(X)$.

Cette notion est équivalente à celle de comparaison uniforme des VaR puisque toute mesure de risque de Wang peut s'écrire comme somme de VaR (cf. le paragraphe 1.1.2.4) :

$$\begin{aligned} X \prec_{st} Y &\Leftrightarrow \rho_g(X) \leq \rho_g(Y) \text{ pour toute fonction de distorsion } g \\ &\Leftrightarrow \text{Var}(X, \alpha) \leq \text{Var}(Y, \alpha) \forall \alpha \in [0; 1]. \end{aligned}$$

La relation $\prec_{st}$ est un ordre partiel sur l'ensemble des lois de probabilités.

L'ordre stochastique ne permet pas de comparer toutes les variables aléatoires. En effet, il est possible d'avoir simultanément :

$$\text{Var}(X, \alpha) \leq \text{Var}(Y, \alpha),$$

et

$$\text{Var}(X, \beta) > \text{Var}(Y, \beta).$$

Pour s'en convaincre, une simple expérience Bernoulli peut être envisagée. Considérons les deux variables aléatoires :

$$\begin{cases} P[X = 0] = 0,2 \\ P[X = 3] = 0,8 \end{cases} \text{ et } \begin{cases} P[Y = 1] = 0,5 \\ P[Y = 2] = 0,5 \end{cases}$$

$$\text{On a } \text{Var}(X; 0,1) \leq \text{Var}(Y; 0,1) \text{ et } \text{Var}(X; 0,4) > \text{Var}(Y; 0,4).$$

En revanche, on a l'équivalence suivante (pour autant que les espérances existent) :

$$X \prec_{st} Y \Leftrightarrow E[\varphi(X)] \leq E[\varphi(Y)] \text{ pour toute fonction } \varphi \text{ croissante.}$$

En particulier, on a l'implication :

$$X \prec_{st} Y \Rightarrow E[X] \leq E[Y],$$

qui signifie intuitivement que le risque X est plus « petit » que le risque Y . Cette conséquence explique le fait que l'on parle de dominance stochastique au premier ordre pour désigner la relation $\prec_{st}$.

1.2.3. Ordre convexe

Définition 13. *Ordre convexe croissant*

On dit que X est moins dangereux que Y sur la base de l'ordre convexe croissant ($\prec_{icx}$) et l'on note $X \prec_{icx} Y$ si, pour toute fonction de distorsion g concave, on a : $\rho_g(X) \leq \rho_g(Y)$.

Cette notion est équivalente à celle de comparaison uniforme des TVaR, puisque :

$$\begin{aligned} X \prec_{icx} Y &\Leftrightarrow \rho_g(X) \leq \rho_g(Y) \text{ pour toute fonction de distorsion } g \text{ concave} \\ &\Leftrightarrow \text{TVaR}(X, \alpha) \leq \text{TVaR}(Y, \alpha) \text{ pour tout } \alpha \in [0; 1]. \end{aligned}$$

On supposera dans la suite que les risques ont des primes pures finies ce qui garantit l'existence des TVaR. Cette relation est également connue sous les

noms de dominance stochastique du deuxième ordre, d'ordre Stop-Loss et d'ordre sur les TVaR.

Définition 14. *Ordre convexe*

On dira que le risque X est moins dangereux que le risque Y de même prime pure au sens de l'ordre convexe ($\prec_{cx}$), s'il est moins dangereux au sens de l'ordre convexe croissant, i. e. si $X \prec_{cx} Y \Leftrightarrow X \prec_{icx} Y$ et $E[X] = E[Y]$.

Pour toute fonction φ convexe croissante et pour autant que les variances existent, on a : $X \prec_{cx} Y \Leftrightarrow E[\varphi(X)] \leq E[\varphi(Y)]$.

En particulier pour $\varphi : x \mapsto x^2$ pour $x \geq 0$, on a : $X \prec_{cx} Y \Rightarrow \text{Var}[X] \leq \text{Var}[Y]$ pour autant que les variances existent. Intuitivement cette propriété signifie que si $X \prec_{cx} Y$, le risque X est moins « variable » que le risque Y .

L'ordre convexe permet de comparer des variables aléatoires de même espérance, ce que ne permettait pas l'ordre stochastique puisque

$$X \prec_{st} Y \text{ et } E[X] = E[Y] \Leftrightarrow X =_{loi} Y.$$

Proposition 1. *Théorème de séparation*

On a l'équivalence :

$$X \prec_{icx} Y \Leftrightarrow \exists Z \text{ tel que } X \prec_{st} Z \prec_{cx} Y.$$

Le théorème de séparation permet d'établir que si X est moins dangereux que Y selon l'ordre convexe croissant, X est à la fois plus « petit » ($\prec_{st}$) et moins « variable » ($\prec_{cx}$) que Y .

1.2.4. Bornes comonotones d'une somme de variables aléatoires

L'ordre convexe nous permet de disposer de bornes pour une somme de variables aléatoires.

Proposition 2. *Bornes comonotones d'une somme de variables aléatoires*

Pour tout vecteur $X = (X_1, \dots, X_n)$ et pour toute variable aléatoire Λ , on

a les inégalités : $\sum_{i=1}^n E[X_i | \Lambda] \prec_{cx} \sum_{i=1}^n X_i \prec_{cx} \sum_{i=1}^n F_{X_i}^{-1}(U)$, où U est une variable aléatoire de loi uniforme sur $[0;1]$.

Ce résultat est démontré dans Kaas et al. (2000).

Le majorant de cette double inégalité est appelé la contrepartie comonotone du vecteur X . En effet ces deux vecteurs ont les mêmes marginales mais le vecteur $X^c = (F_{X_1}^{-1}, \dots, F_{X_n}^{-1})$ est comonotone.

Grâce aux propriétés de l'ordre convexe, cette majoration nous permet de disposer d'un maximum pour les primes Stop-Loss de la somme de risques $\sum_{i=1}^n X_i$. Concernant la minorant, lorsque c'est possible, on choisira Λ de manière à ce que $X^l = (E[X_1|\Lambda], \dots, E[X_n|\Lambda])$ soit un vecteur comonotone ce qui permettra parfois d'explicitier analytiquement $\sum_{i=1}^n E[X_i|\Lambda]$. On trouvera dans Dhaene, Vanduffel et al. (2004) une procédure pour déterminer Λ de manière à ce que X^l soit comonotone lorsque les marginales de X sont log-normales.

De plus, lorsque l'on évaluera le risque associé à $\sum_{i=1}^n X_i$ par une mesure de risque de distorsion, on disposera d'un encadrement dont les bornes seront simples à déterminer puisque les mesures de risque de distorsion sont comonotones additives (cf. le sous-paragraphe 1.1.2.4).

2. Un traitement différencié du risque

Bien que participant d'un même mouvement de rationalisation et d'homogénéisation, les trois référentiels en évolution (valorisation économique : *embedded value*, comptabilités : normes IFRS et contrôle prudentiel : Solvabilité 2) n'arrivent pas à converger complètement sur le traitement qui doit être fait du risque. En effet, si ces trois approches s'accordent sur le fait que les entreprises d'assurance doivent utiliser la meilleure information disponible et, en particulier, les données de marché dès que c'est possible, les objectifs poursuivis sont différents et conduisent à des approches distinctes de traitement du risque.

En effet, si l'approche prudentielle vise à déterminer des niveaux de provisions techniques et de fonds propres minimaux qui doivent permettre à la compagnie de *supporter* le risque, les approches de valorisation économique et comptable ont un objectif similaire, celui de donner une valeur à un portefeuille de contrats ou à une entreprise. Ces deux démarches sont très différentes.

2.1. L'approche assurantielle

La démarche prudentielle adoptée dans Solvabilité 2 est à rapprocher de la philosophie de transfert de risque sur laquelle s'est développée l'activité d'assurance.

Rappelons que l'opération d'assurance consiste en un transfert de risque de l'assuré vers l'assureur moyennant une prime. Pour que cette opération soit profitable pour les deux parties, il faut que l'aversion au risque des deux parties soit différente. C'est la situation de base en assurance puisque la compagnie d'assurance, couvrant un ensemble de risques, bénéficie d'effets de mutualisation qui ont pour conséquence de réduire le risque relatif de

l'entreprise². Si l'opération ainsi réalisée peut sembler déséquilibrée en espérance, puisque si le tarif est bien établi, la prime payée par l'assuré est supérieure à la prime pure du risque transféré, elle s'équilibre du fait de la différence d'aversion relative au risque.

Considérons un exemple très simple : deux assurés sont soumis à un risque de même nature (mais touche les assurés de manière indépendante) qui, avec une probabilité de 50 % leur causera une perte de 1. La situation individuelle d'un assuré peut être résumée par

$$X_1 = \begin{cases} 0, & p = 0,5, \\ 1, & 1-p = 0,5. \end{cases}$$

La prime pure (l'espérance mathématique) pour s'assurer contre un tel risque s'élève à 0,5.

Supposons que nos deux assurés s'associent pour supporter les éventuelles pertes causées par ce risque. Leur association peut être résumée par :

$$X_1 + X_2 = \begin{cases} 0, & p^2 = 0,25, \\ 1, & 2p(1-p) = 0,5, \\ 2, & (1-p)^2 = 0,25. \end{cases}$$

Cette association n'a pas d'impact sur le niveau de la prime pure (toujours 0,5). En revanche, on peut remarquer que si les deux assurés ont contribué à cette association en apportant la prime pure, cela leur suffira pour se prémunir de ce risque dans 75 % des cas, à comparer avec un cas sur deux lorsqu'ils gardent ce risque individuellement.

Ainsi l'opération d'assurance repose sur la mutualisation décrite mathématiquement par la loi des grands nombres. Rappelons que celle-ci énonce le fait que si on considère ensemble un très grand nombre de risques de même nature mais indépendants, alors le montant que l'on aura à payer sera proportionnel au nombre de ces risques et certain (au sens contraire d'aléatoire).

En pratique les assureurs se trouvent confrontés à deux écueils :

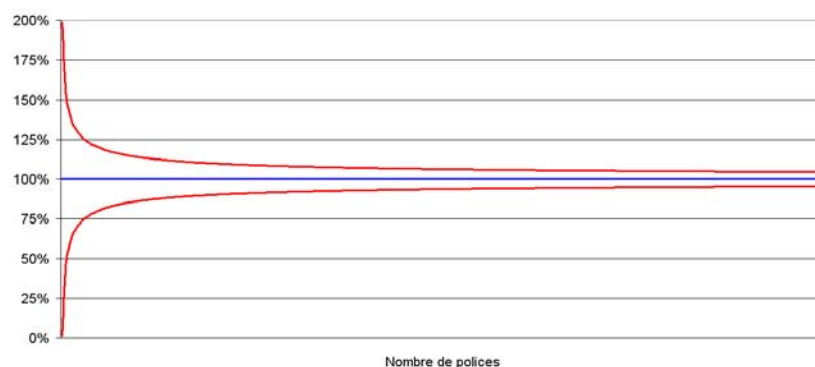
- les portefeuilles de risques similaires ne sont pas de taille infinie,
- les risques qui constituent ces portefeuilles ne sont pas systématiquement indépendants les uns des autres.

Dans ce contexte et dans un souci de protection des assurés, les pouvoirs publics ont imposé aux assureurs de disposer, en plus des provisions techniques destinées à couvrir les prestations futures espérées, de ressources propres destinées à pallier l'absence de mutualisation totale. À ce sujet, il est intéressant que le dispositif actuel (issu notamment de la Directive Solvabilité 1) prévoit des exigences en matière de capitaux proportionnelles à la masse des provisions en assurance vie et proportionnelles à la somme des primes ou des sinistres en

². Tout au moins dans la situation de référence d'indépendance entre les différents risques. La prise en compte de la possible dépendance entre ceux-ci est étudiée dans le Chapitre 5.

assurance non-vie. De tels dispositions n'intègrent pas³ l'effet de mutualisation qui va de pair avec la taille du groupe comme le montre la Figure 3 qui montre le rapprochement relatif des bornes de l'intervalle de confiance à 90 % autour de l'espérance.

Figure 3 - Évolution des bornes de l'intervalle de confiance à 90 % en fonction du nombre de polices assurées



Le dispositif Solvabilité 2 a pour vocation de rendre les entreprises d'assurance plus solvables en fixant notamment des exigences quantitatives de fonds propres en référence au risque réel supporté par ces sociétés. Incidemment, cette référence explicite au risque passe nécessairement par le choix d'une mesure de risque (cf. le paragraphe 1.1).

2.2. L'approche économique

L'approche économique ne répond pas aux mêmes impératifs que l'approche prudentielle. En effet, qu'il s'agisse de problématiques d'Embedded Value ou de normes IFRS, l'objectif est similaire : donner une valeur à un portefeuille de contrats ou à une compagnie. L'idée est ici de donner la vision la plus juste possible de ce que « vaut » l'entreprise ou le portefeuille.

Par exemple, les travaux de l'IASB sur la phase II de la norme IFRS dédiée au contrat d'assurance définissent le montant à inscrire en provisions au titre d'un ensemble de contrats à sa *Current Exit Value* (CEV) ou valeur actuelle de sortie, à savoir, le prix que l'assureur qui détient cet ensemble de contrats pourrait donner à un autre assureur pour lui transférer tous les droits et toutes les obligations résultants de ces contrats.

Cette définition renvoie donc à la notion de prix de marché alors même qu'un tel marché secondaire des contrats d'assurance n'existe pas *stricto sensu*. Puisque l'assureur ne peut lire ces prix sur un marché organisé, les travaux du Board évoquent les approches financières de valorisation qui ont connu un essor formidable dans les années soixante-dix avec notamment les travaux de Black, Scholes et Merton. Ces travaux s'appuient notamment sur l'hypothèse d'absence d'opportunité d'arbitrage.

³. Ce point est à nuancer en assurance non-vie par l'application de deux coefficients différents en fonction de la taille du portefeuille (cf. l'exemple du Chapitre 3).

2.2.1. Introduction binomiale

Considérons une action cotée S_0 en 0 et une option d'achat sur cette action dont la valeur actuelle est C_0 . L'option arrive à échéance en 1, date à laquelle nous supposons que l'action vaut $S_u = (1+u) \times S_0$ avec une probabilité p et $S_d = (1+d) \times S_0$ avec une probabilité $1-p$. Son prix d'exercice est K , le flux de l'option peut donc être résumé par :

$$\begin{array}{c} \nearrow \\ C_0 \\ \searrow \\ C_u = [S_u - K]^+ = [(1+u) \times S_0 - K]^+ \\ C_d = [S_d - K]^+ = [(1+d) \times S_0 - K]^+ \end{array}$$

On remarque qu'en prenant une position longue (achat) sur une action et courte (vendeuse) sur un certain nombre n d'options, il est possible de constituer un portefeuille dont la valeur en 1 est certaine. En effet, il suffit pour cela de déterminer n tel que :

$$S_u - nC_u = S_d - nC_d,$$

soit

$$n = \frac{[S_u - K]^+ - [S_d - K]^+}{(u-d)S}.$$

Un portefeuille constitué dans ces proportions procure un montant certain en 1, il est donc sans risque. En l'absence d'opportunité d'arbitrage⁴, un tel portefeuille procure le taux sans risque r . Sa valeur actuelle est donc donnée par sa valeur future (certaine) actualisée au taux sans risque, soit :

$$S_0 - n \times C_0 = \left\{ (1+d)S_0 - n \times [(1+d)S_0 - K]^+ \right\} e^{-r}.$$

On en déduit, la valeur actuelle de l'option d'achat :

$$C_0 = \frac{1}{n} \left\{ \left(1 - (1+d)e^{-r} \right) S_0 - ne^{-r} \times [(1+d)S_0 - K]^+ \right\}.$$

La valeur de cette option ne dépend pas directement de la probabilité p de hausse du cours de l'action sur la période. Ce résultat peut intuitif s'explique par le fait que cette information est déjà contenue dans le niveau actuel du cours de l'action S_0 . En effet, potentiels et probabilités de hausse et de baisse

⁴. Un marché présente des opportunités d'arbitrage lorsqu'il permet de mettre en œuvre des stratégies systématiquement meilleures que d'autres. L'hypothèse d'absence d'opportunité d'arbitrage repose sur l'hypothèse de réactivité du marché : si une opportunité d'arbitrage existe, le marché va l'utiliser ce qui va modifier les prix pour les mettre au niveau d'équilibre pour lesquels il n'y aura plus cette opportunité.

Incidentement, cette hypothèse contraint à ce que $1+d \leq e^r \leq 1+u$.

du cours du titre ont déjà été incorporés par le marché dans la valeur initiale de l'action.

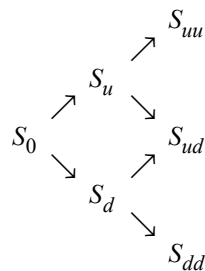
Si l'on cherche à exprimer la valeur du call en fonction de ses flux possibles en fin de période, après quelques développements et réductions, on obtient :

$$C_0 = \{qC_u + (1-q)C_d\}e^{-r},$$

où $q = \frac{(e^r - 1) - d}{u - d}$.

Cette réécriture nous permet d'observer que le prix du call peut s'exprimer comme l'espérance de ses flux futurs actualisés au taux sans risque dans un univers dans lequel la probabilité de hausse du cours de l'action vaut q . Dans cet univers, le prix d'un actif est égal à l'espérance de ses flux futurs actualisés⁵. Cette propriété se traduit par le fait que, dans cet univers, les agents économiques donnent le même prix à des investissements, de niveau de risque différent, mais de même espérance de gain d'où la dénomination d'« univers risque-neutre ».

Cet exemple binomial s'étend naturellement au cas multi-périodique :



Il s'agit alors de mettre en œuvre la même technique que précédemment mais à rebours (valorisation de C_u et C_d puis de C) avec l'hypothèse supplémentaire qu'il n'y a pas de frais de transaction (recomposition du portefeuille d'arbitrage à la fin de chaque période).

Le lecteur intéressé pourra se référer à Hull (1999) pour une présentation plus complète de l'approche binomiale et plus généralement sur les méthodes de valorisation des options et des dérivés financiers.

2.2.2. Retour sur les hypothèses

Le modèle binomial présenté *supra* et plus généralement le modèle de Black et Scholes (1993) et ses dérivés (cf. Hull (1999)) reposent sur deux hypothèses principales relatives au marché : celui-ci est parfait et complet.

On qualifie de parfait un marché idéalisé dans lequel :

- les titres sont divisibles,
- les ventes à découvert sont autorisées sans limite,

⁵. On parle de propriété martingale des flux futurs actualisés.

- les agents économiques disposent de la même information et aucun d'entre eux n'a la capacité de faire bouger les prix par son seul comportement,
- il n'y a pas de friction (frais, fiscalité) sur les achats et les cessions de titres,
- les taux de prêts et d'emprunts sont identiques.

Dans un tel marché, les méthodes de valorisation des actifs financiers, et notamment des dérivés, reposent sur l'hypothèse d'absence d'opportunité d'arbitrage. On dit qu'il y a absence d'opportunité d'arbitrage lorsque aucune stratégie financière n'est systématiquement préférable à une autre.

Sur les marchés financiers, ces opportunités peuvent exister, néanmoins le fonctionnement du marché conduit à ce qu'elles soient repérées et que les arbitragistes prennent position ce qui conduit mécaniquement à équilibrer les prix.

Par ailleurs, un marché est complet lorsque tous les états du monde d'arrivée sont atteignables par une certaine composition du portefeuille des actifs du marché. C'est cette dernière hypothèse qui assure l'unicité de la mesure de probabilité sous laquelle les prix actualisés sont des martingales. Lorsque ce n'est pas le cas, il existe une infinité de probabilités équivalentes pour lesquelles la propriété martingale est vérifiée. L'obtention d'un prix est alors conditionnée par le choix d'une de ces mesures de probabilité. Ballotta (2005) étudie les prix obtenus sous différentes probabilités dans le cas d'un contrat participatif.

2.3. Confrontation des deux approches en assurance

L'analyse des contrats d'épargne toujours plus sophistiqués montre que certaines garanties (taux garantis, participation aux bénéfices financiers ou encore les garanties plancher sur les contrats en unités de comptes), pouvaient être exprimées formellement comme des flux d'options et elles sont alors valorisées comme l'espérance actualisée en probabilité risque neutre des flux qu'elles engendrent. On est conduit à mettre en avant deux points en particulier :

- pour un engagement d'assurance, l'évaluation financière est menée conditionnellement à une réalisation de l'aléa démographique (pour une garantie plancher, par le théorème des probabilités totales, on conditionne par la connaissance de la date du décès, par exemple), le résultat final consistant à prendre la valeur moyenne des résultats obtenus. Compte tenu de la mutualisation imparfaite d'un portefeuille réel, le « coût » de la garantie optionnelle n'est *in fine* qu'une approximation de sa valeur théorique ;
- quelle est la légitimité de l'évaluation d'un engagement via la détermination d'un prix sur un marché ? La question pour l'assureur n'est pas tant de donner un prix à un portefeuille (sauf pour l'*embedded value*) que de provisionner à un niveau suffisant pour garantir une maîtrise satisfaisante des risques gérés.

Concrètement, l'assureur doit faire face à deux types de risques : des risques mutualisables, dont la prise en charge fonde l'activité d'assurance et des risques non mutualisables, comme ceux que l'on vient d'évoquer.

Les risques non mutualisables sont souvent financiers mais il peuvent être directement associés aux engagements d'assurance : mortalité stochastique et, plus généralement, tous les phénomènes qui remettent en cause l'indépendance entre les assurés (attentats, risques environnementaux, etc.). Il importe donc de fixer une méthode techniquement fondée de calcul des provisions pour les risques non mutualisables. En effet, sur ces risques systématiques (autres que risques de marché), il n'est pas possible de tenir le raisonnement introduit en 2.2 qui permet de constituer une stratégie qui élimine le risque. En particulier, on ne dispose pas d'actifs financiers ou de dérivés d'assurance qui permettent de mettre en place de telles stratégies⁶. Pour ces risques d'assurance systématiques, la distribution du passif n'a donc sens qu'en probabilité historique.

C'est cette logique qu'adopte le projet Solvabilité 2 en fixant comme critère de référence le contrôle de la probabilité de ruine à un an.

Par ailleurs, un engagement (ou une marge de risque) déterminé par les techniques financières n'a pas de sens autre que conventionnel (i.e. imposé par la norme) puisque la gestion du portefeuille d'arbitrage n'est pas mise en oeuvre et, quand bien même elle le serait, les imperfections et la durée des engagements rendent la couverture approximative.

Il convient donc de distinguer la fixation d'un prix qui, en l'absence d'un marché organisé, relève d'une démarche conventionnelle et le contrôle des risques pour la gestion technique.

Si, dans le premier cas, les techniques de couverture, pour lesquelles la probabilité risque neutre fournit un moyen de calcul simple, sont acceptables, il en va différemment pour quantifier les risques portés par l'assureur, dont l'évaluation se fait dans le monde réel. La probabilité risque neutre peut être utilisée lors d'étapes intermédiaires de valorisation, mais la mesure du risque nécessite le recours à la probabilité historique. Les deux approches ne sont pas exclusives l'une de l'autre, elles traitent des problématiques différentes.

Il faut garder à l'esprit que l'assurance ne consiste pas à neutraliser des risques via des techniques de couverture mais à les gérer en contrôlant le niveau d'incertitude qui leur est attaché. Si le recours à des méthodes financières basées sur une condition d'absence d'opportunité d'arbitrage apparaît légitime pour fixer un prix dans un cadre normatif (IFRS et MCEV), l'analyse de la solvabilité doit être effectuée avec les lois de probabilité réelles. L'application des méthodes financières en assurance doit être conduite avec discernement.

⁶. Le développement de la titrisation de risques d'assurance permettra, peut-être, dans le futur de multiplier le nombre de risques pour lesquels on disposera de dérivés d'assurance.

3. Contrats en unités de compte : les garanties plancher

La garantie plancher en cas de décès sur les contrats en unités de comptes est l'exemple type d'un engagement qui dépend simultanément du risque de marché et du risque de mortalité.

Après avoir présenté les contrats en unités de comptes tels qu'ils sont commercialisés en France et les risques qu'ils peuvent présenter pour les assureurs, nous revenons plus particulièrement sur les problématiques de tarification et de gestion des risques induits par les garanties plancher en cas de décès.

3.1. Contrats en unités de compte

Les contrats en unités de comptes sont des contrats d'assurance vie pour lesquels les primes versées sont investies sur des supports choisis par l'assuré. L'épargne évolue en fonction de ces supports et l'assureur ne garantit que le nombre de parts et pas la valeur de celles-ci. C'est donc l'assuré qui supporte le risque de placement. Ces contrats peuvent être mono-support ou multi-supports et, dans ce cas, un des supports peut être un fonds en euros, voire l'actif général de la compagnie.

Les supports sont généralement profilés (cf. les noms des supports des assureurs : prudence, dynamique, performance) de manière à permettre à l'assuré de choisir entre des investissements plus ou moins risqués et selon leur localisation géographique (France, Europe, États-Unis, Asie, pays émergents, etc.)

Ces produits en unités de compte se sont particulièrement développés ces dernières années. En France en 2005, le chiffre d'affaires des contrats d'épargne représente 116,6 milliards d'euros⁷ dont près d'un quart est affecté aux unités de comptes.

3.2. Principaux risques

Le risque de placement étant supporté par l'assuré dans ce type de contrats, les deux principaux risques proviennent de la présence de garanties plancher en cas de vie ou en cas de décès et du risque de renonciation résultant d'une clause contractuelle ou d'une mauvaise information transmise à l'assuré.

3.2.1. Garantie plancher

Dans le cas le plus fréquent, une garantie plancher sur un contrat en unités de comptes permet à l'assuré ou à ses ayants-droits de toucher, au moment du dénouement du contrat, le maximum entre les primes versées et l'épargne valorisée.

On distingue deux types de garanties plancher : en cas de vie (cf. Bergonzat et Cavals (2006)) et en cas de décès (cf. Merlus et Pequeux (2000)).

Dans le premier cas, l'assureur garantit à l'assuré, qu'au terme du contrat, il percevra l'épargne constituée sauf si elle s'avère inférieure aux primes versées auquel cas, ce sont ces primes lui seront reversées. Aujourd'hui très peu

⁷. Cf. le rapport annuel de la FFSA 2005, p. 16.

d'acteurs du marché prévoient encore ce type de clause dans les contrats qu'ils commercialisent.

Dans le second cas, ce sont les bénéficiaires du contrat qui bénéficient de cette disposition à la date de décès de l'assuré. Alors l'incertitude porte sur la valeur de l'épargne et sur le décès de l'assuré. Généralement, les contrats qui comportent ce type de garantie prévoient le prélèvement d'une partie de l'épargne pour la financer. Selon que ce prélèvement est effectué comme une prime de risque ou de manière forfaitaire en pourcentage des encours, il conviendra ou non de constituer une provision pour garantie plancher.

3.2.2. Renonciation

La renonciation consiste, pour un assuré, à annuler le contrat et réclamer les primes qu'il a versées à l'assureur. Ce risque peut résulter :

- de l'application normale de l'art. L132-5-1 qui prévoit que dans un délai de 30 jours à compter de la date de réception de son certificat d'adhésion, l'assuré peut renoncer à son adhésion (cette clause est généralement reprise dans les conditions générales des contrats) ;
- d'un manque d'information des assurés qui a pour effet de perpétuer le droit à renonciation.

Sur le deuxième point, les arrêtés de mars 2006 précisent les informations qui doivent être communiqués aux clients lors de la commercialisation des contrats.

L'arrêté du 8 mars 2006 précise qu'en tête de contrat, un encadré doit permettre de répondre aux questions que se posent le plus souvent les assurés, à savoir : la nature individuelle ou de groupe du contrat, la caractère garanti ou non des sommes investies, la participation aux bénéfices, la disponibilité des sommes et les délais de versement, les frais prélevés par l'assureur, les enjeux du choix de la durée de placement et les modalités de désignation des bénéficiaires.

L'arrêté du 1^{er} mars 2006 impose aux assureurs de faire figurer dans un tableau les valeurs de rachat des 8 premières années et d'expliquer clairement leur mécanisme de calcul lorsque celles-ci ne peuvent être établies dès la souscription (cas des contrats en unités de compte).

Ces deux arrêtés viennent à la suite de la loi 2005-1564 du 15 décembre 2005 qui a créé l'article L132-5-2 qui prévoit l'obligation de remise d'une note d'information distincte des conditions générales du contrat de manière à clarifier l'exhaustivité de l'information dont doit disposer l'assuré pour que le droit à renonciation ne soit pas prorogé. Cependant cette loi n'étant pas rétroactive, le contentieux pour les contrats souscrits antérieurement perdure.

Le risque de renonciation porte essentiellement sur les produits avec supports en unités de comptes pour lesquels la propension d'une diminution forte de la valeur du support sur un horizon même réduit (30 jours) est vraisemblable.

On distingue donc deux problématiques dans le cadre de la gestion de ce risque : celle liée à la faculté contractuelle et celle résultant de l'éventuelle incomplétude de l'information aux assurés.

D'une manière générale, pour parer au risque de renonciation sur la période des 30 jours suivant la souscription, les conditions générales des contrats en unités de comptes spécifient que l'épargne ne sera investie de manière effective sur les supports UC qu'à expiration de ce délai. En attendant, elle est le plus souvent investie sur un fonds monétaire, voire l'actif général de l'assureur.

Concernant le risque lié à l'information aux assurés, si les services juridiques des assureurs ont intégré les exigences résultant de la loi du 15 décembre 2005 et des arrêtés de mars 2006 aux nouveaux contrats, le risque existe souvent pour les contrats plus anciens. En particulier, l'arrêt du 7 mars 2006 de la deuxième chambre civile de la Cour de cassation a confirmé la mauvaise rédaction de nombreux contrats d'assurance vie.

La gestion de ce risque peut être envisagée de diverses manières :

- mise à plat de l'information aux assurés avec mise en exergue du droit de renonciation de 30 jours ;
- clôture des contrats par décision de l'assureur (lorsque c'est possible) avec proposition de basculement de l'épargne vers un nouveau contrat.

3.3. Valorisation et gestion des garanties plancher

Les travaux de Frantz, Chenut et Walhin (2003) illustrent parfaitement les différences entre les méthodes présentées précédemment dans le paragraphe 2 et insistent notamment, en matière de gestion du risque, sur le lien entre le procédé de valorisation et les stratégies mises en œuvre pour gérer le risque.

Ces travaux s'inscrivent dans une problématique de tarification d'un traité de réassurance prévoyant la cession pleine et entière du risque induit par une clause de garantie plancher en cas de décès sur un contrat en unités de comptes mono-support.

Les auteurs étudient les deux approches évoquées précédemment pour associer un prix à un risque :

- l'approche « financière » qui repose sur l'hypothèse de complétude du marché et donc sur la possibilité de mettre en place une stratégie permettant de se couvrir complètement contre le risque (ou de minimiser le risque supporté dans le cas de marchés incomplets) ;
- l'approche « actuarielle » qui repose sur la loi des grands nombres et qui conduit à déterminer la prime comme somme de la prime pure (moyenne des prestations futures) et d'une marge pour risque déterminée en faisant référence au risque supporté par le biais d'une mesure de risque (cf. le sous-paragraphe 1.1).

La garantie plancher étudiée prévoit qu'en cas de décès, les bénéficiaires du contrat perçoivent le maximum entre la valeur des unités de comptes du contrat et la somme des primes versées.

Formellement, si l'on note $S = (S_t)_{t \geq 0}$ le processus décrivant l'évolution de la valeur des unités de comptes au cours du temps et K la somme des primes versées, le réassureur cherche à donner un prix à l'engagement de verser

$$\left[K - S_t \right]^+,$$

en cas de décès de l'assuré en t .

3.3.1. Modélisation

Frantz et al. (2003) modélisent l'évolution de la valeur de l'épargne en UC par un mouvement brownien géométrique :

$$\frac{dS_t}{S_t} = \mu dt + \sigma dB_t,$$

où $B = (B_t)_{t \geq 0}$ est un mouvement brownien sous la probabilité historique P .

De plus, il est supposé que le marché financier est complet, ne présente pas d'opportunité d'arbitrage et que le taux sans risque r est constant.

Dans ce contexte, il existe une unique mesure de probabilité Q équivalente à P sous laquelle les prix actualisés sont des martingales et, d'après le théorème de Girsanov, le prix de l'épargne suit le processus suivant sous cette probabilité :

$$\frac{dS_t}{S_t} = r dt + \sigma d\tilde{B}_t,$$

où $\tilde{B} = (\tilde{B}_t)_{t \geq 0}$ est un mouvement brownien sous la probabilité risque-neutre Q .

Par ailleurs, les assurés sont censés décéder tous selon une même table de mortalité. Dans la suite, les notations actuarielles classiques sont employées :

- T_x la durée résiduelle de survie (aléatoire) sachant que l'on est en vie à l'âge x ;
- ${}_n p_x$, la probabilité de vivre n années supplémentaires sachant que l'on a atteint l'âge x ;
- q_x , la probabilité de décéder dans l'année sachant que l'on a atteint l'âge x .

Enfin, l'évolution des marchés et la mortalité des assurés sont supposées indépendantes.

3.3.2. Primes pures

En faisant l'hypothèse implicite que les prestations sont payées en fin d'année et avec l'aide du théorème des probabilités totales, en conditionnant par la survenance du décès, le flux actualisé auquel doit faire face le réassureur est donné, pour un assuré d'âge x , par :

$$e^{-rT_x} \left[K - S_{T_x} \right]^+ \mathbf{1}_{T_x \leq \tau},$$

où τ désigne le terme du contrat.

S'il est possible de mettre en place une couverture qui permet de se prémunir (complètement) du risque, le prix associé à ce risque ne saurait être différent de celui de la couverture.

Comme on le verra *infra*, le marché n'est pas parfait, les coûts de friction existent et la mise en place d'un portefeuille répliquant nécessite un arbitrage entre rebalancements fréquents (et donc coûts de friction élevés) et rebalancements moins fréquents (et donc coûts de réajustement élevés).

Il est néanmoins possible d'établir une prime pure théorique, obtenue par le théorème des probabilités totales en conditionnant par la probabilité de décès de l'assuré, soit :

$$\begin{aligned}
 SPP &= E_{P \times Q} \left[e^{-rT_x} \left[K - S_{T_x} \right]^+ \mathbf{1}_{T_x \leq \tau} \right] \\
 &= \sum_{n=1}^{\tau} p_{x+n-1} \times q_{x+n-1} \times E_Q \left[e^{-rn} \left[K - S_n \right]^+ \right] \\
 &= \sum_{n=1}^{\tau} p_{x+n-1} \times q_{x+n-1} \times \left[K e^{-rn} \Phi(-d_2(n)) - S_0 e^{-(\mu-r)n} \Phi(-d_1(n)) \right]
 \end{aligned}$$

où $d_2(t) = \frac{\ln S_0/K + (r - \sigma^2/2)t}{\sigma\sqrt{t}}$ et $d_1(t) = d_2^{\hat{f}}(t) + \sigma\sqrt{t}$.

S'agissant d'un risque mixte (mortalité et financier), il convient de préciser ce que l'on entend par prime pure. En effet, si prime pure désigne en assurance le montant moyen que devra déboursier l'assureur qui assure le risque, *SPP* n'est pas réellement une prime pure puisque la valorisation risque-neutre du dérivé financier inclut implicitement une prime de risque (cf. le sous-paragraphe 2.2) et conduit donc à un montant supérieur à celui qui aurait été obtenu par un calcul de valeur actuelle probable traditionnellement mené sur les risques d'assurance.

3.3.3. Gestion du risque

Donner un prix au risque et le gérer de manière effective sont deux choses différentes. Par exemple, dans le paragraphe précédent, les hypothèses financières classiques de perfection du marché et d'absence d'opportunité d'arbitrage, ont permis de donner un prix à la garantie plancher. Cependant, même si ce prix est retenu pour la communication financière, par exemple, le gestionnaire peut ou non mettre en place de manière effective la stratégie de couverture sur laquelle il repose.

En pratique, il est confronté au fait que les hypothèses idéales dans lesquelles on se place pour utiliser les méthodes de valorisation risque-neutre (cf. le sous-paragraphe 2.2.2) ne sont pas vérifiées. En particulier, les coûts de friction ne sont pas nuls et donc recomposer son portefeuille répliquant en temps continu n'est pas envisageable. Si l'on souhaite mettre en place une couverture qui minimise le risque financier, il faut donc choisir une périodicité pour les recompositions du portefeuille répliquant. Comme évoqué précédemment, le gestionnaire est confrontés à deux sources de coût résiduel :

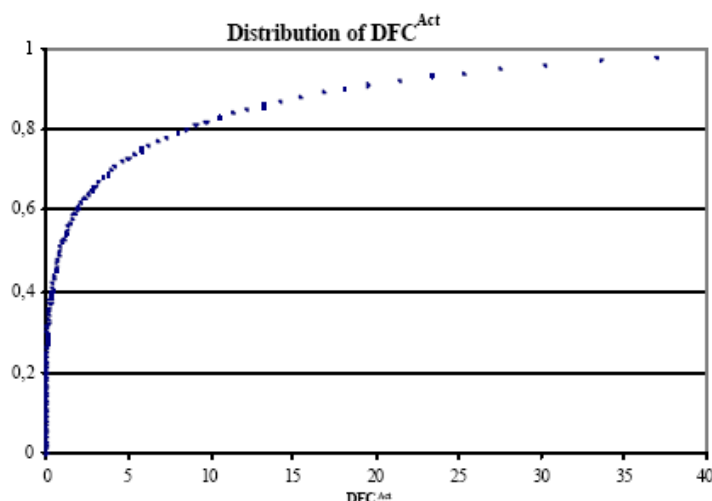
- les frais liés aux cessions et aux achats de titre ;
- les coûts liés au caractère non-continu des reconstitution de portefeuille : entre les instants t et $t+1$, l'actif a peut-être évolué défavorablement auquel cas, le gestionnaire doit faire face à un coût supplémentaire.

Cet aspect est détaillé dans Planchet, Thérond et al. (2005).

En matière de gestion du risque, Frantz et al. (2003) remarquent qu'en tant que gestionnaire du risque, les deux situations peuvent être envisagées : mettre en œuvre ou pas la couverture financière (même si elle est imparfaite du fait des frottements évoqués précédemment). Les deux graphiques suivants sont repris de Frantz et al. (2003).

La Figure 4 reprend la distribution de la valeur actuelle des flux futurs lorsque la couverture financière n'est pas mise en place.

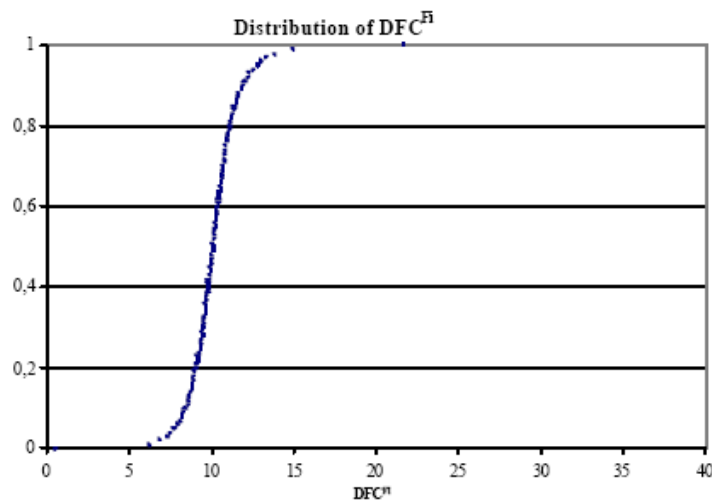
Figure 4 - Distribution des coûts futurs en absence de gestion active de la couverture



De manière alternative, il est possible de mettre en œuvre la stratégie de couverture de manière à minimiser le risque. Il s'agit alors de déterminer le coût de la couverture et de mesurer le risque résiduel.

Dans leurs travaux, Frantz et al. (2003) font l'hypothèse que les frais liés aux achats et aux cessions de titre sont nuls, que les reconstitutions de portefeuille sont annuelles et que donc, seuls les coûts de reconstitution liés au caractère non-continu de ceux-ci présentent une charge pour le gestionnaire. La Figure 5 reprend la distribution de la valeur actuelle des flux futurs dans l'approche lorsque la gestion active de la couverture est mise en place.

Figure 5 - Distribution des coûts futurs lorsque la gestion active de la couverture est mise en place



La mise en œuvre effective de la couverture a pour effet de réduire considérablement le risque supporté par l'assureur, même si elle ne conduit pas à l'éliminer complètement du fait du caractère discret des recompositions du portefeuille. Il subsiste donc un risque résiduel malgré la mise en place de la couverture. Par ailleurs, à ce risque vient s'ajouter celui résultant de la mutualisation imparfaite des décès : les montants couverts sont déterminés à partir de l'estimation des décès futurs du portefeuille.

4. Conclusion

L'objet de ce premier chapitre était de présenter les outils mathématiques les plus utilisés en assurance pour mesurer et comparer des risques, puis d'insister sur le fait que l'approche généralement retenue en termes d'évaluation dépend de la nature des risques eux-mêmes.

Ainsi pour les risques financiers, les modèles de valorisation économique commencent à s'imposer qu'il s'agisse d'établir les états comptable (normes IFRS), les éléments de la communication financière à destination du marché (*Market Consistent Embedded Value*) ou même de déterminer des provisions et les exigences de capitaux de solvabilité (Solvabilité 2). En particulier, ces modèles s'imposent même dans des cas où la stratégie de couverture sur laquelle ils reposent ne peut être mise en œuvre de manière effective. De plus, même dans les cas où il est possible de mettre en œuvre la couverture, l'utilisation de ces méthodes ne signifie pas que la gestion du risque est effectuée par la couverture.

Ce dernier point pose un problème conceptuel dans la mesure où deux sociétés identiques quant à leurs engagements et leurs actifs mais avec une gestion différente, mise en œuvre ou pas de la couverture, obtiendront les mêmes résultats, les mêmes valeurs la plupart du temps. Il n'en demeure pas moins que les profils de risque de ces deux compagnies sont très différents. On

l'a vu dans le cas des garanties plancher en cas de décès sur les contrats en unités de compte en comparant la Figure 4 et la Figure 5. Incidemment, le bon gestionnaire n'est pas uniquement celui qui produit de bons reporting comptables, de communication financière et prudentiels. En effet, bien qu'en constante évolution, ceux-ci ne restent que des images, des projections de l'entreprise d'assurance ; il s'agit pour le gestionnaire d'analyser le risque et de mettre en place les bonnes mesures de gestion de manière à protéger les assurés et corrélativement à pérenniser la rentabilité de l'entreprise, dans l'intérêt notamment de ses actionnaires ou sociétaires.

La prise en compte de ces règles de gestion dans les nouveaux référentiels est un problème à part entière du fait de l'étendu du champ des possibles et de la difficulté d'appréciation de la capacité de l'entreprise à mettre en œuvre dans les faits les règles annoncées aux horizons de gestion annoncés. Nous revenons sur la prise en compte des règles de gestion avec des exemples concrets dans le Chapitre 2.

Bibliographie

- Artzner Ph., Delbaen F., Eber J.M., Heath D. (1999) « Coherent measures of risk », *Mathematical Finance* 9, 203-28.
- Ballotta L. (2005) « A Lévy process-based framework for the fair valuation of participating life insurance contracts », *Insurance: Mathematics and Economics* 37 (2), 173-96.
- Bergonzat A., Cavals F. (2006) *Garanties en cas de vie sur contrats multisupports*, Mémoire d'actuariat, ENSAE.
- Black F., Scholes M. (1973) « The pricing of options and corporate liabilities », *Journal of Political Economy* 81 (3), 637-54.
- Denuit M., Charpentier A. (2004) *Mathématiques de l'assurance non-vie. Tome 1 : principes fondamentaux de théorie du risque*, Paris : Economica.
- Denuit M. (2004) « Actuarial theory for dependent risks », *8th international Congress on Insurance: Mathematics and Economics*, Rome, 14-16 juin 2004.
- Dhaene J., Laeven R.J.A., Vanduffel S., Darkiewicz G., Goovaerts M.J. (2004) « Can a coherent risk measure be too subadditive? », Research Report OR 0431, Department of Applied Economics, Katholieke Universiteit Leuven.
- Dhaene J., Vanduffel S., Tang Q.H., Goovaerts M., Kaas R., Vyncke D. (2004) « Solvency capital, risk measures and comonotonicity: a review », Research Report OR 0416, Department of Applied Economics, Katholieke Universiteit Leuven.
- Dhaene J., Wang S., Young V., Goovaerts M. (2000) « Comonotonicity and maximal Stop-Loss premiums », *Bulletin of the Swiss Association of Actuaries* 2, 99-113.
- Embrechts P. (2000) « Actuarial versus financial pricing of insurance », *Risk Finance* 1 (4), 17-26.
- Embrechts P., Furrer H., Kaufmann R. (2007) « Different Kinds of Risk » in *Handbook of Financial Time Series*, Eds. Andersen, Davis, Kreiss, and Mikosch.
- Frantz C., Chenut X., Walhin J.F. (2003) « Pricing and capital allocation for unit-linked life insurance contracts with minimum death guarantee », *Proceedings of the 13th AFIR Colloquium*, Maastricht.
- Gouriéroux Ch. (1999) *Statistique de l'assurance*, Paris : Economica.
- Hull J.C. (1999) *Options, futures and other derivatives*, 4th edition, Prentice-Hall.
- Hürlimann W. (1999) « On risk and price : Stochastic orderings and measures », *Actes du 27^e colloque ASTIN*, Cancun, 22 mars 2002.

- Kass R., Dhaene J., Goovaerts M. (2000) « Upper and lower bounds for sums of random variables », *Insurance: Mathematics and Economics* 27, 151-68.
- Merlus S., Pequeux O. (2000) *Les garanties plancher des contrats d'assurance vie en UC: tarification et couverture*, Mémoire d'actuariat, ENSAE.
- Partrat Ch., Besson J.L. (2005) *Assurance non-vie. Modélisation, simulation*, Paris : Economica.
- Planchet F., Thérond P.E., Jacquemin J. (2005) *Modèles financiers en assurance. Analyses de risque dynamiques*, Paris : Economica.
- Wang S. (2002) « A risk measure that goes beyond coherence », *Actes du 12^e colloque AFIR*, Cancun, 18 mars 2002.

Chapitre 2

Traitement spécifique du risque : aspects pratiques

Dans le premier chapitre de cette thèse, nous avons étudié la question du traitement du risque en assurance, sous un angle général, en mettant en exergue les différents modèles que l'on rencontre en pratique et notamment leurs limites provenant des conditions de leur emploi.

Ce deuxième chapitre traite du même problème mais sous un angle d'attaque différent et en faisant référence de manière plus explicite aux nouveaux référentiels dans lesquels vont s'inscrire les sociétés d'assurance : normes IFRS, MCEV et Solvabilité 2. Ainsi dans un premier temps, nous reviendrons sur ces trois cadres et leur finalité, et verront notamment qu'ils reposent sur un socle de principes communs.

De manière opérationnelle, ces principes vont conduire à considérer différemment de ce qui a été fait jusqu'à présent des éléments-clé de l'évaluation de l'activité d'assurance tels que les hypothèses actuarielles. Par ailleurs, la divergence des finalités conduit à ne pas considérer de la même manière des événements futurs tels que ceux découlant du comportement des assurés ou de celui de l'assureur.

Enfin, nous verrons, sur un exemple de contrats d'assurance vie de type épargne en euros, l'impact de ces différences en termes de valorisation et de modélisation du risque.

1. De nouveaux référentiels distincts

L'objectif de ce premier paragraphe est de rappeler les grandes lignes des nouveaux référentiels comptables, de communication financière et de solvabilité puis de mettre en évidence les bases communes sur lesquelles ils ont été construits et enfin de revenir plus particulièrement sur ce qui les distingue les uns des autres.

1.1. Contexte

Le système actuel de solvabilité, régit en partie par la Directive Solvabilité 1, repose sur les comptes sociaux des entreprises d'assurance. Comme n'importe quelle entreprise commerciale, les sociétés d'assurance établissent des comptes en vue notamment de déterminer un résultat sur lequel l'État pourra prélever l'impôt. En revanche, une des spécificités de l'activité d'assurance est le contrôle qui en est fait par les autorités. En effet, compte tenu du rôle de l'assurance dans le développement de l'économie et dans la sécurisation de la société, les sociétés d'assurance font l'objet d'un contrôle spécifique de manière à éviter leurs défaillances et ses conséquences pour les

assurés. Ce contrôle repose notamment sur des exigences en matières de fonds propres caractérisées par l'exigence de marge de solvabilité, par des règles sur les provisions techniques (comptables !) et la composition du portefeuille d'actifs. Le superviseur bénéficie également d'états réglementaires qui lui permettent de juger de la solidité financière des assureurs et, le cas échéant, de réclamer des capitaux propres supplémentaires. Ces états réglementaires sont, pour la-plupart, des déclinaisons des éléments comptables traditionnels et pour d'autres des stress-tests ou tests de sensibilité à des scénarios le plus souvent adverses : états C6bis, T3, C8, C9, etc.

La situation actuelle est donc celle de l'enchevêtrement des reportings comptable et prudentiel.

1.1.1. Normes IFRS

Depuis 2005, les sociétés européennes cotées ou faisant appel public à l'épargne doivent publier leurs comptes consolidés selon les normes IFRS pour *International Financial Reporting Standards*. Le passage à ce nouveau référentiel constitue un « choc des cultures » pour nombre de sociétés françaises puisque cette information comptable n'est pas principalement destinée à l'État mais aux marchés financiers. Ce contenu est sensiblement différent puisque l'objectif de ces normes est de donner la meilleure vision économique possible de l'entreprise, de fournir des outils d'aide à la décision à la veuve de Carpentras et de mesurer la richesse créée pour l'actionnaire. En conséquence, les principes sur lesquels repose ce cadre normatif ne sauraient coïncider avec ceux du Plan Comptable Général.

Le principe directeur de ces normes est celui de juste valeur ou *fair value*. Bien que la définition exacte de ce concept soit en cours de refonte dans le contexte notamment de la convergence avec les normes comptables américaines (les FAS), on peut définir la juste valeur d'un actif ou d'un passif comme étant le montant auquel deux parties intéressées et également informées s'échangeraient cet actif ou ce passif.

C'est pour cela que les détracteurs des normes IFRS parlent de « valeur à la casse » de la compagnie : la plupart de ses éléments comptables sont en effet valorisés en juste valeur, c'est à dire au prix auquel ils pourraient être cédés. On se rend compte que si cette notion est assez naturelle s'agissant d'instruments financiers cotés dont on observe le prix sur un marché liquide, il est plus délicat de donner la valeur d'un droit immatériel ou, pour ce qui nous concerne, d'un portefeuille de polices d'assurance.

On pourra consulter Thérond (2003) pour plus de renseignements sur les normes IFRS en général et notamment leurs processus d'élaboration et d'adoption au sein de l'Union européenne.

1.1.2. Solvabilité 2

À cette première évolution qui touche l'ensemble des sociétés mentionnées précédemment, s'en est ajoutée une seconde, propre au monde de l'assurance : Solvabilité 2. Il s'agit en fait d'un nouveau référentiel prudentiel destiné à remplacer celui issu en partie de la Directive Solvabilité 1. Son objectif est double : harmoniser les règles et le contrôle au sein de l'Union européenne, et

se doter d'un système plus adapté au risque effectif supporté par les assureurs. Pour illustration de ce deuxième point, l'exigence en matière de fonds propres ne fera plus l'objet d'un calcul de proportion en fonction des provisions mathématiques en assurance vie ou des primes et des sinistres en non-vie, mais devra permettre à l'entreprise de ne pas être en ruine à la fin de l'exercice avec une très forte probabilité (99,5 %). Il s'agit en fait de passer du système actuel dans lequel les sociétés d'assurance disposent de marges de sécurité implicites à plusieurs niveaux (provisions techniques, plus-values latentes, marge de solvabilité, etc.) à un système dans lequel la meilleure information possible est utilisée et les marges de prudence sont explicitées (marge pour risque dans les provisions techniques, capital de solvabilité).

On pourra consulter Planchet, Thérond et al. (2005) et Thérond et Girier (2005) pour un descriptif plus complet des principes de Solvabilité 2 et de son mécanisme d'élaboration.

1.1.3. Market Consistent Embedded Value

À ces deux référentiels « obligatoires » pour les entreprises concernées, s'ajoute un troisième référentiel vers lequel se tournent de plus en plus de sociétés d'assurance : celui de la communication financière par le biais de l'*embedded value*. Utilisée depuis les années 1990, l'*embedded value* ou valeur intrinsèque devient un outil de communication de plus en plus prisé et de plus en plus sophistiqué. Il s'agit en fait de déterminer une valorisation de l'entreprise à partir de la projection de son portefeuille et de la rapprocher de sa capitalisation boursière le cas échéant. Ces techniques ont beaucoup évolué depuis quelques années et notamment depuis la publication en 2004 des principes de l'*European Embedded Value* (EEV) par le CFO Forum¹. Il s'agit en fait d'un recueil des principes que doit respecter une évaluation d'*embedded value* pour avoir le label EEV, l'objectif étant d'uniformiser au maximum ces principes pour accroître la comparabilité des chiffres produits.

Parmi les approches qui répondent aux principes de l'EEV, l'approche dite *market-consistent* tend à s'affirmer comme le standard. Il s'agit en fait d'une méthode qui valorise les risques répliquables selon l'approche économique décrite dans le Chapitre 1 et qui n'affecte pas de prime de risque aux risques d'assurance mutualisables. À l'instar de ce qui a été fait dans le cadre plus général de l'EEV, un recueil des principes de la *Market Consistent Embedded Value* (MCEV) devrait voir le jour en 2007.

1.2. Des finalités différentes

Pour comprendre les différences et les similitudes entre les trois référentiels précédemment évoqués, il convient de revenir sur leurs objectifs. Si IFRS et MCEV relèvent de l'empire de la communication financière, Solvabilité 2 n'est pas destinée à l'information des marchés financiers mais à la détermination d'exigences de solvabilité auxquelles doit répondre l'assureur. En cela les principes rappelés dans la troisième étude d'impact quantitatif QIS 3 menée par le CEIOPS (cf. l'annexe) sont assez parlant puisqu'il est précisé que :

¹. Le CFO Forum est une association d'assureurs européens significatifs en termes de part de marché.

- la formule standard doit permettre d'estimer le niveau de capital que doit posséder l'assureur aujourd'hui pour ne pas être en ruine dans un an avec une probabilité de 99,5 % ;
- les provisions techniques doivent comporter une marge pour risque qui, en cas de défaillance de l'assureur, permettrait leur transfert vers un autre acteur du marché sans que, pour cela, celui-ci n'ait besoin de se financer.

Pour comparaison, rappelons que les normes IFRS et une MCEV visent à donner une information sur la valeur de la société dans le cadre d'une transaction dans des conditions normales. Si ces deux référentiels ont le même but, ils n'ont pas la même portée : alors que la MCEV est une valorisation globale de la compagnie, les comptes IFRS recèlent d'informations plus détaillées, le niveau de calcul étant celui du portefeuille.

Le Tableau 2 propose un comparatif des finalités et des principes de ces trois cadres d'évaluation.

Tableau 2 - Comparatif des référentiels IFRS, Solvabilité 2 et MCEV

	MCEV	Solvabilité 2	IFRS phase 2
Objectifs	Valorisation économique (transactions) Information financière	Contrôle prudentiel	Information financière
Information fournie	Valeur de la compagnie Rentabilité des affaires en portefeuille (et des affaires nouvelles)	Provisions techniques Exigence de fonds propres	Provisions techniques, Capitaux propres (variation des éléments d'actifs et de passifs)
Niveau d'application	Social puis consolidé	Social puis consolidé	Consolidé
Valorisation des engagements d'assurance	Espérance complétée de la valeur de l'ensemble des options et des garanties financières	Espérance + marge pour risque	<i>Current Exit Value</i> : montant qui serait exigé en contrepartie du transfert de l'engagement sur un marché (espérance + marge pour risque + marge pour service)
Valorisation des risques	Risques valorisés à travers les cash-flows futurs et les options & garanties	marge pour risque explicite (CoC ou VaR) au niveau du portefeuille + capital de solvabilité (SCR) qui doit contrôler le risque global de la compagnie	marge pour risque, au niveau du portefeuille, qui correspond à une prime de risque normalement disponible sur les marchés
Actualisation	Taux sans risque	Taux sans risque	Taux sans risque

1.3. Un socle de principes communs

Il n'en demeure pas moins que, s'ils ont des finalités différentes, les trois référentiels usent de moyens semblables en s'appuyant sur un socle de principes communs.

Ainsi, il est acquis dans les trois contextes, que ce sont les meilleures hypothèses possibles qui doivent être utilisées dans les projections. Si cela

semble naturel, il convient de se rappeler que les calculs de provision en normes françaises sont effectués sur la base d'hypothèses prudentes (lois de mortalité prudentes, taux d'actualisation nuls ou positifs mais prudents, etc.) Nous revenons plus précisément sur ce point dans le sous-paragraphe 2.1.

1.3.1. Traitement des risques financiers

Par ailleurs un autre élément significatif est celui du traitement du risque financier, les trois référentiels convergent pour le valoriser selon les approches économiques présentées dans le Chapitre 1. Il s'agit donc de valoriser les risques financiers au prix de leur portefeuille répliquant ou de couverture. L'objectif étant de coller le plus possible aux prix observés sur le marché.

La limite de cet exercice réside dans le fait que deux portefeuilles identiques (mêmes clauses contractuelles, réglementaires, mêmes assurés et même actif en représentation) comportant des risques financiers dans deux entreprises d'assurances différentes se verront attribuer la même valeur, obtenue selon des modèles de type Black et Scholes, quand bien même l'une déciderait de mettre en place la couverture de manière à éliminer le risque et pas l'autre. C'est la situation étudiée au Chapitre 1 dans le cas de la garantie plancher en cas de décès sur un contrat en unités de compte.

Dans le cadre de Solvabilité 2, on obtiendrait donc le même niveau de provision, néanmoins la différence d'exposition au risque financier se retrouverait au niveau de l'exigence de capital de solvabilité (le SCR, cf. l'annexe). Pour ce qui est des comptes IFRS, on obtiendrait également un même niveau de provision au titre du fait que, dans le cas du transfert du portefeuille, la gestion de ce risque revient au nouvel assureur qui est libre de mettre ou pas en place cette stratégie. Aussi, il n'y a pas de raison que le prix du portefeuille soit différent selon la politique de gestion mise en place par l'assureur qui détient actuellement le portefeuille. Ce point précis est une des difficultés rencontrées par le projet de norme IFRS phase 2 dédiée aux contrats d'assurance (cf. également le sous-paragraphe 2.2.2).

1.3.2. Traitement des risques non-financiers

Concernant les risques non-financiers, normes IFRS et Solvabilité 2 convergent sur le fait qu'il faut les valoriser selon une décomposition *best estimate* + marge pour risque.

Si la notion de *best estimate* est proche entre les référentiels comptables et prudentiels (modulo les différences sur certaines hypothèses, cf. le sous-paragraphe 2.1), celles de marge pour risque ne répondent pas à la même définition.

Dans Solvabilité 2, la marge pour risque doit permettre le transfert du portefeuille en cas de faillite de l'assureur. QIS 3 (cf. l'annexe) adopte ainsi la méthode du coût du capital pour la définir comme étant le coût d'immobilisation du capital relatif à cet engagement.

Le projet de norme IFRS phase 2 dédié aux contrats d'assurance propose une définition différente : la marge pour risque est la prime de risque, au-delà du *best estimate*, que requiert un intervenant du marché pour accepter de se voir transférer le risque.

Bien que distinctes, ces deux approches ne sont pas très éloignées, si l'on considère, par exemple, que la prime de risque requise par le marché correspond au coût d'immobilisation du capital de solvabilité.

Les calculs de MCEV font une distinction supplémentaire parmi les risques non-financiers : ces risques sont-ils diversifiables ou pas ? Le principe étant de ne pas accorder de prime de risque aux risques diversifiables puisque l'actionnaire, conformément à la loi des grands nombres, pourrait l'éliminer. Par exemple, le risque de fluctuation d'échantillonnage de la mortalité au sein du portefeuille d'assurance ne conduira pas à constituer une marge pour risque. En revanche, le risque systématique de mortalité (cf. Planchet et Thérond (2007a)) qui n'est pas diversifiable en nécessitera une : l'actionnaire ne peut pas s'en prémunir en achetant des titres des différents assureurs de la place puisqu'ils sont tous touchés par ce même risque.

2. Des incidences opérationnelles

Sur le plan opérationnel, les trois nouveaux référentiels amènent à réviser les techniques et méthodes mises en œuvre jusqu'à présent lorsqu'il s'agissait notamment d'établir des comptes en normes françaises.

Parmi les éléments les plus importants, le recours aux meilleures hypothèses possibles nécessite une attention particulière d'autant plus que c'est un principe commun aux différents cadres précédemment évoqués.

Par ailleurs, les projections stochastiques des portefeuilles d'assurance en général et d'assurance vie en particulier nécessitent la modélisation du comportement de l'assureur et des assurés dès lors que des choix sont offerts à l'un ou aux autres. C'est notamment le cas de toutes les décisions d'ordre discrétionnaire. Elles sont particulièrement nombreuses dans les contrats d'épargne. Parmi elles, on peut citer :

- pour les assurés, la possibilité de racheter totalement ou partiellement leur contrat, la possibilité de ne plus payer ses primes périodiques (mixtes), la possibilité d'effectuer des versements libres, etc.
- pour les assureurs, la politique de gestion des actifs, la revalorisation discrétionnaire des contrats et corrélativement la gestion de la provision pour participation aux bénéfices (PPB), etc.

Au sujet de ces futures décisions de gestion, la lecture des textes de référence sur les différents projets laisse augurer des différences assez nettes.

2.1. Hypothèses actuarielles

L'évaluation des engagements d'assurance repose sur des hypothèses de projection. Dans le contexte français d'établissement des comptes et des états réglementaires à destination des autorités de contrôle, le Code des assurances prescrit une bonne partie des hypothèses dites actuarielles : tables de mortalité, de maintien en incapacité temporaire, taux d'actualisation, etc.

La plupart de ces hypothèses sont volontairement prudentes de manière à obtenir des niveaux de provisions techniques relativement prudents. Cela conduit en particulier à l'introduction de marges implicites de prudence. Elles sont implicites dans la mesure où elles résultent de l'application d'hypothèses

prudentes. Ainsi elles ne sont pas quantifiées, par exemple, par analyse de l'écart sur les provisions par rapport à des hypothèses plus réalistes.

Les nouveaux référentiels économiques, comptables et prudentiels exigent pour leur part de mesurer explicitement le risque et de fixer la prudence en référence à ce risque. C'est notamment l'objet de la marge pour risque.

Dans ce contexte, les assureurs doivent utiliser la meilleure information possible pour évaluer leurs engagements. En particulier, les hypothèses dites actuarielles doivent, lorsque c'est possible, être fixées en référence au marché ou à l'expérience de l'entreprise.

2.1.1. Références internes et externes

Concernant les hypothèses, une des principales distinctions que l'on peut faire entre les normes IFRS d'une part et les autres référentiels d'autre part provient du calibrage de ces hypothèses sur des données internes ou externes.

En effet, de par la définition de la *Current Exit Value*, il s'agit, dans le référentiel IFRS, de déterminer une valeur de transfert du portefeuille sur laquelle s'entendraient deux assureurs. Sur ces bases, l'IASB précise que les paramètres de projection doivent être ceux du marché et non ceux de l'entreprise qui détient le portefeuille en question (i.e. celle dont on établit les comptes). Par exemple, si l'assureur A a des frais de gestion de ses actifs de l'ordre de 0,5 % alors que, pour le reste du marché, ils ne sont que de l'ordre de 0,3 %, le portefeuille doit être valorisé sur les bases des paramètres de marché puisque le 0,5 % n'est propre qu'à l'assureur A et que le transfert du portefeuille d'implique pas le transfert de ces frais de gestion des actifs alourdis.

Si ce principe semble cohérent avec le modèle d'évaluation retenu (la *Current Exit Value*), il n'en demeure pas moins que ces « paramètres de marché » sont rarement observables, ce qui pose une difficulté supplémentaire à l'utilisation de ces références externes.

Pour des hypothèses telles que la mortalité en revanche, c'est bien l'expérience issue du portefeuille qui doit être utilisée dans la mesure où le transfert de celui-ci ne va pas conduire à modifier la mortalité de ses assurés.

Cette question des références externes ou internes à l'entreprise dans le choix des paramètres de projection amène également celle des règles de gestion future. En effet, il n'est pas possible de projeter les flux futurs d'une société d'assurance sans modéliser les règles de gestion futures en matière d'allocation d'actifs, de revalorisation de contrats, de gestion des provisions pour participation aux bénéfices, etc. Or chaque entreprise a ses propres règles et la valeur du portefeuille peut significativement varier selon les règles utilisées. Ce point est plus particulièrement développé dans le sous-paragraphe 2.2 *infra*.

2.1.2. Le cas de la mortalité

Que ce soit dans le cadre de l'évaluation des engagements d'un régime de retraite, la valorisation d'un portefeuille d'épargne ou encore le suivi technique de contrats en cas de décès, la mortalité constitue un paramètre déterminant du

résultat des valorisations. La question se pose donc du choix pertinent de l'hypothèse à retenir.

D'une manière générale, les normes comptables IFRS conduisent à privilégier, au travers d'une approche « économique » de la valorisation de l'entreprise, le choix d'hypothèses « réalistes », tenant compte de l'expérience du portefeuille (IFRS 4 et le projet de norme assurance en phase II) ou de l'entreprise (IAS 19, cf. Planchet et Thérond (2004b)). Dans le cas particulier de la mortalité, cela conduit naturellement à vouloir se tourner vers des « tables d'expérience » en lieu et place de références exogènes (tables nationales, tables réglementaires dans le cas des assureurs, etc.) pas nécessairement en phase avec la réalité du risque.

Mais, dès lors que le risque est viager, et notamment dans les problématiques de rentes, la table de mortalité utilisée se doit d'être prospective afin de prendre en compte l'évolution future des taux de décès, ce qui induit de fortes contraintes en terme de volume de données si l'on souhaite construire une « surface de mortalité » spécifique de la population concernée anticipant correctement les évolutions à venir.

Ainsi, lors de la construction des nouvelles tables de mortalité réglementaires utilisées par les assureurs pour le provisionnement de leur engagements viagers, les quelques centaines de milliers de têtes observées sur une dizaine d'années n'ont pas suffi à construire une table autonome (ou « endogène ») et il est apparu nécessaire de s'appuyer sur des tables construites préalablement sur l'ensemble de la population française pour dégager des tendances de long terme.

En pratique, la taille des groupes existants et le nombre d'années disponibles sont souvent insuffisants pour espérer réaliser une construction robuste d'une telle table « endogène », obtenue uniquement à partir des observations issues du groupe considéré. Ceci est valable tant sur des portefeuilles de rentiers que dans le cas d'entreprises.

Il n'en demeure pas moins qu'il n'apparaît pas efficient de ne pas tenir compte de l'information apportée par les observations. En effet, les déterminants de la mortalité sont nombreux : le sexe, le mode de vie, le niveau de revenu, la région d'habitation figurent parmi les plus importants. Une population particulière donnée (portefeuille d'assureur ou rentiers futurs et en cours dans le cas d'un régime supplémentaire d'entreprise) présente donc *a priori* une mortalité différente de celle décrite par des références nationales ou de place.

Dans ce contexte, dès lors que l'on abandonne une évaluation prudentielle des engagements au profit d'une évaluation « réaliste » avec une quantification qui se veut explicite de la marge de risque, il est indispensable de prendre en compte l'information apportée par les données disponibles, sous peine de déformer « l'image économique fidèle » de l'entreprise que l'on entend donner.

Se pose toutefois la question du moyen d'y parvenir, compte tenu des difficultés techniques exposées précédemment.

Un examen plus attentif de la structure d'une table de mortalité prospective conduit alors à distinguer le niveau de la mortalité (âge par âge) d'une part et

son évolution anticipée dans le futur d'autre part. Déterminer le niveau de la mortalité à un moment donné revient à fixer une table du moment, proposer une tendance pour l'évolution future conduit à lui ajouter une dimension prospective pour aboutir à une surface de mortalité. On parle alors de table prospective.

Les données disponibles fournissent en général une détermination suffisamment riche pour la construction d'une table du moment, quitte à procéder à des extrapolations en dehors de la plage où des observations sont disponibles. Ces extrapolations sont aisées dans le cadre des modèles paramétriques de mortalité, comme par exemple le modèle classique de Makeham ou les modèles de régression logistique (cf. Planchet et Thérond (2006)).

C'est au moment de déterminer la tendance d'évolution des taux de mortalité dans le futur que les données s'avèrent insuffisantes en pratique : historiques trop faibles (souvent moins de 10 ans) et effectifs insuffisants (quelques dizaines de milliers) pour supprimer le bruit issu des fluctuations d'échantillonnage rendent la démarche prospective délicate à appliquer directement, sauf à prendre des risques importants sur l'appréciation de la tendance. Il existe par exemple des versions paramétriques du célèbre modèle de Lee-Carter qui permettent d'estimer une surface de mortalité à partir d'un nombre réduit de paramètres : la diminution du nombre de paramètres permet une estimation plus fiable et augmente le pouvoir prédictif du modèle, mais elle augmente en parallèle le risque d'inadéquation du modèle à la réalité.

Toutefois, on peut noter à ce stade que la question n'est pas tant fondamentalement ici de construire une table de mortalité que d'être en mesure de sélectionner une hypothèse bien adaptée au groupe pour le risque considéré, et à tout le moins mieux adaptée qu'une référence réglementaire utilisée de manière arbitraire. En effet, si l'objectif de construire une table spécifique pour la population étudiée s'avère délicat à atteindre, une démarche pragmatique consiste à se tourner vers les différentes références existantes et à rechercher celle qui, parmi elle, représente le mieux le comportement de la population en terme de mortalité. Les outils statistiques d'analyse de l'adéquation d'une loi de mortalité donnée a priori à des observations issues de l'expérience sont classiques et peuvent ici être employés avec succès.

Ainsi, il est donc important de noter qu'en pratique l'analyse de la mortalité d'expérience du groupe peut conduire à retenir comme « table d'expérience » une table exogène au groupe, par exemple la TPG 1993, une table INSEE, ou plus généralement toute table dont on aurait des raisons de penser qu'elle peut raisonnablement représenter la mortalité de la population considérée ; d'une manière générale, la justification de l'adéquation d'une table à un groupe donné pour un risque donné (vie ou décès) relève d'une analyse différente de la construction d'une table de mortalité propre au groupe. Et cette analyse est plus robuste et plus simple à mettre en œuvre que la construction proprement dite.

Face à ces contraintes, l'obtention d'une information raisonnablement fiable sur la mortalité d'expérience du groupe passe donc in fine par le positionnement de cette mortalité par rapport à une référence, que celle-ci soit une table nationale, une table de place ou une table d'expérience construite sur

une population plus large présentant des similitudes de comportement sur ce registre.

Une table de mortalité est, on l'a vu, essentiellement décrite par le niveau de la mortalité à un moment donné et la tendance future d'évolution des taux de décès (ou de tout autre grandeur décrivant la survie, comme par exemple l'indicateur classique d'espérance de vie). Cette structure conduit à observer que le positionnement par rapport à une référence peut, en fonction de la qualité des données à disposition, s'envisager de deux manières :

- un positionnement sur le niveau uniquement, la tendance reprenant strictement la tendance de la référence proposée ;
- un positionnement conjoint sur le niveau et la tendance. Plus délicat à mettre en oeuvre techniquement, il permet une appréciation plus fine du risque de longévité lorsque le volume de données est suffisant.

Pour l'une ou l'autre de ces deux approches, les outils techniques existent : on peut citer notamment les régressions logistiques ou le modèle de Cox et ses dérivés.

On peut d'ailleurs noter que les nouvelles tables réglementaires TGH et TGF 05 légitiment d'une certaine manière cette démarche. En effet, la relative petite taille des portefeuilles utilisés pour la construction a nécessité comme on l'a rappelé infra le calage des tendances sur des tables prospectives préalablement construites pour l'occasion sur la base de données INSEE. Ainsi, les tables réglementaires elles-mêmes sont en un certain sens le résultat d'un positionnement de la mortalité des assurés par rapport à la mortalité générale.

Au global si la construction d'une table d'expérience propre au groupe et élaborée à partir de ses seules données n'est en pratique pas une solution systématiquement envisageable, les techniques de positionnement de la mortalité d'expérience par rapport à une référence externe fournissent une palette d'outils opérationnels conduisant à une appréciation plus réaliste du risque porté.

C'est dans ce contexte que s'inscrivent les travaux de Hardy et Panjer (1998). Ceux-ci consistent à mettre en place un modèle de crédibilité² qui offre aux différentes entreprises du marché canadien de calculer, à partir de leurs propres données, un taux d'abattement relatif à une table de mortalité standard. Grâce à ce modèle, qui repose sur une variante du modèle de Bühlmann et Straub (1970), chaque entreprise peut donc calculer son propre taux d'abattement. Incidemment, ce type de modèle permet aux entreprises qui ne disposent pas de données suffisantes pour établir une table d'expérience uniquement à partir de celles-ci, d'utiliser des hypothèses de mortalité plus adaptées à leur propre risque que la table standard.

Ce genre de démarche est typiquement ce qui peut être mise en œuvre par des associations d'assureurs de manière à disposer d'un maximum d'informations pour apprécier la structure du risque interne aux assureurs et sa variabilité entre les assureurs. Incidemment, en se regroupant pour mener de telles études, des assureurs avec des portefeuilles modestes peuvent déduire des

². Cf. l'ouvrage de référence de Bühlmann et Gisler (2005) pour plus de détail sur les modèles de crédibilité en général et celui dit de Bühlmann-Straub en particulier.

lois d'expérience qui leur sont propres, ce qu'ils n'auraient pu faire sur la base de leurs seules données.

2.2. Futures décisions de gestion

Les calculs prospectifs préconisés par Solvabilité 2, les principes de l'EEV ou encore les normes IFRS ne peuvent faire l'économie de la modélisation des décisions de gestion future. Il s'agit de se donner un modèle qui permettra de simuler de manière automatique le comportement des agents économiques (ici l'assureur et les assurés) et de l'adapter aux conditions simulées (état de l'assuré, conditions de marché, etc.)

Dans la suite de ce paragraphe, nous dressons la liste des comportements à modéliser pour les assurés comme pour l'assureur et analysons ce que prévoient les différents référentiels à leur sujet.

2.2.1. Comportement des assurés

Les principales options qui sont offertes à un assuré pendant la vie de son contrat relèvent :

- de la gestion du contrat : rachats partiels ou total, transfert, arbitrages entre des supports en unités de compte en euros pour un contrat d'épargne multi-support ;
- du paiement (ou non) des primes futures : versements libres, versements programmés et primes périodiques ;
- des options qui lui sont offertes : conversion de l'épargne en rente, prorogation ou modification des garanties du contrat dans des conditions favorables, etc.

Bien qu'ils fassent partie de la vie effective des contrats, l'analyse de ces comportements n'est pas aisée. En effet, si la plupart des assureurs ont construit des lois d'expérience sur le rachat de contrats, celles-ci sont le plus souvent segmentées en fonction de l'ancienneté du contrat (cruciale en raison des dates fiscales des quatrième et huitième anniversaires) sans distinction de l'âge de l'assuré ou plus encore du moment auquel est intervenu le rachat et donc des conditions de marché correspondantes.

La transposition de ces lois dans un contexte de projection stochastique des actifs notamment est encore plus ardue. Pour le rachat par exemple, il s'agit, le plus souvent, de déformer les lois standard obtenues sans distinction des conditions de marché pour les alourdir ou au contraire les alléger en fonction de l'évolution des taux d'intérêt par exemple. Cet exercice se heurte à un écueil qu'il est difficile de contourner autrement que par une décision arbitraire : dans l'économie actuelle, avec des taux d'intérêts à deux chiffres, quel serait le comportement des assurés bénéficiant de contrats à taux garantis de l'ordre de 2,5 % ? Ce type de scénario peut se produire dans des projections stochastiques et les lois de comportement *ad hoc* n'ont pas pu être calibrées pour celui-ci puisqu'il n'existe pas de données statistiques.

De manière générale, sur la modélisation du comportement des assurés, les référentiels peuvent diverger. C'est notamment le cas sur les primes futures.

Le projet de norme assurance IFRS phase II prévoit des critères relativement stricts pour inclure les primes futures dans l'évaluation des provisions des contrats d'assurance. Ainsi, le Board prévoit que les flux projetés doivent tenir compte des primes futures :

- si l'assureur dispose d'un droit contractuel lui permettant d'exiger leur paiement ; ou
- si l'assureur ne peut refuser de percevoir ces primes et que la valeur actuelle de ces primes est inférieure à la valeur actuelle des prestations additionnelles qui en résultent ; ou
- si le paiement de ces primes permet à l'assuré de conserver son assurabilité.

Ainsi des primes futures probables mais qui ne correspondraient pas à l'un de ces trois critères ne devront pas être prises en compte pour l'évaluation de la provision IFRS alors qu'elles devront l'être pour un calcul de MCEV par exemple (cf. principe n° 6 de l'EEV Principes).

2.2.2. Comportement de l'assureur

La modélisation du comportement de l'assureur est essentielle puisqu'elle va avoir de lourds impacts sur l'évolution des contrats. C'est surtout le cas pour les contrats d'épargne en assurance vie dont l'évolution dépend grandement de :

- la gestion du portefeuille financier : réalisation des plus ou moins-values latentes qui participe à la constitution du taux de rendement comptable sur lequel est déterminée la participation aux bénéfices réglementaire et, le plus souvent, contractuelle également, recomposition de la composition du portefeuille, gestion de la réserve de capitalisation, etc.
- la revalorisation des contrats : incorporation, dotation et reprise à la PPB.

Il s'agit donc, pour l'assureur, de mettre en évidence les décisions de gestion qu'il prendrait dans différentes conditions de marché de manière à établir une règle générique qui s'adaptent à celles-ci. Les résultats des évaluations seront dépendants de ces règles. C'est d'ailleurs une des difficultés auxquelles se trouve confronté le projet de norme IFRS phase 2 sur les contrats d'assurance : la méthode retenue pour valoriser un engagement d'assurance est celle de la *Current Exit Value* (CEV) ou valeur de sortie du portefeuille. Dès lors, on peut se demander si cette valeur de sortie devrait dépendre des règles de gestion de l'assureur dans la mesure où le transfert effectif du portefeuille vers un autre assureur conduirait à ce que ce soit cet autre assureur qui le gère avec ses propres règles de gestion, pas nécessairement identiques à celle du premier (et même, le plus souvent, différentes). Il n'en demeure pas moins qu'il n'est pas possible de se passer de la modélisation de ces règles pour projeter les flux futurs. De plus ces comportements sont étroitement liés³ : une politique de

³. Un modèle de rachat dépendant des taux servis et des taux de marché est présenté dans le dernier paragraphe du Chapitre 5.

revalorisation des contrats qui se contente des taux minimums garantis conduira à des rachats plus massifs qu'une politique d'attribution de PB plus en phase avec le marché.

2.2.3. Exemple : gestion du portefeuille d'actifs

Pour illustrer l'impact de la modélisation des décisions de gestion sur les valeurs de portefeuille, nous avons modélisé deux politiques de gestion d'actifs à partir du portefeuille d'actif à la date d'évaluation. Nous avons observé les valeurs obtenues pour un contrat de type épargne qui comporte une garantie de taux et une clause de participation aux bénéfices.

De manière à isoler le phénomène que l'on souhaite observer, nous avons fait abstraction des risques biologiques tels que le décès ou l'arrêt de travail et des risques liés au comportement des assurés (rachats, prorogation, etc.), ainsi l'assuré bénéficie de son épargne au terme prévu contractuellement.

2.2.3.1 Présentation du produit

Considérons un contrat d'épargne de durée 8 ans qui comporte une garantie de taux annuelle⁴ de 2,5 % et une clause de participation bénéficiaire de 90 %. Si l'on note PM_k , la valeur de l'épargne à la date anniversaire k , le processus de revalorisation, à chaque date anniversaire est le suivant :

$$PM_{k+1} = \left(1 + \max\{2,5\%; 90\% \times \rho_{k+1}\}\right) PM_k,$$

où ρ_k désigne le rendement du portefeuille financier sur la k -ème période.

Ainsi on suppose que l'assureur n'attribue pas de participation aux bénéfices au-delà de son engagement contractuel. Dans un cas réel, cette hypothèse doit être levée du fait de la prise en compte des rachats (cf. notamment le modèle de rachat présenté en fin du Chapitre 5 et le paragraphe 3).

2.2.3.2 Modélisation du rendement financier

Dans la suite nous supposons que l'assureur dispose aujourd'hui d'un actif composé d'actions dans une proportion $\theta = 30\%$ et d'un bon de capitalisation pour les 70 % restants. De plus, nous supposons que le bon capitalise à un taux d'intérêt certain et constant au cours des huit années du contrat égal à 4 %. Par ailleurs, le cours de la poche action est modélisé par un mouvement brownien géométrique avec une volatilité de 20 %.

Étudions à présent les deux stratégies envisagées.

Dans la première, l'assureur ne touche plus à ses actifs jusqu'au terme du contrat. Dans ce cas de figure, le rendement (discret) pour la k -ème période est donné par l'expression suivante :

⁴. Il s'agit d'une garantie cliquet comme présente dans les contrats français : chaque année les 2,5 % sont acquis. Pour mémoire, les contrats anglais comportent plus souvent des *terminal bonus* pour lesquels la revalorisation n'intervient effectivement qu'au terme.

$$\rho(k) = \frac{\theta S_k + (1-\theta) B_k}{\theta S_{k-1} + (1-\theta) B_{k-1}} - 1,$$

où S_k et B_k sont respectivement les prix des actions et du bon de capitalisation à la date k . On supposera, sans perte de généralité, que $S_0 = B_0 = 1$.

La deuxième stratégie consiste à recomposer, chaque début d'année, le portefeuille financier de manière à ce que la poche action représente 30 % de la valeur totale du portefeuille. Dans ce cas, le rendement (discret) pour la k -ème période est donné par l'expression suivante :

$$\rho(k) = \theta \frac{S_k}{S_{k-1}} + (1-\theta) \frac{B_k}{B_{k-1}} - 1.$$

Le risque que l'on cherche à analyser étant purement financier, conformément aux principes évoqués précédemment dans le sous-paragraphe 1.3, ces rendements ont été modélisés dans l'univers risque-neutre. Sous la mesure risque-neutre Q , le cours des actions évolue selon le processus suivant :

$$\frac{dS_t}{S_t} = rdt + \sigma d\hat{B}_t,$$

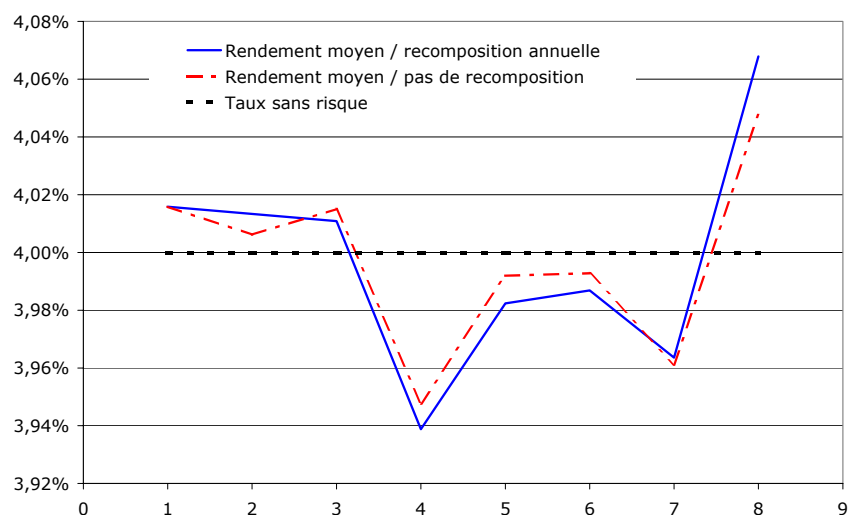
où $\hat{B}$ est un mouvement brownien sous Q et $r = \ln(B_{k+1}/B_k) = \ln(1 + 4\%)$ est le taux sans risque instantané.

Dans l'univers risque-neutre, les deux processus de rendement ont, sur une période, une même espérance égale au rendement sans risque.

Les distributions des rendements ont pu être obtenues grâce aux méthodes de Monte Carlo⁵.

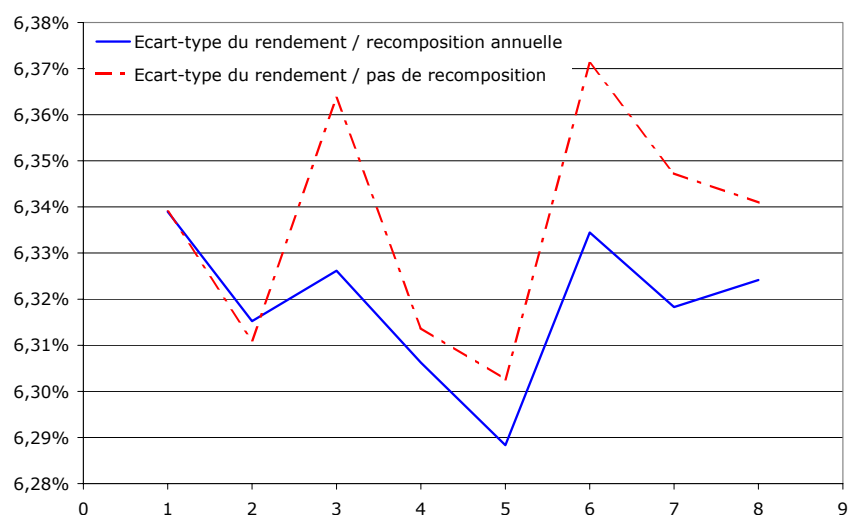
Comme le montrent les graphiques suivants, les deux stratégies mènent à des résultats sensiblement différents. La Figure 6 représente l'évolution du rendement moyen au cours des années selon la règle de gestion utilisée et le positionne par rapport au rendement sans risque.

⁵. Les résultats ont été obtenus à partir de 30 000 simulations. Le mouvement brownien géométrique a fait l'objet d'une discrétisation exacte et les variables aléatoires gaussiennes ont été simulées à partir de réalisations uniformes simulées par l'algorithme du tore mélangé et transformées par l'approximation de Moro. Cf. le Chapitre 6 pour le détail de ces techniques.

Figure 6 - Évolution du rendement moyen de l'actif

Malgré le nombre important de simulations effectuées, il réside un biais qui atteint un maximum de 1,7 % pour le rendement moyen de la huitième période. Néanmoins, le rendement moyen de ces deux stratégies est comparable, ce qui est normal puisqu'elles ont le même rendement espéré.

La Figure 7 représente l'évolution de l'écart-type du rendement au cours des huit périodes dans les deux stratégies.

Figure 7 - Évolution de l'écart-type du rendement

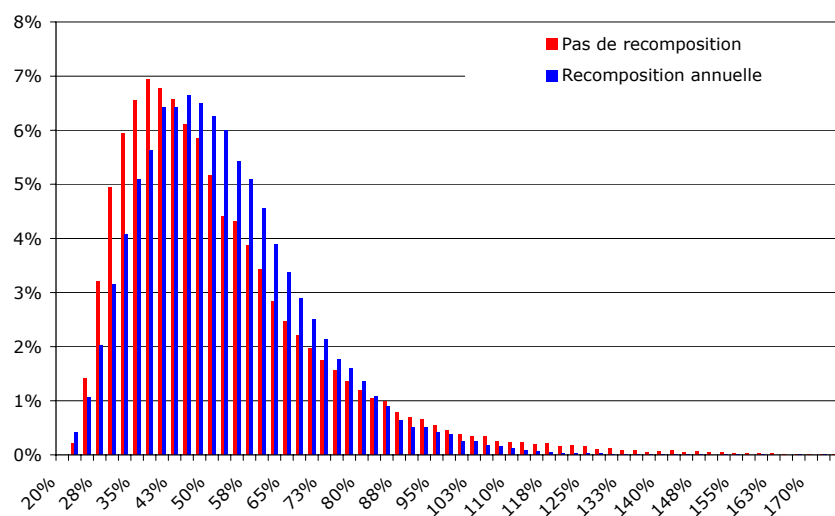
La stratégie de reconstitution périodique du portefeuille conduit à des rendements moins variables que celle qui consiste à acheter et ne plus toucher au portefeuille financier jusqu'au terme.

Nous allons voir dans le paragraphe suivant les incidences de ces différences sur le contrat d'assurance vie considéré.

2.2.3.3 Résultats

La Figure 8 représente la distribution de la revalorisation de l'épargne au terme pour l'assuré selon les deux règles de gestion des actifs.

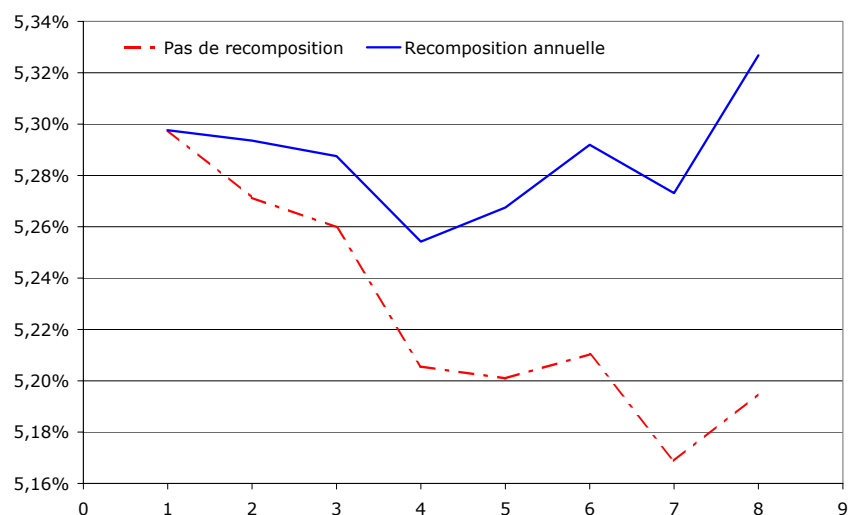
Figure 8 - Distribution empirique de la revalorisation l'épargne au terme



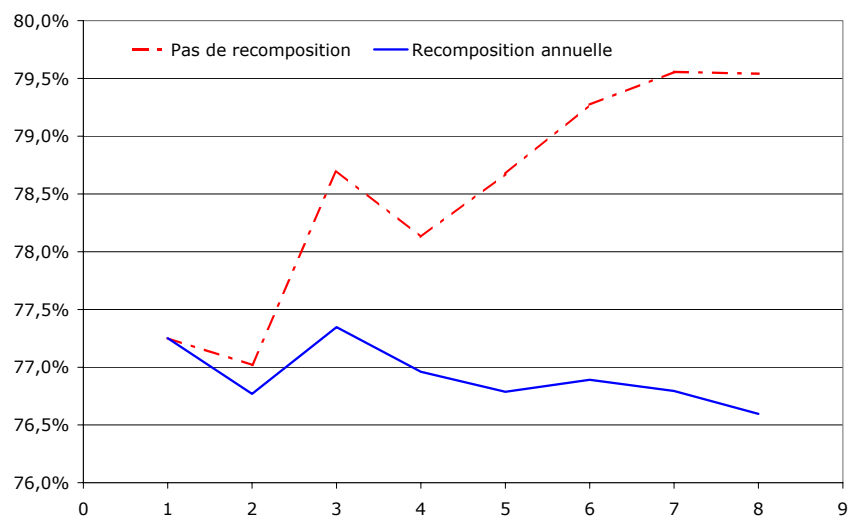
On peut remarquer que la politique de reconstitution annuelle du portefeuille conduit à une distribution plus resserrée autour de la valeur moyenne. À l'inverse, l'autre stratégie a une queue de distribution plus lourde vers les montants d'épargne élevés mais son cœur se situe davantage en deçà de la valeur moyenne. Si les deux distributions sont différentes, l'espérance de la revalorisation au terme est, dans les deux cas, de l'ordre de 51 %, ce qui représente un gain de 10 % sur 8 ans par rapport au sans risque. Incidemment, si l'assureur choisit de reconstituer son portefeuille en tout début de contrat pour n'investir que sur les bons de capitalisation, la revalorisation au terme (certaine dans ce cas précis) n'atteint que 32,7 %. La gestion du portefeuille financier a donc une incidence importante sur la valorisation du contrat.

Si les valeurs de l'épargne à chaque date d'inventaire ont des moyennes très proches dans les deux politiques de gestion envisagées, au niveau des rendements actuels, les profils de rendement sont différents. La Figure 9 reprend l'évolution du taux de revalorisation moyen de l'épargne au cours de la vie du contrat dans les deux situations étudiées tandis que la Figure 10 présente son coefficient de variation⁶.

⁶. Le coefficient de variation d'une variable aléatoire est le rapport de son écart-type sur son espérance (lorsque les deux premiers moments existent). C'est un indicateur de variabilité couramment utilisé pour comparer des risques qui n'ont pas la même taille (espérance).

Figure 9 - Évolution du taux de revalorisation moyen

Le rendement moyen avec reconstitution annuelle s'avère supérieur à celui sans. Cependant, sa variabilité étant plus faible (cf. la Figure 10), les deux stratégies aboutissent à des montants espérés d'épargne au terme très proches.

Figure 10 - Évolution du coefficient de variation du taux de revalorisation

Conformément, à ce que l'on avait déjà observé au sujet des rendements du portefeuille d'actifs, la stratégie de reconstitution annuelle conduit des revalorisations de l'épargne moins variable que l'autre stratégie étudiée.

3. Cas pratique : portefeuille d'assurance vie

Dans ce paragraphe, nous allons voir comment les principes précédemment présentés s'appliquent dans le cas de la valorisation d'un produit d'assurance vie individuel (contrat d'épargne en euros) commercialisé par une société d'assurance.

Dans ce paragraphe, sauf mention contraire, toutes les références juridiques sont relatives au code des assurances.

3.1. Présentation du produit

Le produit modélisé est un contrat d'épargne en euros commercialisé par une société d'assurance qui prévoit, contre le paiement à la souscription d'une prime unique, le versement d'un capital :

- à l'assuré en cas de survie au terme ;
- à ses ayants-droits en cas de décès de l'assuré avant le terme.

Le montant du capital versé en cas de décès ou au terme est garanti à la souscription. Il n'est pas possible d'effectuer de versements libres sur ce contrat. En revanche, l'assuré peut, à tout moment, demander le rachat de son contrat.

3.1.1. Caractéristiques techniques

L'objet de ce sous-paragraphe est de détailler les caractéristiques techniques de ce contrat afin d'identifier les risques qu'il fait courir à l'assureur et de déterminer la modélisation qui peut en être faite.

3.1.1.1 Caractéristiques générales

La durée du produit est fixe et de 8 ans. Le contrat peut cependant être prorogé par tacite reconduction à un taux minimum garanti de 60 % du TME.

La prime nette initialement versée est investie dans l'actif général de l'assureur. Il n'y a pas de canton spécifique pour ce produit.

3.1.1.2 Revalorisation de l'épargne

Ce type de contrat est principalement souscrit par les assurés en tant qu'outil de placement permettant de bénéficier des avantages fiscaux que confère la réglementation des contrats d'assurance.

De ce fait, sa principale raison d'être de ce produit est de procurer un rendement financier à l'assuré. Ainsi l'épargne constituée par la prime unique (diminuée des frais sur versement) est revalorisée chaque année civile. Le processus de revalorisation de ces contrats est complexe. En effet, il dépend de la politique de gestion de l'assureur sous les contraintes contractuelles (taux minimum garanti), réglementaire (participation aux bénéfices) et économiques (taux affichés par la concurrence).

Contractuellement, chaque police se voit attribuer à sa souscription un taux minimum garanti (TMG) qui dépend, pour ce contrat, des conditions de marché à la souscription (référence au Taux Moyen des Emprunts de l'État Français, le

TME). Ce taux est un plancher pour le taux de revalorisation effectif de l'épargne.

Par ailleurs, le Code des assurances prévoit qu'au moins 90 % des bénéfices techniques et 85 % des bénéfices financiers doivent être redistribués aux assurés. En particulier, le taux de rendement des actifs est calculé selon les références françaises : le rendement comptable n'intègre pas les éventuelles plus-values latentes. Les éventuelles moins-values latentes sont quant à elles, au moins en partie, constatée via l'inscription dans les comptes d'éventuelles provisions pour dépréciation durable⁷ (PDD) (provision ligne à ligne) ou provision pour risque d'exigibilité⁸ (PRE).

Enfin, si l'assureur veut éviter des vagues massives de rachats qui peuvent lui poser des problèmes de liquidité de ses actifs et ne souhaite pas donner de signe négatif aux éventuels futurs clients, sa politique de revalorisation ne peut faire l'économie de la prise en compte des revalorisations accordées par ses concurrents. Cette nécessité de communication sur les taux de revalorisation s'observe également sur les taux futurs au travers d'offres à fenêtre par exemple. En effet, l'article A132-3 autorise de fixer et de communiquer un taux minimum de revalorisation pour l'année suivante pouvant s'élever jusqu'à 85 % de la moyenne sur deux ans des taux de rendement de l'actif constatés.

La revalorisation des contrats intervient au 31/12 de chaque année. Les contrats de l'année sont revalorisés au pro-rata de leur présence au cours de l'année.

Concernant les contrats sortis au cours d'une année. Compte-tenu du fait que le taux de revalorisation n'est calculé qu'en fin d'année, ils sont revalorisés à un taux fixé par l'assureur (supérieur à leur TMG) au prorata de leur présence sur l'année.

3.1.1.3 Rachats

Le contrat étudié prévoit deux types de rachat :

- Le rachat total qui permet à l'assuré de se voir verser une somme égale à sa provision mathématique (PM) diminuée d'une indemnité de 0,3 % par semestre entier restant à courir jusqu'au huitième anniversaire du contrat. Cette clause respecte les dispositions du code des assurances qui prévoient, qu'en cas de rachat, la pénalité demandée par l'assureur ne peut être supérieure à 5 % de la PM.
- Les retraits, possibles à partir du premier anniversaire, qui permettent à l'assuré de retirer une partie de son épargne avec les mêmes

⁷. Lorsqu'un titre est durablement en moins-value latente, il est prévu que les compagnies d'assurance inscrivent, dans leurs comptes sociaux, une PDD pour constater comptablement cette situation. Il s'agit donc d'une provision ligne à ligne dont les modalités pratiques sont notamment régies par l'avis n° 2002-F du 18 décembre 2002 du Conseil National de la Comptabilité (CNC).

⁸. La PRE est une provision globale qui est constituée lorsque une moins-value latente nette est constatée sur l'ensemble des placements R332-20 (titres non amortissables). Cf. art. R331-5-1.

pénalités qu'en cas de rachat total. Un retrait ne peut conduire à ce que l'épargne restante soit inférieure à un certain seuil.

3.1.1.4 Frais

Ce contrat prévoit deux types de frais :

- des prélèvements sur encours qui sont effectués chaque fin d'année au moment de la revalorisation ;
- des frais sur les versements effectués lors du paiement de la prime unique à la souscription.

3.1.2. Identification des risques

L'objet de ce sous-paragraphe est de cartographier l'ensemble des risques qui influent sur le résultat engendré par ce contrat.

De manière à identifier quels sont les conséquences de tel ou tel risque, nous allons les comparer à ce que nous appellerons dans la suite la situation de référence.

3.1.2.1 Situation de référence

La situation de référence correspond au cas le plus fréquent : l'assuré souscrit son contrat et le mène à son terme (ne sort ni pour raison de décès, ni pour cause de rachat total) dans son intégralité (pas de retrait). Dans le même temps, l'assureur réalise chaque année un taux de rendement financier comparable à celui du marché qui lui permet de revaloriser les contrats à des taux du même ordre de grandeur que ceux de la concurrence.

Cette situation, qui peut paraître idéale, est la plus probable comme on peut le constater. L'hypothèse sur le décès est la plus vraisemblable compte tenu de l'âge moyen des assurés (autour de 50 ans) et de la durée du contrat (8 ans).

Concernant le rachat, qui est étroitement lié aux taux effectifs de revalorisation, il est relativement faible notamment grâce au fait que la durée du contrat (8 ans) est collée sur une durée fiscale et que les taux de rendement de la société sont comparables à ceux observés sur le marché (la composition de l'actif n'est pas atypique et le portefeuille financier recèle des plus-values latentes pour amortir d'éventuels chocs adverses).

Par rapport à ce scénario de référence, plusieurs alternatives peuvent survenir :

- décès de l'assuré avant le terme ;
- rachat partiel ou total du contrat ;
- sous-performance financière.

Ces différentes situations sont détaillées ci-après.

3.1.2.2 Décès de l'assuré

En cas de décès de l'assuré avant le terme du contrat, l'assureur verse aux ayants-droits l'épargne acquise.

Par rapport à la situation standard présentée *supra*, il prélèvera pas les frais sur encours futurs. Dans le même temps, il n'aura plus à supporter le risque, essentiellement financier, associé au contrat.

3.1.2.3 Rachat du contrat

La situation du rachat total est semblable, du point de vue de l'assureur, à celle du décès de l'assuré. La seule différence réside dans la pénalité que va pouvoir conserver l'assureur.

3.1.2.4 Performance financière

Comme nous le verrons par la suite, la situation la plus défavorable pour l'assureur est celle où il réalise des contre-performances financières.

Contractuellement, s'il ne réalise pas le TMG, il va devoir constater des pertes pour revaloriser les contrats et ainsi puiser dans ses fonds propres. Incidemment cela conduira à détériorer sa marge de solvabilité.

De plus, même si le rendement financier s'avère supérieur au TMG mais ressort inférieur aux attentes du marché (et des assurés), l'assureur peut choisir de réaliser des pertes pour afficher des taux de revalorisation en ligne avec la concurrence de manière à endiguer une éventuelle vague de rachats massive et à ne pas décourager d'éventuels futurs clients.

3.2. Le portefeuille des assurés

Le portefeuille est constitué de 160 assurés d'âge moyen 55 ans. A la date d'évaluation (31/12/2006), l'épargne moyenne accumulée est de l'ordre de 14 k€. S'agissant de la production de l'année 2006, le terme résiduel moyen est de l'ordre de 7, 5 ans.

3.3. Modélisation retenue

De manière à être aussi exhaustif que possible avec les différents risques en jeu, nous avons choisi de modéliser les risques engendrés

3.3.1. Mortalité des assurés

Le risque de mortalité des assurés a été modélisé grâce aux techniques de simulation à partir d'une loi de mortalité nationale.

Compte tenu de la taille du portefeuille, le risque systématique de mortalité (cf. Planchet et Thérond (2007a)) n'a pas été modélisé du fait :

- de la difficulté de sa modélisation sur un portefeuille de cette taille,
- de son faible impact sur une durée si courte (inférieure à 8 ans en moyenne) et sur un portefeuille de cette taille.

3.3.2. Rendement financier

C'est le rendement comptable de l'actif général qui a été modélisé directement. Il est supposé corrélé avec un rendement « de marché » de l'actif des autres sociétés d'assurance.

Ces deux rendements sont modélisés dans l'univers risque-neutre par des mouvements browniens arithmétiques d'espérance le taux sans risque supposé constant sur l'horizon de projection.

3.3.3. Rachat de contrats

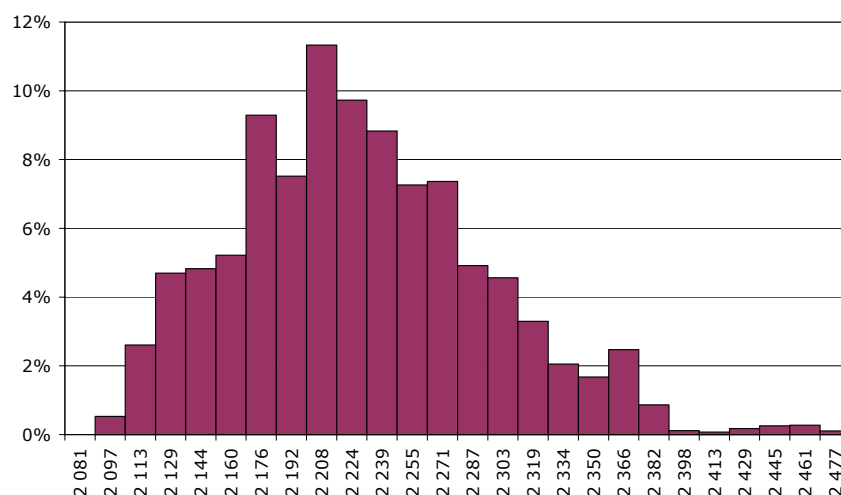
On dispose d'une loi centrale de rachat fonction de l'ancienneté du contrat. Cette loi est déformée en fonction des revalorisations accordées à l'épargne des assurés : si cette revalorisation s'éloigne trop de la revalorisation moyenne accordée par le marché (déduite du rendement de l'actif général du marché), la probabilité de rachat augmente proportionnellement à l'écart de revalorisation constaté.

3.4. Analyse des résultats

Chaque simulation a donné lieu à un engagement obtenu comme étant la valeur des prestations futures actualisées au taux sans risque.

L'ensemble de ces valeurs permet de tracer une distribution empirique de la valeur actuelle des prestations futures.

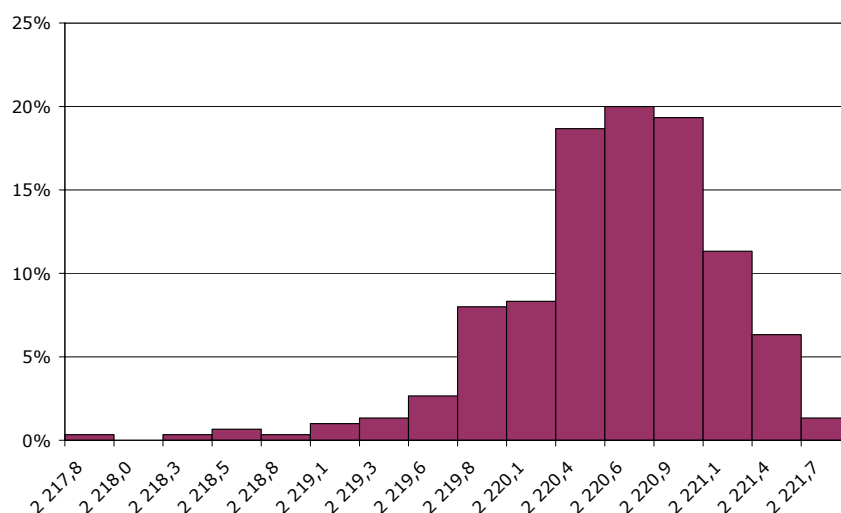
Figure 11 - Densité empirique de la valeur actuelle des prestations futures



Si l'on mesure le risque associé à cette distribution empirique grâce à la variance (cf. le Chapitre 1 pour une présentation des différentes mesures de risque usuelles en assurance), une décomposition de variance nous permet de décomposer le risque global entre risque financier et risque non-financier. Pour ce contrat et ce portefeuille, le risque financier représente 99 % du risque global. Comme les référentiels IFRS, Solvabilité 2 et d'EEV prévoient de valoriser les risques financiers à leur prix de marché ou de manière équivalente au prix de leur portefeuille de réplcation, il est naturel de s'intéresser à la distribution de la valeur actuelle des prestations futures lorsque le risque financier est couvert. Il s'agit en fait de reproduire, dans notre situation, la démarche utilisée par Frantz, Chenut et Walhin (2003) dans le cas des contrats

en unités de compte (cf. le paragraphe 3 du Chapitre 1). La distribution ainsi obtenue est représentée par sa densité empirique dans la Figure 12

Figure 12 - Densité empirique des prestations futures actualisées



Comme annoncé précédemment, on remarque que, même sur un portefeuille de taille aussi restreinte, le risque non-financier, et en particulier le risque de mortalité et la partie structurelle du risque de rachat, ont un très faible impact sur la variabilité de l'engagement.

3.5. Conclusion

Compte tenu de la part non-significative de la part du risque non-financier dans l'engagement de l'assureur, on peut considérer que la valorisation de ce type de contrat peut s'épargner la simulation des fluctuations d'échantillonnage de risques tels que la mortalité. Ces risques doivent néanmoins être pris en compte, en moyenne, via leur probabilité de survenance.

De manière opérationnelle, cette conclusion est intéressante dans la mesure où la mortalité des assurés augmente considérablement les temps de simulation.

Il n'est en revanche pas sûr que l'on tire les mêmes conclusions pour des contrats de prévoyance à prime périodique disposant d'une composante épargne. Pour de tels contrats, des garanties d'exonération des primes prévoient qu'en cas d'incapacité temporaire, l'assureur verse les primes périodiques à la place de l'assuré. Dans de tels cas, les fluctuations d'échantillonnage, du risque arrêt de travail par exemple, peuvent conduire à une variabilité accrue de l'engagement de l'assureur.

Bibliographie

- Bühlmann H., Straub E. (1970) « Glaubwürdigkeit für Schadensätze », *Mitteilungen der Vereinigung Schweizerischer Versicherungsmathematiker* 70.
- Bühlmann H., Gisler A. (2005) *A course in credibility theory*, Berlin : Springer.
- Frantz C., Chenut X., Walhin J.F. (2003) « Pricing and capital allocation for unit-linked life insurance contracts with minimum death guarantee », *Proceedings of the 13th AFIR Colloquium*, Maastricht.
- Hardy M.R., Panjer H.H. (1998) « A credibility approach to mortality risk », *ASTIN Bulletin* 28 (2), 269-83.
- Planchet F., Thérond P.E. (2004b) « Les principes de valorisation des engagements sociaux », *Revue fiduciaire comptable*, n° 310, 18-25.
- Planchet F., Thérond P.E. (2006) *Modèles de durée. Applications actuarielles*, Paris : Economica.
- Planchet F., Thérond P.E. (2007a) *Pilotage technique d'un régime de rentes viagères*, Paris : Economica.
- Planchet F., Thérond P.E., Jacquemin J. (2005) *Modèles financiers en assurance. Analyses de risque dynamiques*, Paris : Economica.
- Thérond P.E. (2003) *Impact des futures normes IFRS sur la tarification et le provisionnement des contrats d'assurance vie : mise en oeuvre de méthodes par simulation*, Mémoire d'actuariat, ISFA.
- Thérond P.E. (2005) « Contrôle de la solvabilité des compagnies d'assurance : évolutions récentes », Séminaire *Gestion des risques et assurance*, Hanoï, 28 février 2005.
- Thérond P.E., Girier G. (2005) « Solvabilité 2 en construction », *L'Argus de l'assurance*, Hors-série de décembre.

Chapitre 3

Incidence sur la gestion technique d'un assureur

Jusqu'alors dans le dispositif prudentiel français et, plus largement, dans de nombreux dispositifs européens, la solvabilité des sociétés d'assurance est assurée par une série de provisions¹ au passif, des contraintes sur les supports admissibles pour les placements à l'actif et la marge de solvabilité. En effet, les placements sont soumis à diverses règles prudentielles ayant pour but d'assurer leur sécurité. En particulier, leur répartition géographique et entre les différentes classes d'actifs, leur dispersion et leur congruence (cohérence de la monnaie dans laquelle ils sont libellés avec celle dans laquelle seront payées les prestations) répondent à des règles strictes². Ce mode de fonctionnement a pour conséquence directe que les problématiques de détermination des fonds propres, d'une part, et d'allocation d'actifs, d'autre part, sont distinctes. En effet, dans la cadre de la réglementation européenne actuelle, le niveau minimal des fonds propres (l'exigence de marge de solvabilité) dont doit disposer un assureur dépend uniquement³ du niveau des provisions techniques en assurance vie et du niveau des prestations et des primes en assurance non-vie.

Le projet Solvabilité 2 en cours d'élaboration modifie profondément ces règles en introduisant comme critère explicite de détermination du niveau des fonds propres le contrôle du risque global supporté par la société. Ce risque devra notamment être quantifié au travers de la probabilité de ruine. Ainsi, la détermination de l'allocation d'actifs se trouve de fait intégrée dans la démarche de fixation du niveau des fonds propres, la structure de l'actif impactant directement la solvabilité de l'assureur.

Un certain nombre de travaux s'attachent à définir des approches standards pour la détermination du capital de solvabilité (cf. AAI (2004), Djehiche et Hörfelt (2004) et la formule publiée par le CEIOPS dans le cadre de l'étude d'impact quantitatif QIS 3⁴). Dans ces approches la structure de l'actif est une donnée et les auteurs s'attachent à déterminer le niveau minimal de capital qui contrôle le risque global de la société.

Nous proposons ici un point de vue alternatif dans lequel les interactions entre le niveau du capital cible requis et l'allocation sont explicitement prises en compte et où l'on cherche à déterminer directement un couple capital / allocation d'actifs. La motivation de cette approche est qu'il nous apparaît

¹ Cf. art. R331-6 du Code des assurances.

² Cf. art. R332-1 et suiv. du Code des assurances.

³ Le calcul de l'exigence de marge de solvabilité intègre également la réassurance des sociétés.

⁴ Cf. l'annexe.

préférable dans ce nouveau contexte de déterminer de manière conjointe le niveau du capital et la manière d'allouer ce dernier, du fait des fortes interactions que Solvabilité 2 induit à ce niveau.

Pour cela nous allons utiliser le critère de maximisation des fonds propres économiques (MFPE) initialement élaboré dans des problématiques d'assurance vie dans Planchet et Thérond (2004b). Ce critère revient à déterminer le couple capital / allocation d'actifs qui maximise la valeur initiale de la firme (exprimée en pourcentage des fonds propres réglementaires) lorsque celle-ci est mesurée sous l'opérateur espérance. Ce critère permet d'obtenir une allocation dont la détermination ne dépend pas d'un critère subjectif tel que le niveau de la probabilité de ruine. L'allocation déterminée est ensuite confrontée *ex post* à de tels indicateurs de manière à la calibrer et à valider son respect des contraintes réglementaires.

Après avoir présenté un modèle simplifié de société d'assurance sur laquelle nous allons travailler, nous reprenons la définition d'allocation optimale au sens du critère de MFPE et l'explicitons, en particulier, d'abord dans le cadre de la réglementation française puis d'un référentiel de type Solvabilité 2.

Nous illustrons ensuite la mise en œuvre de ce critère dans le cas d'une société d'assurance couvrant deux types de risques dépendants et devant composer son portefeuille financier parmi un actif risqué (dont le rendement est modélisé par un processus de Lévy simple) et un actif sans risque. Provisions techniques et niveaux de fonds propres minimaux sont déterminés dans les deux référentiels prudentiels ; puis le critère de MFPE conduit à une allocation d'actifs dans le cadre de la réglementation française et un couple capital cible / allocation d'actifs dans le référentiel de type Solvabilité 2.

Ce chapitre s'appuie sur Planchet et Thérond (2005a) et (2007b).

1. Modélisation de la société d'assurance

Une société couvre sur une période⁵ n risques qui engendreront sur cette période les montants de sinistres $S_1, \dots, S_n$ où S_i correspond à la charge de sinistres de l'ensemble des polices de la branche i dont la fonction de répartition sera notée F_i . Cette modélisation n'est pas restrictive dans la mesure où la variable aléatoire (v.a.) S_i peut représenter les prestations effectivement versées sur la période ou leur valeur actualisée en fin de période. Nous ferons l'hypothèse que le versement des prestations a lieu en fin de période. Nous supposerons également que l'assureur ne souscrit pas de nouveau contrat en cours de période et que toutes les survenances de sinistres sont connues en fin de période.

En début de période, conformément à la réglementation, l'assureur a doté ses provisions techniques d'un montant L_0 . Parallèlement l'assureur dispose d'un niveau de fonds propres E_0 qui doit être supérieur au niveau minimum de

⁵. Les documents de travail de la Commission européenne privilégient une approche mono-périodique pour l'appréciation de la solvabilité.

fonds propres réglementaires E_0^R . Nous supposons que l'assureur place le montant $E_0 + L_0$ dans m actifs financiers $A = (A_1, \dots, A_m)$ avec les proportions⁶ $\omega = (\omega_1, \dots, \omega_m)$. Pour simplifier les notations, nous poserons, sans perte de généralité, pour tout $j \in \{1, \dots, m\}$, $A_j = 1$ et $\sum_{j=1}^m \omega_j = 1$. Nous noterons Ω l'ensemble des choix de portefeuille admissibles au regard de la réglementation.

Dans le cadre de la législation française actuelle, le calcul de L_0 et E_0^R dépend uniquement⁷ de $S = (S_1, \dots, S_n)$, alors que sous un référentiel du type Solvabilité 2, E_0^R est également fonction de $\omega, A_1, \dots, A_m$.

En fin de période, l'assureur doit payer le montant de prestations $\hat{S} = \sum_{i=1}^n S_i$. Il dispose comme ressource de $(L_0 + E_0) \sum_{j=1}^m \omega_j A_j$. Dans la suite, nous ferons l'hypothèse que les vecteurs aléatoires S et A sont indépendants.

Dans la suite Φ désignera la fonction de répartition de la loi normale centrée réduite $\mathcal{N}(0;1)$.

2. Critère de maximisation des fonds propres économiques

Après avoir introduit de manière générale le critère de maximisation des fonds propres économiques, nous étudions sa mise en œuvre dans le cadre de la réglementation française actuelle d'une part, puis d'un référentiel de type Solvabilité 2 d'autre part.

2.1. Présentation générale

En espérance, l'assureur devra déboursier $E[\hat{S}] = E \sum_{i=1}^n S_i = \sum_{i=1}^n E[S_i]$ en fin de période. Par ailleurs et à la même date, ses ressources seront constituées par ses provisions techniques, ses fonds propres et les produits financiers qu'ils ont engendrés, soit $(L_0 + E_0) \sum_{j=1}^m \omega_j A_j$. En espérance, l'actif de l'assureur en fin de période atteindra donc le montant :

⁶. On suppose donc implicitement que les règles de placement sont identiques pour les actifs en représentation des provisions techniques et pour les actifs associés aux fonds propres.

⁷. Sous réserve de l'utilisation d'un principe de primes ne dépendant que des lois des S_i .

$$\mathbb{E} \left[(L_0 + E_0) \sum_{j=1}^m \omega_j A_j \right] = (L_0 + E_0) \sum_{j=1}^m \omega_j \mathbb{E} [A_j].$$

Dans la suite nous noterons $\Lambda_0^\omega = \left(\sum_{i=1}^n S_i \right) \left(\sum_{j=1}^m \omega_j A_j \right)^{-1}$ la charge des prestations actualisée au taux de rendement du portefeuille financier et $\Sigma_0^\omega = (L_0 + E_0) - \Lambda_0^\omega$. La variable aléatoire Σ_0^ω peut s'interpréter comme la valeur (aléatoire) du surplus au terme actualisée au taux de rendement de l'actif.

Définition 15. *Fonds propres et provisions économiques*

Nous appellerons « fonds propres économiques » l'espérance de ce surplus actualisé $\mathbb{E} [\Sigma_0^\omega] = \mathbb{E} [E_0 + L_0 - \Lambda_0^\omega]$ et « provision économique » la quantité $\mathbb{E} [\Lambda_0^\omega]$.

Pour tout $\omega \in \Omega$, placer en début de période la quantité $\mathbb{E} [\Lambda_0^\omega]$ selon l'allocation ω permet d'être, en espérance, capable de payer les prestations en fin de période. En effet, la fonction inverse étant convexe sur $]0; +\infty[$, l'inégalité de Jensen assure que pour toute variable aléatoire X à valeurs strictement positives, $\mathbb{E} [X^{-1}] > (\mathbb{E} [X])^{-1}$ et donc

$$\mathbb{E} [\Lambda_0^\omega] \times \mathbb{E} \left[\sum_{j=1}^m \omega_j A_j \right] \geq \mathbb{E} \left[\sum_{i=1}^n S_i \right],$$

puisque A et S ont été supposés indépendants.

Aussi les fonds propres économiques $\mathbb{E} [\Sigma_0^\omega]$ peuvent s'interpréter comme une valorisation, sous l'opérateur espérance, de la compagnie en 0. L'assureur peut rechercher à maximiser cette valorisation en référence aux capitaux de la société par le biais de son allocation d'actifs, *i. e.* à choisir le portefeuille ω qui maximise la quantité

$$\phi(\omega) = \frac{\mathbb{E} [\Sigma_0^\omega]}{E_0},$$

i. e. le rapport entre la « valeur économique » de la société et ses fonds propres comptables.

Dans la suite, on dira qu'une allocation ω^* est optimale au sens du critère de maximisation des fonds propres économiques (MFPE) si ω^* est solution du programme d'optimisation $\sup_{\omega \in \Omega} \{\phi(\omega)\}$.

On note que l'ensemble des solutions ω^* de $\sup_{\omega \in \Omega} \{\phi(\omega)\}$ coïncide avec l'ensemble des solutions du programme $\sup_{\omega \in \Omega} \left\{ \left(L_0 - E \left[\Lambda_0^\omega \right] \right) / E_0 \right\}$. Le terme E_0^{-1} ne peut quant à lui être éliminé puisque l'assureur doit disposer d'un niveau de fonds propres $E_0 \geq E_0^R$ et que E_0^R peut dépendre de ω (ce sera notamment le cas dans le référentiel de type Solvabilité 2 présenté *infra*).

On démontre de manière immédiate que si les variables aléatoires $\sum_{i=1}^n S_i$ et $\left(\sum_{j=1}^m \omega_j A_j \right)^{-1}$ sont indépendantes, ω^* est optimale au sens du critère MFPE si, et seulement si, ω^* est solution du programme d'optimisation :

$$\sup_{\omega \in \Omega} \left\{ E_0^{-1} \left(L_0 - E \left[\sum_{i=1}^n S_i \right] E \left[\left(\sum_{j=1}^m \omega_j A_j \right)^{-1} \right] \right) \right\}.$$

Dans la suite, on supposera que l'assureur dispose du niveau minimal de fonds propres réglementaires et donc que $E_0 = E_0^R$.

2.2. Mise en œuvre dans le cadre réglementaire français

La réglementation française impose d'évaluer les provisions techniques branche par branche. Elle précise⁸ que ces provisions correspondent à la « valeur estimative des dépenses (...) nécessaires au règlement de tous les sinistres survenus et non payés ». Aussi le montant à provisionner en début de période au titre de la branche i sera l'espérance de S_i . Par ailleurs l'escompte des provisions est prohibé⁹ de sorte que l'on a $\forall i \in \{1, \dots, n\}, L_0^i = E[S_i]$ et

$$L_0 = \sum_{i=1}^n L_0^i.$$

Dans le cadre de la réglementation actuellement en vigueur, E_0^R est appelé « exigence de marge de solvabilité¹⁰ » et est indépendant de $\omega, A_1, \dots, A_m$. En effet, le niveau de l'exigence de marge de solvabilité est uniquement fonction des sinistres passés et des primes encaissées¹¹ : l'assureur calcule donc séparément l'exigence de marge selon les méthodes à partir des primes d'une

⁸. Cf. art. R331-6 du Code des assurances.

⁹. À l'exception de branches à très long développement telles que l'assurance responsabilité civile en assurance construction par exemple.

¹⁰. Cf. art. L334-1 du Code des assurances.

¹¹. Cf. Directive européenne 2002/13/CE.

part et des sinistres d'autre part, puis retient la plus grande des deux valeurs. Nous supposons ici que l'exigence de marge de solvabilité est calculée à partir des primes. En l'absence de réassurance, la méthode à partir des primes consiste à retenir le montant

$$18\% \times \min \{ \Pi_S, 50 \text{ M€} \} + 16\% \times \max \{ \Pi_S - 50 \text{ M€}, 0 \},$$

où Π_S désigne le montant total des primes commerciales encaissées. Dans la suite, nous supposons ainsi que

$$E_0^R = 18\% \times (1 + \gamma) E[\tilde{S}],$$

où γ est le taux de chargement des primes, que nous supposons commun à toutes les branches.

Dans ce contexte, le passif de la société n'intègre ni la dépendance qui peut exister entre les branches, ni les risques liés aux placements et, *a fortiori*, les risques actif-passif. L'indépendance de E_0^R par rapport à $\omega, A_1, \dots, A_m$ permet d'établir trivialement que, dans le cadre réglementaire français, une allocation ω^* est optimale au sens du critère de maximisation des fonds propres économiques si ω^* est solution du programme d'optimisation :

$$\inf_{\omega \in \Omega} E \left[\left(\sum_{j=1}^m \omega_j A_j \right)^{-1} \right].$$

Dans ce modèle mono-périodique et en appliquant la réglementation actuelle en vigueur, l'allocation optimale au sens du critère de MFPE ne dépend donc que des caractéristiques des placements financiers. Rappelons que lorsqu'il est mis en oeuvre dans un modèle multi-périodique tel que dans le cas d'un régime de rentiers (cf. Planchet et Thérond (2004b)), la structure du passif est intégrée par le biais du profil espéré des flux futurs. Le problème d'optimisation à résoudre s'écrit alors

$$\inf_{\omega \in \Omega} E \left[\sum_{k=1}^a \left\{ E[Flux(k)] \left(\sum_{j=1}^m \omega_j A_j(k) \right)^{-1} \right\} \right],$$

qui est une généralisation naturelle du critère présenté *supra*.

Dans la situation traitée ici, le fait que le critère de MFPE ne dépende que de l'actif est donc la conséquence d'une part du fait que la marge de solvabilité ne dépend pas de ω , et d'autre part que le modèle considéré est mono-périodique.

2.3. Mise en œuvre dans le référentiel Solvabilité 2

Le projet Solvabilité 2 vise à introduire des outils de gestion de la solvabilité *globale* de la société d'assurance. En termes quantitatifs, il s'agira essentiellement :

- de déterminer, pour chaque branche, un niveau de provisions techniques qui intègre la dangerosité du risque appréciée par une mesure de risque ;
- de déterminer un niveau de « capital cible » qui contrôle, avec une forte probabilité, tous les risques supportés par la société sur un horizon fixé.

Dans la suite, nous supposons que pour chaque branche, l'assureur doit provisionner, en début de période, le montant qui lui permet de faire face à ses engagements de payer en fin de période les prestations engendrées par cette branche dans 75 % des cas. Cela revient à provisionner, pour chaque branche i , la Value-at-Risk (VaR) à 75 % du risque S_i escomptée au taux d'intérêt sans risque sur la période r . Nous supposons, sans perte de généralité, que la période est de durée 1.

On rappelle que la Value-at-Risk (VaR) de niveau α associée au risque X est donnée par $\text{VaR}(X, \alpha) = \inf \{x \mid \Pr[X \leq x] \geq \alpha\}$.

Pour chaque branche i , on aura donc :

$$L_0^i = \text{VaR}(S_i, 75\%)e^{-r}.$$

Par ailleurs, nous supposons que le « capital cible » E_0^R est déterminé de telle manière que l'entreprise soit capable, en fin de période, de faire face à ses engagements envers ses assurés avec une probabilité de 99,5 %. Formellement cela signifie que E_0^R est le plus petit montant qui remplit la condition :

$$\Pr \left[\left(L_0 + E_0^R \right) \sum_{j=1}^m \omega_j A_j \geq \sum_{i=1}^n S_i \right] \geq 0,995.$$

Donc $E_0^R + L_0$ est la VaR à 99,5 % de Λ_0^ω . Le niveau du capital cible est donc égal à :

$$E_0^R = \inf \left\{ E_0 \geq 0 \mid \Pr \left[\Lambda_0^\omega - \sum_{i=1}^n \text{VaR}(S_i, 75\%)e^{-r} \leq E_0 \right] \geq 0,995 \right\}.$$

Cette expression nous permet d'observer qu'à la différence de la marge de solvabilité, le capital cible est fonction de la dépendance stochastique entre les différentes branches et des risques liés aux placements financiers.

3. Recherche de l'allocation optimale

Dans cette partie, nous allons mettre en œuvre le critère de MFPE dans le cas d'une société d'assurance qui couvre deux risques S_1 et S_2 et qui doit composer son portefeuille financier parmi deux actifs A_1 et A_2 .

Après avoir comparé les bilans obtenus dans les deux systèmes prudentiels, nous comparons les allocations obtenues par le critère de MFPE et leur sensibilité à l'évolution de différents paramètres.

3.1. Modélisation des risques

La compagnie supporte deux types de risque : les risques de passif liés à la sinistralité engendrée par son portefeuille de contrats d'assurance et les risques de placements. Par nature ces deux types de risques sont différents aussi bien dans la manière dont ils affectent l'assureur que dans le pilotage qui peut être mis en œuvre pour les contrôler. Ainsi, à la différence des risques de sinistres, les risques de placements ne se mutualisent pas mais l'assureur peut les contrôler par sa politique d'investissement.

3.1.1. Risques des sinistres

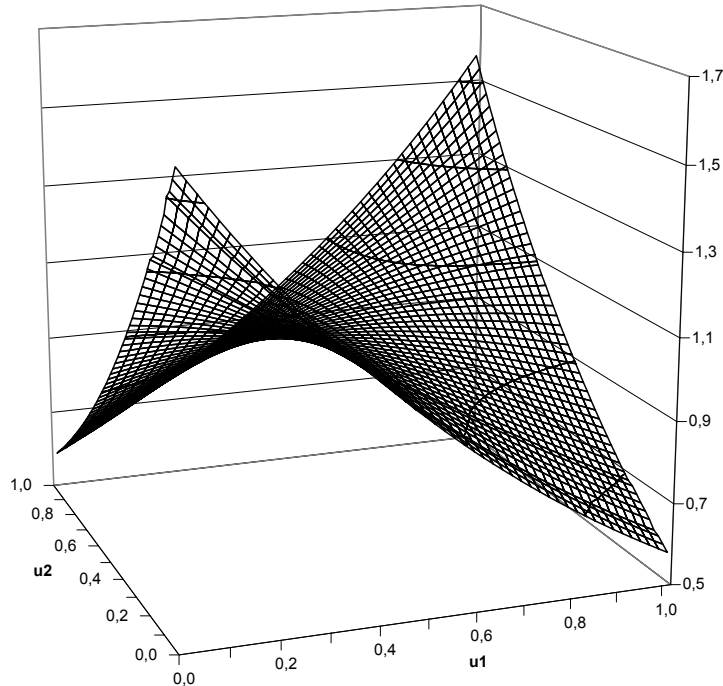
Nous supposons que les deux risques S_1 et S_2 couverts par la société d'assurance sont distribués selon des lois log-normales $\mathcal{LN}(\mu_1, \sigma_1)$ et $\mathcal{LN}(\mu_2, \sigma_2)$. Notons C_α la copule de Franck qui sera supposée modéliser la dépendance entre ces deux risques. Le théorème de Sklar (cf. Chapitre 5 pour une présentation de la théorie des copules) nous indique que la loi du couple (S_1, S_2) peut s'écrire, de manière unique,

$$\Pr[S_1 \leq s_1, S_2 \leq s_2] = C_\alpha(F_1(s_1), F_2(s_2)),$$

où :

- F_1 et F_2 , les fonctions de répartition respectivement de S_1 et S_2 , sont telles que $\Pr[\ln S_i \leq s] = \Phi\left(\frac{s - \mu_i}{\sigma_i}\right)$ pour $i \in \{1; 2\}$,
- $C_\alpha(u_1, u_2) = -\frac{1}{\alpha} \ln \left\{ 1 + \frac{(e^{-\alpha u_1} - 1)(e^{-\alpha u_2} - 1)}{e^{-\alpha} - 1} \right\}$.

La copule de Franck permet de disposer d'une structure de dépendance qui permet, selon la valeur de α , de modéliser des risques indépendants ($\alpha \rightarrow 0$), avec une dépendance négative ($\alpha \rightarrow -\infty$ conduit à la copule de la borne inférieure de Fréchet) ou avec une dépendance positive ($\alpha \rightarrow +\infty$ conduit à la copule de la borne supérieure de Fréchet). La structure de dépendance induite par cette copule est illustrée par le graphe suivant.

Figure 13 - Densité de la copule de Franck de paramètre 1

Pour les applications numériques, nous utiliserons les paramètres suivants¹² :

$$\begin{array}{lll} \mu_1 = 5,0099 & \sigma_1 = 0,0377 & \alpha = 1 \\ \mu_2 = 3,8421 & \sigma_2 = 0,3740 & \gamma = 0,15 \end{array}$$

Notons que $\alpha = 1$ correspond à des risques dont la dépendance est positive.

Comme il n'est pas possible d'expliciter par une formule directement utilisable la loi d'une somme de v.a. log-normales, nous utiliserons les techniques de Monte Carlo pour obtenir la fonction de répartition empirique de cette somme. Cependant, en pratique, il apparaît souvent souhaitable de diminuer le nombre de variables à simuler, dès lors il peut être intéressant d'utiliser des approximations de leur loi. Par exemple, Dhaene et al. (2004b) proposent d'approximer $\hat{S}$ par la *comonotonic upper bound* $\hat{S}^c$ de $\hat{S}$ définie par

$$\hat{S}^c = \sum_{i=1}^2 \exp\{\mu_i + \sigma_i \Phi^{-1}(U)\},$$

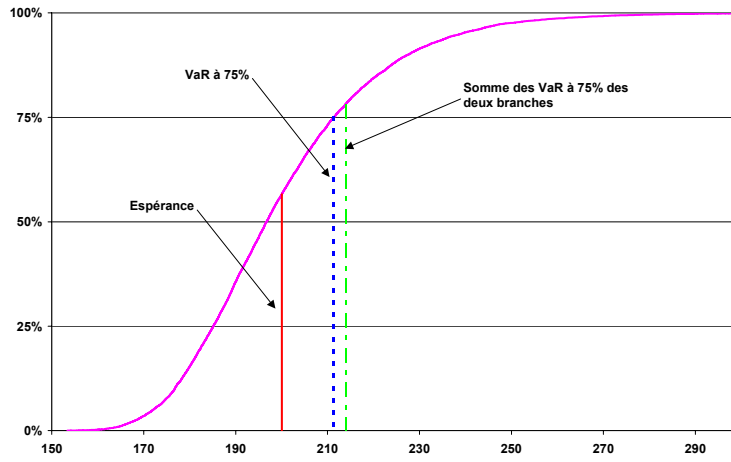
où U est une v.a. de loi uniforme sur $[0;1]$.

¹². Les paramètres sont en fait choisis pour que l'espérance de la charge sinistre pour le risque n° 1 soit égale à 150 et à 50 pour le risque n° 2.

Cette approximation permet notamment de disposer de formules fermées pour calculer des Value-at-Risk ou encore des Tail-VaR. Son utilisation requiert néanmoins de mesurer l'erreur d'approximation commise. Ce point ne sera pas développé plus avant dans le présent travail.

Le graphique suivant présente la distribution de la charge totale de sinistres obtenue pour cette modélisation du passif avec les paramètres indiqués *supra*.

Figure 14 - Distribution de la charge totale de sinistres



On remarque en particulier que la somme des VaR à 75 % est supérieure à la VaR à 75 % de la somme des risques bien que la dépendance entre les risques soit positive. On rappelle en effet que la Value-at-Risk n'est pas une mesure de risque cohérente au sens de Artzner et al. (1999) car elle n'est pas sous-additive (cf. le Chapitre 1 pour un contre-exemple).

3.1.2. Risques des placements

Nous supposons que l'assureur doit composer son portefeuille financier parmi deux actifs A_1 et A_2 . Pour fixer les idées, nous supposons que A_1 est un actif risqué et A_2 un bon de capitalisation. Nous supposons que le cours de l'actif risqué suit le modèle de Merton, et est donc un processus à saut défini par

$$A_1(t) = \exp \left\{ \left(\mu - \frac{\sigma^2}{2} \right) t + \sigma B_t + \sum_{k=1}^{N_t} U_k \right\},$$

où :

- $B = (B_t)_{t \geq 0}$ est un mouvement brownien,
- $N = (N_t)_{t \geq 0}$ est un processus de Poisson d'intensité λ ,
- $U = (U_k)_{k \geq 1}$ est une suite de variables aléatoires indépendantes identiquement distribuées de loi une loi normale $\mathcal{N}(0, \sigma_u)$,

– les processus B , N et U sont mutuellement indépendants.

Les sauts sont ici, dans un souci de simplicité, supposés symétriques et en moyenne nuls ; des modèles plus élaborés à sauts dissymétriques peuvent également être proposés (cf. Ramezani et Zeng (1998)).

Notons Ψ_t la tribu engendrée par les B_s, N_s pour $s \leq t$ et $U_k 1_{\{k \leq N_t\}}$ pour $j \geq 1$; B est un mouvement brownien standard par rapport à la filtration Ψ , N est un processus adapté à cette même filtration. De plus pour tout $t > s$, $N_t - N_s$ est indépendant de la tribu Ψ_s .

Proposition 3. *Distribution de $A_1(t)$*

Pour tout $x > 0$,

$$\Pr[A_1(t) \leq x] = \sum_{n=0}^{+\infty} \Phi \left[\frac{\ln x - \left(\mu - \sigma^2/2\right)t}{\sqrt{n\sigma_u^2 + t\sigma^2}} \right] e^{-\lambda t} \frac{(\lambda t)^n}{n!}.$$

Démonstration : Soit $x > 0$, on a :

$$\begin{aligned} \Pr[A_1(t) \leq x] &= \Pr \left[\sum_{k=1}^{N_t} U_k + \left(\mu - \sigma^2/2\right)t + \sigma B_t \leq \ln x \right] \\ &= \sum_{n=0}^{+\infty} \Pr \left[\sum_{k=1}^{N_t} U_k + \left(\mu - \sigma^2/2\right)t + \sigma B_t \leq \ln x, N_t = n \right] \\ &= \sum_{n=0}^{+\infty} \Pr \left[\sum_{k=1}^n U_k + \left(\mu - \sigma^2/2\right)t + \sigma B_t \leq \ln x \right] \Pr[N_t = n], \end{aligned}$$

puisque les processus N , B et U sont mutuellement indépendants. Par ailleurs, les processus $\sum_{k=1}^n U_k$ et σB_t étant indépendants et gaussiens, leur somme est

également gaussienne : $\sum_{k=1}^n U_k + \sigma B_t \sim \mathcal{N} \left(0; \sqrt{n\sigma_u^2 + t\sigma^2} \right)$. Enfin comme N est

un processus de Poisson d'intensité λ , pour tout $t > 0$, la v. a. N_t est distribuée selon une loi de Poisson de paramètre λt et donc

$$\Pr[N_t = n] = e^{-\lambda t} \frac{(\lambda t)^n}{n!}. \quad \square$$

Lorsque $\lambda = 0$ (cas de l'absence de sauts) on retrouve la loi log-normale usuelle du brownien géométrique. Dans le cas général, l'expression de la Proposition 4 permet d'approcher la distribution de l'actif en ne conservant qu'un nombre fini de termes dans la somme.

Par ailleurs, nous supposons qu'à la date t , le bon de capitalisation A_2 vaut $A_2(t) = e^{rt}$ où r est le taux d'intérêt sans risque, supposé constant sur la période étudiée, utilisé pour escompter les provisions dans le paragraphe 2.3.

Pour les applications numériques, nous utiliserons les paramètres suivants.

$$\begin{aligned} \mu &= 0,6 & \sigma &= 0,15 & r &= 0,344 \\ \lambda &= 0,5 & \sigma_u &= 0,2 \end{aligned}$$

Les sauts seront donc d'espérance nulle et, en moyenne, il en surviendra un toutes les deux périodes. Par ailleurs le taux sans risque r a été pris de manière à ce que le taux d'escompte discret soit de 3,5 %.

Enfin nous ferons l'hypothèse que $\Omega = [0;1] \times [0;1]$, ce qui signifie que les ventes à découvert sont interdites à l'assureur.

3.2. MFPE dans la réglementation française

Dans un premier temps, nous allons nous intéresser à ce qui se passe dans la réglementation française dans laquelle le niveau minimal de fonds propres ne dépend pas du choix de portefeuille et donc dans laquelle le problème d'optimisation n'intègre pas, lorsque l'on ne considère qu'une période, la variabilité des flux de sinistres.

3.2.1. Bilan initial

Dans le cadre de la réglementation française, l'assureur va doter ses provisions techniques en début de période du montant $L_0 = E[S_1] + E[S_2]$.

Le montant de ses fonds propres E_0 sera supposé être égal à la marge de solvabilité E_0^R soit : $E_0^R = 18\% \times (1 + \gamma) E(S_1 + S_2)$. Avec les modélisations de sinistres retenues *supra*, le passif de l'assureur est donc déterminé par

$$\begin{cases} L_0 = \exp\left(\mu_1 + \frac{\sigma_1^2}{2}\right) + \exp\left(\mu_2 + \frac{\sigma_2^2}{2}\right) \\ E_0 = E_0^R = 0,18(1 + \gamma) \left(\exp\left(\mu_1 + \frac{\sigma_1^2}{2}\right) + \exp\left(\mu_2 + \frac{\sigma_2^2}{2}\right) \right) \end{cases}$$

Avec les paramètres sélectionnés, le bilan en 0 de la société est résumé dans le tableau *infra*.

Tableau 3 - Bilan initial dans la réglementation française actuelle

BILAN (réglementation française)	
$E_0 + L_0 = 241,4$	$E_0 = 41,4$
	$L_0^1 = 150$
	$L_0^2 = 50$

3.2.2. Allocation optimale

Dans le cadre de la réglementation française, on a vu que l'allocation $\omega = (\omega_1, 1 - \omega_1)$ est optimale au sens du critère de MFPE si ω_1 est solution du programme d'optimisation $\inf_{\omega \in \Omega} E \left[(\omega_1 A_1 + (1 - \omega_1) A_2)^{-1} \right]$. La proposition suivante donne la condition nécessaire et suffisante pour que ce programme d'optimisation ait une solution non triviale¹³.

Proposition 4. *Unicité de la solution du critère MFPE*

Le programme $\inf_{\omega \in [0;1]} E \left[(\omega A_1 + (1 - \omega) A_2)^{-1} \right]$ admet une unique solution ω^* non triviale ($\omega^* \in]0;1[$) si, et seulement si,

$$r < \mu < r + 2\sigma^2 + \lambda \left[\exp(\sigma_u^2) - \exp(\sigma_u^2/2) \right].$$

Ce résultat est démontré en annexe.

On peut vérifier que les paramètres $(r, \mu, \sigma, \sigma_u, \lambda)$ choisis en 3.1.2 sont tels qu'il existe une solution non triviale :

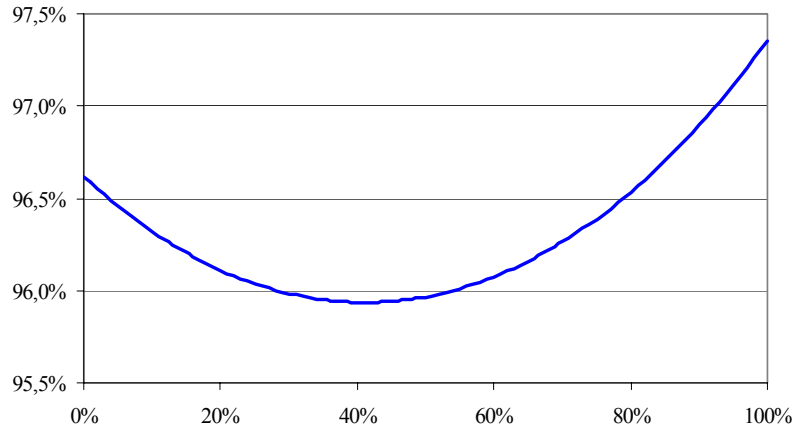
$$\underbrace{\underbrace{r}_{0,0344} < \underbrace{\mu}_{0,06} < r + 2\sigma^2 + \lambda \left[\exp(\sigma_u^2) - \exp(\sigma_u^2/2) \right]}_{0,0897}.$$

Cette proposition permet d'établir que si $\mu < r$, l'assureur consacrera l'intégralité de son portefeuille financier au bon de capitalisation. En revanche si $\mu > r + 2\sigma^2 + \lambda \left(\exp\{\sigma_u^2\} - \exp\{\sigma_u^2/2\} \right)$, provisions techniques et fonds propres seront exclusivement investis dans l'actif risqué.

Le graphique suivant présente l'évolution de la quantité $E \left[(\omega_1 A_1 + (1 - \omega_1) A_2)^{-1} \right]$ en fonction de ω_1 .

¹³. L'expression analytique du minimum est en revanche délicate à obtenir et les calculs numériques seront menés par des techniques de simulation.

Figure 15 - Provision économique exprimée en pourcentage de la provision réglementaire en fonction de ω_1



Cette fonction a un minimum non trivial pour $\omega_1 = 0,391$. La société maximisera donc ses fonds propres économiques si elle place son actif initial pour 39,1 % en actions.

3.2.3. Probabilité de ruine

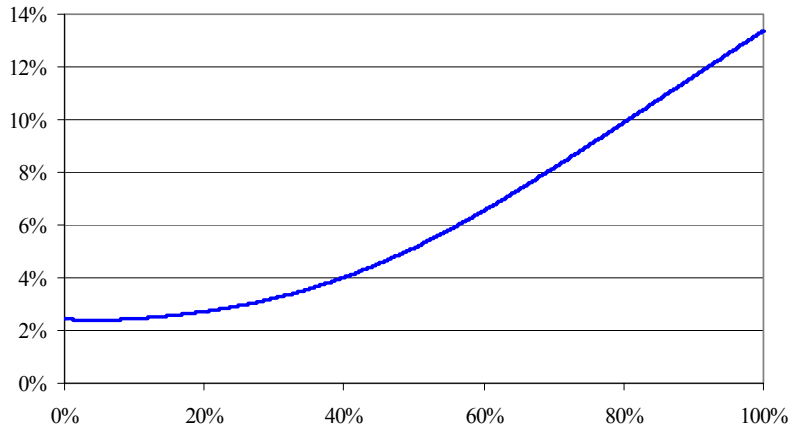
Afin de mesurer le niveau de prudence associé à l'allocation fournie par le critère de MFPE, il est naturel de déterminer la probabilité de ruine qui lui est associée. Aussi nous allons nous intéresser à la probabilité que l'assureur ne puisse payer les prestations en fin de période. Nous supposons que la société

est en faillite en fin de période si $\hat{S} > (E_0 + L_0) \left(\sum_{j=1}^m \omega_j A_j \right)$.

À chaque allocation ω , il est possible d'associer le niveau de probabilité de ruine $\pi(\omega)$ défini par

$$\pi(\omega) = \Pr \left[\hat{S} > (E_0 + L_0) \sum_{j=1}^m \omega_j A_j \right].$$

L'expression analytique de $\pi(\omega)$ est complexe, mais cette grandeur peut être aisément approchée numériquement par des techniques de simulation. Le graphique suivant reprend l'évolution de la probabilité de ruine en fonction de la part d'actifs investie en actions.

Figure 16 - Probabilité de ruine en fonction de ω_1 

Avec un niveau de fonds propres initial égal à la marge de solvabilité, la probabilité de ruine de l'assureur en fin de période est minimale (2,4 %) lorsqu'il a placé ses provisions et ses fonds propres pour 4,3 % en actions. L'allocation optimale au sens du critère de MFPE (39,1 %) correspond quant à elle à une probabilité de ruine de l'ordre de 3,9 %.

3.3. MFPE dans un référentiel du type Solvabilité 2

Dans la cadre d'un référentiel de type Solvabilité 2, comme le niveau minimal de fonds propres dépend de l'allocation d'actifs, la mise en œuvre du critère MFPE passe, dans un premier temps, par la mise en lumière de la relation liant l'allocation et le capital cible ; puis par la détermination du couple capital cible / allocation qui maximise le rapport entre les fonds propres économiques et le capital cible.

3.3.1. Bilan initial

Comme $S_i \sim \mathcal{LN}(\mu_i, \sigma_i)$, $\Pr[S_i \leq x] = \Phi\left(\frac{\ln x - \mu_i}{\sigma_i}\right)$ et donc

$$\text{VaR}(S_i, 75\%) = \exp\{\mu_i + \sigma_i \Phi^{-1}(0,75)\}.$$

Ce qui nous permet au passage de vérifier que si les charges sinistres sont distribuées selon des lois log-normales, l'utilisation de la VaR pour déterminer le niveau des provisions est cohérente avec une approche *market value margin* (telle qu'évoquée dans le DSOP) qui consisterait à prendre $\exp\{\sigma_i \Phi^{-1}(p) - \sigma_i^2/2\} - 1$ comme coefficient de majoration de l'espérance. En effet, la log-normalité de la charge sinistres permet d'exprimer, par le biais d'un coefficient ne dépendant que de σ_i , la VaR en fonction de l'espérance puisque

$$\text{VaR}(S_i, p) = (1 + \exp\{\sigma_i \Phi^{-1}(p) - \sigma_i^2/2\} - 1) E[S_i].$$

Comme $L_0 = \sum_{i=1}^2 \text{VaR}(S_i, 75\%) e^{-r}$, le niveau total des provisions techniques est donné par

$$L_0 = \sum_{i=1}^2 \exp\{\mu_i - r + \sigma_i \Phi^{-1}(0, 75)\} = 148,55 + 57,97 = 206,52.$$

On peut noter que la faible variabilité du risque S_1 conduit, dans ce référentiel, à un niveau de provisions pour ce risque (148,55) inférieur à celui obtenu dans la réglementation française actuelle (150). La situation est inverse pour le risque S_2 . Au global, le changement de référentiel prudentiel conduit à augmenter les provisions techniques de 3,26 %.

Rappelons que le niveau du capital cible E_0^R dépend des risques de passif comme des risques de placement puisqu'il est solution du programme d'optimisation

$$\inf \left\{ E_0 \geq 0 \mid \Pr \left[\Lambda_0^\omega \leq E_0 + L_0 \right] \geq 99,5\% \right\},$$

qui admet une unique solution car la distribution sous-jacente est absolument continue. Ce programme peut se réécrire

$$\inf \left\{ E_0 \geq 0 \mid \Pr \left[\omega_1 A_1 + (1 - \omega_1) A_2 \leq \frac{S_1 + S_2}{E_0 + L_0} \right] \leq 0,5\% \right\}.$$

Comme nous avons supposé que l'évolution des actifs est indépendante de la sinistralité, on a

$$\Pr \left[\omega_1 A_1 + (1 - \omega_1) A_2 \leq \frac{S_1 + S_2}{E_0 + L_0} \right] = \int \Pr \left[\omega_1 A_1 + (1 - \omega_1) A_2 \leq \frac{s}{E_0 + L_0} \right] f_S(s) ds,$$

où f_S désigne la densité de la v. a. $\hat{S} = S_1 + S_2$. On a démontré *supra* que

$$\Pr \left[\omega_1 A_1 + (1 - \omega_1) A_2 \leq \frac{s}{E_0 + L_0} \right] = \sum_{n=0}^{+\infty} \Phi \left[\frac{\ln \rho(s) - (\mu - \sigma^2/2)}{\sqrt{n\sigma_u^2 + \sigma^2}} \right] e^{-\lambda} \frac{\lambda^n}{n!},$$

$$\text{où } \rho(s) = \frac{1}{\omega_1} \left(\frac{s}{E_0 + L_0} - (1 - \omega_1) e^r \right).$$

$$\text{Comme } \int \sum_{n=0}^{+\infty} \Phi \left[\frac{\ln \rho(s) - (\mu - \sigma^2/2)}{\sqrt{n\sigma_u^2 + \sigma^2}} \right] e^{-\lambda} \frac{\lambda^k}{k!} f_S(s) ds \text{ n'est pas simple à}$$

calculer, nous avons choisi de résoudre numériquement ce programme d'optimisation en simulant des réalisations de $\hat{S} = S_1 + S_2$ puis en estimant

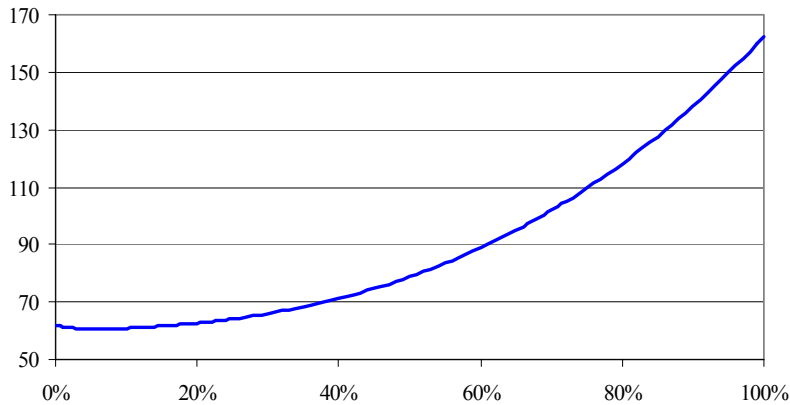
$\Pr\left[\Lambda_0^\omega \leq E_0 + L_0\right]$ à partir de la moyenne empirique des

$$p_k^w = \sum_{n=0}^{+\infty} \Phi\left[\frac{\ln p(s_k) - (\mu - \sigma^2/2)}{\sqrt{n\sigma_u^2 + \sigma^2}}\right] e^{-\lambda} \frac{\lambda^n}{n!} \text{ où } s_k \text{ désigne une réalisation de la}$$

variable aléatoire $\hat{S}$. La méthode d'obtention des réalisations de $\hat{S}$ est décrite en annexe de ce chapitre.

En pratique nous avons développé la somme jusqu'à $n = 7$ bien que le cinquième terme de la suite soit déjà négligeable devant la somme des précédents. Cette méthode nous permet d'obtenir la courbe suivante qui représente le niveau du capital cible en fonction de l'allocation stratégique $\omega = (\omega_1, 1 - \omega_1)$.

Figure 17 - Capital cible en fonction de la part investie en actions ω_1



Cette fonction a un minimum non trivial pour $\omega_1 = 6,1\%$ avec un capital cible de 60,71. Pour cette allocation, le passif de l'assureur s'élève à 267,24 contre 241,40 dans le cadre réglementaire français. Le passif peut atteindre 368,99 pour un actif intégralement composé d'actions.

3.3.2. Allocation optimale

Dans un référentiel de type Solvabilité 2, l'allocation $\omega = (\omega_1, 1 - \omega_1)$ est optimale au sens du critère de MFPE si ω_1 est solution du programme d'optimisation

$$\sup_{\omega_1 \in [0,1]} \varphi(\omega_1),$$

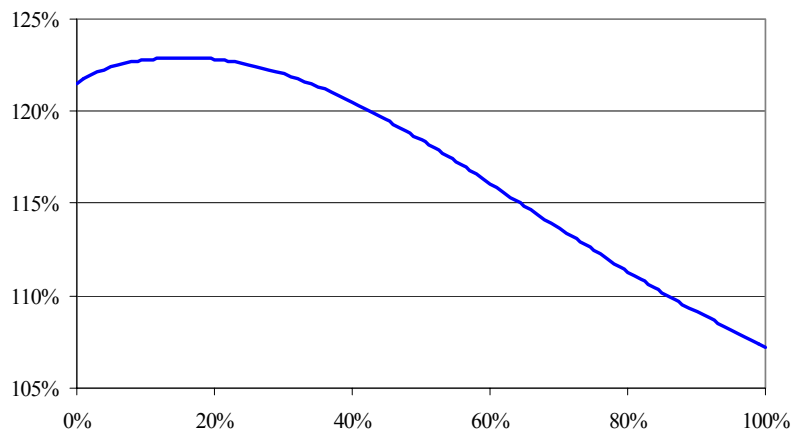
$$\text{où } \varphi(\omega_1) = \frac{1}{E_0^R(\omega_1)} \sum_{i=1}^2 \left\{ e^{\mu_i - r + \sigma_i} \Phi^{-1}(0,75) - e^{\mu_i + \sigma_i^2/2} E\left[\left(\omega_1 A_1 + (1 - \omega_1) e^r\right)^{-1}\right] \right\}.$$

Les conditions du premier ordre de ce programme sont délicates à expliciter du fait de la présence du terme E_0^R au dénominateur de la fonction objectif. Aussi nous avons cherché à résoudre numériquement ce problème qui, avec les paramètres utilisés, possède une unique solution sur $]0;1[$.

Pour une allocation $\omega = (\omega_1, 1 - \omega_1)$ fixée, le montant du capital cible a été déterminé en 3.3.1 ; il reste donc à faire varier l'allocation optimale, puis à calculer pour chaque allocation la valeur de la fonction objectif.

Le graphique suivant reprend l'évolution du rapport entre les fonds propres économiques et le fonds propres réglementaires selon la part ω_1 initialement investie en actif risqué A_1 .

Figure 18 - Graphe de ϕ



On constate empiriquement que la fonction objectif présente un maximum sur $]0;1[$.

Si l'assureur souhaite maximiser ses fonds propres économiques, il devra composer son portefeuille financier de 15,4 % d'actions en début de période. Pour réaliser cette allocation, les actionnaires devront fournir un capital réglementaire de 62,6. Pour cette allocation et ce capital, la société sera valorisée, sous l'opérateur espérance, à 75,7.

La probabilité de ruine définie en 3.2.2 est évidemment égale à 0,5 % puisque le capital cible a été déterminé de manière à contrôler la ruine avec cette probabilité.

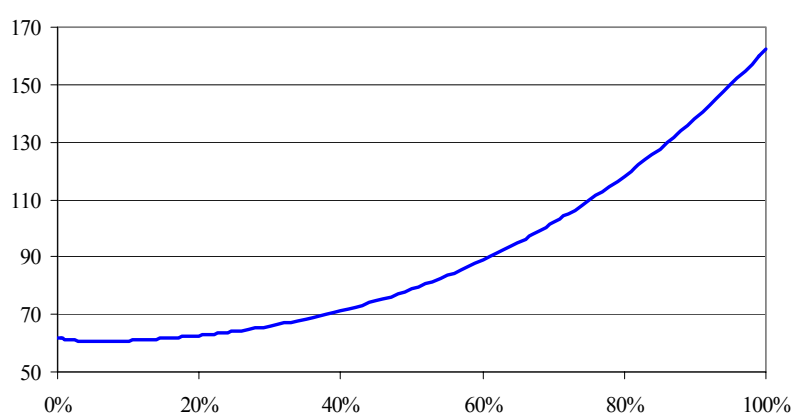
4. Conclusion

L'avènement du nouveau référentiel prudentiel Solvabilité 2 va engendrer une refonte des processus de gestion des compagnies d'assurance. En effet, l'objectif annoncé d'une exigence en matière de fonds propres en référence explicite au risque global supporté constitue une révolution pour le monde de l'assurance dans la mesure où toute décision de gestion qui a un impact sur

l'exposition au risque de l'assureur a, par ricochet, un impact sur l'exigence de fonds propres de celui-ci.

Cette situation a été analysée plus particulièrement ici au sujet des choix d'allocation stratégique. Si l'on étudie la Figure 17 reprise ci-dessous, on se rend compte que le choix d'une allocation stratégique a un impact important sur le niveau minimal de capitaux dont doit disposer l'assureur.

Figure 19 - Capital cible en fonction de la part investie en actions ω_1



Pour comparaison, dans le référentiel actuel, seules les règles de limitation entre les différentes catégories d'actifs se seraient imposées¹⁴ à l'assureur. Par ailleurs, le calcul de l'exigence de marge de solvabilité (en référence aux provisions mathématiques en vie, aux sinistres et aux prestations en non-vie) n'est directement dépendant du risque effectif que dans la mesure où celui-ci a justement un impact sur les provisions en vie, les prestations ou les sinistres en non-vie. L'exigence de marge de solvabilité calculée selon la règle *stricto sensu* est donc, dans la plupart des cas, indépendante des règles de gestion de l'actif, de la gestion actif-passif ou encore des risques opérationnels.

Comme ans Solvabilité 2, toute règle de gestion qui déforme le risque supporté par la compagnie a un impact sur le capital de solvabilité, il va s'agir de mettre en œuvre des processus de décision qui permettront de limiter le risque opérationnel d'une décision prise sans en mesurer les conséquences prudentielles. Ce type de mesure est d'ailleurs du ressort du pilier II de Solvabilité 2 comme plus généralement tous les processus opérationnels et les procédures de contrôles par les superviseurs.

Les sociétés ont commencé à mettre en place ce type de mesures depuis quelques années, sous l'impulsion, entre autres de l'application des normes IFRS depuis 2005. Ainsi, par exemple, de nombreuses sociétés qui avaient inscrits des obligations dans la catégorie *Held-To-Maturity* (HTM) de l'IAS 39

¹⁴. Pour ce qui est des règles automatiques : le superviseur peut demander un relèvement de l'exigence de marge de solvabilité s'il juge que la composition de l'actif expose trop l'assureur au risque action par exemple. En tout état de cause, ce jugement appartient au superviseur et n'est pas automatique.

(cf. Planchet, Thérond et al. (2005) pour un résumé de l'IAS 39) de manière à pouvoir les valoriser selon la méthode du coût amorti plutôt qu'en valeur de marché, ont mis en place des procédures de contrôle strictes de manière à ce qu'un de ces titres ne soit pas cédé inopportunément ce qui aurait pour conséquence de re-qualifier tous les autres titres HTM dans la catégorie *Available For Sale* (AFS) et ainsi de les valoriser en valeur de marché.

Ce type de risque a d'ailleurs incité une majorité d'acteurs français du marché de l'assurance à délaisser totalement la catégorie HTM et à valoriser tous leurs actifs financiers en juste valeur.

Annexe A : Démonstration des résultats mathématiques

Proposition 4. Le programme $\inf_{\omega \in [0;1]} E \left[\left(\omega A_1 + (1 - \omega) A_2 \right)^{-1} \right]$ admet une unique solution ω^* non triviale ($\omega^* \in]0;1[$) si, et seulement si,

$$r < \mu < r + 2\sigma^2 + \lambda \left[\exp(2\sigma_u^2) - \exp(\sigma_u^2/2) \right].$$

Démonstration : Cette démonstration va s'articuler en deux étapes. Dans un premier temps nous allons déterminer $E \left[A_1^p(t) \right]$ pour tout réel p . Ce résultat nous permettra d'expliciter une condition sur les paramètres des modèles d'actifs $(r, \mu, \sigma, \sigma_u, \lambda)$ pour que $\omega^* \in]0;1[$.

Etape 1 : Soit $p \in \mathbb{R}$, on a :

$$E \left[A_1^p(t) \right] = E \left[\exp \left\{ p \left(\mu - \frac{\sigma^2}{2} \right) t + p \sigma B_t + p \sum_{k=1}^{N_t} U_k \right\} \right].$$

Les termes aléatoires de l'exponentielle sont indépendants ce qui ramène le calcul au produit de $E \left[\exp \{ p \sigma B_t \} \right]$ et $E \left[\exp \left\{ p \sum_{k=1}^{N_t} U_k \right\} \right]$.

Pour tout $t > 0$, B_t est une v. a. de loi $\mathcal{N}(0; \sqrt{t})$ donc

$$E \left[\exp \{ p \sigma B_t \} \right] = \exp \left\{ \frac{p^2 \sigma^2}{2} t \right\}.$$

Par ailleurs on a vu (cf. le sous-paragraphe 2.1) que

$$E \left[\exp \left\{ p \sum_{k=1}^{N_t} U_k \right\} \right] = e^{-\lambda t} \sum_{n=0}^{+\infty} \frac{(\lambda t)^n}{n!} E \left[\exp \left\{ p \sum_{k=1}^n U_k \right\} \right].$$

Comme les sauts sont gaussiens et centrés,

$$E \left[\exp \left\{ p \sum_{k=1}^n U_k \right\} \right] = \exp \left\{ \frac{n p^2 \sigma_u^2}{2} \right\},$$

et donc

$$E \left[\exp \left\{ p \sum_{k=1}^{N_t} U_k \right\} \right] = e^{-\lambda t} \sum_{n=0}^{+\infty} \frac{(\lambda t)^n}{n!} \exp \left\{ \frac{n p^2 \sigma_u^2}{2} \right\} = \exp \left\{ \lambda t \left[\exp(p^2 \sigma_u^2/2) - 1 \right] \right\}.$$

Ce qui permet d'obtenir

$$\mathbb{E}\left[A_1^p(t)\right] = \exp\left\{p\left(\mu - \frac{\sigma^2}{2}\right)t + \frac{p^2\sigma^2}{2}t + \lambda t\left[\exp\left(p^2\sigma_u^2/2\right) - 1\right]\right\}.$$

Etape 2 : La fonction objectif dans le cas réglementaire français s'écrit

$$\varphi(\omega) = \mathbb{E}\left[\left(e^r + \omega X\right)^{-1}\right] \text{ avec } X = A_1(1) - e^r. \text{ On a donc}$$

$$\frac{\partial \varphi}{\partial \omega}(\omega) = -\mathbb{E}\left[X\left(e^r + \omega X\right)^{-2}\right],$$

et

$$\frac{\partial^2 \varphi}{\partial \omega^2}(\omega) = 2\mathbb{E}\left[X^2\left(e^r + \omega X\right)^{-3}\right].$$

On en déduit trivialement que la fonction objectif est convexe : $\frac{\partial^2 \varphi}{\partial \omega^2}(\omega) > 0$.

Par ailleurs

$$\frac{\partial \varphi}{\partial \omega}(0) = -e^{-2r} \mathbb{E}(X) = -e^{-2r} (e^\mu - e^r),$$

et

$$\frac{\partial \varphi}{\partial \omega}(1) = -\mathbb{E}\left[A_1^{-2} (A_1 - e^r)\right] = \mathbb{E}\left[A_1^{-2} e^r\right] - \mathbb{E}\left[A_1^{-1}\right].$$

Cette dernière expression se calcule à l'aide du résultat de l'étape 1 puisque

$$\mathbb{E}\left[A_1^{-2} e^r\right] = \exp\left\{r - 2\mu + 3\sigma^2 + \lambda\left(e^{2\sigma_u^2} - 1\right)\right\},$$

et

$$\mathbb{E}\left[A_1^{-1}\right] = \exp\left\{-\mu + \sigma^2 + \lambda\left(e^{\sigma_u^2/2} - 1\right)\right\}.$$

Ce qui nous permet d'écrire

$$\frac{\partial \varphi}{\partial \omega}(1) = \exp\left\{r - 2\mu + 3\sigma^2 + \lambda\left(e^{2\sigma_u^2} - 1\right)\right\} - \exp\left\{-\mu + \sigma^2 + \lambda\left(e^{\sigma_u^2/2} - 1\right)\right\}.$$

La fonction φ est strictement convexe sur $[0;1]$, elle atteint donc un unique

minimum sur $]0;1[$ si $\frac{\partial \varphi}{\partial \omega}(0) < 0$ et $\frac{\partial \varphi}{\partial \omega}(1) > 1$, i. e. si

$$r < \mu < r + 2\sigma^2 + \lambda\left(\exp\left\{2\sigma_u^2\right\} - \exp\left\{\sigma_u^2/2\right\}\right). \quad \square$$

Annexe B : Simulation des réalisations de la charge de sinistres

Les réalisations de la charge sinistres $\hat{S}$ ont été obtenues par les techniques de Monte Carlo, à partir de la méthode des distributions conditionnelles qui permet de simuler des v. a. dont la dépendances est modélisée par une copule.

Cette méthode des distributions conditionnelles consiste à simuler¹⁵ indépendamment deux réalisations v_1, v_2 de v. a. de loi uniforme sur $[0;1]$ puis d'utiliser la transformation suivante :

$$\begin{cases} u_1 = v_1 \\ u_2 = C_{u_1}^{-1}(v_2) \end{cases}$$

$$\text{où } C_{u_1}(u_2) = \Pr[F_2(S_2) \leq u_2 | F_1(S_1) = u_1] = \frac{\partial C}{\partial u_1} C(u_1, u_2).$$

Dans le cas de la copule de Franck utilisée dans ce travail, cette dernière expression se calcule analytiquement :

$$C_{u_1}(u_2) = \frac{\exp(-\alpha u_1)(\exp(-\alpha u_2) - 1)}{\exp(-\alpha) - 1 + (\exp(-\alpha u_1) - 1)(\exp(-\alpha u_2) - 1)}.$$

De plus cette fonction s'inverse analytiquement puisque :

$$C_{u_1}^{-1}(u_2) = -\frac{1}{\alpha} \ln \left\{ 1 + \frac{(\exp(-\alpha) - 1)u_2}{u_2 + (1 - u_2)\exp(-\alpha u_1)} \right\}.$$

Enfin la charge totale de sinistres simulée est obtenue par

$$s = F_1^{-1}(u_1) + F_2^{-1}(u_2).$$

¹⁵. Pour les illustrations numériques, les simulations de réalisations de v.a. ont été obtenues à partir du générateur du tore mélangé présenté dans Planchet et Thérond (2004a) et au Chapitre 6.

Bibliographie

- AAI (2004) *A global framework for insurer solvency assessment*, <http://www.actuaires.org>.
- Artzner Ph., Delbaen F., Eber J.M., Heath D. (1999) « Coherent measures of risk », *Mathematical Finance* 9, 203-28.
- Commission européenne (2003) « Conception d'un futur système de contrôle prudentiel applicable dans l'Union européenne – Recommandation des services de la Commission », MARKT/2509/03.
- Commission européenne (2004) « Solvency II – Organisation of work, discussion on pillar I work areas and suggestions of further work on pillar II for CEIOPS », MARKT/2543/03.
- Dhaene J., Vanduffel S., Tang Q.H., Goovaerts M., Kaas R., Vyncke D. (2004) « Solvency capital, risk measures and comonotonicity: a review », Research Report OR 0416, Department of Applied Economics, K.U.Leuven.
- Dhaene J., Vanduffel S., Goovaerts M., Kaas R., Vyncke D. (2005) « Comonotonic approximations for optimal portfolio selection problems », *Journal of Risk and Insurance* 72, n° 2.
- Djehiche B., Hörfelt P. (2005) « Standard approaches to asset & liability risk », *Scandinavian Actuarial Journal* 2005, n° 5, 377-400.
- Merton R.C. (1976) « Option pricing when underlying stock returns are discontinuous », *Journal of Financial Economics* 3, 125-44.
- Planchet F., Thérond P.E. (2004a) « Simulation de trajectoires de processus continus », *Belgian Actuarial Bulletin* 5, 1-13.
- Planchet F., Thérond P.E. (2004b) « Allocation d'actifs d'un régime de rentes en cours de service », *Proceedings of the 14th AFIR Colloquium*, Boston, 111-34.
- Planchet F., Thérond P.E. (2005a) « Asset allocation : new constraints induced by the Solvency 2 project », *Proceedings of the 36th ASTIN Colloquium*, Zürich.
- Planchet F., Thérond P.E. (2007b) « Allocation d'actifs selon le critère de maximisation des fonds propres économiques en assurance non-vie : présentation et mise en oeuvre dans la réglementation française et dans un référentiel de type Solvabilité 2 », *Bulletin Français d'Actuariat* 7 (13), 10-38.
- Planchet F., Thérond P.E., Jacquemin J. (2005) *Modèles financiers en assurance. Analyses de risque dynamiques*, Paris : Economica.
- Ramezani C.A., Zeng Y. (1998) « Maximum likelihood estimation of asymmetric jump-diffusion processes: application to security prices », Working paper.

PARTIE II

MODÉLISATIONS AVANCÉES EN ASSURANCE

La deuxième partie de cette thèse est consacrée aux techniques avancées de gestion du risque d'une compagnie d'assurance.

L'avènement du référentiel prudentiel Solvabilité 2 et, dans une moindre mesure, du nouveau cadre comptable induit par la phase II de la norme IFRS dédiée aux contrats d'assurance, va systématiser l'emploi de la Value-at-Risk (VaR) en assurance. Particulièrement utilisées en finance de marché, les techniques de calcul de VaR exigent une adaptation spécifique dans un contexte assurantiel de par la nature des risques qui sont ainsi mesurés. Schématiquement on distingue deux contextes, qui imposent des approches différentes :

- La mesure du risque lié à la sinistralité au moyen d'une VaR dans le corps de la distribution : la provision¹ devra suffire à payer les sinistres dans 75 % des cas ;
- la mesure de risque liée à la ruine de la compagnie par le biais d'une VaR d'ordre très élevé : le capital de solvabilité devra permettre à la compagnie de ne pas être en ruine, à la fin de l'exercice, avec une probabilité supérieure à 99,5 %.

Dans la première situation, que l'on adopte une approche VaR historique ou que l'on cherche à modéliser la sinistralité, on reste dans un cadre dans lequel on dispose d'un matériel statistique de taille généralement suffisante pour estimer une VaR dans le corps de la distribution. Dans le second cas, on est en revanche confronté à l'absence d'observations et il convient alors, dans un premier temps, de modéliser les variables de base qui influent sur la solvabilité de la compagnie, dans un deuxième temps, de simuler la ruine de la compagnie et enfin d'estimer une VaR d'ordre élevé. Cette dernière étape nécessite le recours à la théorie des valeurs extrêmes. Le cinquième chapitre s'attache à présenter les contextes de calcul de VaR en assurance, les résultats probabilistes sous-jacents et les limites de ces critères.

Par ailleurs un autre phénomène à intégrer est celui de la dépendance entre les risques. L'activité d'assurance, qui repose sur la mutualisation et la

¹. Les travaux les plus récents du CEIOPS sur Solvabilité 2 (cf. l'annexe) donnent la préférence à une approche coût du capital plutôt qu'à la VaR à 75 % pour le calibrage de la marge pour risque.

convergence en probabilité énoncée dans la loi des grands nombres, s'est développée sur une des hypothèses principales de cette loi : l'indépendance entre les risques. Les exigences de solvabilité comme les référentiels économiques de valorisation ne font aujourd'hui plus l'économie de la non prise en compte des risques inhérents à la dépendance. Ainsi le capital de solvabilité fixé en référence au risque global supporté par la société ne saurait s'analyser comme la somme des capitaux qui permettraient de couvrir chaque risque individuellement. Par ailleurs, dans une approche de valorisation, si les marchés financiers n'accordent pas de prime à un risque mutualisable, il n'en est pas de même pour les risques non-mutualisables (ou systématiques) dont les sociétés d'assurance cherchent à se défaire au maximum que ce soit au travers de traités de réassurance ou, de plus en plus, grâce aux opérations de titrisation. Par exemple, AXA a récemment titrisé une partie du risque non-mutualisable sur son portefeuille automobile en France et du risque lié à l'évolution à moyen terme (5 ans) de la mortalité en Europe, aux États-Unis et au Japon.

Le sixième chapitre aborde cette problématique en présentant, dans un premier temps, les aspects mathématiques de la mesure et de la modélisation de la dépendance ; puis en analysant l'impact de la prise en compte de celle-ci dans les contextes assurantiels de détermination d'un capital de solvabilité et de valorisation du contrat d'épargne présenté dans le deuxième chapitre.

Cet exemple n'a pu être mis en œuvre que grâce à l'utilisation de méthodes de Monte Carlo. En effet, la complexité des contrats d'assurance et l'interdépendance intrinsèque et temporelle des phénomènes qui concourent à leur dénouement obligent les assureurs à avoir un recours de plus en plus systématique aux techniques de simulation. Par exemple, l'alternative à l'utilisation de la *formule standard* (cf. l'annexe) pour déterminer le capital de solvabilité passera par la mise en place d'un modèle interne. Un tel modèle revient à modéliser l'ensemble des variables qui influent sur le risque global supporté par la compagnie (cf. la Figure 20 au début du Chapitre 4). La mise en œuvre effective d'un tel modèle ne pourra faire l'économie du recours aux techniques de simulation. C'est justement l'objet du septième chapitre qui revient notamment sur les techniques de discrétisation temporelle des processus stochastiques continus, l'estimation des paramètres et la génération de nombres aléatoires. Les illustrations des différentes techniques en jeu s'appuient sur un modèle de taux d'intérêts.

Chapitre 4

Limites opérationnelles : la prise en compte des extrêmes

À l'instar de ce qu'a connu le monde bancaire ces dernières années avec les accords dits de Bâle II, le monde européen de l'assurance se prépare à se doter d'un nouveau référentiel prudentiel qui résultera du projet Solvabilité 2. Construit également sur une structure à trois piliers, le futur système de solvabilité vise à ce que les entreprises d'assurance mesurent et gèrent mieux les risques auxquels elles sont soumises. Une bonne gestion des risques pourra avoir pour conséquence une moindre exigence en terme de capitaux propres. Le premier pilier, consacré aux exigences quantitatives, fait largement référence à des mesures de risques bien connues des financiers, les Value-at-Risk (VaR). Les travaux en cours sur Solvabilité 2 prévoient que les assureurs devront disposer d'un montant de provisions techniques¹ leur permettant d'honorer leurs engagements de payer les prestations avec une probabilité de 75 % (à ce stade, une approche *cost of capital* est également envisagée en assurance-vie). Ils devront de plus disposer d'un niveau de fonds propres leur permettant de ne pas être en ruine, à horizon un an, avec une très forte probabilité (*a priori* 99,5 %). Pour mémoire, dans le système actuel, les montants de provisions sont, le plus souvent, déterminés par des méthodes déterministes prudentes (méthodes et hypothèses conservatrices d'estimation de la charge ultime de sinistres et absence d'actualisation ou actualisation à un taux résolument prudent). L'exigence minimale de fonds propres est quant à elle régie par le système de la marge de solvabilité : les assureurs doivent disposer d'un niveau de fonds propres et de plus-values latentes sur les placements financiers supérieur à une exigence de marge de solvabilité qui est exprimée en pourcentage des provisions mathématiques en assurance vie, et en pourcentage des primes perçues ou des sinistres payés en assurance non-vie. Ce système est donc entièrement déterministe et ne fait pas explicitement référence au risque réellement supporté par la compagnie. En particulier, deux entreprises qui ont des profils de risque très différents mais les mêmes éléments comptables se voient exiger un même niveau de capitaux propres². Il convient néanmoins de préciser que, d'un point de vue prudentiel, ce système s'avère relativement

¹. Le passif d'une société d'assurance est essentiellement composé des provisions techniques d'une part, et des fonds propres d'autre part.

². Ce constat est néanmoins à modérer de par l'intervention de l'autorité de contrôle (l'ACAM) qui peut relever les exigences de capitaux propres des assureurs.

efficace au vu du faible nombre de sociétés d'assurance européennes qui ont fait faillite au cours des dernières années.

Si elles font toutes les deux références à une VaR, les deux nouvelles exigences quantitatives prévues par Solvabilité 2 sont très différentes de par leur nature. L'exigence portant sur les provisions techniques ne va pas poser de problème pratique majeur. En effet, les assureurs disposent d'un nombre conséquent d'observations et de méthodes stochastiques de provisionnement robustes. De plus le quantile à estimer n'est pas très élevé, on se situe donc dans le cœur de la distribution pour lequel on dispose d'une information de qualité. Il en va différemment pour la détermination du niveau souhaitable de fonds propres ou *Solvency Capital Requirement* (SCR). En effet, l'assureur n'observe pas directement la variable d'intérêt (le résultat), il ne dispose donc pas de matériel statistique pour estimer directement la VaR.

Une formule standard³ sera proposée et aura pour vocation d'approcher autant que possible le critère de probabilité de ruine. Le projet prévoit, qu'en parallèle, les sociétés pourront construire un modèle interne qui, sur la base de la modélisation de l'ensemble des variables influant la solvabilité de la compagnie, permettra de simuler la situation de la société à horizon un an et ainsi de mesurer le niveau de fonds propres dont a besoin aujourd'hui la compagnie pour ne pas être en ruine, un an plus tard, avec une probabilité de 99,5 %. La construction d'un tel modèle est une problématique à part entière : elle nécessite la modélisation fine des variables de base et de leurs interactions. Un tel modèle doit de plus intégrer des contraintes techniques liées à la puissance de calcul informatique. Par ailleurs, le montant à estimer est un quantile extrême de la distribution. Une estimation empirique naïve robuste par le quantile empirique à 99,5 % nécessiterait un nombre de simulations extrêmement important, chacune de ces réalisations résultant de la simulation de différents phénomènes tels que la sinistralité, l'écoulement des provisions techniques, l'évolution des actifs financiers, etc. L'estimation d'un tel quantile devra donc passer par l'utilisation des techniques issues de la théorie des extrêmes qui s'est particulièrement développée depuis le milieu des années 70 et les travaux de Pickands (1975) et Hill (1975), et plus récemment de Smith (1987), Dekkers et de Haan (1989) ou encore Dekkers et al. (1989) et a été rapidement appliquée aux problématiques financières et assurantielles (Embrechts et al. (1997)). Or ces résultats utilisent les données de queue ; cela pose un nouveau problème, celui de la robustesse de ces données qui, rappelons-le, ne sont pas observées mais simulées. En effet, les modèles internes passeront, le plus souvent, par une modélisation paramétrique des variables de base (évolution du cours des actions par exemple), or ces modèles ne collent pas parfaitement aux observations. En particulier les paramètres sont estimés sur l'ensemble de la distribution au risque de ne pas bien représenter les queues de distribution. Dans la problématique qui nous intéresse, ce sont justement ces queues de distribution qui vont engendrer les situations les plus

³. Dans le cadre de l'étude d'impact quantitative QIS2 portant que les exigences de fonds propres auxquelles pourrait conduire le projet Solvabilité II, le Comité Européen des Contrôleurs d'Assurance et de Pensions Professionnelles (CEIOPS) a publié un modèle qui vise à la mesure de chaque risque et à l'agrégation des capitaux correspondants par une formule du type de celle du RBC américain.

catastrophiques en termes de solvabilité, i.e. les valeurs extrêmes sur lesquelles vont reposer l'estimation de la VaR à 99,5 %.

Ceci pose le problème de la robustesse de ce critère de VaR à 99,5 %, avec en ligne de mire le souci de la sécurité des assurés : les sociétés d'assurance pourraient être tentées, en jouant sur les paramètres ou sur les modèles des variables de base, de sous-estimer le SCR. Après avoir décrit les spécificités du calcul de VaR en assurance, nous présentons différentes méthodes d'estimation de quantiles extrêmes, puis nous explicitons et illustrons les limites de la mesure de risque proposée pour déterminer le SCR dans le projet Solvabilité 2. Dans un souci de clarté, les principaux résultats de la théorie des extrêmes utilisés dans ce travail sont présentés en annexe.

Ce chapitre s'appuie sur Thérond et Planchet (2007).

1. Calcul de VaR en assurance

La transposition de la VaR (et de la TVaR) aux problématiques d'assurance est récente et impose une approche radicalement différente de l'approche bancaire. En effet, la situation de référence du monde bancaire consiste à estimer la VaR sur un échantillon important de gains / pertes sur un portefeuille ou une position. Les données sont disponibles en quantité, avec une fréquence importante. Les approches de type VaR historique sont ainsi le socle sur lequel sont construits de nombreux modèles bancaires (cf. Christoffersen et al. (2001)). L'adaptation des techniques bancaires au portefeuille d'actifs d'un assureur nécessite un certain nombre d'aménagements, essentiellement pour tenir compte de la durée de détention plus importante. On pourra par exemple se reporter à Fedor et Morel (2006) qui présentent des analyses quantitatives sur ce sujet. Dans le cadre de la détermination du capital de solvabilité, de nouvelles difficultés apparaissent ; la nature des données disponibles invalidant l'approche historique. Il convient ici de revenir sur les situations d'assurance dans lesquelles on est amené à évaluer des VaR ; en pratique on peut en distinguer principalement deux, dont on va voir qu'elle imposent des approches différentes :

- le calcul d'une provision via un quantile à 75 % de la distribution des sinistres ;
- la détermination du niveau du capital de solvabilité (SCR) pour contrôler une probabilité de ruine en imposant que celle-ci soit inférieure à 0,5 % à horizon un an.

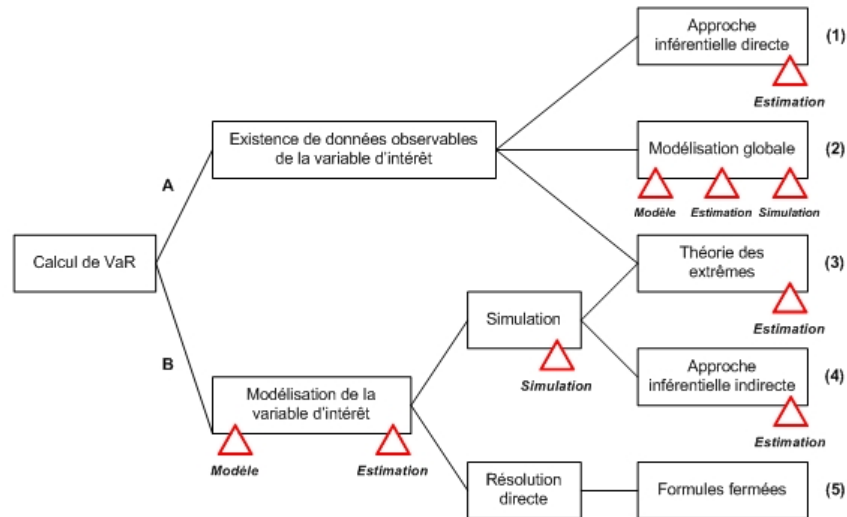
La première situation constitue une simple évolution par rapport à la situation qui prévaut actuellement ; en effet, les provisions sont aujourd'hui calculées sur la base d'une espérance (approche *best estimate*, cf. Blum et Otto (1998)). On conçoit que, du point de vue de la technique statistique, les deux calculs (calcul d'une espérance et calcul d'une VaR à 75 %) ne soient pas fondamentalement distincts. En particulier on dispose dans les deux cas de données permettant de mettre en oeuvre les techniques classiques de statistique inférentielle. En pratique, on propose souvent une modélisation de la charge des sinistres à l'aide d'un modèle paramétrique et tout se ramène alors à l'estimation des paramètres de la loi considérée (cf. Partrat et Besson (2005) pour une revue des lois les plus utilisées en pratique).

Pour ce qui concerne la détermination du capital de solvabilité (SCR), la situation est radicalement différente ; en effet, on ne dispose que de très peu de données directes (quelques années au plus d'observation du résultat par exemple) et d'aucune donnée dans la zone de la distribution concernée (si on en avait, l'assureur aurait fait faillite...). Alors que dans le cas précédent la variable d'intérêt était directement observable, la VaR à calculer est maintenant la résultante d'un modèle, souvent complexe, mettant en jeu les différents postes du bilan de l'assureur : prestations, provisions, actifs financiers, etc. La démarche va donc consister à construire un modèle en décrivant chaque poste du bilan ; en pratique l'approche consiste à modéliser le passif, dont le poste le plus important est constitué des provisions techniques, puis l'actif et enfin les éventuelles interactions actif / passif. Après avoir estimé les paramètres de ces modèles, on peut obtenir, en général par des techniques de simulation, une estimation de la loi du résultat ; de cette estimation, on déduira enfin le niveau de la VaR cherchée. La mise en pratique d'une telle démarche comporte à chaque étape des risques qu'il importe *a minima* d'identifier :

- un risque de modèle : le modèle utilisé n'offre qu'une représentation imparfaite de la réalité ; au surplus, les modèles usuellement employés tant à l'actif qu'au passif conduisent à sous-estimer les situations extrêmes ; nous reviendrons sur ce point potentiellement très pénalisant dans l'approche Solvabilité 2 ;
- un risque d'estimation : les paramètres estimés qui alimenteront le modèle sont entachés d'une erreur. Les conséquences de cette erreur peuvent être lourdes dans le cas d'un modèle peu robuste (voir par exemple Windcliff et Boyle (2004) pour une illustration dans le cas du modèle de Markowitz) ;
- un risque de simulation : la distribution du résultat sera, en général, obtenue par simulation, et ne sera donc qu'approchée.

De plus, s'agissant de l'estimation d'un quantile d'ordre élevé (VaR à 99,5 %), et la forme de la loi du résultat n'étant en général pas aisée à ajuster globalement à un modèle paramétrique, il faudra se tourner vers les techniques de valeurs extrêmes pour calculer finalement la mesure de risque, avec l'apparition ici d'un nouveau risque d'estimation, sur la queue de la distribution simulée. La Figure 20 propose une synthèse de ces différents risques et des situations dans lesquelles on les rencontre.

Figure 20 - Typologie des différents risques rencontrés



Typiquement les études de variations extrêmes de variables financières classiques s'inscrivent dans le cadre des branches A3 ou A1 : on dispose d'un nombre important d'observations et les variables sont relativement régulières (elles possèdent généralement des moments d'ordre 3 ou 4). En assurance, l'estimation d'une VaR à 75 % dans le cadre de la détermination d'une provision Solvabilité 2 par exemple, correspondra, dans le meilleur cas, à la branche A1 pour les risques pour lesquels on dispose de beaucoup d'observations (risques non-vie à forte fréquence) et le plus fréquemment aux branches B4 ou B5 (on se donne un modèle, paramétrique le plus souvent, puis selon le cas on estime les paramètres et l'on en déduit la VaR ou l'on procède par les méthodes de Monte Carlo).

L'estimation du SCR procède, quant à elle, de la branche B3 qui cumule les risques de modèle, de simulation et d'estimation (à deux reprises). Les risques d'estimation et de simulation ne sont pas spécifiques à la problématique du présent travail, et les techniques propres à les contrôler sont robustes et accessibles.

En revanche, le risque de modèle prend ici une importance particulière ; de fait, et comme on l'illustre infra, la mesure correcte du capital de solvabilité impose de refondre l'ensemble de la modélisation des postes du bilan pour prendre en compte correctement les événements extrêmes et éviter ainsi une sous-estimation du capital de solvabilité.

2. Notations

Considérons un n -échantillon de variables aléatoires (v.a.) $X_1, \dots, X_n$ indépendantes et identiquement distribuées (i.i.d.) de fonction de répartition F et de fonction de survie $\bar{F} : x \mapsto 1 - F(x)$. La statistique d'ordre associée sera notée $X_{n,n}, \dots, X_{1,n}$ et est définie par

$$\min\{X_1, \dots, X_n\} = X_{n,n} \leq X_{n-1,n} \leq \dots \leq X_{1,n} = \max\{X_1, \dots, X_n\}.$$

Par ailleurs, on notera F_u la fonction de répartition de la v.a. ${}_uX$ qui représente le surplus au-delà du seuil u de X , lorsque X dépasse le seuil u , soit

$$F_u(x) := \Pr[{}_uX \leq x] = \Pr[X - u \leq x \mid X > u].$$

On notera de plus x_F la borne supérieure du support de X de fonction de répartition F , soit

$$x_F := \sup\{x \in \mathbb{R} \mid F(x) < 1\}.$$

Enfin, nous noterons N_u le nombre d'observations qui dépassent le seuil u , soit

$$N_u = \sum_{i=1}^n 1_{X_i > u}.$$

3. Estimation de quantiles extrêmes

Plaçons-nous dans la situation standard où l'on observe directement des réalisations du phénomène dont on souhaite déterminer la valeur qu'il ne dépassera qu'avec une faible probabilité. L'objet de ce paragraphe est de préciser les différentes méthodes d'estimation possibles et de mesurer les erreurs d'estimation commises relativement à la quantité de données disponibles.

Les différentes méthodes sont illustrées à partir d'un échantillon simulé de réalisations d'une variable aléatoire de loi de Pareto de première espèce.

3.1. Estimation naturelle

Comme $X_{k,n}$ est un estimateur naturel du quantile d'ordre $1 - (k-1)/n$, un estimateur naturel de $F^{-1}(p)$ ($p \in]0; 1[$) est

$$([pn] - pn + 1)X_{n-[pn]-1,n} + (pn - [pn])X_{n-[pn],n},$$

où $[.]$ désigne l'opérateur partie entière.

Cette méthode nécessite pour être efficace en pratique un volume de données jamais disponible dans notre contexte.

3.2. Ajustement à une loi paramétrique

Une méthode naturelle consiste à ajuster l'ensemble des données à une loi paramétrique puis à estimer la fonction quantile au niveau de probabilité souhaité. En effet, on rappelle que si $\hat{\theta}$ est l'estimateur du maximum de vraisemblance (e.m.v.) du paramètre θ de la loi de X , alors $F_{\hat{\theta}}^{-1}(p)$ est l'e.m.v. de $F_{\theta}^{-1}(p)$. De plus, les méthodes quantiles bootstrap BC (BCa) (cf.

Efron et Tibshirani (1993) et le paragraphe 4) permettent d'obtenir des intervalles de confiance de l'estimation.

Cette méthode d'estimation se décompose en deux étapes :

- ajustement à une loi paramétrique : choix d'une ou plusieurs lois, estimation des paramètres, tests d'adéquation ;
- inversion de la fonction de répartition (analytiquement quand c'est possible, numériquement sinon).

La principale difficulté de cette méthode consiste en le choix des lois paramétriques qui vont être utilisées. Elles doivent répondre à la double contrainte de permettre une évaluation numérique de leur fonction quantile aux ordres souhaités et doivent être en adéquation avec les données. Ce dernier point est particulièrement problématique dans la mesure où l'estimation des paramètres est effectuée sur l'ensemble de la distribution observée et représente rarement bien les valeurs extrêmes. À moins d'être assuré que la loi choisie pour l'ajustement est la véritable loi d'où sont issues les données, ce qui est exceptionnel en pratique, l'intérêt de cette méthode apparaît très limité dans un contexte d'estimation de quantiles extrêmes.

3.3. Approximation GPD

Connue sous le nom de méthode *Peaks Over Threshold* (POT), cette technique repose sur le théorème de Pickands-Balkema-de Haari (cf. l'annexe B.3) qui établit que, pour u assez grand, F_u peut être approximée par une distribution Pareto généralisée (GPD). Une fois les paramètres de la GPD estimés (cf. l'annexe A3), comme $\bar{F}(x) = \bar{F}(u) \times \bar{F}_u(x-u)$, on dispose de l'approximation suivante :

$$\bar{F}(x) \approx \frac{N_u}{n} \left(1 + \frac{\hat{\xi}}{\hat{\beta}} (x-u) \right)^{-1/\hat{\xi}},$$

pour $x > u$. L'inversion de cette fonction de répartition nous fournit un estimateur de la fonction quantile F^{-1} aux ordres élevés :

$$\hat{x}_p = u + \frac{\hat{\xi}}{\hat{\beta}} \left(\left(\frac{n}{N_u} (1-p) \right)^{-\hat{\xi}} - 1 \right).$$

N.B. Si l'on choisit comme seuil une des observations $X_{k+1,n}$ (par exemple le quantile empirique à 99 % pour estimer le quantile à 99,5 %) alors $N_u = k$ et l'estimateur se réécrit de la manière suivante :

$$\hat{x}_p = X_{k+1,n} + \frac{\hat{\xi}}{\hat{\beta}} \left(\left(\frac{n}{k} (1-p) \right)^{-\hat{\xi}} - 1 \right),$$

pour $p > 1 - k/n$.

Cette méthode présente en pratique la difficulté du choix du seuil u . En effet, u doit être assez grand pour que l'approximation GPD soit bonne mais pas trop proche du quantile à estimer de manière à ce que l'on dispose de suffisamment de matériel statistique pour que l'estimation de la probabilité d'être entre u et ce quantile (estimation qui fait apparaître N_u) soit fiable. La problématique posée par le choix de u est similaire à celle du choix du nombre d'observations à utiliser dans le cadre de l'estimation de Hill de l'épaisseur de la queue de distribution (cf. l'annexe C.2.4). Une approche pratique visant à choisir le seuil u consiste à tracer

$$e_n(u) := \frac{1}{N_u} \sum_{i=1}^n [X_i - u]^+,$$

l'estimateur empirique de l'espérance résiduelle de X et à choisir u de manière à ce que $e_n(x)$ soit approximativement linéaire pour $x \geq u$. En effet, comme la fonction espérance résiduelle d'une loi GPD de paramètre $\xi < 1$ est linéaire :

$$e(v) := E[vX] = \frac{\beta + \xi v}{1 - \xi},$$

on cherchera un u aussi petit possible sous la contrainte que l'estimateur empirique de l'espérance résiduelle des excès au delà de u soit approximativement linéaire.

Il ne faut pas attendre de cette technique la bonne valeur de u , elle fournit néanmoins une aide précieuse. En pratique, plusieurs valeurs de u doivent être testées.

3.4. Estimateur de Hill

Considérons les fonctions de répartition appartenant au domaine d'attraction maximum (DAM) de Fréchet ($\xi > 0$). Le théorème taubérien (cf. l'annexe B.2) nous indique que $\bar{F}(x) = x^{-1/\xi} \mathcal{L}_F(x)$, où $\mathcal{L}_F$ est une fonction à variation lente. Ainsi, on a

$$\frac{\bar{F}(x)}{\bar{F}(X_{k+1,n})} = \frac{\mathcal{L}_F(x)}{\mathcal{L}_F(X_{k+1,n})} \left(\frac{x}{X_{k+1,n}} \right)^{-1/\xi}.$$

Pour k suffisamment petit et $x > X_{k+1,n}$, on peut négliger (au premier ordre) le rapport des fonctions à variation lente ; il vient alors

$$\bar{F}(x) \approx \bar{F}(X_{k+1,n}) \left(\frac{x}{X_{k+1,n}} \right)^{-1/\xi}.$$

D'où l'on déduit l'estimateur de F :

$$\hat{F}(x) = 1 - \frac{k}{n} \left(\frac{x}{X_{k+1,n}} \right)^{-1/\xi_k^H},$$

pour $x > X_{k+1,n}$. L'inversion de cette fonction nous donne l'estimateur de Hill de la fonction quantile pour des ordres élevés ($p > 1 - k/n$) :

$$\hat{x}_p^H = X_{k+1,n} \left(\frac{n}{k} (1-p) \right)^{-\frac{1}{\xi_p^H}}.$$

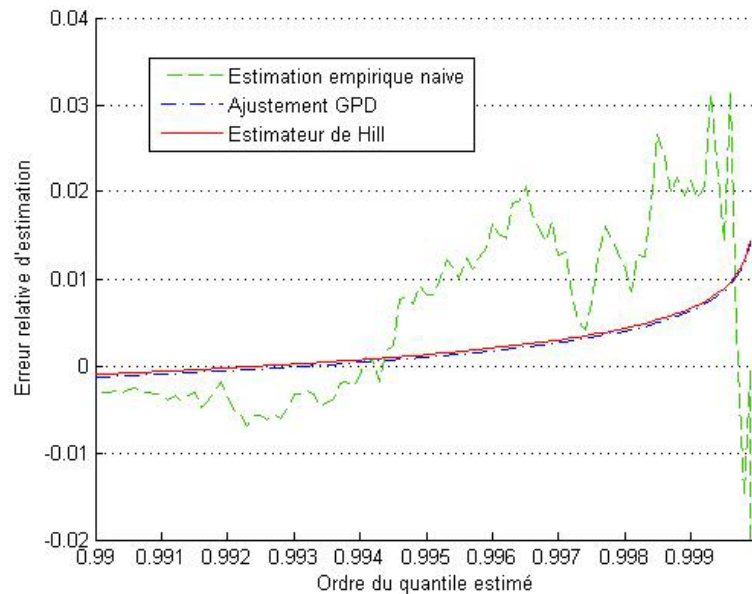
Le problème du choix de k est similaire à celui rencontré dans l'estimation du paramètre de queue ξ (cf. l'annexe C.2.4).

3.5. Illustration

La Figure 21 illustre les différentes techniques dans le cadre de l'estimation d'un quantile extrême d'une loi de Pareto de première espèce, i.e. de fonction de répartition

$$F(x) = 1 - \left(\frac{x_0}{x} \right)^\alpha, \text{ pour } x > x_0.$$

Figure 21 - Estimation d'un quantile extrême : erreur relative d'estimation



Une telle loi appartient au DAM de Fréchet. La méthode POT (exacte dans le cas d'une loi de Pareto de première espèce) et l'estimateur de Hill donnent des résultats comparables et nettement plus précis que l'estimateur naturel.

4. Application du bootstrap

La méthode du bootstrap proposée initialement par Efron (1979) est devenue d'application courante en assurance (cf. Partrat et Jal (2002) et Verall

et England (1999)). Plus généralement, le lecteur intéressé trouvera une revue des applications de la méthode bootstrap en économétrie dans Horowitz (1997).

Dans le contexte d'estimation de valeurs extrêmes, son application nous permet de disposer d'intervalles de confiance utiles à l'appréciation de la robustesse des estimations.

4.1. Présentation

Le principe de la méthode classique consiste à remarquer que, pour un échantillon de taille suffisante, la fonction de répartition F de la loi sous-jacente peut être approchée par la fonction de répartition empirique :

$$F_n(x) = \frac{1}{n} \sum_{i=1}^n 1_{\{X_n \leq x\}}.$$

En pratique ce principe très général conduit à décliner des méthodes bootstrap adaptées à différents contextes : séries chronologiques, modèles linéaires, modèles de régression non linéaires, etc. On en présente ci-après le principe dans le cas simple de la construction d'un intervalle de confiance pour un estimateur de tendance centrale.

Pour évaluer une fonctionnelle de F de la forme $I(g) = \int g(x) dF(x)$, l'approche par simulation classique suggère de calculer

$$I_n(g) = \int g(x) dF_n(x) = \frac{1}{n} \sum_{i=1}^n g(X_i),$$

qui constitue une approximation de $I(g)$. La théorie asymptotique fournie des informations sur la loi de la statistique $I_n(g)$: le théorème central limite permet en effet de prouver que la loi asymptotique est normale et d'en déduire, par exemple, des intervalles de confiance de la forme

$$J_\alpha = \left[I_n(g) - \frac{\sigma_g}{\sqrt{n}} \varphi^{-1} \left(1 - \frac{\alpha}{2} \right), I_n(g) + \frac{\sigma_g}{\sqrt{n}} \varphi^{-1} \left(1 - \frac{\alpha}{2} \right) \right],$$

avec $\sigma_g^2 = \text{Var}(g(X))$ qu'il s'agit d'estimer. Toutefois dans certaines situations l'approximation asymptotique n'est pas suffisamment précise, et la méthode bootstrap fournit une approche alternative pour obtenir des informations sur la loi de $I_n(g)$.

Cette technique consiste à considérer des réalisations $I_n^b(g)$, pour $b = 1, \dots, B$ où $B \leq n^n$ obtenues en remplaçant l'échantillon initial $(X_1, \dots, X_n)$ par les échantillons « bootstrapés » $(X_1^b, \dots, X_n^b)$ obtenus en effectuant des tirages avec remise dans $(X_1, \dots, X_n)$. On dit que B est la taille de l'échantillon bootstrap.

En effet, évaluer des statistiques par simulation se ramène alors à générer des échantillons à l'aide de la distribution empirique. Or un tirage dans la distribution empirique s'obtient simplement par un tirage avec remise des n valeurs dans l'échantillon initial. On obtient ainsi au plus n^n échantillons « bootstrapés » à partir desquels on va calculer les estimateurs empiriques des statistiques d'intérêt.

Le bootstrap permet ainsi, à partir d'un tirage aléatoire au sein de l'échantillon d'origine, de le perturber afin d'obtenir de nouveaux estimateurs des paramètres.

Une fois que l'on dispose de B estimateurs $I_n^b(g)$ de $I(g)$, en posant

$$\hat{\mu}[I_n^1(g), \dots, I_n^B(g)] = \frac{1}{B} \sum_{b=1}^B I_n^b(g),$$

et

$$\hat{\sigma}^2[I_n^1(g), \dots, I_n^B(g)] = \frac{1}{B-1} \sum_{b=1}^B [I_n^b(g) - \hat{\mu}]^2,$$

on obtient un intervalle de confiance de la forme

$$J_\alpha = \left[\hat{\mu} - \frac{\hat{\sigma}}{\sqrt{B}} \Phi^{-1}\left(1 - \frac{\alpha}{2}\right), \hat{\mu} + \frac{\hat{\sigma}}{\sqrt{B}} \Phi^{-1}\left(1 - \frac{\alpha}{2}\right) \right].$$

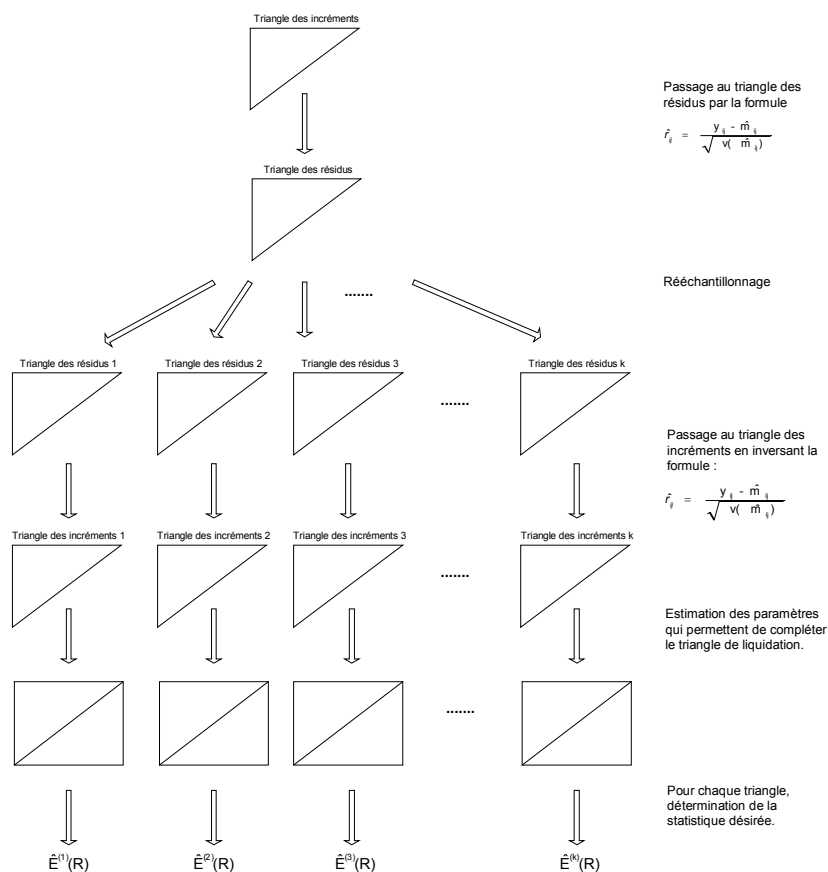
Deux techniques alternatives permettent d'affiner les résultats : le bootstrap percentile et la méthode BCa (cf. Efron et Tibshirani (1993)). Ces techniques sont décrites *infra* dans le cas particulier de l'estimation de quantiles.

La méthode bootstrap n'est toutefois pas pertinente dans tous les cas de figures : elle ne fonctionne par exemple pas lorsque l'on s'intéresse à la valeur maximum de l'échantillon. En effet, par construction, il est impossible de constituer des échantillons contenant des valeurs supérieures à la valeur maximale de l'échantillon de départ. Pour ce type de problème, d'autres méthodes existent telle que celle proposée par Zelterman (1993) qui consiste à rééchantillonner les intervalles entre deux valeurs adjacentes des données de départ plutôt que les données elles-mêmes. Cette méthode semi-paramétrique est illustrée dans le cas de l'estimation par intervalle de l'intensité maximale des séismes dans différentes régions par Bottard (1996).

Tel quel, le bootstrap est en revanche très efficace dans toutes les statistiques basées de près ou de loin sur la moyenne. Les calculs de capital de solvabilité ne correspondent pas à cette configuration.

La technique du bootstrap est communément utilisée en assurance pour analyser la variabilité des montants de sinistres et obtenir des erreurs de prédiction pour différentes méthodes de provisionnement, et notamment pour les méthodes basées sur Chain Ladder et sur les modèles linéaires généralisés (cf. Partrat et Jal (2002)). Le schéma général (repris de Planchet, Thérond et al. (2005)) est résumé par la Figure 22.

Figure 22 - Méthode bootstrap en provisionnement non-vie



4.2. Calcul d'un intervalle de confiance pour une VaR

Les différentes techniques de bootstrap sont présentées dans ce paragraphe dans le contexte spécifique de l'estimation par intervalle d'un quantile.

Dans le cas classique, où l'on suppose que les observations sont des réalisations de variables aléatoires dont la loi sous-jacente appartient à une famille paramétrique (gaussienne, log-normale, Pareto, Benktander, etc.), l'estimation d'un quantile d'ordre quelconque peut être envisagée de la manière suivante :

- estimation $\hat{\theta}$ du paramètre θ (par maximum du vraisemblance quand c'est possible ou par une autre méthode) ;
- inversion de la fonction de répartition ;
- estimation ponctuelle de la VaR par $\hat{VaR}(q) = F_{\hat{\theta}}^{-1}(q)$;
- recherche d'un intervalle de confiance.

On s'intéresse à la précision de l'estimateur ainsi obtenu. Dans le cas particulier où le paramètre est estimé par l'estimateur du maximum de vraisemblance (e.m.v.), on déduit des propriétés générales de cet estimateur que $\hat{VaR}(q)$ est l'e.m.v. de la VaR (comme fonction d'un e.m.v.). Il est donc asymptotiquement sans biais et gaussien. En obtenant une estimation de sa variance asymptotique, il sera donc possible de construire un intervalle de confiance.

La loi de la statistique $\hat{VaR}(q)$ est toutefois difficile à déterminer. Dans ce contexte, il est conseillé de se tourner vers la méthode bootstrap.

Dans le cas de l'estimation d'une VaR, l'échantillon initial est constitué par les n observations de X utilisées pour estimer les paramètres du modèle. La statistique d'intérêt est :

$$\hat{VaR}(q) = F_{\hat{\theta}}^{-1}(q).$$

On cherche à construire un intervalle de confiance de la forme :

$$\Pr(b_1 \leq \hat{VaR}(q) \leq b_2) = 1 - \alpha.$$

On distingue trois techniques bootstrap permettant de construire l'intervalle de confiance recherché :

- bootstrap classique : estimation de la variance bootstrapée de $\hat{VaR}(q)$ que l'on utilise ensuite avec l'hypothèse de normalité asymptotique ;
- bootstrap percentile : classement des estimations ponctuelles obtenues pour chaque échantillon bootstrap et constitution des bornes de l'intervalle de confiance par les estimations correspondantes ;
- estimation directe de l'intervalle de confiance bootstrapé via la méthode Bias corrected and accelerated (BCa) (cf. Efron et Tibshirani (1993)).

Dans la première approche, l'estimation de la variance bootstrapée consiste simplement à calculer la variance empirique de l'échantillon des $\hat{VaR}_b(q) = F_{\hat{\theta}_b}^{-1}(q)$ avec $\hat{\theta}_b$ le paramètre estimé sur le b -ème échantillon bootstrap. On obtient alors aisément un intervalle de confiance en utilisant la normalité asymptotique de l'estimateur du maximum de vraisemblance. En pratique, cette méthode n'est pas recommandée pour les problèmes non-paramétriques (pour lesquels l'estimation, pour chaque échantillon, du quantile reposerait sur son estimateur empirique).

Une méthode plus performante et plus adaptée aux problèmes non-paramétriques est le bootstrap percentile. Cette technique consiste à classer les $\hat{VaR}_b(q)$ par ordre croissant et à prendre comme bornes de l'intervalle de confiance les $B \times \alpha/2$ et $B \times (1 - \alpha/2)$ -èmes plus grandes valeurs. Cette technique est bien adaptée aux problèmes non-paramétriques car elle ne recourt pas aux propriétés asymptotiques de l'e.m.v.

La méthode BCa (cf. Efron (1987)) se propose d'optimiser la construction de l'intervalle de confiance, comme la méthode percentile sans s'appuyer sur les propriétés asymptotiques de l'e.m.v. Cette méthode consiste à déterminer les bornes de l'intervalle de confiance en prenant pour b_i la valeur d'ordre $B \times \beta_i$ de l'échantillon des $\widehat{VaR}_b(q)$, où

$$\beta_1 = \Phi \left(z_0 + \frac{z_0 + u_{\alpha/2}}{1 - \gamma(z_0 + u_{\alpha/2})} \right) \text{ et } \beta_2 = \Phi \left(z_0 + \frac{z_0 + u_{1-\alpha/2}}{1 - \gamma(z_0 + u_{1-\alpha/2})} \right),$$

avec Φ la fonction de répartition de la loi normale centrée réduite, u_c le quantile d'ordre c de cette même loi, le paramètre de correction de biais $z_0 = \Phi^{-1}(k)$ où k est la proportion des échantillons bootstrapés pour lesquels la VaR estimée est inférieure à la VaR estimée sur l'ensemble des échantillons. Enfin le paramètre d'accélération γ est estimé par :

$$\gamma = \frac{\sum_{i=1}^n (\widehat{VaR}_i - \widehat{VaR})^3}{6 \left[\sum_{i=1}^n (\widehat{VaR}_i - \widehat{VaR})^2 \right]^{3/2}},$$

où $\widehat{VaR}_i$ est l'estimation de la VaR provenant de l'échantillon Jackknife n° i (échantillon observé auquel on a retiré l'observation i) et $\widehat{VaR}$ est la moyenne de ces n estimations (cf. Quenouille (1949)).

4.3. Illustration numérique

On illustre numériquement ces méthodes dans le cas d'une loi log-normale. Dans ce cas il est immédiat de vérifier que :

$$VaR_p(X) = F^{-1}(p) = \exp\left(m + \sigma\Phi^{-1}(p)\right),$$

et l'on obtient donc comme estimateur du maximum de vraisemblance de la VaR d'ordre p :

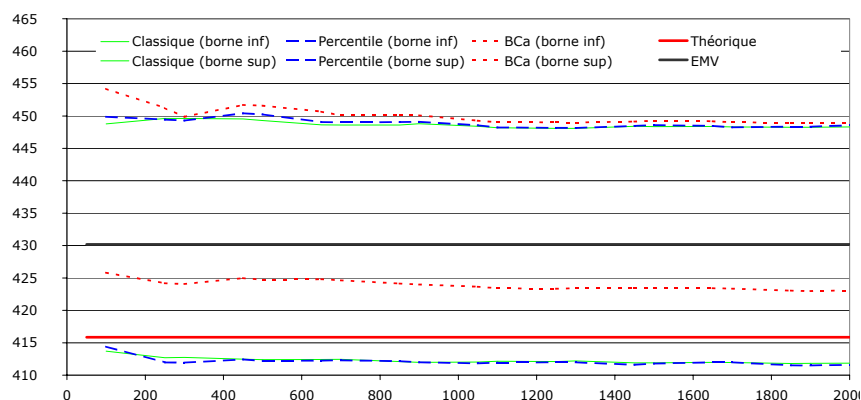
$$\widehat{VaR}_p(X) = \exp\left(\hat{m} + \hat{\sigma}\Phi^{-1}(p)\right),$$

$$\text{où } \hat{m} = \frac{1}{n} \sum_{i=1}^n X_i \text{ et } \hat{\sigma} = \sqrt{\frac{1}{n} \sum_{i=1}^n (X_i - \hat{m})^2}.$$

On retient les valeurs numériques suivantes : $m = 5$, $\sigma = 0,4$, ce qui conduit à une VaR théorique à 99,5 % de

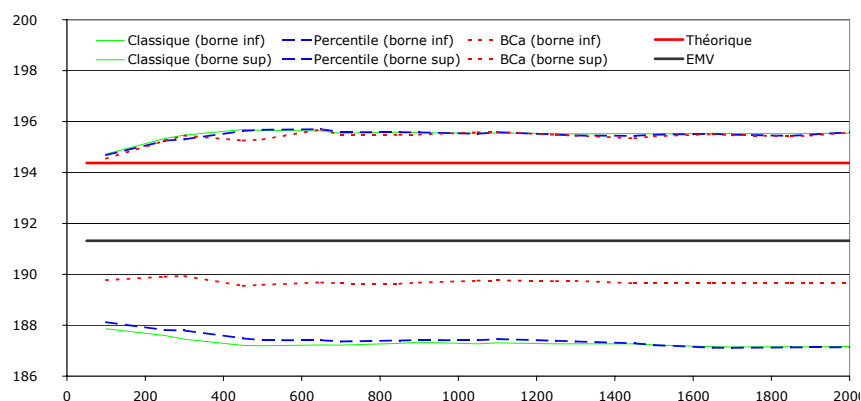
$$VaR_{99,5\%}(X) = \exp\left(m + \sigma\Phi^{-1}(p)\right) = 415,85.$$

On estime ce quantile sur la base d'un échantillon de taille 1 000 et on détermine les intervalles de confiance à 90 % avec les méthodes classique, percentile et BCA. On obtient en fonction de la taille de l'échantillon bootstrap les résultats typiques suivants :

Figure 23 - Intervalles de confiance (VaR à 99,5 %)

On note que les méthodes classique et percentile sont très proches. La méthode Bca fournit un intervalle dont la borne supérieure est proche des deux autres méthodes, mais avec une borne inférieure plus élevée. Dans l'exemple ci-dessus on constate par exemple que si la vraie valeur est bien dans les intervalles de confiance « classique » et « percentile », elle est inférieure à la borne inférieure de l'intervalle fournie par la méthode BCa.

Les résultats obtenus par les méthode bootstrap restent peu robustes pour un quantile élevé. Pour un quantile d'ordre inférieur, comme par exemple le quantile à 75 %, les estimations deviennent nettement plus fiables :

Figure 24 - Intervalles de confiance (VaR à 75 %)

5. Robustesse du SCR

L'objet de ce paragraphe est de préciser les limites d'un critère de type VaR extrême pour fixer le capital de solvabilité d'une société d'assurance. En effet la complexité du dispositif à mettre en place rend la pertinence des résultats issus du modèle interne toute relative. Deux points sont plus particulièrement développés et illustrés à l'aide d'un exemple simple : l'erreur d'estimation des

paramètres des variables de base puis l'erreur de spécification de modèle de ces variables.

5.1. Estimation des paramètres des variables de base

Considérons le modèle interne simplifié inspiré des travaux de Deelstra et Janssen (1998) dans lequel les sinistres de l'année sont modélisés par une variable aléatoire de distribution log-normale et payés en fin d'année. Supposons de plus que le rendement des actifs financiers sur la période soit gaussien et indépendant de la sinistralité. Formellement, notons $B \sim \mathcal{LN}(m, s)$ la v.a. modélisant le montant à payer en fin de période et $\rho \sim \mathcal{N}(\mu, \sigma)$ le rendement financier sur la période.

Déterminons le montant minimal d'actif a_0 dont doit disposer l'assureur en début de période pour ne pas être en ruine, en fin de période, avec une probabilité supérieure à $1 - \alpha$. Formellement, a_0 est la solution d'un problème d'optimisation :

$$a_0 = \min \left\{ a > 0 \mid \Pr[ae^\rho \geq B] > 1 - \alpha \right\}.$$

Comme $Be^{-\rho} \sim \mathcal{LN}(m - \mu, \sqrt{s^2 + \sigma^2})$, on dispose d'une expression analytique pour la valeur de a_0 :

$$a_0 = \exp \left\{ m - \mu + \Phi^{-1}(1 - \alpha) \sqrt{s^2 + \sigma^2} \right\}.$$

Étudions à présent le sensibilité de a_0 aux paramètres des modèles de base. Par exemple pour le passif, on a :

$$\frac{1}{a_0} \frac{\partial a_0}{\partial m} = 1,$$

et

$$\frac{1}{a_0} \frac{\partial a_0}{\partial s} = \frac{\Phi^{-1}(1 - \alpha)}{\sqrt{1 + \sigma^2/s^2}}.$$

Ainsi une erreur relative d'estimation du paramètre m conduit à la même erreur relative sur a_0 . Par ailleurs, une erreur relative de 1 % sur s conduit à une erreur relative de $\Phi^{-1}(1 - \alpha) / \sqrt{1 + \sigma^2/s^2}$ sur a_0 . Cela correspond pour une VaR à 99,5 %, lorsque $s \approx \sigma$, à une erreur 1,82 fois plus grande.

5.2. Simulation

La complexité inhérente à la modélisation du résultat d'une activité d'assurance rend incontournable le recours aux techniques de simulation pour obtenir des résultats numériques. Si le principe de base de ces méthodes est simple et universel (utilisation de la convergence forte de la Loi des grands

nombres), une mise en œuvre efficace nécessite de prendre quelques précautions. En effet, l'utilisation de résultats approchés par simulation est source de différentes erreurs :

- fluctuations d'échantillonnage liées au nombre fini de tirages effectués ;
- biais de discrétisation lors de la transformation d'un modèle continu dans sa version discrète ;
- erreurs associées aux approximations utilisées par certaines techniques d'inversion ;
- biais induits par un choix mal approprié du générateur de nombres aléatoires.

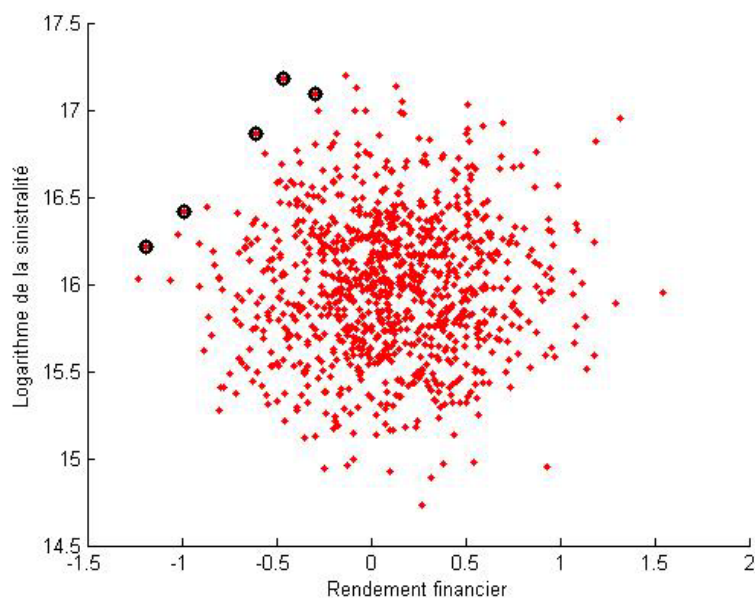
Aussi, les algorithmes utilisés devront permettre un contrôle quantitatif de l'ensemble de ces sources d'erreurs afin de pouvoir calibrer le nombre de tirages nécessaire pour le degré de précision (un niveau de confiance étant fixé *a priori*) requis.

5.3. Spécification du modèle

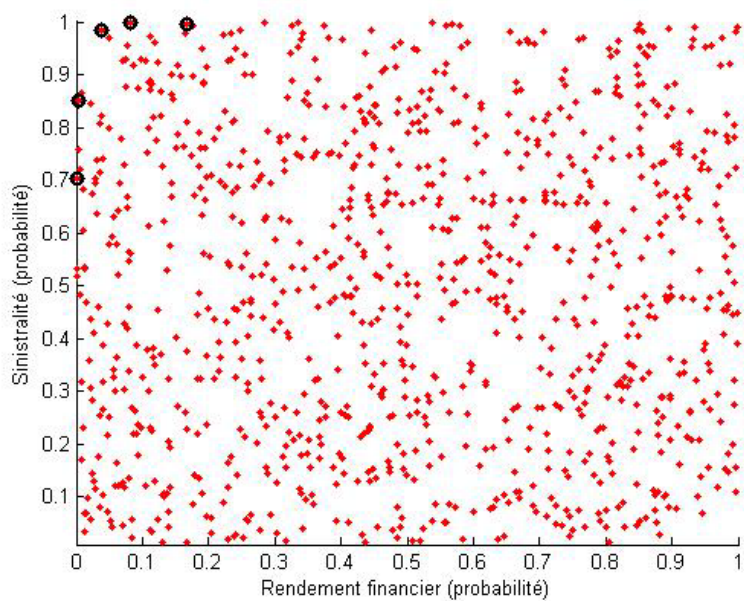
5.3.1. Contexte et motivation

La construction d'un modèle interne nécessite la modélisation des différentes variables influant sur la solvabilité de l'assureur. Pour ce faire, la première étape consiste en la modélisation de ces différentes variables. Dans le cas d'ajustements paramétriques, il est naturel d'effectuer l'estimation des paramètres et les tests d'adéquation sur l'intégralité des données disponibles. Cependant, dans notre problématique, ce sont les queues de distribution qui vont influencer sur le niveau du SCR. Or celles-ci sont souvent mal décrites par l'approche globale utilisant une loi paramétrique simple unique. En particulier, la plupart des modèles usuels (log-normal, Benktander, etc.) conduisent à des lois dans le DAM de Gumbel alors que les observations tendent à opter pour des queues de distribution dans le DAM de Fréchet. En effet, les modèles stochastiques qui fonctionnent actuellement chez les assureurs ont initialement été créés pour effectuer des études de rentabilité ou calculer des provisions techniques ; les résultats étant, le plus souvent, appréciés par des critères de type espérance-variance (ou VaR dans le cœur de la distribution – 75 % par exemple), la modélisation des événements extrêmes a dans ce contexte une influence toute relative. Dans l'approche Solvabilité 2, on ne s'intéresse quasiment plus qu'aux valeurs extrêmes de la variable d'intérêt : le capital dont on doit disposer aujourd'hui pour ne pas être en ruine dans un an. Cette remarque est à nuancer par le fait que la plupart des assureurs qui investissent dans le développement de modèles internes escomptent, qu'à quelques aménagements près, celui-ci leur permettra d'affiner leur plan stratégique et de répondre aux futures exigences comptables résultant du passage à la phase II de la norme dédiée aux contrats d'assurance.

Considérons le modèle présenté ci-dessus. La Figure 25 présente un échantillon de 1000 réalisations du montant d'actif dont doit disposer la société en 0 pour ne pas être en ruine en 1.

Figure 25 - Modèle interne simplifié : identification des valeurs extrêmes

Les cinq points soulignés correspondent aux cinq scénarios qui ont conduits aux valeurs maximales.

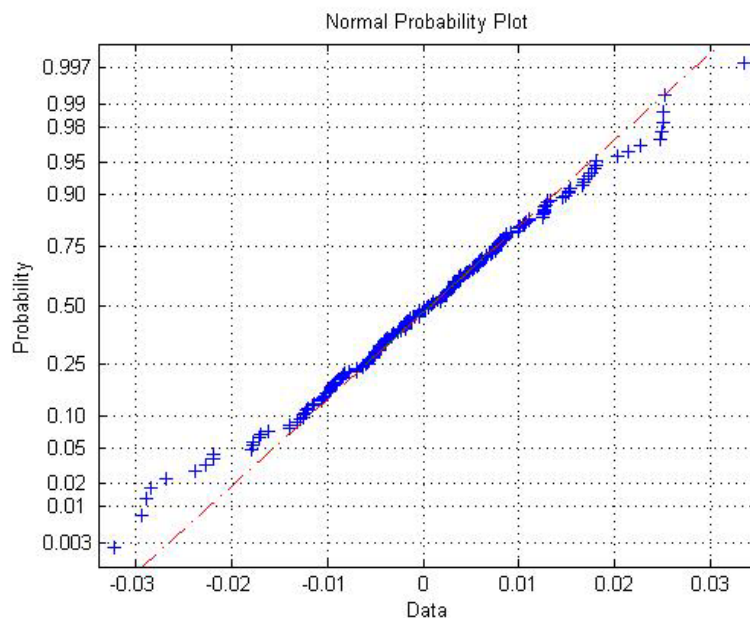
Figure 26 - Modèle interne simplifié : identification des valeurs extrêmes (probabilités)

On remarque sur la Figure 26, qui représente les probabilités associées à chaque grandeur, que ces points se situent dans la queue de distribution d'au moins une des deux variables de base. Or justement aux queues de distribution, l'adéquation des données au modèle retenu est souvent imparfaite.

5.3.2. Modélisations avancées

Considérons la Figure 27 qui représente le QQ-plot normal du rendement journalier du titre TOTAL (de juillet 2005 à avril 2006).

Figure 27 - Rendement journalier du titre TOTAL : QQ-plot loi empirique vs loi normale



On peut d'ores et déjà observer graphiquement que l'ajustement semble globalement satisfaisant mais que les queues de distribution du modèle (en particulier celle des rendements négatifs pour ce qui nous intéresse) sont trop fines. Les tests statistiques de Jarque-Béra et de Lilliefors au seuil 5 % conduisent en effet à ne pas rejeter l'hypothèse de normalité du rendement. Cependant, dans le cadre de la détermination d'un SCR sur le fondement d'un critère VaR à 99,5 %, le modèle gaussien tendrait à minimiser le risque pris en investissant sur ce placement. Par exemple, sur les données du titre TOTAL, le quantile à 0,5 % du rendement quotidien observé vaut $-0,0309$ alors que le quantile du même ordre du modèle gaussien ajusté vaut $-0,0286$, ce qui représente une erreur de l'ordre de 7,5 %. Pour remédier à cela, dans les situations où l'on dispose d'un nombre conséquent d'observations, on pourrait être tenté d'adopter une approche non-paramétrique en utilisant la distribution empirique. Cela n'est toutefois pas satisfaisant en pratique, car de telles approches sont techniquement plus difficiles à implémenter, nécessitent des temps de simulation plus longs et pénalisent une compréhension simplifiée du

modèle. Dans ce contexte, il est naturel de conserver une approche paramétrique et de rechercher un modèle qui représente mieux les queues de distributions.

Le modèle suivant est une extension naturelle du rendement gaussien.

Supposons que le rendement de l'actif sur la période soit régi par le processus suivant

$$r = \mu + \sigma_0 \varepsilon_0 + \sigma_u \sum_{i=1}^N \varepsilon_i ,$$

où les $\varepsilon_0, \varepsilon_1, \dots$ sont des v.a.i.i.d. de loi $\mathcal{N}(0;1)$ et indépendantes de $N \sim \mathcal{P}(\lambda)$ ($\lambda > 0$).

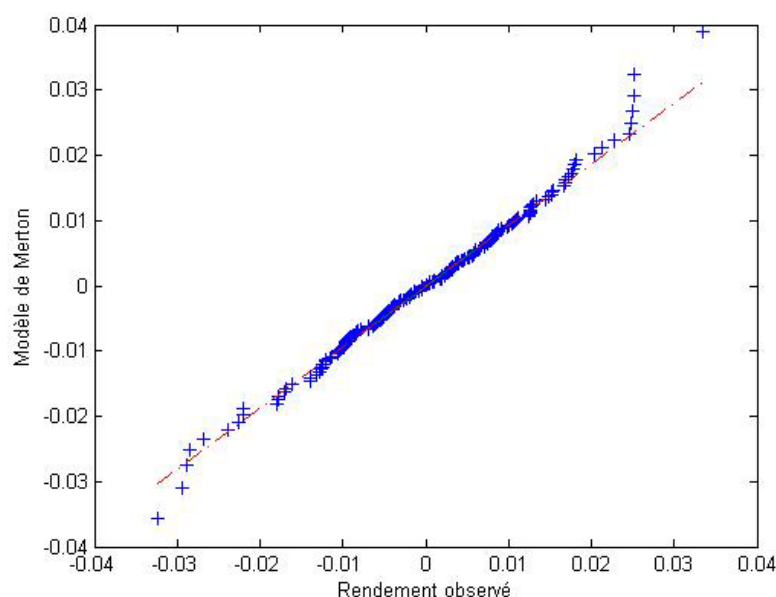
Ce processus est une version mono-périodique du processus de Merton (1976). L'estimation des paramètres par maximum de vraisemblance nécessite l'utilisation de techniques numériques (cf. Planchet et Thérond (2005b)).

Dans notre problématique, on peut toutefois souhaiter calibrer manuellement les paramètres de manière à alourdir les queues de distribution. Par exemple, en fixant l'espérance et la variance globale du rendement (identiques à celle du modèle gaussien), on peut pondérer la variabilité représentée par la composante à sauts de manière à ce que la queue de distribution du modèle soit aussi lourde que celle des observations. Formellement, cela revient à choisir λ, σ_0 et σ_u sous la contrainte :

$$\sigma_0^2 + \lambda \sigma_u^2 = \sigma^2 .$$

La Figure 28 représente cet ajustement. En la comparant avec la Figure 27, on observe que les queues de distributions ont été alourdies. D'une manière générale, la démarche illustrée supra dans le cas de la modélisation du rendement financier devra être reconduite pour l'ensemble des variables de base à l'actif et au passif sous peine de sous-estimer le SCR.

Figure 28 - Rendement journalier du titre TOTAL : QQ-plot loi empirique vs modèle de Merton



6. Conclusion

Le projet Solvabilité 2 a pour objectif d'établir un référentiel prudentiel harmonisé et cohérent à l'ensemble des sociétés d'assurance européennes. Pour cela, il prévoit notamment d'inciter les sociétés à mieux modéliser leurs risques en les autorisant à construire des modèles internes qui aboutissent à un capital de solvabilité inférieur à celui fourni par la "formule standard". Cependant, nous avons vu que la robustesse du critère VaR à 99,5 % sur des données non-observées mais générées par le modèle interne est toute relative. En particulier, les modèles stochastiques actuellement en place chez les assureurs (tant à l'actif qu'au passif) ne sont pas orientés vers l'estimation de valeurs extrêmes et sont, le plus souvent, construits à partir de modélisations des "variables de base" qui sous-estiment les queues de distribution. L'utilisation en l'état de tels modèles conduirait à sous-estimer le SCR. Aussi les sociétés d'assurance qui souhaiteront développer un modèle interne devront modéliser plus finement les queues de distribution. Parallèlement, les autorités qui auditeront les modèles internes devront porter une attention particulière à celles-ci dans le processus de validation. En particulier, une harmonisation au niveau européen de ces procédures est indispensable sous peine d'introduire une distorsion de concurrence entre les différents intervenants. Au global, la prise en compte du critère de contrôle d'une probabilité de ruine qui sous-tend le dispositif Solvabilité 2 implique donc une profonde refonte des modèles utilisés jusqu'alors dans les sociétés d'assurance.

Annexe A : Loi de Pareto généralisée (GPD)

A.1 Définition

Définition 16. *Loi de Pareto généralisée*

Soit $H_{\xi,\beta}$ la fonction de répartition définie pour $\beta > 0$ par

$$H_{\xi,\beta}(x) = \begin{cases} 1 - \left(1 + \frac{\xi}{\beta}x\right)^{-1/\xi}, & \text{si } \xi \neq 0, \\ 1 - \exp(-x/\beta), & \text{si } \xi = 0. \end{cases}$$

Cette fonction de répartition correspond à la loi de Pareto généralisée (GPD) de paramètres ξ et β . Elle est définie pour $x > 0$ si $\xi > 0$ et pour $0 \leq x \leq -\beta/\xi$ si $\xi < 0$.

On notera dans la suite $D(\xi, \beta)$ le domaine de définition de $H_{\xi,\beta}$.

On montre (cf. Denuit et Charpentier (2005)) que la loi de Pareto généralisée peut être vue comme une log-gamma ou encore, dans le cas où $\xi > 0$, comme un mélange de lois exponentielles dont le paramètre suit une loi gamma.

A.2 Quelques propriétés

Les résultats suivants sont énoncés pour une variable aléatoire Y distribuée selon une GPD de paramètre (ξ, β) .

Propriété 1. Si $\xi < 1$, on a

$$\mathbb{E} \left(1 + \frac{\xi}{\beta} Y \right)^{-r} = \frac{1}{1 + \xi r}, \text{ pour } r > -1/\xi,$$

$$\mathbb{E} \left[\ln \left(1 + \frac{\xi}{\beta} Y \right) \right]^k = \xi^k k!, \text{ pour } k \in \mathbb{N},$$

$$\mathbb{E} \left[Y \left(\bar{H}_{\xi,\beta}(Y) \right)^r \right] = \frac{\beta}{(r+1-\xi)(r+1)}, \text{ pour } (r+1)/\xi > 0.$$

Propriété 2. La variable aléatoire Y admet des moments jusqu'à l'ordre $\lceil \xi^{-1} \rceil$ et l'on a

$$\mathbb{E} [Y^r] = \frac{\beta^r \Gamma(\xi^{-1} - r)}{\xi^{r+1} \Gamma(\xi^{-1} + 1)} r!, \text{ pour } r \leq \lceil \xi^{-1} \rceil.$$

Propriété 3. (Stabilité) La variable aléatoire $Y_u = [Y - u | Y > u]$ est distribuée selon une GPD de paramètre $(\xi, \beta + \xi u)$. On en déduit que si $\xi < 1$, alors pour tout $u < y_F$,

$$E[Y - u | Y > u] = \frac{\beta + \xi u}{1 - \xi}, \text{ pour } \beta + u\xi > 0.$$

On rappelle que y_F est la borne supérieure du support de Y , soit

$$y_F = \sup\{y \in \mathbb{R}, F(y) < 1\}.$$

A.3 Estimation des paramètres

Considérons un n -échantillon $(Y_1, \dots, Y_n)$ de la variable aléatoire Y de fonction de répartition $H_{\xi, \beta}$.

A.3.1 Méthode du maximum de vraisemblance

La densité f de Y vaut

$$f(y) = \frac{1}{\beta} \left(1 + \frac{\xi}{\beta} y\right)^{-1/\xi - 1}, \text{ pour } y \in D(\xi, \beta).$$

On en déduit la log-vraisemblance

$$\ln L(\xi, \beta; Y_1, \dots, Y_n) = -n \ln \beta - \left(1 + \frac{1}{\xi}\right) \sum_{i=1}^n \ln \left(1 + \frac{\xi}{\beta} Y_i\right).$$

En utilisant la reparamétrisation $\tau = \xi/\beta$, l'annulation des dérivées partielles de la log-vraisemblance conduit au système

$$\begin{cases} \xi = \frac{1}{n} \sum_{i=1}^n \ln(1 + \tau Y_i) =: \hat{\xi}(\tau), \\ \frac{1}{\tau} = \frac{1}{n} \left(\frac{1}{\xi} + 1\right) \sum_{i=1}^n \frac{Y_i}{1 + \tau Y_i}. \end{cases}$$

L'estimateur du maximum de vraisemblance de (ξ, τ) est $(\hat{\xi} = \hat{\xi}(\hat{\tau}), \hat{\tau})$ où $\hat{\tau}$ est solution de

$$\frac{1}{\tau} = \frac{1}{n} \left(\frac{1}{\hat{\xi}(\tau)} + 1\right) \sum_{i=1}^n \frac{Y_i}{1 + \tau Y_i}.$$

Cette dernière équation se résout numériquement de manière itérative pour autant que l'on dispose d'une valeur initiale τ_0 pas trop éloignée de τ . En pratique, cette valeur initiale pourra être obtenue par la méthode des moments (pour autant que ceux-ci existent jusqu'à l'ordre 2) ou par la méthode des quantiles.

Lorsque $\xi > -1/2$, Hosking et Wallis (1987) ont montré la normalité asymptotique des estimateurs du maximum de vraisemblance :

$$n^{1/2} \left(\hat{\xi}_n - \xi, \frac{\hat{\beta}_n}{\beta} - 1 \right) \xrightarrow[n \rightarrow \infty]{\mathcal{L}} \mathcal{N} \left[0, (1 + \xi) \begin{pmatrix} 1 + \xi & -1 \\ -1 & 2 \end{pmatrix} \right].$$

Ce résultat permet en particulier de calculer les erreurs approximatives d'estimation commises par les estimateurs du maximum de vraisemblance.

A.3.2 Méthode des moments

D'après les résultats du paragraphe A.2, si $\xi < 1/2$, les deux premiers moments de Y existent et l'on a

$$\mu_1 := E[Y] = \frac{\beta}{1 - \xi},$$

et

$$\mu_2 := \text{Var}[Y] = \frac{\beta^2}{(1 - \xi)^2 (1 - 2\xi)}.$$

On en déduit que

$$\xi = \frac{1}{2} - \frac{\mu_1^2}{2\mu_2} \quad \text{et} \quad \beta = \frac{\mu_1}{2} \left(1 + \frac{\mu_1^2}{\mu_2} \right).$$

En remplaçant μ_1 et μ_2 par leurs estimateurs empiriques, on obtient les estimateurs de la méthode des moments $\hat{\xi}_{MM}$ et $\hat{\beta}_{MM}$. Ces estimateurs sont simples à mettre en oeuvre, mais ne fonctionnant que pour $\xi < 1/2$, ils nécessitent la connaissance *a priori* de cette information.

A.3.3 Méthode des moments pondérés par les probabilités

Une alternative intéressante à la méthode des moments a été proposée par Hosking et Wallis (1987). En remarquant que

$$\omega_r := E \left[Y \bar{H}_{\xi, \beta}^r(Y) \right] = \frac{\beta}{(r+1)(r+1-\xi)}, \text{ pour } r = 0, 1,$$

on obtient

$$\beta = \frac{2\omega_0\omega_1}{\omega_0 - 2\omega_1} \quad \text{et} \quad \xi = 2 - \frac{\omega_0}{\omega_0 - 2\omega_1}.$$

En remplaçant ω_0 et ω_1 par leurs estimateurs empiriques, on obtient les estimateurs des moments pondérés par les probabilités $\hat{\xi}_{PWM}$ et $\hat{\beta}_{PWM}$.

Dans le cas d'échantillons de taille réduite, Hosking et Wallis (1987) montrent que, lorsque $\xi < 1/2$, ces estimateurs sont plus efficaces que ceux du maximum de vraisemblance. Le domaine de validité de cette méthode est sa

principale limite à une utilisation en assurance : si la plupart des risques admettent des moments d'ordre 2, ce n'est pas une généralité (tempêtes, tremblements de terre, risques industriels, responsabilité civile, etc.). Des généralisations (Diebolt et al. (2005b)) permettent d'étendre cette méthode à $\xi < 3/2$.

A.3.4 Méthode des quantiles

La fonction quantile d'une loi GPD de paramètre (ξ, β) est donnée par

$$q_p := H_{\xi, \beta}^{-1}(p) = \frac{\beta}{\xi} \left[(1-p)^{-\xi} - 1 \right].$$

En remarquant que

$$\frac{q_{p_2}}{q_{p_1}} = \frac{(1-p_2)^{-\xi} - 1}{(1-p_1)^{-\xi} - 1}, \text{ pour } p_1, p_2 \in]0; 1[,$$

la solution $\hat{\xi}_{MQ}$ de

$$\frac{Q_{p_2}}{Q_{p_1}} = \frac{(1-p_2)^{-\xi} - 1}{(1-p_1)^{-\xi} - 1}$$

où Q_{p_1} et Q_{p_2} sont les quantiles empiriques d'ordres p_1 et p_2 , est l'estimateur des quantiles de ξ .

A.3.5 Méthodes bayésiennes

Des développements récents (cf. Coles et Powell (1996) ou Diebolt et al (2005a)) proposent des approches bayésiennes pour estimer les paramètres de la GPD. Utilisant des algorithmes *Markov Chain Monte Carlo* (MCMC), ces méthodes permettent d'intégrer une information *a priori* (avis d'expert) dans des contextes où l'on dispose d'un nombre réduit d'observations.

Annexe B : Résultats probabilistes

Sont successivement rappelés dans ce paragraphe des résultats de calcul des probabilités concernant la loi du maximum, l'épaisseur des queues de distribution et la loi des excès au-delà d'un seuil. Les différents résultats sont démontrés dans Embrechts et al. (1997).

B.1 Loi du maximum

Définition 17. *Lois de même type*

Deux v.a. X et Y sont dites de même type s'il existe deux constantes $a \in \mathbb{R}$ et $b > 0$ telles que $Y =_d a + bX$.

Théorème 1. Fisher-Tippett-Gnedenko

Considérons une suite $X_1, X_2, \dots$ de v.a.i.i.d. S'il existe une suite de réels (a_n) , une suite positive (b_n) et une loi non-dégénérée H telles que

$$\frac{X_{1,n} - a_n}{b_n} \rightarrow_d H,$$

alors H est du même type qu'une des lois suivantes :

$$\begin{aligned} G(x) &= \exp\left\{-\left(1 + \xi x\right)^{-1/\xi}\right\}, \quad \text{pour } 1 + \xi x \geq 0, \\ \Lambda(x) &= \exp\left\{-\exp(-x)\right\}, \quad \text{pour } x \in \mathbb{R}. \end{aligned}$$

Les fonctions de répartition $G_+(\xi > 0)$, $\Lambda(\xi = 0)$ et $G_-(\xi < 0)$ correspondent respectivement aux lois de Fréchet, de Gumbel et de Weibull.

Définition 18. Domaine d'attraction maximum

On dira d'une fonction de répartition F qui répond aux hypothèses du théorème de Fisher-Tippett-Gnedenko qu'elle appartient au domaine d'attraction maximum (DAM) de H .

La représentation de Jenkinson-Von Mises fournit une caractérisation synthétique des lois extrêmes : la distribution généralisée des valeurs extrêmes (Generalized Extreme Value, ou GEV). La fonction de répartition GEV a la forme suivante :

$$H_{\xi, \mu, \sigma}(x) = \begin{cases} \exp\left\{-\left(1 + \xi \frac{x - \mu}{\sigma}\right)^{-\frac{1}{\xi}}\right\} & \text{si } 1 + \xi \frac{x - \mu}{\sigma} > 0, \xi \neq 0, \\ \exp\left\{-\exp\left(-\frac{x - \mu}{\sigma}\right)\right\} & \text{si } \xi = 0, \end{cases}$$

où le paramètre de localisation μ est directement lié à la valeur la plus probable de la loi, il donne donc une information approximative sur le cœur de la distribution alors que σ est le paramètre de dispersion qui indique l'étalement des valeurs extrêmes. Enfin ξ est l'indice de queue précédemment introduit.

B.2 Épaisseur des queues de distribution**Définition 19. Fonction à variation régulière**

Une fonction $h : \mathbb{R}_+^* \rightarrow \mathbb{R}_+$ est dite à variation régulière (en $+\infty$) d'indice α si h vérifie

$$\lim_{x \rightarrow +\infty} \frac{h(tx)}{h(x)} = t^\alpha \quad \text{pour tout } t > 0.$$

Si $\alpha = 0$, on parlera de variation lente ; si $\alpha = \infty$, de variation rapide.

Théorème 2. Théorème taubérien

Une v.a. de fonction de répartition F et de transformée de Laplace L_F est à variation régulière d'indice $-\alpha$ ($\alpha \leq 0$) si les conditions équivalentes suivantes sont vérifiées ($\mathcal{L}_F$, $\mathcal{L}_{F^{-1}}$, $\mathcal{L}_L$ et $\mathcal{L}_f$ désignent des fonctions à variation lente) :

(i) $\bar{F}$ est à variation régulière d'indice $-\alpha$, i.e. $\bar{F}(x) = x^{-\alpha} \mathcal{L}_F(x)$.

(ii) La fonction quantile est à variation régulière :

$$F^{-1}(1-1/x) = x^{1/\alpha} \mathcal{L}_{F^{-1}}(x).$$

(iii) La transformée de Laplace de F vérifie $L_F(t) = t^\alpha \mathcal{L}_L(1/t)$.

(iv) Si la densité existe et vérifie $xf(x)/\bar{F}(x) \rightarrow \alpha$ quand $x \rightarrow +\infty$, alors la densité est à variation régulière d'indice $-(1+\alpha)$, i.e. $f(x) = x^{-(1+\alpha)} \mathcal{L}_f(x)$.

La condition relative à la transformée de Laplace permet d'établir que la propriété de variation régulière à paramètre fixé est stable par convolution.

B.3 Loi des excès au-delà d'un seuil**Théorème 3. Pickands-Balkema-de Haari**

Une fonction de répartition F appartient au domaine maximum d'attraction de G_ξ si, et seulement si, il existe une fonction positive $\beta(\cdot)$ telle que

$$\limsup_{u \rightarrow x} \left\{ \left| u F(x) - H_{\xi, \beta(u)}(x) \right| \right\} = 0.$$

Ce théorème établit le lien entre le paramètre de la loi du domaine d'attraction maximum et le comportement limite des excès au-delà d'un seuil grand. En particulier, l'indice de queue ξ est identique au paramètre de la loi GPD qui décrit le coût résiduel des sinistres dépassant un seuil suffisamment élevé.

Ceci permet notamment de distinguer les lois à queue épaisse qui appartiennent au DAM de Fréchet ($\xi > 0$) des lois à queue fine qui appartiennent au DAM de Gumbel ($\xi = 0$). Le tableau suivant indique le comportement limite de certaines lois usuelles.

Tableau 4 - Comportement extrême de lois usuelles en assurance

Lois à queue épaisse $\xi > 0$	Lois à queue fine $\xi = 0$	Lois bornées à droite $\xi < 0$
Cauchy Pareto log-gamma Student α -stable ($\alpha < 2$)	gamma normale log-normale Weibull Benktander	uniforme beta

La propriété suivante concerne le nombre N_u de dépassements d'un seuil u assez élevé.

Proposition 5. *Nombre de dépassements d'un seuil élevé*

Le nombre de dépassements du seuil u_n dans un échantillon de taille n est asymptotiquement distribué selon une loi de Poisson pour autant que la probabilité de dépasser u_n diminue proportionnellement à l'inverse de la taille de l'échantillon. Formellement, on a

$$\lim_{n \rightarrow \infty} n\bar{F}(u_n) = \tau \Rightarrow N_{u_n} \rightarrow_d \mathcal{P}(\tau).$$

En effet, comme $\Pr[N_{u_n} = k] = \Pr\left[\sum_{i=1}^n 1_{X_i > u_n} = k\right]$ et que les $X_1, X_2, \dots$ sont indépendants et identiquement distribués, on a :

$$\Pr[N_{u_n} = k] = C_n^k \bar{F}^k(u_n) (1 - \bar{F}(u_n))^{n-k}, \text{ pour } n \geq k.$$

Cette expression peut se réécrire sous la forme :

$$\begin{aligned} \Pr[N_{u_n} = k] &= C_n^k \bar{F}^k(u_n) (1 - \bar{F}(u_n))^{n-k} \\ &= \frac{1}{(1 - \bar{F}(u_n))^k} \left(1 - \frac{1}{n}\right) \dots \left(1 - \frac{k-1}{n}\right) \bar{F}^k(u_n) (1 - \bar{F}(u_n))^n. \end{aligned}$$

Lorsque $n \rightarrow \infty$, $n\bar{F}(u_n) \rightarrow \tau$, $\bar{F}(u_n) \rightarrow 0$ et $j/n \rightarrow 0$ pour $j \in \{0, \dots, k-1\}$, donc

$$\lim_{n \rightarrow \infty} \Pr[N_{u_n} = k] = \frac{\tau^k}{k!} \lim_{n \rightarrow \infty} \left(1 - \frac{n\bar{F}(u_n)}{n}\right)^n = \frac{\tau^k}{k!} e^{-\tau},$$

la dernière égalité pouvant se montrer par un développement limité de $\ln(1 - n\bar{F}(u_n)/n)^n$.

Annexe C : Estimation du paramètre de queue

L'épaisseur de la queue de la fonction de répartition F est résumée par le paramètre ξ de la loi dont elle fait partie du DAM. Ainsi les lois appartenant au DAM de la loi de Fréchet ($\xi > 0$) voient leur queue décroître en fonction puissance tandis que celles appartenant au DAM de Gumbel ($\xi = 0$) ont des queues qui décroissent de manière exponentielle.

Les estimateurs de quantiles extrêmes faisant appel à ce paramètre de queue ξ , son estimation doit faire l'objet d'une attention particulière.

L'objet de cette section est donc de présenter les principales méthodes d'estimation de l'épaisseur de la queue de distribution. Sont ainsi notamment présentés les estimateurs de Pickands (1975), de Hill (1975) et de Dekkers et al. (1989).

C.1 Méthodes paramétriques

C.1.1 Ajustement à la loi du maximum

Le théorème de Fisher-Tippett-Gnedenko nous donne la loi asymptotique de $X_{1,n}$. A supposer que l'on dispose de réalisations de cette v.a., i.e. de m échantillons permettant d'observer m réalisations $x_{1,n}^{(1)}, \dots, x_{1,n}^{(m)}$ de $X_{1,n}$, la méthode du maximum de vraisemblance permettrait d'estimer les paramètres de la loi limite et en particulier ξ .

C.1.2 Ajustement de la loi limite des excès

Considérons une distribution F appartenant au domaine d'attraction maximum de G_ξ . D'après le théorème de Pickands-Balkema-de Haari, il existe une fonction positive $\beta(\cdot)$ telle que

$$\lim_{u \rightarrow x} \sup_{x > 0} \left\| u F(x) - H_{\xi, \beta(u)}(x) \right\| = 0.$$

En particulier, le paramètre ξ de la distribution Pareto généralisée $H_{\xi, \beta}$ est identique à celui de G_ξ . Ainsi pour un seuil u assez élevé, la distribution $H_{\xi, \beta(u)}$ est une bonne approximation de F_u dont on dispose de N_u observations $X_{N_u, n} - u, \dots, X_{1, n} - u$. Les techniques présentées au paragraphe A.3 permettent alors d'estimer ξ .

En particulier, si $\xi > -1/2$ est estimé par la méthode du maximum de vraisemblance, Smith (1987) montre que la variance asymptotique de cet estimateur vaut $(1 + \xi)^2 / N_u$.

C.2 Méthodes semi-paramétriques

C.2.1 Estimateur de Pickands

Pour $k/n \rightarrow 0$, l'estimateur de Pickands, défini par

$$\xi_k^P = \frac{1}{\ln 2} \ln \frac{X_{k,n} - X_{2k,n}}{X_{2k,n} - X_{4k,n}},$$

est un estimateur convergent de ξ . De plus, sous certaines conditions supplémentaires portant sur k et F (cf. Dekkers et de Haan (1989)), il est asymptotiquement gaussien :

$$\sqrt{k}(\xi_k^P - \xi) \rightarrow_d \mathcal{N}\left(0, \frac{\xi^2(2^{\xi+1} + 1)}{(2(2^\xi - 1)\ln 2)^2}\right).$$

C.2.2 Estimateur de Hill

Pour $k/n \rightarrow 0$, l'estimateur de Hill défini par

$$\xi_k^H = \frac{1}{k} \sum_{j=1}^k \ln \frac{X_{j,n}}{X_{k,n}},$$

est un estimateur convergent de ξ . De plus, sous certaines conditions sur k et F (cf. de Haan et Peng (1998)), il est asymptotiquement gaussien :

$$\sqrt{k}(\xi_k^H - \xi) \rightarrow_d \mathcal{N}(0, \xi^2).$$

Bien que plus performant que l'estimateur de Pickands (cf. le rapport des variances asymptotiques), l'estimateur de Hill n'est utilisable que pour les distributions de Fréchet ($\xi > 0$).

C.2.3 Estimateur de Dekkers-Einmahl-de Haan

L'estimateur de Dekkers-Einmahl-de Haan est une généralisation de l'estimateur de Hill à l'ensemble des lois extrêmes ($\xi \in \mathbb{R}$).

Pour $k/n \rightarrow 0$, l'estimateur de Dekkers-Einmahl-de Haan défini par

$$\xi_k^{DEdH} = 1 + \xi_k^{(1)} - \frac{1}{2} \left(1 - \frac{(\xi_k^{(1)})^2}{\xi_k^{(2)}} \right)^{-1},$$

avec $\xi_k^{(i)} = \frac{1}{k} \sum_{j=1}^k \left(\ln \frac{X_{j,n}}{X_{k+1,n}} \right)^i$, est un estimateur convergent de ξ . De plus,

sous certaines conditions sur k et F (cf. Dekkers et al. (1989)), il est asymptotiquement gaussien :

$$\sqrt{k}(\xi_k^{DEdH} - \xi) \rightarrow_d \mathcal{N}(0, 1 + \xi^2).$$

Cet estimateur est également connu sous l'appellation d'estimateur des moments, les $\xi_k^{(i)}$ pouvant s'interpréter comme des moments empiriques.

C.2.4 Nombre d'observations à utiliser

Les résultats concernant les estimateurs de l'indice de queue énoncés précédemment sont asymptotiques : ils sont obtenus lorsque $k \rightarrow \infty$ et $k/n \rightarrow 0$. Comme en pratique, on ne dispose que d'un nombre d'observations n fini, il s'agit de choisir k de manière à ce que l'on dispose de suffisamment de matériel statistique (les $X_{k,n}, \dots, X_{1,n}$) tout en restant dans la queue de distribution ($k \ll n$).

En particulier, l'estimateur de Hill satisfaisant la propriété asymptotique

$$\sqrt{k}(\xi_k^H - \xi) \rightarrow_d \mathcal{N}(0, \xi^2),$$

pour $k \rightarrow \infty$ avec un *certain taux de croissance* en fonction de n , on pourrait être tenté de choisir k aussi grand que possible de manière à minimiser l'erreur quadratique moyenne commise par ξ_k^H . Cependant, le comportement au deuxième ordre de la fonction à variation lente $\mathcal{L}_F$ introduite dans le théorème taubérien induit un biais lorsque k est trop grand. Des solutions pour fixer k de manière à disposer d'un estimateur asymptotiquement sans biais ont notamment été proposées par Goldie et Smith (1987) puis de Haan et Peng (1998).

Concernant l'estimateur de Hill pour des fonctions appartenant au DAM de Fréchet, de Haan et Peng (1998) ont proposé de retenir le nombre d'observation k^* qui réduit l'erreur quadratique moyenne de l'estimateur de Hill, i.e.

$$k^*(n) = \begin{cases} 1 + n^{2\xi/(2\xi+1)} \left(\frac{(1+\xi)^2}{2\xi} \right)^{1/(2\xi+1)}, & \text{si } \xi \in]0; 1[\\ 2n^{2/3}, & \text{si } \xi > 1. \end{cases}$$

Toutefois, on remarque que k^* s'exprime en fonction de ξ qui n'est pas observé. De plus ce critère ne concerne que l'estimateur de Hill.

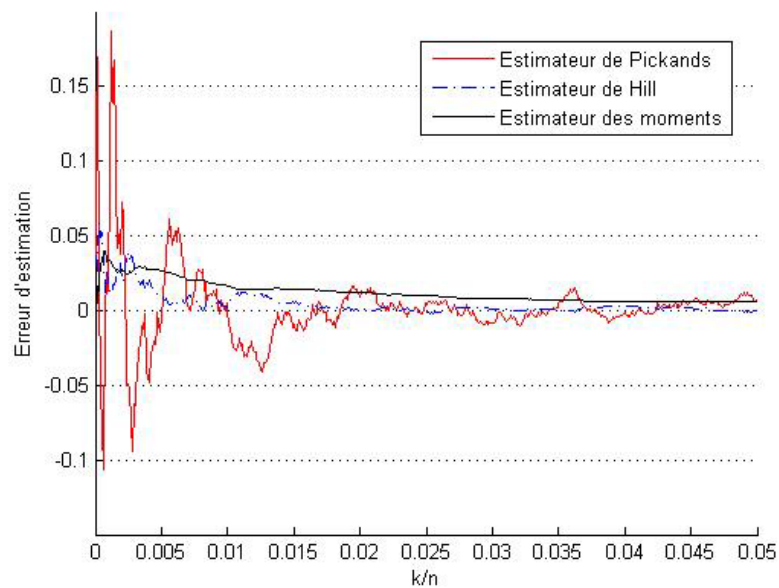
En pratique, il sera préférable de suivre la démarche proposée par Embrechts et al. (1997) qui suggèrent de tracer les estimateurs en fonction de k/n et de choisir k dans un intervalle où la droite des estimations est approximativement horizontale (cf. la Figure 29 détaillée infra).

C.2.5 Illustration

La Figure 29 illustre les différents estimateurs du paramètre de queue d'une loi de Pareto de première espèce (de fonction de répartition $F(x) = 1 - (x_0/x)^\alpha$, pour $x > x_0$). Cette famille de distributions faisant partie

du DAM de Fréchet, les estimations de Pickands, de Hill et de Dekkers-Einmahl-de Haan ont pu être tracées en fonction de k .

Figure 29 - Estimation de l'épaisseur de la queue d'une distribution de Pareto



On observe la suprématie de l'estimateur de Hill sur ceux de Pickands et de Dekkers-Einmahl-de Haan. Par ailleurs, pour $k < 0,02 \times n$, l'estimateur de Hill est relativement volatile. On serait donc amené à utiliser de l'ordre de 2,5 % des données les plus extrêmes pour estimer l'épaisseur de la queue de distribution.

Bibliographie

- Blum K. A., Otto D. J. (1998) « Best estimate loss reserving : an actuarial perspective », *CAS Forum Fall* 1, 55-101.
- Bottard S. (1996) « Application de la méthode du Bootstrap pour l'estimation des valeurs extrêmes dans les distributions de l'intensité des séismes », *Revue de statistique appliquée* 44 (4), 5-17.
- Christoffersen P., Hahn J., Inoue, A. (2001) « Testing and comparing value-at-risk measures », *Journal of Empirical Finance* 8 (3), 325-42.
- Coles S., Powell E. (1996) « Bayesian methods in extreme value modelling : a review and new developments », *Internat. Statist. Rev.* 64, 119-36.
- Davidson R., MacKinnon J.G. (2004) « Bootstrap Methods in Econometrics », working paper.
- de Haan L., Peng L. (1998) « Comparison of tail index estimators », *Statistica Neerlandica* 52 (1), 60-70.
- Deelstra G., Janssen J. (1998) « Interaction between asset liability management and risk theory », *Applied Stochastic Models and Data Analysis* 14, 295-307.
- Dekkers A., de Haan L. (1989) « On the estimation of the extreme-value index and large quantile estimation », *Annals of Statistics* 17, 1795-832.
- Dekkers A., Einmahl J., de Haan L. (1989) « A moment estimator for the index of an extreme-value distribution », *Annals of Statistics* 17, 1833-55.
- Denuit M., Charpentier A. (2005) *Mathématiques de l'assurance non-vie. Tome 2 : tarification et provisionnement*, Paris : Economica.
- Diebolt J., El-Aroui M., Garrido S., Girard S. (2005a) « Quasi-conjugate bayes estimates for gpd parameters and application to heavy tails modelling », *Extremes* 8, 57-78.
- Diebolt J., Guillou A., Rached I. (2005b) « Approximation of the distribution of excesses using a generalized probability weighted moment method », *C. R. Acad. Sci. Paris, Ser. I* 340 (5), 383-8.
- Efron B. (1979) « Bootstrap methods: Another look at the Jackknife », *Ann. Statist.*, 71-26.
- Efron B. (1987) « Better bootstrap confidence intervals », *Journal of the American Statistical Association* 82, 171-200.
- Efron B., Tibshirani R. J. (1993) *An introduction to the bootstrap*, Chapman & Hall.
- Embrechts P., Klüppelberg C., Mikosch T. (1997) *Modelling Extremal Events for Insurance and Finance*, Berlin: Springer Verlag.

- England P.D., Verall R.J. (1999) « Analytic and bootstrap estimates of prediction errors in claim reserving », *Insurance: mathematics and economics* 25, 281-93.
- Fedor M., Morel J. (2006) « Value-at-risk en assurance : recherche d'une méthodologie à long terme », *Actes du 28e congrès international des actuaires*, Paris.
- Goldie C., Smith R. (1987) « Slow variation with remainder: a survey of the theory and its applications », *Quarterly Journal of Mathematics Oxford* 38 (2), 45-71.
- Hill B. (1975) « A simple general approach to inference about the tail of a distribution », *Annals of Statistics* 3, 1163-74.
- Horowitz J.L. (1997) « Bootstrap methods in econometrics: theory and numerical performance », in *Advances in Economics and Econometrics: Theory and Application*, volume 3, 188-222, D. M. Kreps, K. F. Wallis (eds), Cambridge: Cambridge University Press.
- Hosking J. R., Wallis J. R. (1987) « Parameter and quantile estimation for the generalized pareto distribution », *Technometrics* 29, 339-49.
- Jal P., Partrat Ch. (2004) « Evaluation stochastique de la provision pour sinistres », Conférence scientifique de l'Institut des Actuaires, Paris, 20 janvier 2004.
- Merton R. (1976) « Option pricing when underlying stock returns are discontinuous », *Journal of Financial Economics* 3, 125-44.
- Partrat C., Besson J.L. (2005) *Assurance non-vie. Modélisation, simulation*, Paris : Economica.
- Pickands J. (1975) « Statistical inference using extreme orders statistics », *Annals of Statistics* 3, 119-31.
- Planchet F., Thérond P.E. (2005b) « L'impact de la prise en compte des sauts boursiers dans les problématiques d'assurance », *Proceedings of the 15th AFIR Colloquium*, Zürich.
- Quenouille M. (1949) « Approximate tests of correlation in time series », *Journal of the Royal Statistical Society, Soc. Series B*, 11, 18-84.
- Smith R. L. (1987) « Estimating tails of probability distributions », *Annals of Statistics* 15, 1174-207.
- Thérond P.E., Planchet F. (2007) « Provisions techniques et capital de solvabilité d'une compagnie d'assurance : méthodologie d'utilisation de Value-at-Risk », *Assurances et gestion des risques* 74 (4), 533-63 et in *Proceedings of the 37th ASTIN Colloquium*, Orlando.
- Windcliff H. Boyle, P. (2004) « The 1/n pension investment puzzle », *North American Actuarial Journal* 8 (3), 32-45.
- Zajdenweber D. (2000) *Économie des extrêmes*, Paris : Flammarion.

Zelterman D. (1993) « A semiparametric Bootstrap technique for simulating extreme order statistics », *Journal of American Statistical Association* 88 (422), 477-85.

Chapitre 5

Prise en compte de la dépendance

Le cœur de l'activité d'assurance repose sur le principe de mutualisation découlant de la loi des grands nombres.

Théorème 4. *Loi faible des grands nombres*

Soit $X_1, X_2, \dots$ une suite de variables aléatoires indépendantes d'espérances $m_1, m_2, \dots$ finies et de variances $\sigma_1^2, \sigma_2^2, \dots$ finies. Si

$$\frac{1}{n} \sum_{i=1}^n m_i \rightarrow m \text{ et si } \frac{1}{n^2} \sum_{i=1}^n \sigma_i^2 \rightarrow 0, \text{ alors } \frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{\text{Pr}} m.$$

La loi des grands nombres requiert l'indépendance des risques. Si les risques ne sont pas indépendants, on a toujours la propriété limite $E\left[\frac{1}{n} \sum_{i=1}^n X_i\right] = m$, mais on ne dispose plus de la convergence en probabilité qui assure la mutualisation.

Intéressons-nous par exemple à deux variables aléatoires X_1 et X_2 de loi exponentielle de paramètre 1, i. e. telles que $\Pr[X_i < x] = 1 - e^{-x}$ pour $x \in \mathbf{R}_+$.

Supposons que ces deux variables modélisent deux risques supportés par un assureur. En faisant l'hypothèse d'indépendance, la charge globale de sinistres $X_1 + X_2$ est distribuée selon une loi gamma $\mathcal{G}(2;1)$ alors qu'elle sera distribuée selon une loi exponentielle de paramètre 1/2 lorsque la dépendance est parfaitement positive. Si les espérances de ces deux lois sont égales, il en sera autrement des primes Stop-Loss ou encore des Value-at-Risk.

Aussi la prise en compte de la dépendance stochastique est primordiale dans une approche de contrôle du risque. Notons enfin la part croissante prise par la théorie des copules dans la cadre de la modélisation des risques d'assurance.

Après avoir examiné les différents types de dépendance et les outils pour quantifier le degré de dépendance, on s'attache dans cette section à décrire les outils usuels de modélisation des structures de dépendance.

On distingue d'une part la dépendance linéaire (dont l'intérêt est lié de manière assez intime au caractère gaussien des phénomènes sous-jacents) et la dépendance non linéaire, au travers de la théorie des copules.

La dernière partie de ce chapitre est consacrée à deux exemples de modélisation de la dépendance en assurance : le rachat de contrats d'épargne en euros et le calcul d'un capital de solvabilité par une formule modulaire.

1. Analyse mathématique de la dépendance

L'étude mathématique de la dépendance entre variables aléatoires nous permet de disposer d'outils pour quantifier la force de cette dépendance : les mesures de dépendance, et de structures de dépendance.

Ces deux aspects sont abordés successivement.

1.1. Mesures de dépendance

L'objet de ce paragraphe est de définir ce qu'est une mesure de dépendance et de présenter les mesures les plus usuelles. On se restreint tout d'abord au cas de deux risques avant de généraliser.

1.1.1. Définition

Définition 20. Mesure de dépendance

Une mesure de dépendance est une application qui associe à deux variables aléatoires un réel permettant de quantifier la force de la dépendance qui lie ces deux variables aléatoires.

Cette définition n'est pas très restrictive mais en pratique, on demandera souvent à une mesure de risque de posséder un certain nombre de « bonnes propriétés », i. e. d'être une mesure de concordance définie comme suit.

Définition 21. Mesure de concordance

Une mesure de dépendance δ est une mesure de concordance si elle possède les propriétés suivantes :

$$(P1) \quad \delta(X_1, X_2) = \delta(X_2, X_1) \text{ (symétrie) ;}$$

$$(P2) \quad -1 \leq \delta(X_1, X_2) \leq 1 \text{ (normalisation) ;}$$

$$(P3) \quad \delta(X_1, X_2) = 1 \Leftrightarrow (X_1, X_2) =_{loi} (F_1^{-1}(U), F_2^{-1}(U)) \text{ où } U \sim \mathcal{U}[0;1] ;$$

$$(P4) \quad \delta(X_1, X_2) = -1 \Leftrightarrow (X_1, X_2) =_{loi} (F_1^{-1}(U), F_2^{-1}(1-U)) \text{ où } U \sim \mathcal{U}[0;1] ;$$

(P5) Si f est strictement monotone,

$$\delta(f(X_1), X_2) = \begin{cases} \delta(X_1, X_2) & \text{si } f \uparrow \\ -\delta(X_1, X_2) & \text{si } f \downarrow. \end{cases}$$

Ainsi une mesure de dépendance δ qui a la propriété $\delta(X_1, X_2) = 0 \Leftrightarrow X_1 \text{ et } X_2 \text{ indépendantes}$ n'est pas une mesure de concordance.

En effet si l'on considère le couple $(X_1, X_2) =_{loi} (\cos Z, \sin Z)$ où $Z \sim \mathcal{U}[0; 2\pi]$. Puisque $(X_1, X_2) =_{loi} (-X_1, X_2)$, on a :

$\delta(X_1, X_2) = \delta(-X_1, X_2)$. Par ailleurs $f : x \mapsto -x$ est strictement décroissante donc $\delta(X_1, X_2) = \delta(-X_1, X_2) = -\delta(X_1, X_2) = 0$ alors que X_1 et X_2 ne sont pas indépendantes.

En revanche, nous verrons que la plupart des mesures de dépendance utilisées en pratique vérifient la propriété :

$$X_1 \text{ et } X_2 \text{ indépendantes} \Rightarrow \delta(X_1, X_2) = 0.$$

1.1.2. Coefficient de corrélation de Pearson

Le coefficient de corrélation de Pearson, également appelé coefficient de corrélation linéaire, est la première mesure de dépendance à avoir été utilisée. Il repose sur la propriété suivante :

$$\text{Var}[X_1 + X_2] = \text{Var}[X_1] + \text{Var}[X_2] + 2\text{Cov}[X_1, X_2].$$

La covariance mesure donc le surcroît de variabilité (éventuellement négatif) de la somme de deux variables aléatoires par rapport à la somme de la variabilité de chacune de ces variables aléatoires. La covariance permet donc d'apprécier le sens de la covariation de deux variables aléatoires. Le coefficient de corrélation de Pearson en est une version normée sans dimension.

Soit deux variables aléatoires X_1 et X_2 admettant des moments jusqu'à l'ordre 2. Le coefficient de corrélation de Pearson entre X_1 et X_2 que l'on notera $r(X_1, X_2)$ est défini par

$$r(X_1, X_2) = \frac{\text{Cov}[X_1, X_2]}{\sqrt{\text{Var}[X_1]\text{Var}[X_2]}}.$$

Plus le coefficient de corrélation de Pearson sera grand en valeur absolue, plus la dépendance entre les variables aléatoires sera forte.

Par ailleurs, quelles que soient les variables aléatoires X_1 et X_2 admettant des moments jusqu'à l'ordre 2, on a toujours $-1 \leq r(X_1, X_2) \leq 1$.

Notons que considérant deux variables aléatoires X_1 et X_2 de fonction de répartition F_1 et F_2 , il n'est pas toujours possible de trouver une distribution jointe telle que les bornes de l'inégalité soient atteintes.

Propriétés : Le coefficient de corrélation possède les deux propriétés suivantes :

- $|r(X_1, X_2)| = 1 \Leftrightarrow X_2 =_{loi} a + bX_1$ avec $b \neq 0$. De plus le signe de r est identique au signe de b .
- X_1 et X_2 indépendantes $\Rightarrow r(X_1, X_2) = 0$. La réciproque est fausse.

La démonstration de la première propriété est triviale. Pour la deuxième, il suffit de considérer $X \sim \mathbf{N}(0;1)$. X est symétrique donc $E[X^3] = 0$. Posons $Y = X^2$, $E[XY] = 0$, donc $r(X, Y) = 0$ alors que X et Y ne sont pas indépendants.

En revanche le coefficient de corrélation linéaire n'est pas une mesure de concordance. En effet, il n'a pas les propriétés (P3) et (P4) de la définition d'une mesure de concordance présentée au paragraphe 1.1.1.

1.1.3. Coefficient de corrélation des rangs de Kendall

Le coefficient de corrélation des rangs de Kendall, ou τ de Kendall, utilise les concepts de concordance et de discordance définis comme suit.

Définition 22. *Concordance de couples*

Les couples (X_1, X_2) et (Y_1, Y_2) sont en concordance si

$$(X_1 < Y_1 \text{ et } X_2 < Y_2) \text{ ou } (X_1 > Y_1 \text{ et } X_2 > Y_2).$$

Définition 23. *Discordance de couples*

Les couples (X_1, X_2) et (Y_1, Y_2) sont en discordance si

$$(X_1 < Y_1 \text{ et } X_2 > Y_2) \text{ ou } (X_1 > Y_1 \text{ et } X_2 < Y_2).$$

Il faut noter que ces deux couples peuvent n'être ni en concordance ni en discordance.

Définition 24. *Coefficient de corrélation des rangs de Kendall*

Soit $X = (X_1, X_2)$ un vecteur aléatoire dont les fonctions de répartition marginales sont continues. Le coefficient de corrélation des rangs de Kendall du vecteur X , que l'on notera $\tau(X_1, X_2)$ est défini par $\tau(X_1, X_2) = \Pr[(X_1 - Y_1)(X_2 - Y_2) > 0] - \Pr[(X_1 - Y_1)(X_2 - Y_2) < 0]$, où (Y_1, Y_2) est un couple indépendant de (X_1, X_2) et de même loi.

Le τ de Kendall compare ainsi la probabilité que deux couples indépendants mais de même loi soient en concordance et la probabilité qu'ils soient en discordance. Aussi un τ positif signifiera qu'en probabilité, les couples considérés sont plus souvent en concordance qu'en discordance.

Proposition 6. *Pour tout couple $X = (X_1, X_2)$ de fonctions de répartition marginales continues, $\tau(X_1, X_2) = 4 E[F_X(X_1, X_2)] - 1$.*

Démonstration : Soit (Y_1, Y_2) est un couple indépendant de (X_1, X_2) et de même loi. Comme $\Pr[(X_1 - Y_1)(X_2 - Y_2) > 0] + \Pr[(X_1 - Y_1)(X_2 - Y_2) < 0] = 1$, d'après la définition du τ de Kendall, il vient $\tau(X_1, X_2) = 2 \Pr[(X_1 - Y_1)(X_2 - Y_2) > 0] - 1$. Comme :

$$\Pr[(X_1 - Y_1)(X_2 - Y_2) > 0] = \Pr[X_1 < Y_1, X_2 < Y_2] + \Pr[X_1 > Y_1, X_2 > Y_2],$$

on peut écrire que $\Pr[(X_1 - Y_1)(X_2 - Y_2) > 0] = E[F_X(Y_1, Y_2) + F_Y(X_1, X_2)]$. Par ailleurs, les couples X et Y étant indépendants et distribués identiquement, on a $E[F_X(Y_1, Y_2)] = E[F_Y(X_1, X_2)]$, ce qui assure le résultat. $\square$

Quelles que soient les fonctions f et g toutes deux croissantes ou toutes deux décroissantes, on a $\tau(f(X_1), g(X_2)) = \tau(X_1, X_2)$.

Par ailleurs on dispose de l'implication :

$$X_1 \text{ et } X_2 \text{ indépendantes} \Rightarrow \tau(X_1, X_2) = 0.$$

En effet, si X_1 et X_2 sont indépendantes, alors $F_X(x_1, x_2) = \Pr[X_1 \leq x_1] \Pr[X_2 \leq x_2]$ donc $E[F_X(X_1, X_2)] = 1/4$ ce qui assure le résultat grâce à l'écriture de la Proposition 6.

On en déduit que le τ de Kendall est une mesure de concordance.

1.1.4. Coefficient de corrélation des rangs de Spearman

Le coefficient de corrélation des rangs de Spearman¹, ou ρ de Spearman, fait également appel aux notions de concordance et de discordance en comparant le couple dont on veut mesurer la dépendance à une version indépendante de ce couple.

Définition 25. Coefficient de corrélation des rangs de Spearman

Considérons le couple $X = (X_1, X_2)$ de fonctions de répartition marginales continues F_1 et F_2 , le coefficient de corrélation des rangs de Spearman de ce couple, noté $\rho(X_1, X_2)$ est défini par

$$\rho(X_1, X_2) = 3 \left\{ \Pr[(X_1 - X_1^\perp)(X_2 - X_2^\perp) > 0] - \Pr[(X_1 - X_1^\perp)(X_2 - X_2^\perp) < 0] \right\},$$

où le couple $X^\perp = (X_1^\perp, X_2^\perp)$ a les mêmes marginales que X mais est indépendant.

Proposition 7. Pour tout couple $X = (X_1, X_2)$ de fonctions de répartition marginales continues, $\rho(X_1, X_2) = r(F_1(X_1), F_2(X_2))$.

Le ρ de Spearman ne dépend donc pas des marginales F_1 et F_2 , puisque $F_1(X_1)$ et $F_2(X_2)$ sont des variables aléatoires de loi uniforme sur $[0; 1]$.

Par ailleurs, quelles que soient les fonctions f et g toutes deux croissantes ou toutes deux décroissantes, on a $\rho(f(X_1), g(X_2)) = \rho(X_1, X_2)$.

Enfin le ρ de Spearman est une mesure de concordance.

1.1.5. Liens entre le τ de Kendall et le ρ de Spearman

Plusieurs relations lient les deux mesures de concordance. Le lecteur trouvera les démonstrations des deux premières dans Denuit et Charpentier (2004). La troisième se déduit des deux premières. Lorsqu'il n'y a pas d'ambiguïté possible, on omettra de préciser à quel couple s'appliquent les mesures de dépendance.

¹. Pour mémoire, Spearman était psychologue.

Propriétés : Pour tout couple $X = (X_1, X_2)$ de fonctions de répartition marginales continues, on a :

$$-1 \leq 3\tau - 2\rho \leq 1 ;$$

$$-1 + 2\left(\frac{1+\tau}{2}\right)^2 \leq \rho \leq 1 - 2\left(\frac{1-\tau}{2}\right)^2 ;$$

$$\frac{-1+3\tau}{2} \leq \rho \leq \frac{1+2\tau-\tau^2}{2} \text{ si } \tau \geq 0 ,$$

et

$$\frac{-1+2\tau+\tau^2}{2} \leq \rho \leq \frac{1+3\tau}{2} \text{ si } \tau \leq 0 .$$

1.2. Structures de dépendance

Si deux variables aléatoires n'ont qu'une manière d'être indépendantes, la dépendance peut quant à elle revêtir de nombreuses formes. Ce sont ces formes que tentent de capter les structures de dépendance. L'objet de ce paragraphe est de présenter trois types de structure de dépendance qui seront étudiées dans des espaces de Fréchet.

1.2.1. Espaces de Fréchet

Définition 26. *Classe de Fréchet*

La classe de Fréchet associée aux fonctions de répartition F_1 et F_2 , notée $\mathbf{F}(F_1, F_2)$, est constituée de l'ensemble des distributions de probabilités des couples $X = (X_1, X_2)$ dont les marginales sont respectivement F_1 et F_2 .

Les classes de Fréchet constituent un cadre idéal pour analyser la dépendance puisqu'elles ne regroupent que des lois de probabilité de couples qui ne diffèrent entre elles que par leur structure de dépendance et non par leurs marginales.

Définition 27. *Borne supérieure de Fréchet*

La fonction de répartition W définie par

$$W(x_1, x_2) = \min\{F_1(x_1), F_2(x_2)\} ,$$

est appelée borne supérieure de Fréchet de la classe $\mathbf{F}(F_1, F_2)$.

Définition 28. *Borne inférieure de Fréchet*

La fonction de répartition M définie par

$$M(x_1, x_2) = \max\{F_1(x_1) + F_2(x_2) - 1, 0\} ,$$

est appelée borne inférieure de Fréchet de la classe $\mathbf{F}(F_1, F_2)$.

Ces deux fonctions de répartition tirent leur nom de la propriété suivante.

Propriété : Pour tout $F_X \in \mathbf{F}(F_1, F_2)$,

$$M(x) \leq F_X(x) \leq W(x) \text{ pour tout } x \in \mathbf{R}^2.$$

Cette propriété traduit le fait que, pour tout couple $X = (X_1, X_2)$, la classe de Fréchet correspondante est bornée par les fonctions de répartition W et M .

Dans la suite, on parlera parfois de couple $X = (X_1, X_2)$ appartenant à la classe de Fréchet $\mathbf{F}(F_1, F_2)$, il s'agit d'un abus de langage car une classe de Fréchet est composée de fonctions de répartition et non de vecteurs aléatoires : c'est en fait la fonction de répartition F_X du couple X qui appartient à la classe de Fréchet.

1.2.2. Comparaison supermodulaire

La comparaison supermodulaire permet de comparer la force de la dépendance existant entre les composantes de couples de risques. Cette comparaison utilise les fonctions supermodulaires, également appelée fonctions superadditives.

Rappelons qu'une fonction $g : \mathbf{R}^2 \rightarrow \mathbf{R}$ est dite supermodulaire lorsque pour tous $\varepsilon > 0, \delta > 0, x \in \mathbf{R}^2$, l'inégalité suivante est vérifiée :

$$g(x_1 + \varepsilon, x_2 + \delta) + g(x_1, x_2) \geq g(x_1 + \varepsilon, x_2) + g(x_1, x_2 + \delta).$$

En particulier lorsque g est dérivable par rapport à chacun de ces arguments en un point x , comme on a :

$$g(x_1 + \varepsilon, x_2 + \delta) - g(x_1, x_2 + \delta) \geq g(x_1 + \varepsilon, x_2) - g(x_1, x_2),$$

lorsque $\varepsilon \rightarrow 0$, il vient

$$\frac{\partial}{\partial x_1} g(x_1, x_2 + \delta) \geq \frac{\partial}{\partial x_1} g(x_1, x_2).$$

Puis lorsque $\delta \rightarrow 0$, on a enfin

$$\frac{\partial^2}{\partial x_1 \partial x_2} g(x_1, x_2) \geq 0.$$

Ce qui nous conduit naturellement à la caractérisation suivante pour les fonctions dérivables par rapport à chaque argument.

Propriété : Une fonction $g \in C^{1,1}(\mathbf{R}^2, \mathbf{R})$ est supermodulaire si, et seulement

si, $\frac{\partial^2}{\partial x_1 \partial x_2} g(x) \geq 0$ pour tout $x \in \mathbf{R}^2$.

Cette caractérisation en amène une autre.

Propriété : Soit $g \in C^{1,1}(\mathbf{R}^2, \mathbf{R})$ une fonction supermodulaire. Pour toutes fonctions f_1 et f_2 toutes deux croissantes ou toutes deux décroissantes, la fonction $\Psi : x \mapsto g(f_1(x), f_2(x))$ est également supermodulaire.

Définition 29. Comparaison supermodulaire

Le couple X est dit inférieur au sens supermodulaire au couple Y , et l'on note $X \prec_{sm} Y$, lorsque pour toute fonction supermodulaire g telle que $E[g(X)]$ et $E[g(Y)]$ existent, on a l'inégalité $E[g(X)] \leq E[g(Y)]$.

Lorsque $X \prec_{sm} Y$, on pourra dire que l'intensité de la dépendance qui lie les variables aléatoires X_1 et X_2 est moins forte que celle qui lie les risques Y_1 et Y_2 .

De plus, pour tous couples X et Y dont les fonctions de répartition appartiennent à $\mathbf{F}(F_1, F_2)$, quelles que soient les fonctions f_1 et f_2 croissantes, on a : $X \prec_{sm} Y \Rightarrow (f_1(X_1), f_2(X_2)) \leq (f_1(Y_1), f_2(Y_2))$.

Pour pouvoir être comparées les fonctions de répartition de deux couples doivent faire partie de la même classe de Fréchet. En effet,

$$X \prec_{sm} Y \Rightarrow X_i =_{loi} Y_i \text{ pour } i = 1, 2 \Leftrightarrow F_X, F_Y \in \mathbf{F}(F_1, F_2).$$

Ce résultat se démontre en utilisant les fonctions supermodulaires $g_1 : z \mapsto \mathbf{1}_{\{z > x\}}$ et $g_2 : z \mapsto \mathbf{1}_{\{z \leq x\}}$.

L'utilisation de ces mêmes fonctions g_1 et g_2 permet de démontrer la caractérisation suivante de l'ordre supermodulaire.

Proposition 8. Soient X et Y deux couples dont les fonctions de répartition appartiennent à une même classe de Fréchet. On a

$$X \prec_{sm} Y \Leftrightarrow F_X(x) \leq F_Y(x) \text{ pour tout } x \in \mathbf{R}^2.$$

Le résultat suivant traduit l'intuition selon laquelle la somme de deux risques est d'autant plus dangereuse que les risques sont dépendants.

Proposition 9. Pour toutes $F_X, F_Y \in \mathbf{F}(F_1, F_2)$, on a :
 $X \prec_{sm} Y \Leftrightarrow X_1 + X_2 \prec_{cx} Y_1 + Y_2$.

Démonstration : Soient X et Y deux couples tels que $F_X, F_Y \in \mathbf{F}(F_1, F_2)$. D'après la Proposition 8, on a :

$$X \prec_{sm} Y \Leftrightarrow \forall t \geq 0, \forall s \leq t, F_X(t, s - t) \leq F_Y(t, s - t)$$

En remarquant que $\int_0^t \mathbf{1}_{\{x_1 \leq s, x_2 \leq t-s\}} ds = (t - x_1 - x_2)^+$ et en utilisant le théorème de Fubini, on obtient l'équivalence :

$$\begin{aligned} X \prec_{sm} Y &\Leftrightarrow \forall t \geq 0, E[(X_1 + X_2 - t)^+] \leq E[(Y_1 + Y_2 - t)^+] \\ &\Leftrightarrow X_1 + X_2 \prec_{cx} Y_1 + Y_2. \end{aligned}$$

Par ailleurs, comme

$$\begin{aligned}
\mathbb{E}\left[(X_1 + X_2 - t)^+\right] - \mathbb{E}\left[(t - X_1 - X_2)^+\right] &= t + \mathbb{E}[X_1 + X_2] \\
&= t + \mathbb{E}[Y_1 + Y_2] \\
&= \mathbb{E}\left[(Y_1 + Y_2 - t)^+\right] - \mathbb{E}\left[(t - Y_1 - Y_2)^+\right]
\end{aligned}$$

puisque $X_i \stackrel{loi}{=} Y$ pour $i = 1, 2$, on a :

$$\mathbb{E}\left[(X_1 + X_2 - t)^+\right] \leq \mathbb{E}\left[(Y_1 + Y_2 - t)^+\right] \Leftrightarrow \mathbb{E}\left[(t - X_1 - X_2)^+\right] \leq \mathbb{E}\left[(t - Y_1 - Y_2)^+\right],$$

qui assure que $X \prec_{sm} Y \Leftrightarrow \forall t \geq 0, \mathbb{E}\left[(X_1 + X_2 - t)^+\right] \leq \mathbb{E}\left[(Y_1 + Y_2 - t)^+\right]$. $\square$

On dispose de l'implication suivante qui relie la comparaison supermodulaire avec les mesures de dépendance étudiées dans le paragraphe 1.1 :

Propriété : Quels que soient les couples X et Y , on a

$$X \prec_{sm} Y \Rightarrow \begin{cases} r(X) \leq r(Y) \\ \tau(X) \leq \tau(Y) \\ \rho(X) \leq \rho(Y) \end{cases}.$$

1.2.3. Comonotonie et antimonotonie

Les concepts de comonotonie et d'antimonotonie sont équivalents à celui de dépendance parfaite.

Définition 30. Comonotonie

Le couple $X = (X_1, X_2)$ est comonotone s'il existe des fonctions croissantes g_1 et g_2 telles que $X \stackrel{loi}{=} (g_1(Z), g_2(Z))$.

Définition 31. Antimonotonie

Le couple $X = (X_1, X_2)$ est antimonotone s'il existe une fonction croissante g_1 et une fonction décroissante g_2 telles que $X \stackrel{loi}{=} (g_1(Z), g_2(Z))$.

Ces notions correspondent à la plus forte dépendance possible, ce qui se traduit, quand on se place dans l'espace de Fréchet correspondant, à l'équivalence avec les bornes de la classe de Fréchet.

Ainsi, le couple $X = (X_1, X_2)$ est comonotone si, et seulement si, il admet W comme fonction de répartition.

De la même manière, le couple $X = (X_1, X_2)$ est antimonotone si, et seulement si, il admet M comme fonction de répartition.

Rappelons que la VaR de la somme de deux risques comonotones est égale à la somme des VaR de chacun de ces risques, pour autant que les fonctions de

répartitions marginales de ces deux variables aléatoires soient continues (cf. les propriétés des mesures de risques dans le **Chapitre 1**).

1.2.4. Dépendance positive par quadrant

Le concept de dépendance positive par quadrant repose sur la comparaison de la fonction de répartition d'un couple par rapport à celle d'un couple de mêmes marginales mais dont les composantes sont indépendantes.

Définition 32. *Dépendance positive par quadrant*

Le couple $X = (X_1, X_2)$ dont la fonction de répartition appartient à la classe $\mathbf{F}(F_1, F_2)$ est dépendant positivement par quadrant (DPQ) si pour tout $x \in \mathbf{R}^2$, on a : $\bar{F}_X(x) \geq \bar{F}_1(x_1)\bar{F}_2(x_2)$.

Le nom de cette structure de dépendance provient du fait que la probabilité que le vecteur X soit dans le quart de quadrant supérieur droit au point (x_1, x_2) est supérieure à celle que le couple $X^\perp$, version indépendante de X , y soit. Remarquons qu'un couple indépendant répond à la définition de couple DPQ.

Les probabilités conditionnelles nous offrent une caractérisation de la dépendance positive par quadrant.

Caractérisation : Le couple X est DPQ si, et seulement si,

$$\Pr[X_2 > x_2 \mid X_1 > x_1] \geq \Pr[X_2 > x_2] \text{ pour tout } x \in \mathbf{R}^2.$$

On déduit de cette caractérisation que :

- X est DPQ $\Leftrightarrow X^\perp \prec_{sm} X$;
- pour tout couple $X = (X_1, X_2)$ dépendant positivement par quadrant, on a l'équivalence :
 $X_1 \perp X_2 \Leftrightarrow r(X_1, X_2) = 0 \Leftrightarrow \tau(X_1, X_2) = 0 \Leftrightarrow \rho(X_1, X_2) = 0$.

1.2.5. Croissance conditionnelle

Définition 33. *Croissance conditionnelle*

Le couple X est dit conditionnellement croissant lorsque

- quels que soient $x_1 \leq y_1$, $\Pr[X_2 > t \mid X_1 = x_1] \geq \Pr[X_2 > t \mid X_1 = y_1]$ pour tout $t \in \mathbf{R}$,
- quels que soient $x_2 \leq y_2$

$$\Pr[X_1 > t \mid X_2 = x_2] \geq \Pr[X_1 > t \mid X_2 = y_2], \text{ pour tout } t \in \mathbf{R}.$$

Cette notion de croissance conditionnelle peut également s'exprimer à l'aide de l'ordre stochastique (cf. **Chapitre 1**). En effet, on a l'équivalence :

X est conditionnellement croissant

$$\Leftrightarrow \begin{cases} [X_1 | X_2 = x_2] \prec_{st} [X_1 | X_2 = y_2] \\ [X_2 | X_1 = x_1] \prec_{st} [X_2 | X_1 = y_1] \end{cases} \text{ pour tous } x_1 \leq y_1 \text{ et } x_2 \leq y_2.$$

Cette notion de dépendance est plus forte que celle de dépendance positive par quadrant : X est conditionnellement croissant $\Rightarrow X$ est DPQ.

2. Rappels sur la dépendance linéaire

Les catégories d'actifs en portefeuille d'une société d'assurance sont en général corrélées. Cette corrélation peut être matérialisée (cf. Planchet, Thérond et al. (2005)) par le biais de la matrice ci-dessous où l'on a noté z^{Actif} le mouvement brownien associé à l'évolution de l'actif « Actif ».

	$dz_t^{\text{Obligation}}$	dz_t^{Action}	$dz_t^{\text{Immobilier}}$	$dz_t^{\text{Monétaire}}$
$dz_t^{\text{Obligation}}$	dt	$\rho_{12}dt$	$\rho_{13}dt$	$\rho_{14}dt$
dz_t^{Action}	$\rho_{12}dt$	dt	$\rho_{23}dt$	$\rho_{24}dt$
$dz_t^{\text{Immobilier}}$	$\rho_{13}dt$	$\rho_{23}dt$	dt	$\rho_{34}dt$
$dz_t^{\text{Monétaire}}$	$\rho_{14}dt$	$\rho_{24}dt$	$\rho_{34}dt$	dt

En vue d'aboutir à une dé-corrélation des actifs, on utilise la décomposition de Choleski d'un vecteur gaussien. On considère un vecteur aléatoire gaussien $X = (X_1, X_2, \dots, X_n)$ dont la distribution est normale non dégénérée. Sa densité s'écrit de la manière suivante :

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

où l'on a noté $\mu = (\mu_1, \mu_2, \dots, \mu_n)$ le vecteur espérance des lois marginales et Σ la matrice de variance – covariance. Plus précisément, cette dernière matrice s'écrit :

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

Comme Σ est symétrique, il est possible de trouver une matrice T définie comme suit :

$$\mathbf{T} = \begin{bmatrix} b_{11} & 0 & 0 & 0 \\ b_{21} & b_{22} & 0 & 0 \\ b_{31} & b_{32} & b_{33} & 0 \\ b_{41} & b_{42} & b_{43} & b_{44} \end{bmatrix},$$

et vérifiant $\Sigma = \mathbf{T}\mathbf{T}'$. D'autre part, on sait que si $\mathbf{Z}$ est un vecteur i.i.d. de loi commune $\mathcal{N}(0;1)$, alors le vecteur $\mathbf{X} = \mathbf{T}\mathbf{Z} + \boldsymbol{\mu}$ a pour matrice de variance-covariance Σ . Cette dernière relation permet de retrouver de proche en proche les valeurs des coefficients de la matrice triangulaire $\mathbf{T}$:

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

En considérant la variance de chacun des termes de cette égalité, on trouve $b_{11} = 1$.

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

En considérant la variance de chacun des termes de cette égalité, on trouve $b_{11} + b_{22} = 1$. Si l'on s'intéresse à la covariance, on obtient 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 $b_{22} = (1 - \rho_{12}^2)^{1/2}$. L'application de cette méthode permet d'aboutir aux valeurs des b_{ij} :

$$b_{ij} = \frac{\rho_{ij} - \sum_{k=1}^{j-1} b_{ik}b_{jk}}{\left(\rho_{jj} - \sum_{k=1}^{j-1} b_{jk}^2 \right)^{1/2}}.$$

Et la matrice $\mathbf{T} = (t_{i,j})$ s'écrit donc :

$$\mathbf{T} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ \rho_{12} & \sqrt{1 - \rho_{12}^2} & 0 & 0 \\ \rho_{13} & \frac{\rho_{23} - \rho_{12}\rho_{13}}{\sqrt{1 - \rho_{12}^2}} & \sqrt{1 - \rho_{13}^2 - \frac{(\rho_{23} - \rho_{12}\rho_{13})^2}{1 - \rho_{12}^2}} & 0 \\ \rho_{14} & \frac{\rho_{24} - \rho_{14}\rho_{12}}{\sqrt{1 - \rho_{12}^2}} & \frac{\rho_{34} - \rho_{14}\rho_{13} - \frac{\rho_{24} - \rho_{43}\rho_{12}}{\sqrt{1 - \rho_{12}^2}} \frac{\rho_{23} - \rho_{13}\rho_{12}}{\sqrt{1 - \rho_{12}^2}}}{\sqrt{1 - \rho_{13}^2 - \frac{(\rho_{23} - \rho_{13}\rho_{12})^2}{1 - \rho_{12}^2}}} & \sqrt{1 - t_{4,1}^2 - t_{4,2}^2 - t_{4,3}^2} \end{bmatrix}$$

3. La théorie des copules

Le concept de copule a été introduit par Abe Sklar en 1959, comme solution d'un problème de probabilité énoncé par Maurice Fréchet dans le cadre des espaces métriques aléatoires (travaux réalisés avec Berthold Schweizer).

Ce concept est resté longtemps très peu utilisé en statistique ; on peut toutefois citer les travaux de Kimeldorf et Sampson sur la dépendance (1975) ainsi que les recherches de Paul Deheuvels à la fin des années 70.

L'étude systématique des copules débute dans le milieu des années 1980 avec Christian Genest et son équipe. L'article fondateur est celui de Genest et Mac Kay (1986).

Depuis, de nombreux développements statistiques ont été menés par C. Genest et ses coauteurs (R.J. MacKay, L.P. Rivest, P. Capéraà, A.L. Fougères, K. Ghoudi, B. Rémillard).

Les copules sont devenues un outil de base dans la modélisation des distributions multivaluées en finance et en assurance.

Le présent chapitre est une introduction à la théorie des copules, il en présente de manière succincte les principales définitions et outils.

3.1. Copules, lois conjointes et dépendance

3.1.1. Rappels et notations

Si F désigne la fonction de répartition d'une v.a. X , alors par définition $F(x) = \Pr(X \leq x)$; la fonction F est continue à droite avec des limites à gauche ; on définit alors l'inverse généralisée de F en posant $F^{-1}(y) = \inf(x | F(x) \leq y)$.

La variable aléatoire $U = F(X)$ est distribuée selon une loi uniforme sur $[0,1]$ (car $\Pr(U \leq u) = \Pr(X \leq F^{-1}(u)) = u$).

3.1.2. Définitions et résultats de base

Une copule en dimension n est une fonction C définie sur $[0,1]^n$ et à valeurs dans $[0,1]$ possédant les propriétés suivantes :

$$C_n(u) = C(1, \dots, 1, u, 1, \dots, 1) = u,$$

$$\sum_{k_1=1}^2 \dots \sum_{k_n=1}^2 (-1)^{k_1 + \dots + k_n} C(u_1^{k_1}, \dots, u_n^{k_n}) \geq 0,$$

pour tous $(u_1^1, \dots, u_n^1)$ et $(u_1^2, \dots, u_n^2)$ tels que $u_i^1 \leq u_i^2$. Cette condition s'interprète en disant que le C -volume de l'hypercube $[x, y]$ de $[0,1]^n$ est positif quel que soit l'hypercube, le C -volume étant défini par :

$$V_C([x, y]) = \sum \varepsilon(z) C(z_1, \dots, z_n),$$

la somme étant étendue à tous les sommets de l'hypercube, et la fonction $\varepsilon(z)$ valant 1 si un nombre pair de sommets ont une valeur de x et -1 sinon. Il s'agit d'une propriété de « croissance » de la copule. On dit que la fonction C est « supermodulaire » (voir également le sous-paragraphe 1.2.2).

Il résulte directement de la définition d'une copule que C est une fonction de répartition d'un vecteur dont les composantes sont uniformes. De plus, si on se donne des distributions réelles univariées $F_1, \dots, F_n$, alors $F(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n))$ est la fonction de répartition d'un vecteur aléatoire dont les lois marginales sont les lois $F_1, \dots, F_n$. Cela est la conséquence directe du fait que pour toute v.a. X de fonction de répartition F , $F(X)$ suit une loi uniforme sur $[0, 1]$.

Exemple n° 1 : Soit $F(x, y) = (1 - e^{-x} - e^{-y})^{-1}$ la distribution logistique bivariée de Gumbel ; par intégration on trouve que les lois marginales sont identiques, de fonction de répartition $F_X(x) = (1 - e^{-x})^{-1}$; cette fonction est aisément inversible et on en déduit que :

$$C(u, v) = F(F_X^{-1}(u), F_X^{-1}(v)) = \frac{uv}{u + v - uv}.$$

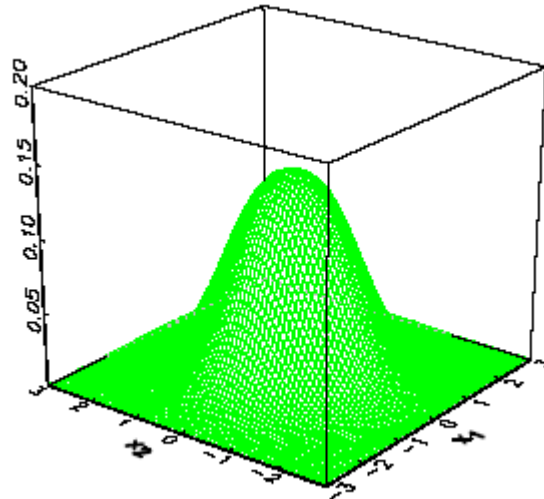
Le calcul explicite de la copule est toutefois rarement possible car elle nécessite l'inversion des fonctions de répartition des marginales.

Exemple n° 2 : Utilisation de la copule cubique avec des marginales gaussiennes.

La copule cubique est définie pour $1 \leq \alpha \leq 2$ par :

$$C(u_1, u_2) = u_1 u_2 + \alpha u_1 (u_1 - 1) (2u_1 - 1) u_2 (u_2 - 1) (2u_2 - 1).$$

Avec des marginales gaussiennes, on obtient la fonction de répartition représentée sur la Figure 30.

Figure 30 - Fonction de répartition de la copule cubique

La copule cubique présente la particularité suivante : si les marginales sont continues et symétriques, leur coefficient de corrélation est nul, alors qu'elles ne sont pas indépendantes. On peut ainsi aisément construire des distributions à marginales non corrélées dépendantes.

Les copules apparaissent ainsi être un outil naturel pour construire des distributions multivariées. De plus, lorsque les distributions marginales sont suffisamment régulières on a le théorème de représentation suivant :

Théorème 5. Théorème de Sklar

Si F est une distribution de dimension n dont les lois marginales $F_1, \dots, F_n$ sont continues, alors il existe une copule unique telle que :

$$F(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n)).$$

Remarque : Lorsque les marginales ne sont pas continues, il est toujours possible de définir une copule, mais celle-ci n'est plus unique et de ce fait perd beaucoup de son intérêt. En effet, on peut toujours poser :

$$C(u_1, \dots, u_n) = F(F_1^{-1}(u_1), \dots, F_n^{-1}(u_n)),$$

avec les « inverses généralisées » des marginales : $F_j^{-1}(u) = \inf \{t \mid F_j(t) \geq u\}$ (ces fonctions sont également appelées fonctions quantiles).

Ce résultat est important en pratique pour les applications à l'assurance et à la finance, car il indique que l'analyse d'une problématique multivariée peut être décomposée en deux étapes, tout d'abord l'identification des distributions marginales, puis l'analyse de la structure de dépendance entre les composantes.

Le théorème de Sklar fournit ainsi une représentation canonique d'une distribution multivariée, via la donnée des distributions marginales et de la structure de dépendance.

Deux fonctions particulières jouent un rôle important dans l'analyse des copules, les bornes de Fréchet ; on pose par définition :

$$C^-(u_1, \dots, u_n) = \max \left(\sum_{i=1}^n u_i - n + 1, 0 \right) \text{ et } C^+(u_1, \dots, u_n) = \min(u_1, \dots, u_n).$$

On a pour toute copule C et tout u dans $[0,1]^n$:

$$C^-(u) \leq C(u) \leq C^+(u).$$

La majoration est la conséquence directe de $C_n(u) = C(1, \dots, 1, u, 1, \dots, 1) = u$; la borne supérieure est atteinte avec le vecteur $(U, \dots, U)$ où la v.a. U suit une loi uniforme sur $[0,1]$. Cela prouve que la fonction $C^+(\cdot)$ est une copule. Ce résultat reste valable à une transformation strictement croissante des composantes près ; on parle alors de variables comonotoniques.

En dimension 2, la borne inférieure est atteinte avec le couple $(U, 1 - U)$; en dimension supérieure à 2 la fonction $C^-(\cdot)$ n'est pas une copule.

La situation d'indépendance des v.a. est associée à une copule particulière, la copule $C^\perp(\cdot)$, définie par $C^\perp(u_1, \dots, u_n) = \prod_{i=1}^n u_i$.

Propriété : Si $X = (X_1, \dots, X_n)$ est un vecteur aléatoire et $h_1, \dots, h_n$ sont des fonctions strictement croissantes, alors $h(X) = (h_1(X_1), \dots, h_n(X_n))$ a la même copule que X .

En effet, comme les fonctions h_i sont strictement croissantes, elles sont inversibles et

$$P(h_1(X_1) \leq x_1, \dots, h_n(X_n) \leq x_n) = P(X_1 \leq h_1^{-1}(x_1), \dots, X_n \leq h_n^{-1}(x_n)),$$

donc

$$P(h_1(X_1) \leq x_1, \dots, h_n(X_n) \leq x_n) = C(F_{X_1}(h_1^{-1}(x_1)), \dots, F_{X_n}(h_n^{-1}(x_n))).$$

De par l'unicité de la décomposition copule (en remarquant que la transformation affecte uniquement les marginales dont la loi devient $F_{h(X_i)} = F_{X_i} \circ h_i^{-1}$), on conclut que $h(X) = (h_1(X_1), \dots, h_n(X_n))$ a la même copule que X .

3.1.3. Exemples de copules

Il existe un grand nombre de copules adaptées à différentes situations ; toute répartition associée à un vecteur dont les marginales sont uniformes sur $[0,1]$ définit une copule. Toutefois, quelques formes particulières sont souvent

utilisées en pratique du fait de leur simplicité de mise en œuvre. On peut notamment citer les exemples ci-après :

Tableau 5 - Exemples de copules

Nom	Copule
Gaussienne ²	$C(u_1, \dots, u_n, \rho) = \Phi_\rho \left(N^{-1}(u_1), \dots, N^{-1}(u_n) \right)$
Clayton	$C(u, v, \theta) = (u^{-\theta} + v^{-\theta} - 1)^{-1/\theta}$
Franck	$C(u, v, \theta) = -\frac{1}{\theta} \ln \left(1 + \frac{(\exp(-\theta u) - 1)(\exp(-\theta v) - 1)}{\exp(-\theta) - 1} \right)$
Gumbel	$C(u, v, \theta) = \exp \left(- \left[(-\ln(u))^\theta + (-\ln(v))^\theta \right]^{1/\theta} \right)$

On peut illustrer avec une copule simple le fait (classique) que la non corrélation peut être observée dans des situations de forte dépendance entre les variables.

Il suffit de considérer la copule suivante :

$$C(u_1, u_2) = \begin{cases} u_1, & 0 \leq u_1 \leq \frac{u_2}{2} \leq \frac{1}{2}, \\ \frac{u_2}{2}, & 0 \leq \frac{u_2}{2} \leq u_1 \leq 1 - \frac{u_2}{2}, \\ u_1 + u_2 - 1, & \frac{1}{2} \leq 1 - \frac{u_2}{2} \leq u_1 \leq 1, \end{cases}$$

alors on vérifie que $\text{Cov}(U_1, U_2) = 0$ et pourtant $\Pr(U_2 = 1 - |2U_1 - 1|) = 1$, ce qui signifie que si on connaît la valeur prise par U_1 alors on connaît presque sûrement celle prise par U_2 .

3.1.4. Les copules archimédiennes

La notion de copule archimédienne regroupe un certain nombre de copules ci-dessus (Clayton, Gumbel, Franck); l'idée d'une copule archimédienne de générateur φ est que la transformation $\omega(u) = \exp(-\varphi(u))$ appliquée aux marginales « rend les composantes indépendantes » :

$$\omega(C(u_1, \dots, u_n)) = \prod_{i=1}^n \omega(u_i).$$

². Φ_ρ est la distribution normale multivariée de matrice de corrélations ρ .

Définition 34. Copule archimédienne

La copule archimédienne de générateur φ est définie par l'égalité

$$C(u_1, \dots, u_n) = \varphi^{-1}(\varphi(u_1) + \dots + \varphi(u_n)), \text{ dès lors que } \sum_{i=1}^n \varphi(u_i) \leq \varphi(0), \text{ et}$$

$C(u_1, \dots, u_n) = 0$ sinon. Le générateur doit être choisi de classe C^2 de sorte que $\varphi(1) = 0$, $\varphi'(u) \leq 0$ et $\varphi''(u) > 0$.

Les copules archimédiennes peuvent être également définies via un conditionnement par une « variable de structure », Z ; à cet effet, on se donne une v.a. positive Z telle que :

- la transformée de Laplace de Z est $\psi = \varphi^{-1}$,
- conditionnellement à Z les v.a. X_i sont indépendantes,
- $\Pr(X_i \leq x | Z) = \omega(F_{X_i}(x))^Z$,

alors la copule du vecteur $(X_1, \dots, X_n)$ est

$C(u_1, \dots, u_n) = \varphi^{-1}(\varphi(u_1) + \dots + \varphi(u_n))$. L'idée importante ici est que, via un conditionnement, on retrouve la situation d'indépendance.

En effet, par définition de la copule :

$$\begin{aligned} C(u_1, \dots, u_n) &= \Pr(F_{X_1}(X_1) \leq u_1, \dots, F_{X_n}(X_n) \leq u_n) \\ &= E \left[\Pr \left[(F_{X_1}(X_1) \leq u_1, \dots, F_{X_n}(X_n) \leq u_n) | Z \right] \right] \end{aligned}$$

mais on a l'identité $E[\exp(-\varphi(u)Z)] = \varphi^{-1}(\varphi(u)) = u$ et, conditionnellement à Z les v.a. X_i sont indépendantes et de part la forme des distributions conditionnelles :

$$C(u_1, \dots, u_n) = E \prod_{i=1}^n \omega(u_i^Z) = E \exp \left(-Z \sum_{i=1}^n \varphi(u_i) \right) = \varphi^{-1}(\varphi(u_1) + \dots + \varphi(u_n)).$$

En particulier, la copule $C^\perp(\cdot)$ est archimédienne, de générateur $-\ln u$, qui est un cas limite de copule de Gumbel ; en effet, la copule de Gumbel est archimédienne de générateur $(-\ln u)^\theta$, avec $\theta \geq 1$. Les copules de Clayton et de Franck sont également archimédiennes.

³. Cette transformation définit une fonction de répartition car la fonction φ est décroissante.

Tableau 6 - Générateurs de copules archimédiennes

Copule	Générateur
Indépendance	$-\ln u$
Clayton	$u^{-\theta} - 1, \theta \geq 0$
Franck	$-\ln \left(\frac{e^{-\theta u} - 1}{e^{-\theta} - 1} \right), \theta \neq 0$
Gumbel	$(-\ln u)^\theta, \theta \geq 1$

3.1.5. Copules et mesures de dépendance

Comme on l'a vu plus haut dans ce chapitre, il existe en statistique un certain nombre de grandeurs proposées par leurs auteurs comme des mesures de la dépendance entre deux variables aléatoires. Parmi elles, on peut citer le ρ de Spearman (coefficient de corrélation de Pearson, mais calculé sur les rangs des deux variables), le τ de Kendall (probabilité de concordance des rangs moins probabilité de discordance des rangs), l'indice de Gini (mesure d'équirépartition de deux échantillons).

Ces coefficients sont en effet associés à la structure de dépendance du couple, et on peut donc logiquement les exprimer en fonction de la copule associée.

Tableau 7 - Mesures de dépendance et copules associées

Indice	Expression usuelle	Expression copule
Rho de Spearman	$\rho = 1 - \frac{6 \sum_{i=1}^n d_i^2}{n(n^2 - 1)}$	$\rho = 12 \int_0^1 \int_0^1 u_1 u_2 dC(u_1, u_2) - 3$
Tau de Kendall	$\tau = \frac{4R}{n(n-1)} - 1$	$\tau = 4 \int_0^1 \int_0^1 C(u_1, u_2) dC(u_1, u_2) - 1$
Indice de Gini	$\gamma = \frac{\sum_{i,j=1}^n x_i - x_j }{2\mu n^2}$	$\gamma = 2 \int_0^1 \int_0^1 (u_1 + u_2 - 1 - u_1 - u_2) dC(u_1, u_2)$

N.B. Dans le Tableau 7, les notations suivantes ont été employées :

- le d_i de la formule du ρ de Spearman représente la différence des rangs de l'individu i dans chacun des classements ;

- dans l'expression du τ de Kendall, R désigne la somme sur tous les individus du nombre de paires concordantes ;
- dans celle de l'indice de Gini, μ représente l'espérance de la distribution.

On pourra vérifier que le coefficient de corrélation de Pearson ne peut se mettre sous une forme ne faisant intervenir que la copule, les marginales restent présentes dans l'expression : ce coefficient n'est pas, comme on l'a vu, une mesure de dépendance.

3.2. Inférence statistique

En pratique lors de la mise en œuvre d'un modèle intégrant des copules il convient d'être capable d'estimer la structure de dépendance à partir des données disponibles.

Les méthodes d'inférence statistique utilisées sont des déclinaisons de méthodes standards, paramétriques ou non paramétriques selon le cas.

3.2.1. Estimation non paramétrique

3.2.1.1 Rappel : fonction de répartition empirique

Le cadre général de l'estimation non paramétrique d'une loi marginale s'appuie sur la fonction de répartition empirique, définie par :

$$F_K(x) = \frac{1}{K} \sum_{i=1}^K 1_{(x_i \leq x)}$$

pour un n -échantillon $(x_1, \dots, x_K)$ de la loi F . En dimension n , si l'on se donne $(x_1^k, \dots, x_n^k)_{1 \leq k \leq K}$ un K -échantillon du vecteur (de dimension n) X , on peut généraliser l'expression de la fonction de répartition empirique en posant :

$$F_K(x_1, \dots, x_n) = \frac{1}{K} \sum_{i=1}^K 1_{(x_1^i \leq x_1, \dots, x_n^i \leq x_n)}.$$

Cet estimateur conduit à un estimateur non paramétrique naturel d'une copule présenté *infra*.

3.2.1.2 Copule empirique

Soit comme ci-dessus $(x_1^k, \dots, x_n^k)_{1 \leq k \leq K}$ un K -échantillon du vecteur (de dimension n) X . En observant que $C(u_1, \dots, u_n) = F(F_1^{-1}(u_1), \dots, F_n^{-1}(u_n))$ on introduit la notion de copule empirique à partir des version empiriques des fonctions de répartition de cette expression :

$$C_K(u_1, \dots, u_n) = F_K(F_{1,K}^{-1}(u_1), \dots, F_{n,K}^{-1}(u_n)),$$

avec $F_{j,K}^{-1}(u) = \inf \{t \mid F_{j,K}(t) \geq u\}$. Cette notion a été introduite par Deheuvels (1979).

Remarque : C est la fonction de répartition d'un vecteur dont les marginales sont uniformes sur $[0,1]$; cette propriété n'est plus vraie pour C_K qui est, conditionnellement aux observations, la distribution d'un vecteur dont les marges sont réparties uniformément sur l'ensemble discret $\left\{1, \frac{1}{K}, \dots, \frac{K-1}{K}\right\}$. Il suffit donc de connaître les valeurs de la copule empirique en ces points discrets ; on a en particulier le lien avec les statistiques d'ordre suivant :

$$C_K\left(\frac{k_1}{K}, \dots, \frac{k_n}{K}\right) = \frac{1}{K} \sum_{k=1}^K 1\left\{x_1^k \leq x_1^{(k_1)}, \dots, x_n^k \leq x_n^{(k_n)}\right\}.$$

Cette expression peut être également écrite à partir des rangs des observations (ce qui est intuitif, puisque la copule est invariante par toute transformation croissante des marginales) :

$$C_K\left(\frac{k_1}{K}, \dots, \frac{k_n}{K}\right) = \frac{1}{K} \sum_{k=1}^K 1\left\{r_1^k \leq k_1, \dots, r_n^k \leq k_n\right\}.$$

Deheuvels (1979) a obtenu des résultats asymptotiques sur C_K et $K^{-1/2}(C_K - C)$, généralisant les résultats classiques sur la fonction de répartition empirique.

Les copules empiriques sont notamment utiles pour fournir des estimateurs non paramétriques de mesures de dépendance.

Par exemple on peut ainsi proposer l'estimateur suivant du ρ de Spearman :

$$\hat{\rho} = \frac{12}{K^2 - 1} \sum_{k_1=1}^K \sum_{k_2=1}^K \left(\hat{C}\left(\frac{k_1}{K}, \frac{k_2}{K}\right) - \frac{k_1 k_2}{K^2} \right).$$

On pourra se référer à Deheuvels, Peccati et Yor (2004) pour des résultats asymptotiques sur les copules empiriques.

3.2.1.3 Identification d'une copule archimédienne

Dans le cas d'une copule archimédienne, la copule est entièrement définie par son générateur ; on remarque alors que si $K(u) = \Pr(C(U_1, U_2) \leq u)$, on a :

$$K(u) = u - \frac{\varphi(u)}{\varphi'(u)}.$$

En effet, par définition d'une copule archimédienne, on a : $K(u) = P(\varphi^{-1}(\varphi(U_1) + \varphi(U_2)) \leq u)$; or la loi conditionnelle se calcule à partir de la copule : $P(U_2 \leq u_2 | U_1 = u_1) = \frac{\partial C}{\partial u_1}(u_1, u_2)$.

La fonction K est liée au τ de Kendall par $\tau = 4 \int_0^1 (1 - K(u)) du - 1$.

En effet, comme $\tau = 4 E(C(U_1, U_2)) - 1$, d'après la définition de K qui est la fonction de répartition de $C(U_1, U_2)$, on en déduit la relation liant K et τ .

3.2.2. Estimation paramétrique et semi-paramétrique

On se place ici dans le cas où la distribution conjointe dépendrait d'un paramètre, que l'on cherche à estimer.

3.2.2.1 Méthode des moments

Cette méthode est notamment utilisée pour les mesures de dépendance ; l'estimateur des moments de la mesure de dépendance considérée est alors simplement obtenu en égalant l'expression paramétrique (analytique) de la mesure avec un estimateur non paramétrique de cette même mesure.

Par exemple, pour le tau de Kendall, on a $\hat{\tau} = \frac{c - d}{c + d}$ avec c (respectivement d) le nombre de paires concordantes (respectivement discordantes) dans l'échantillon. Donc pour une copule de Gumbel l'expression $\tau = 1 - \frac{1}{\theta}$ conduit à l'estimateur des moments $\hat{\theta} = \frac{1}{1 - \hat{\tau}}$.

3.2.2.2 Maximum de vraisemblance

Il s'agit de la méthode classique du maximum de vraisemblance. De l'égalité $F(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n))$ on déduit par dérivation l'expression de la densité du vecteur $(X_1, \dots, X_n)$:

$$f(x_1, \dots, x_n) = c(F_1(x_1), \dots, F_n(x_n)) \prod_{i=1}^n f_i(x_i),$$

où $c(u_1, \dots, u_n) = \frac{\partial^n C(u_1, \dots, u_n)}{\partial u_1 \dots \partial u_n}$ désigne la densité de la copule. L'expression

de la log-vraisemblance de l'échantillon $(x_1^k, \dots, x_n^k)_{1 \leq k \leq K}$ s'en déduit immédiatement :

$$l(\theta) = \sum_{k=1}^K \ln \left(c \left(F_1(x_1^k, \theta), \dots, F_n(x_n^k, \theta), \theta \right) \right) + \sum_{k=1}^K \sum_{i=1}^n f_i(x_i^k, \theta).$$

Il reste maximiser cette expression en θ , ce qui peut s'avérer en pratique fastidieux. La méthode IFM, présentée ci-après, propose une résolution par étapes du problème de maximisation.

3.2.2.3 Méthode IFM

Cette méthode, proposée par Louis et Shih (1995) consiste à « découper » le problème d'estimation en deux étapes successives : l'estimation des paramètres $\theta_1, \dots, \theta_n$ des marginales, puis l'estimation du paramètre θ_C de la copule.

Le paramètre θ est donc décomposé sous la forme $\theta = (\theta_1, \dots, \theta_n, \theta_C)$ et on commence par déterminer les estimateurs du maximum de vraisemblance des paramètres des marginales, soit :

$$\hat{\theta}_i = \arg \max_{\theta_i} \sum_{k=1}^K f_i(x_i^k, \theta_i).$$

On « injecte » alors ces estimateurs dans la partie « copule » de la log-vraisemblance, ce qui conduit à :

$$\hat{\theta}_C = \arg \max_{\theta_C} \sum_{k=1}^K \ln \left(c \left(F_1(x_1^k, \hat{\theta}_1), \dots, F_n(x_n^k, \hat{\theta}_n), \theta_C \right) \right).$$

D'autres procédures peuvent être imaginées, comme par exemple l'estimation non paramétrique des marginales suivi d'un maximum de vraisemblance pour le paramètre de la copule (procédure « omnibus » de Genest et al. (1995) ou Louis et Shih (1995)).

3.3. Copules et valeurs extrêmes

Une copule de valeurs extrêmes est par définition une copule vérifiant, pour tout $\lambda > 0$:

$$C(u_1^\lambda, \dots, u_n^\lambda) = [C(u_1, \dots, u_n)]^\lambda.$$

On peut vérifier par exemple qu'une copule de Gumbel possède cette propriété. Caperaa, Fougères et Genest (1997) ont introduit la notion de « copule archimax » qui englobe les copules archimédiennes et les copules EV.

Le lien entre la théorie des copules et les valeurs extrêmes est donné par le théorème suivant.

Théorème 6. Si $X_p^{(p)} = \max_{1 \leq i \leq p} X_i$ désigne le maximum d'un échantillon de taille p , et si on a n échantillons et des constantes de normalisation telles que :

$$\lim_{p \rightarrow \infty} \Pr \left(\frac{X_{p,1}^{(p)} - a_{p,1}}{b_{p,1}} \leq x_1, \dots, \frac{X_{p,n}^{(p)} - a_{p,n}}{b_{p,n}} \leq x_n \right) = F(x_1, \dots, x_n),$$

alors la copule associée à la représentation canonique de F est une copule de valeurs extrêmes.

Par ailleurs, on dit qu'une copule bivariable présente de la dépendance de queue si $\exists \lambda \in]0,1]$ tel que :

$$\lambda(u) = \frac{\bar{C}(u, u)}{1 - u} \xrightarrow{u \rightarrow 1} \lambda,$$

avec la fonction $\bar{C}(u, v) = 1 - u - v + C(u, v)$ de survie conjointe. L'idée de ce quotient est de mesurer la probabilité que U dépasse le seuil u sachant que V l'a dépassé :

$$\lambda(u) = \Pr(U > u | V > u).$$

La modélisation de la dépendance des extrêmes est particulièrement importante dans les problématiques de détermination d'un capital de solvabilité. En effet, lorsque celui-ci est censé prémunir la société de la ruine avec une très forte probabilité (99,5 % dans les documents de travail du projet Solvabilité 2), il s'agit de modéliser non seulement les queues de distribution des risques qui concourent à la solvabilité de la compagnie mais également les dépendances entre ces queues.

En particulier, des événements extrêmes telles que les tempêtes peuvent toucher simultanément les queues de distribution de branches d'assurance dont la dépendance des queues de distribution sont faibles.

Ces considérations se heurtent aux difficultés pratiques de modélisation de tels dépendances sur les queues de distribution du fait notamment du peu d'information disponibles pour les estimer.

3.4. Copules et simulation

Les copules fournissent un cadre d'analyse bien adapté à la « fabrication » de distributions conjointes possédant une structure de dépendance donnée.

En effet, soit $X = (X_1, \dots, X_n)$ est un vecteur gaussien centré de matrice de variance-covariance fixée ; on peut remplacer les lois marginales gaussiennes par d'autres lois (par exemple pour tenir compte d'un caractère leptocurtique affirmé des observations) tout en conservant la même structure de dépendance.

Pour simuler des réalisations d'un vecteur $X = (X_1, \dots, X_n)$ à partir de la décomposition « copule » de F , il faut être capable de simuler des réalisations

d'un vecteur U ayant des marginales uniformes et dont la distribution est la copule C , puisqu'on a l'égalité en loi:

$$X = (F_1^{-1}(U_1), \dots, F_n^{-1}(U_n)).$$

La suite de ce paragraphe illustre quelques méthodes qui permettent de mener à bien cette tâche. Par ailleurs, la génération de nombres aléatoires et de simulation de réalisations de variables aléatoires est abordée dans le Chapitre 6.

3.4.1. La méthode des distributions

On suppose que l'on se trouve dans une situation où la loi conjointe de X est plus facile à simuler directement que la copule C ; c'est par exemple le cas de la copule gaussienne : un vecteur gaussien de dimension n est aisé à simuler (via la décomposition de Cholesky de la matrice de variance-covariance), alors que la copule gaussienne n'est pas simple à simuler directement.

La méthode consiste à simuler des réalisations de X (de loi F) et à appliquer ensuite la transformation $U = (F_1(X_1), \dots, F_n(X_n))$.

3.4.2. La méthode des distributions conditionnelles

3.4.2.1 Présentation générale

On se donne $(U_1, U_2)^4$ un vecteur dont la distribution est C ; on simule (v_1, v_2) uniformes et indépendantes.

On pose alors $u_1 = v_1$, puis

$$u_2 = C_{u_1}^{-1}(v_2),$$

avec $C_{u_1}(u_2) = P(U_2 \leq u_2 | U_1 = u_1)$. Cette expression se calcule à partir de la copule C :

$$C_{2|1}(u_1, u_2) = P(U_2 \leq u_2 | U_1 = u_1) = \frac{\partial C}{\partial u_1}(u_1, u_2).$$

En effet, si l'on écrit

$$P(U_2 \leq u_2 | u_1 \leq U_1 \leq u_1 + h) = \frac{P(U_1 \leq u_1 + h, U_2 \leq u_2) - P(U_1 \leq u_1, U_2 \leq u_2)}{P(U_1 \leq u_1 + h) - P(U_1 \leq u_1)},$$

(qui est simplement une application de la formule $P(B)P(A|B) = P(A \cap B)$); cette expression s'écrit à l'aide de la copule :

$$P(U_2 \leq u_2 | u_1 \leq U_1 \leq u_1 + h) = \frac{(C(u_1 + h, u_2) - C(u_1, u_2))/h}{(C(u_1 + h, 1) - C(u_1, 1))/h}.$$

⁴. Cette méthode s'étend sans difficulté en dimension n .

On remarque alors que puisque les marginales sont uniformes alors $C(u,1)=u$ et donc le dénominateur de la fraction est égal à 1 ; on fait alors tendre h vers 0 pour obtenir le résultat.

Toutefois en pratique il n'est en général pas aisé d'inverser la fonction de répartition conditionnelle, sauf dans quelques cas particuliers.

Exemple : Copule de Franck

Dans le cas de la copule de Franck la méthode des distributions conditionnelles peut être mise en œuvre simplement. L'expression de cette copule est :

$$C(u, v, \theta) = -\frac{1}{\theta} \ln \left(1 + \frac{(\exp(-\theta u) - 1)(\exp(-\theta v) - 1)}{\exp(-\theta) - 1} \right).$$

On en déduit que $C_{2|1}(u, v) = \frac{\exp(-\theta u)(\exp(-\theta v) - 1)}{\exp(-\theta) - 1 + (\exp(-\theta u) - 1)(\exp(-\theta v) - 1)}$ et on peut alors inverser cette équation (en résolvant l'équation en x $C_{2|1}(u, x) = v$) pour obtenir :

$$C_{2|1}^{-1}(u, v) = -\frac{1}{\theta} \ln \left(1 + \frac{v(\exp(-\theta) - 1)}{v + (1 - v)\exp(-\theta u)} \right).$$

La simulation peut donc être mise en œuvre simplement en simulant indépendamment deux réalisations v_1 et v_2 de variables uniformes sur $[0,1]$ puis en utilisant la transformation :

$$\begin{cases} u = v_1, \\ v = C_{2|1}^{-1}(u, v_2). \end{cases}$$

3.4.2.2 Cas particulier des copules archimédiennes

Dans le cas particulier d'une copule archimédienne, la distribution conditionnelle s'exprime à l'aide du générateur de la copule et conduit à :

$$u_2 = \varphi^{-1} \left(\varphi \left(\varphi^{-1} \left(\varphi \left(\frac{v_1}{v_2} \right) \right) \right) \right) - \varphi(v_1).$$

3.4.3. Les méthodes adaptées spécifiquement à une copule

Comme pour la simulation de loi univariées, des méthodes *ad hoc* adaptées à une copule particulière peuvent être élaborées; par exemple pour la copule de Clayton. L'algorithme décrit ici est proposé par Devroye (1986).

Considérons la copule définie par :

$$C(u, v, \theta) = \left(u^{-\theta} + v^{-\theta} - 1 \right)^{-1/\theta},$$

Sa simulation peut passer par l'utilisations de l'algorithme suivant :

- simulation de x_1 et x_2 deux réalisations indépendantes de v.a. de loi exponentielle de paramètre 1 ($\mathcal{E}(1)$),
- simulation d'une réalisation x issue d'une v.a. de loi $\Gamma(1, \theta)$,
- calcul de $u_2 = \left(1 + \frac{x_2}{x}\right)^{-\theta}$ et $u_1 = \left(1 + \frac{x_1}{x}\right)^{-\theta}$.

Ainsi (u_1, u_2) est distribué selon la copule de Clayton avec des marginales uniformes. Le passage à la réalisation d'un couple (Y_1, Y_2) avec la même structure de dépendance mais des marginales quelconques pourra, par exemple, passer la méthode d'inversion de la fonction de répartition :

$$(y_1, y_2) = (F_1^{-1}(u_1), F_2^{-1}(u_2)).$$

3.5. La méthode NORTA

La théorie des copules fournit des moyens efficaces d'intégrer dans les modèles dynamiques la dépendance entre branches. Une situation que l'on rencontre en pratique souvent est le besoin de simuler des montants de sinistre dans des branches dépendantes, en connaissant d'une part la loi des marginales (c'est à dire les lois de coût de chaque branche) et d'autre part la corrélation entre les branches. On suppose que l'on est dans une situation qui rend délicate l'estimation directe d'une copule, par exemple du fait d'un nombre restreint de données.

L'objectif est donc de simuler un vecteur $(X_1, \dots, X_d)$ dont on connaît les marginales, et en respectant des corrélations fixées. On suppose sans perte de généralité que les v.a. X_i sont centrées et réduites.

L'idée de la méthode NORTA (*normal to anything*) est de considérer que la structure de dépendance du vecteur $(X_1, \dots, X_d)$ est décrite par une copule normale, c'est à dire que l'on fait l'hypothèse que la loi conjointe de $(X_1, \dots, X_d)$ est décrite par l'égalité en loi

$$(X_1, \dots, X_d) = (F_1^{-1}(\Phi(Y_1)), \dots, F_d^{-1}(\Phi(Y_d))),$$

avec $(Y_1, \dots, Y_d)$ un vecteur gaussien centré de matrice de variance-covariance Σ_Y et Φ la fonction de répartition de la loi normale centrée réduite. Chacune des variables Y_i est également réduite, de sorte que Σ_Y est en fait une matrice de corrélation avec des un sur la diagonale et définie entièrement par la donnée des $\frac{d(d-1)}{2}$ termes en dessous.

La solution du problème consiste alors à choisir Σ_Y de sorte que l'on retrouve la matrice Σ_X en calculant les corrélations des variables

$(X_1, \dots, X_d)$. Plus précisément, si on note ρ_{ij} le terme générique de Σ_X , comme les variables sont centrées et réduites on a $\rho_{ij} = E[X_i X_j]$. De même, en notant r_{ij} le terme générique de Σ_Y , r_{ij} est le coefficient de corrélation entre Y_i et Y_j , pour $i < j$.

On a alors une relation de la forme $\rho_{ij} = \omega_{ij}(r_{ij})$ avec

$$\omega_{ij}(r) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} F_i^{-1} \circ \Phi(y_i) F_j^{-1} \circ \Phi(y_j) h(r, y_i, y_j) dy_i dy_j,$$

la fonction h étant la densité d'un vecteur gaussien bi-dimensionnel :

$$h(r, x, y) = \frac{1}{2\pi(1-r)} \exp\left[\frac{x^2 + y^2 - 2rxy}{2(1-r)}\right]. \quad \text{On obtient ainsi } \frac{d(d-1)}{2}$$

équations indépendantes les unes des autres en r_{ij} , que l'on peut résoudre numériquement.

On notera qu'il est préférable d'imposer comme donnée de base la matrice de corrélation des rangs, c'est à dire la matrice de corrélation du vecteur $(F_1(X_1), \dots, F_d(X_d))$ car on obtient alors une expression analytique de r_{ij} : $r_{ij} = 2 \sin(\pi \rho_{ij}^u / 6)$. Cela découle directement de ce que, dans l'intégrale définissant ω_{ij} , les marginales F_i sont remplacées par l'identité.

Dès que $d > 2$ il se peut que le problème n'ait pas de solution, c'est à dire que la matrice de coefficients r_{ij} ainsi obtenue ne soit pas semi-définie positive ; on peut alors chercher la matrice semi-définie positive la plus proche possible de Σ_Y , au sens d'une norme sur l'espace des matrices.

3.6. Exemples d'application des copules

Outre les questions de modélisation de la charge sinistre évoquées ci-dessus, la théorie des copules s'applique directement dans les problématiques de solvabilité (détermination du capital économique). Elle peut également être utilisée dans certains modèles d'actifs.

3.6.1. Mesure du capital sous risque

Les copules interviennent directement dans l'agrégation de risques lorsque la mesure de risque choisie est la VaR. En effet :

$$VaR(X_i, \alpha) = F_{X_i}^{-1}(\alpha).$$

Donc, dans le cas de deux risques, comme par définition d'une copule :

$$F_{X_1 + X_2}(x) = \int_{x_1 + x_2 \leq x} dC(F_{X_1}(x_1), F_{X_2}(x_2)),$$

on voit que la VaR de la somme dépend des distributions marginales et de la copule. Les VaR s'ajoutent si la copule est la borne supérieure de Fréchet.

Cet exemple pose le problème de l'agrégation de risques. Ce thème a notamment été abordé par Meyers, Klinker et Lalonde (2003) dont les travaux ont servi de fondation au modèle de capital économique *Prism* élaboré par l'agence de notation Fitch.

Ce problème est crucial en assurance puisque notamment la formule standard de calcul du SCR en solvabilité 2 reposera, au moins en partie, sur des agrégations de capitaux.

Rappelons que la démarche vise à associer à chaque risque, un capital puis à agréger ces capitaux à l'aide de formules dont les paramètres varient en fonction de la dépendance estimée entre les risques considérés.

À ce stade, il faut d'ailleurs noter que la formule proposée dans le cadre de l'étude quantitative QIS 3 pré-suppose que la dépendance entre les différents risques est linéaire. Cette hypothèse, pratique de manière opérationnelle, pose néanmoins des problèmes conceptuels lorsque le niveau de capital recherché est défini par une VaR d'ordre très élevé (99,5 %) et donc à un niveau de la distribution où les comportements ne sont pas gaussiens.

3.6.2. Application aux processus de diffusion

Les copules sont également utilisées pour analyser les structures de dépendance dans les processus de diffusion.

La copule brownienne a ainsi été introduite par Darsow et al. (1992).

L'intérêt de cette approche est de permettre de construire de nouveaux processus avec une structure de dépendance connue mais des marginales différentes, comme, par exemple, un mouvement brownien à marginales « Student » pour épaissir les queues de distribution. Cela peut constituer une alternative à l'introduction des processus à sauts.

4. Rachat de contrats d'épargne : un modèle ad hoc

Qu'il s'agisse de valoriser un portefeuille en IFRS ou de le projeter pour déterminer une MCEV, les méthodes de valorisation économique des contrats d'assurance vie de type épargne-retraire ne font plus l'économie de la prise en compte de phénomènes de dépendance tels que ceux du comportement de rachat des assurés avec la revalorisation de son épargne et celle observée sur le reste du marché.

Dans ce paragraphe, nous proposons un modèle d'attribution de participation aux bénéfices et de rachat simple qui peut, par exemple, être implémenté pour des calculs de MCEV ou de *Current Exit Value* (cf. le Chapitre 2).

4.1. Notations

Dans la suite, nous considérerons un contrat d'épargne en euros, de taux technique nul, avec les notations suivantes :

- taux minimum garanti : tmg ;

- taux minimum de participation aux bénéfices : taux_pb_min ;
- taux de revalorisation cible pour la période n : $\text{taux_revalo_cible}(n)$
- entrées (primes) de la période n : $\text{entrées}(n)$;
- sorties (rachats, décès) sur la période n : $\text{sorties}(n)$;
- encours en début de période n : $\text{encours}(n)$;
- encours en fin de période n avant revalorisation et prélèvements sur encours : $\text{encours}(n+1)^-$;
- taux de prélèvement sur encours : $\text{taux_prelev_encours}$;
- taux de rendement financier net des charges de placement pour la période n : $\text{taux_rdt_fi}(n)$;
- le volume de plus-values latentes sur les actifs R332-20 en fin de période n : $\text{PVL}(n+1)^-$;
- la contrainte contractuelle ou réglementaire de revalorisation (fonction du taux de rendement financier net des charges de placement de la période) : $\rho(\text{taux_rdt_fi})$.

Dans la suite, on omettra l'argument de la période lorsque cela ne portera pas à confusions.

4.2. Hypothèses simplificatrices

Le modèle proposé repose sur un certain nombre d'hypothèses simplificatrices détaillées comme suit :

- les sorties en cours de période sont revalorisées au taux minimum garanti (tmg) uniquement ;
- les bénéfices techniques sont supposés nuls ;
- les prélèvements sur encours interviennent en fin de période après revalorisation.

On déduit de ces hypothèses, qu'en fin de période, la valeur de l'encours s'élève à :

$$\text{encours}(n+1) = \frac{\left(\text{encours}(n) - \text{sorties}(n) + \frac{\text{entrées}(n)}{2} \right) \times (1 + \text{taux_revalo}(n))}{1 + \text{taux_prelev_encours}}.$$

4.3. Attribution de PB

L'objet de ce paragraphe est de définir un modèle d'attribution de participation aux bénéfices (PB).

Des projections en vue de déterminer une EEV/MCEV doivent intégrer des contraintes d'attribution de PB aussi réalistes que possible. Intuitivement, un modèle d'attribution de PB doit intégrer le taux de rendement des actifs réalisé par la société et un taux cible de revalorisation du contrat, fonction de contraintes juridiques (réglementaires et contractuelles) et de contraintes

économiques (liées aux conditions de marché et aux objectifs commerciaux) de revalorisation.

En effet, le taux de rendement financier des actifs en représentation des contrats (ou le cas échéant de l'actif général de la société) engage la société, au moins au titre des clauses contractuelles et réglementaires d'attribution de bénéfices financiers. Par ailleurs, ce taux indique le taux de revalorisation maximal que peut accorder l'assureur sans toucher à ses plus-values latentes.

Quant au taux cible, il dépend en partie du contexte commercial et des taux pratiqués par le marché indépendamment du taux de rendement des actifs propres à l'assureur.

4.3.1. Définition du taux cible de revalorisation

La société doit déterminer, en fonction de ses contraintes économiques et réglementaires un taux cible de revalorisation.

Le taux cible de revalorisation peut, par exemple, être fonction des taux du marché obligataire sur des horizons semblables à ceux des contrats d'épargne (ex : OAT 10 ans).

Il apparaît naturel que ce taux soit supérieur :

- au taux de prélèvement sur encours,
- au tmg du contrat considéré.

Par exemple, on peut proposer comme taux cible de revalorisation :

$$\text{taux revalo cible} = \max \{ \text{taux OAT}(10 \text{ ans}), \text{taux prelev encours}, \text{tmg} \} .$$

Le mécanisme de détermination de ce taux cible peut être calibré sur les dernières années (dans notre exemple, comparaison du taux cible envisagé par la société et du taux des OAT à 10 ans pendant les mêmes périodes de référence).

4.3.2. Détermination de la revalorisation effective

Deux cas de figure peuvent se produire :

- le taux de rendement financier est suffisant pour servir le taux cible de revalorisation ;
- le taux de rendement financier ne suffit pas à servir le taux cible de revalorisation ; auquel cas, selon les plus-values latentes disponibles, il peut être envisagé de céder une partie des actifs pour réaliser du rendement financier.

Les deux situations sont détaillées *infra*.

4.3.2.1 taux_rdt-fi > taux_revalo_cible

Dans cette situation, le taux de revalorisation effectif doit être au moins égal au taux cible :

$$\text{taux revalo} \geq \text{taux revalo cible} .$$

Dans des situations de marché très favorables, il est possible que les dispositions réglementaires ou contractuelles concernant l'attribution de PB

conduisent à un taux de revalorisation plus favorable que le taux cible, auquel cas

$$\text{taux revalo} = \rho(\text{taux rdt fi}).$$

Au final,

$$\text{taux revalo} = \max(\text{taux revalo cible}, \rho(\text{taux rdt fi})).$$

4.3.2.2 **taux_rdt_fi < taux_revalo_cible**

Dans cette situation, il convient de déterminer si le portefeuille d'actifs recèle des plus-values latentes sur les actifs R332-20.

Schématiquement, trois situations peuvent se produire :

- on constate d'ores et déjà une provision pour risque d'exigibilité (PRE) sur les actifs,
- il existe des plus-values latentes dans un volume insuffisant pour servir le taux cible de revalorisation,
- il existe des plus-values latentes dans un volume suffisant pour servir le taux cible de revalorisation,

1^{er} cas : $PVL < 0$

Dans ce cas, la société ne cède pas d'actif, sert au moins le tmg et peut choisir de verser un taux qui lui est supérieur en réalisant un résultat négatif, soit :

$$\text{taux revalo} = \text{tmg} + \frac{R^-}{\text{encours}(n) - \text{sorties} + \text{entrées}/2},$$

où $R^- \geq 0$ représente la perte supplémentaire que consent à faire l'assureur pour augmenter le taux de revalorisation.

2^e cas : $PVL > 0$ et $PVL < (\text{taux_revalo_cible} - \text{taux_rdt_fi}) \times \text{encours}(n+1)$

La société sert au moins le tmg et peut réaliser une partie de ses plus-values latentes pour servir un taux supérieur au tmg mais ne pouvant égaler le taux cible, soit

$$\text{taux revalo} = \text{tmg} + \alpha \frac{PVL}{\text{encours}(n) - \text{sorties} + \text{entrées}/2} < \text{taux cible},$$

où $\alpha \in [0;1]$ représente la proportion des plus-values latentes sur les actifs R332-20 réalisée pour augmenter le taux de revalorisation.

3^e cas : $PVL > 0$ et $PVL > (\text{taux_revalo_cible} - \text{taux_rdt_fi}) \times \text{encours}(n+1)$

La société peut réaliser une partie de ses plus-values latentes pour servir un taux supérieur au tmg et pouvant aller jusqu'au taux de revalorisation cible, soit

$$\text{taux revalo} = \text{tmg} + \alpha \frac{PVL}{\text{encours}(n) - \text{sorties} + \text{entrées}/2},$$

où $\alpha \in [0;1]$ représente la proportion des plus-values latentes sur les actifs R332-20 réalisée pour augmenter le taux de revalorisation.

4.4. Rachat

Le rachat de contrats peut s'expliquer en partie par des arbitrages d'assurés qui se tournent vers des placements plus performants. En l'absence d'étude de ce phénomène sur les portefeuilles considérés, un modèle simple peut être proposé.

Il consiste à découper le rachat :

- en « rachat structurel » indépendant de l'évolution des marchés financiers et de la politique de revalorisation de l'épargne de la société d'assurance ;
- en « rachat conjoncturel » dépendant de l'évolution des marchés financiers et de la politique de revalorisation de l'épargne de la société d'assurance.

Considérons un contrat en particulier (ou de manière équivalente, un groupe de contrats homogène vis-à-vis du comportement des différents assurés quant au rachat), le taux de rachat moyen constaté dans le passé pour ce contrat est taux_rachat_moyen . Ce taux peut, par exemple, être extrait d'une loi de rachat fonction de l'ancienneté du contrat et de l'âge de l'assuré estimée sur les dernières années.

Considérons que dans des conditions « normales » de marché, la société serve le taux de revalorisation cible. Si le taux de rachat pour ce contrat a été estimé sur une période présentant de telles conditions, on peut supposer que ce taux de rachat moyen constaté par le passé correspond au rachat structurel.

Le taux de rachat conjoncturel peut être modélisé par une fonction croissante de l'écart de taux entre le taux de revalorisation cible de la période précédente et le taux de revalorisation effectif de la période précédente.

Par exemple, le modèle de rachat peut-être le suivant :

$$\text{taux rachat}(n) = \text{taux rachat struct} \times [1 + \alpha (\text{taux revalo cible}(n-1) - \text{taux revalo}(n-1))],$$

où $\alpha \geq 0$ représente la sensibilité du taux de rachat de la période n à l'écart entre le taux cible et le taux de revalorisation effectif sur la période précédente $(n-1)$.

Une alternative prenant en compte également les différentiels de taux des années précédentes (éventuellement avec une sous-pondération) pourrait également être envisagée :

$$\text{taux rachat}(n) = \text{taux rachat struct} \times \left[1 + \sum \alpha_i (\text{taux revalo cible}(n-i) - \text{taux revalo}(n-i)) \right].$$

5. Capital de solvabilité : les méthodes d'agrégation

La prise en compte de la dépendance des risques auxquels sont soumises les compagnies d'assurance est un thème central de Solvabilité 2. En effet, le principal volet quantitatif de cette directive concernera l'exigence de capitaux propres : les sociétés d'assurance devront disposer de ressources (fonds

propres) suffisantes pour ne pas être en ruine à horizon un an avec une très forte probabilité (99,5 %).

Ce critère est très ambitieux dans la mesure où il s'agit de contrôler le *risque global* supporté par la compagnie avec une *très forte probabilité*. Sur ce deuxième point, le Chapitre 4 revient sur les difficultés opérationnelles de mise en œuvre de ce genre de critère lorsque l'on est confronté à des valeurs extrêmes et le Chapitre 6 sur les techniques de simulation auxquelles ils font recours.

En pratique les assureurs vont pouvoir déterminer cette exigence de capital :

- par le biais du modèle interne qui, sur la base de la modélisation de tous les risques de l'assureur, permettra de simuler des besoins en capitaux à partir desquels on estimera la Value-at-Risk à 99,5 % ;
- grâce à l'utilisation d'une formule standard commune à tous les assureurs européens.

Le calibrage de cette formule standard est un des enjeux des deuxième et troisième études d'impact quantitatif menées par le CEIOPS (cf. l'annexe consacrée à QIS 3).

L'élaboration de cette formule standard n'est pas aisée car elle doit permettre d'approcher un quantile extrême tout en restant simple à mettre en œuvre et accessible aux sociétés les plus modestes.

C'est une approche modulaire qui a été retenue par le CEIOPS. Elle consiste, pour chaque risque identifié, à déterminer les besoins en capitaux nécessaires pour se prémunir de chacun de ces risques pris isolément selon le critère précédemment évoqué, puis en l'agrégation des capitaux ainsi obtenus de manière à refléter l'effet de diversification et les éventuelles dépendances entre ces risques.

5.1. Un cadre gaussien

Les méthodes d'agrégation retenues reposent sur des résultats sur les Value-at-Risk dans un cadre gaussien. En effet, si

$$\begin{pmatrix} X \\ Y \end{pmatrix} \approx \mathcal{N} \left(\begin{pmatrix} m_X \\ m_Y \end{pmatrix}, \begin{bmatrix} \sigma_X^2 & \rho\sigma_X\sigma_Y \\ \rho\sigma_X\sigma_Y & \sigma_Y^2 \end{bmatrix} \right),$$

alors $X + Y \approx \mathcal{N} \left(m_X + m_Y, \sigma_X^2 + \sigma_Y^2 + 2\rho\sigma_X\sigma_Y \right)$ et l'on en déduit que

$$VaR_\alpha(X + Y) = m_X + m_Y + \Phi^{-1}(1 - \alpha) \sqrt{\sigma_X^2 + \sigma_Y^2 + 2\rho\sigma_X\sigma_Y}.$$

Dans le cas de la fixation d'un capital de solvabilité à ajouter au montant *best estimate* des provisions, on a donc pour un risque particulier X :

$$SCR_\alpha(X) = VaR_\alpha(X) - m_X = \Phi^{-1}(1 - \alpha) \sigma_X.$$

Pour deux risques on en déduit la formule d'agrégation simple :

$$SCR_\alpha(X + Y) = \sqrt{SCR_\alpha^2(X) + SCR_\alpha^2(Y) + 2\rho SCR_\alpha(X) SCR_\alpha(Y)},$$

qui se généralise naturellement par :

$$SCR = \sqrt{\sum \rho_{ij} SCR_i \times SCR_j}.$$

Incidentement, la VaR est sous-additive dans ce cadre gaussien⁵ et donc cohérente au sens de Artzner et al. (1999). Ce résultat provient du fait que

$$\Phi^{-1}(1-\alpha) \sqrt{\sigma_X^2 + \sigma_Y^2 + 2\rho\sigma_X\sigma_Y} \leq \Phi^{-1}(1-\alpha) [\sigma_X + \sigma_Y],$$

pour tout $\rho \in [-1; 1]$. L'égalité étant obtenue lorsque $\rho = 1$.

5.2. Les limites des modèles retenus

L'approche retenue repose sur une approximation gaussienne, non pas des risques eux-mêmes, mais de leurs effets sur la solvabilité de l'assureur.

Le choix de cette méthode a été dicté par des considérations pratiques : il s'agit de mettre un modèle sous les contraintes opposées de robustesse et de simplicité. Or si l'approximation gaussienne facilite l'agrégation des capitaux, il n'en demeure pas moins qu'elle ne représente pas bien les phénomènes extrêmes que l'on cherche justement à mesurer : le choix du critère de solvabilité étant lui-même un critère extrême (VaR à 99,5 %). En particulier cette approximation gaussienne risque sous-estimer généralement les queue de distribution.

La modélisation retenue risque donc de conduire à une sous-estimation de l'exigence de capital. Aussi le modèle proposé par QIS 3 (cf. l'annexe de cette thèse) retient des coefficients de corrélation jamais négatifs et même relativement élevés dans les calculs d'agrégation de capitaux, ce qui a pour effet de rajouter, *a posteriori*, de la prudence dans le calcul de l'exigence de capitaux.

Incidentement des solutions telles que celle proposée par la méthode NORTA permettraient de se placer dans un cadre de dépendance non linéaire sans qu'il soit besoin de définir d'hypothèses supplémentaires par rapport au cas gaussien.

⁵. Pour autant que $\alpha < 0,5$, ce qui est le cas dans notre contexte puisque $\alpha = 0,005$.

Bibliographie

- Artzner Ph., Delbaen F., Eber J.M., Heath D. (1999) « Coherent measures of risk », *Mathematical Finance* 9, 203-28.
- Caperaa P., Fougères A.L., Genest C. (1997) « A nonparametric estimation procedure for bivariate extreme value copulas », *Biometrika* 84, 567-7.
- Darsow W., Nguyen B., Olsen E. (1992) « Copulas and Markov Processes », *Illinois Journal of Mathematics* 36, n°4, 600-42.
- Deheuvels P. (1978) « Caractérisation complète des lois extrêmes multivariées et de la convergence des types extrêmes », *Publications de l'Institut de Statistique de l'Université de Paris* 23, 1-36
- Deheuvels P. (1979a) « Propriétés d'existence et propriétés topologiques des fonctions de dépendance avec applications à la convergence des types pour des lois multivariées », *C. R. Académie des Sciences de Paris, Série 1*, 288, 145-148
- Deheuvels P. (1979b) « La fonction de dépendance empirique et ses propriétés - Un test non paramétrique d'indépendance », *Académie Royale de Belgique – Bulletin de la Classe des Sciences – 5^e Série*, 65, 274-92
- Denuit M., Charpentier A. (2004) *Mathématiques de l'assurance non-vie. Tome 1 : principes fondamentaux de théorie du risque*, Paris : Economica.
- Embrechts P., McNeil A., Straumann D. (1999) « Correlation: Pitfalls and alternatives », *RISK Magazine*, May, 69-71.
- Genest C., MacKay R. J. (1986) « The joy of copulas: Bivariate distributions with uniform marginals », *The American Statistician* 40, 280-3.
- Joe H. (1997) « Multivariate models and dependence concepts », *Monographs on Statistics and Applied probability* 73, London : Chapman & Hall.
- Louis T., Shi J. (1995) « Inferences on the association parameter in copula models for bivariate survival data », *Biometrics* 51, 1384-99.
- Meyers G.G., Klinker F.L., Lalonde F.A. (2003) « The aggregation and correlation of insurance exposure », *CAS Forum*, 16-82.
- Nelsen R.B. (1999) *An introduction to Copulas*, Lecture Notes in Statistics 139, New-York: Springer Verlag.
- Partrat Ch., Besson J.L. (2005) *Assurance non-vie. Modélisation, simulation*, Paris : Economica.
- Saporta G. (1990) *Probabilités, analyse des données et statistiques*, Paris : Technip.
- Schweizer B., Sklar A. (1958) « Espaces métriques aléatoires », *C. R. Acad. Sci. Paris* 247, 2092-4.

- Schweizer B. (1991) « Thirty years of copula, Advances in probability distributions with given marginals: beyond the copulas », ed. by G. Dall'Aglia, S. Kotz and G. Salinetti, *Mathematics and its applications* v. 67, Kluwer Academic Publishers, 13-50.
- Sklar A. (1959) « Fonctions de répartition à n dimensions et leurs marges », *Publications de l'Institut de Statistique de l'Université de Paris* 8, 229-31.

Chapitre 6

Techniques de simulation

Les processus stochastiques continus sont des outils largement employés en finance et en assurance, notamment pour modéliser taux d'intérêts et cours d'actions. De nombreuses problématiques amènent à la simulation des trajectoires de tels processus. C'est notamment le cas des modèles internes (cf. la Figure 20 au début du Chapitre 4) qui pourront être utilisés pour déroger au niveau de capital fourni par la formule standard dans Solvabilité 2 (cf. l'annexe) ou encore les modèles de projection de l'activité utilisés pour déterminer une MCEV. La mise en œuvre pratique de ces simulations nécessite de discrétiser ces processus, d'estimer les paramètres et de générer des nombres aléatoires en vue de générer les trajectoires voulues. Nous abordons successivement ces trois points que nous illustrons dans quelques situations simples.

Ce chapitre s'appuie en partie sur Planchet et Thérond (2005c).

1. Introduction

La fréquence des variations de cotation des actifs financiers sur les marchés organisés a conduit les financiers à considérer des processus stochastiques continus pour modéliser les évolutions de taux d'intérêt comme de cours d'actions. Si les outils mathématiques sous-jacents peuvent, de prime abord, sembler plus complexes que leurs équivalents discrets, leur usage a notamment permis d'obtenir des formules explicites pour l'évaluation d'actifs contingents (cf. les formules fermées de prix d'options européennes obtenues par Black et Scholes (1973) lorsque l'actif sous-jacent suit un mouvement brownien géométrique). Ces modèles sont aujourd'hui utilisés dans de nombreux domaines, notamment en assurance. Leur mise en œuvre pratique nécessite de les discrétiser, que ce soit pour l'estimation des paramètres ou pour la simulation des trajectoires.

En effet les données disponibles étant discrètes, l'estimation des paramètres des modèles en temps continu n'est pas immédiate. De nombreux processus, tels que le modèle de taux proposé par Cox, Ingersoll et Ross (1985), n'admettant pas de discrétisation exacte, l'estimation des paramètres nécessite souvent une approximation discrétisée. Les estimations directes à partir du modèle discrétisé s'avérant, en général, biaisées lorsque le processus ne dispose pas d'une version discrète exacte, on est alors conduit à utiliser des méthodes indirectes comme la méthode par inférence indirecte (cf. Giet (2003)) ou encore la méthode des moments efficaces pour estimer le modèle.

De plus, de nombreuses problématiques impliquent que l'on soit capable de simuler l'évolution de cours ou de taux modélisés par des processus continus ; la réalisation pratique de telles simulations nécessitera là encore la discrétisation de ces processus. C'est notamment le cas des problématiques de type DFA (*Dynamic Financial Analysis*, cf. Kaufmann, Gadmer et Klett (2001) et Hami (2003)) qui consistent en la modélisation de tous les facteurs ayant un impact sur les comptes d'une société d'assurance dans le but d'étudier sa solvabilité ou de déterminer une allocation d'actif optimale par exemple. Le nombre de variables à modéliser est alors très important : actifs financiers, sinistralité de chaque branche assurée, inflation, réassurance, concurrence et dépendances entre ces variables. Certaines de ces variables seront modélisées par des processus de diffusion ; lorsque l'équation différentielle stochastique (EDS) concernée dispose d'une solution explicite, comme c'est le cas pour le modèle de Vasicek (1977), la discrétisation s'impose à l'utilisateur. Lorsqu'une telle expression n'est pas disponible, l'utilisateur pourra se tourner notamment vers les schémas d'Euler ou de Milstein qui sont des développements d'Itô-Taylor à des ordres plus ou moins importants de l'EDS considérée. L'approximation sera d'autant plus satisfaisante, au sens du critère de convergence forte (cf. Kloeden et Platen (1995)), que l'ordre du développement est élevé.

Par ailleurs, la performance des simulations étant conditionnée par le générateur de nombres aléatoires utilisé, une attention particulière doit être portée à son choix. Nous comparons ici deux générateurs : le générateur pseudo-aléatoire *Rnd* d'Excel / Visual Basic et le générateur quasi-aléatoire obtenu à partir de l'algorithme de la translation irrationnelle du tore présenté par Planchet, Thérond et al. (2005). L'algorithme du tore s'avère nettement plus performant que *Rnd* mais souffre de la dépendance terme à terme des valeurs générées qui induit un biais dans la construction de trajectoires. Pour remédier à ce problème, nous proposons un générateur du *tore mélangé* qui conserve les bonnes propriétés de répartition de l'algorithme du tore et peut être utilisé pour construire des trajectoires de manière efficace.

L'objectif du présent chapitre est de proposer au praticien quelques guides méthodologiques lui permettant d'effectuer de manière efficace des simulations de processus continus, plus particulièrement dans le contexte des problématiques d'assurance. Nous abordons donc les trois étapes clé que sont l'estimation des paramètres, la discrétisation du processus et la génération de nombres aléatoires.

2. Discrétisation de processus continus

Les méthodes de Monte Carlo utilisées, en assurance, lors de la mise en œuvre de modèles du type DFA ou, en finance, lors de l'évaluation d'options font souvent usage de la discrétisation d'équations différentielles stochastiques (EDS). En effet la fréquence de variation du prix des actifs financiers et la commodité mathématique ont souvent conduit les économistes à les modéliser par des processus continus. Toutefois la simulation effective de ces processus requiert la discrétisation du temps et donc la détermination de la loi du processus aux instants de discrétisation. Si pour certains processus tels que le mouvement brownien géométrique, il est possible de déterminer la loi du

processus à n'importe quel instant (on parlera alors de discrétisation exacte), pour les autres, il va falloir les approcher par des processus discrets qui convergent vers les processus que l'on souhaite simuler (on parlera alors de discrétisation approximative).

Même si cette problématique n'est pas abordée ici, remarquons que la discrétisation temporelle est également nécessaire lors de l'estimation des paramètres des modèles ; cette dernière opération reposant évidemment sur des données discrètes. En particulier, Giet (2003) met en lumière l'incidence du choix de la discrétisation sur la qualité de l'estimation des paramètres.

L'objet de ce chapitre est de présenter les différents procédés de discrétisation usuellement employés ainsi que leur efficacité en terme de rapidité de convergence.

Prenons le cas d'un processus défini par l'EDS suivante :

$$\begin{cases} dX_t = \mu(X_t, t)dt + \sigma(X_t, t)dB_t \\ X_0 = x \end{cases}$$

où B est un mouvement brownien standard. Le calcul d'Itô permet de voir cette équation comme une formulation symbolique de :

$$X_t = x + \int_0^t \mu(X_s, s)ds + \int_0^t \sigma(X_s, s)dB_s .$$

Si le processus considéré ne dispose pas d'une discrétisation exacte, un développement d'Itô-Taylor de l'équation sous sa forme intégrale nous permet de disposer d'une version discrétisée approximative. Cette approximation est d'autant plus précise que le développement intervient à un ordre élevé.

Dans la suite, on notera δ le pas de discrétisation, c'est à dire le temps qui s'écoule entre deux instants où l'on va simuler le processus. T désignera l'horizon de projection. D'une manière générale, X sera le processus que l'on souhaite simuler et $(\tilde{X}_{k\delta})_{k \in [1; T/\delta]}$ le processus discret effectivement utilisé pour simuler des réalisations de X aux instants de discrétisations $k\delta$ où $k \in [1; T/\delta]$. Enfin ε désignera une variable aléatoire de loi normale centrée réduite.

2.1. Discrétisation exacte

La simulation d'un processus d'Itô pourra être effectuée directement (sans erreur de discrétisation) dès lors que celui-ci admet une discrétisation exacte. Nous allons voir ici ce qu'il faut entendre par le terme discrétisation exacte qui peut apparaître comme une oxymore et voir que dans le cas des modèles de Black et Scholes d'une part et de Vasicek d'autre part, la simulation des processus pourra passer par une discrétisation exacte.

2.1.1. Définition

Définition 35. *Discrétisation exacte d'un processus continu*

Un processus $(\tilde{X}_{k\delta})_{k \in [1; T/\delta]}$ est une discrétisation exacte du processus X si $\forall \delta > 0, \forall k \in [1; T/\delta] \tilde{X}_{k\delta} \stackrel{L}{\sim} X_{k\delta}$.

Un processus admet une discrétisation exacte dès lors que l'on peut résoudre explicitement l'EDS qui lui est associée.

C'est notamment le cas du mouvement brownien géométrique retenu par Black et Scholes pour modéliser le cours d'une action ou encore celui du processus retenu par Vasicek pour modéliser le taux d'intérêt instantané.

2.1.2. Processus d'Ornstein-Uhlenbeck

Dans le cas du modèle de Vasicek, le taux instantané r est solution de l'EDS

$$dr_t = a(b - r_t)dt + \sigma dB_t.$$

La solution s'écrit dans ce cas :

$$r_t = r_0 e^{-at} + b(1 - e^{-at}) + \sigma e^{-at} \int_0^t e^{as} dB_s.$$

En effet un processus d'Ornstein-Uhlenbeck X est l'unique solution de : $dX_t = -cX_t dt + \sigma dB_t$ avec $X_0 = x$. Si l'on pose $Y_t = X_t e^{ct}$, on a : $dY_t = e^{ct} dX_t + X_t d(e^{ct}) + d\langle X, e^c \rangle_t$. Or $d(e^{ct}) = ce^{ct} dt$ donc $\langle X, e^c \rangle_t = 0$.

Ainsi $dY_t = \sigma e^{ct} dB_t$ et donc $X_t = x e^{-ct} + \sigma e^{-ct} \int_0^t e^{cs} dB_s$. Enfin le résultat est obtenu par : $r_t = X_t + b$ et $a = -c$.

Les propriétés de l'intégrale d'une fonction déterministe par rapport à un mouvement brownien conduisent à la discrétisation exacte :

$$r_{t+\delta} = r_t e^{-a\delta} + b(1 - e^{-a\delta}) + \sigma \sqrt{\frac{1 - e^{-2a\delta}}{2a}} \varepsilon.$$

On montre que si $\sigma(x, t)$ est une fonction du temps, alors X admet une discrétisation exacte (cf. Kloeden et Platen (1995)).

En effet, notons $p \in \mathbb{N}$ et δ le pas de discrétisation retenu et cherchons à simuler la trajectoire d'un mouvement brownien B . Pour obtenir une réalisation de $B_{p\delta}$, il suffit de calculer $B_{p\delta}^\delta = \sqrt{\delta} \sum_{k=1}^p \varepsilon_k$ où les ε_k sont des v.a. mutuellement indépendantes de loi normale centrée réduite. La loi de l'approximation $B_{p\delta}^\delta$ est identique à celle de $B_{p\delta}$. Donc dès lors que $\sigma(x, t)$ est

une fonction qui ne dépend que du temps, $\int_{p\delta}^{(p+1)\delta} \sigma(x, t) dB_t$ reste une variable aléatoire de loi normale centrée et de variance $\int_{p\delta}^{(p+1)\delta} \sigma^2(x, t) dt$.

2.1.3. Mouvement brownien géométrique

Intéressons nous à présent au cas du mouvement brownien géométrique :

$$\frac{dS_t}{S_t} = \mu dt + \sigma dB_t.$$

Ce processus admet une discrétisation exacte :

$$S_t = S_0 \exp \left\{ \left(\mu - \frac{\sigma^2}{2} \right) t + \sigma (B_t - B_0) \right\}.$$

En choisissant un pas de discrétisation δ , on obtient le schéma récursif exact suivant :

$$S_{t+\delta} = S_t \exp \left\{ \left(\mu - \frac{\sigma^2}{2} \right) \delta + \sigma \sqrt{\delta} \varepsilon \right\}.$$

On remarque que, Dans le cas du mouvement brownien géométrique, $\sigma(x, t) = \sigma x$ et n'est donc pas uniquement une fonction du temps. Ceci illustre le fait qu'un processus peut être discrétisé de manière exacte sans que la fonction σ qui lui est associée soit nécessairement uniquement une fonction du temps.

2.2. Discrétisation approximative

Lorsque la discrétisation exacte n'existe pas, il convient de se tourner vers des approximations discrètes du processus continu sous-jacent. Les schémas d'Euler et de Milstein sont les procédés de discrétisation les plus répandus. Tous deux sont des développements d'Itô-Taylor de la forme intégrale de l'équation définissant la diffusion à des ordres différents.

2.2.1. Critère de convergence forte

Dans la suite, nous ferons référence au critère de convergence forte pour classer les procédés de discrétisation. Une discrétisation approximative $\tilde{X}$ converge fortement vers le processus continu X lorsque l'erreur commise sur la valeur finale (à la date T) de la trajectoire obtenue par le processus discrétisé est en moyenne négligeable.

$$\forall T > 0, \quad \lim_{\delta \rightarrow 0} \mathbb{E} \left[\left| \tilde{X}_T^\delta - X_T \right| \right] = 0.$$

La vitesse de convergence de cette équation nous permet d'introduire un ordre entre les procédés de discrétisation. Ainsi le processus discrétisé $\tilde{X}$ converge fortement¹ à l'ordre γ vers le processus X si :

$$\exists K > 0, \exists \delta_0 > 0, \forall \delta \in]0, \delta_0], \quad E \left[\left\| \tilde{X}_T^\delta - X_T \right\| \right] \leq K \delta^\gamma.$$

2.2.2. Développement d'Itô-Taylor

Définissons pour toute fonction g de classe $C^{2,1}$, les opérateurs L_0 et L_1 par

$$\begin{cases} L_0 g(X_t, t) = g_t + g_x \mu + \frac{\sigma^2}{2} g_{yy} \\ L_1 g(X_t, t) = \sigma g_x \end{cases}$$

où les fonctions g , μ , σ et leurs dérivées sont évaluées en (X_t, t) .

L'opérateur L_0 est également appelé Dynkin du processus X .

Rappelons à présent le lemme d'Itô qui, si X est un processus d'Itô, permet de différencier le processus $g(X)$ sous réserve que g soit suffisamment régulière.

Théorème 7. Lemme d'Itô

Si X est un processus d'Itô qui vérifie $dX_t = \mu(X_t, t)dt + \sigma(X_t, t)dB_t$, si g est de classe $C^{2,1}$ et si $Z(t) = g(X_t, t)$, on a :

$$dZ_t = L_0 g(X_t, t)dt + L_1 g(X_t, t)dB_t,$$

ou de manière équivalente :

$$Z_t = g(X_t, t) = g(X_0, 0) + \int_0^t L_0 g(X_s, s)ds + \int_0^t L_1 g(X_s, s)dB_s.$$

Rappelons le principe du développement d'Itô-Taylor. Ce développement repose sur la combinaison du lemme d'Itô et de l'écriture sous forme intégrale d'un processus.

En effet, en considérant que

$$X_t = X_0 + \int_0^t dX_s,$$

il vient

$$X_t = X_0 + \int_0^t (\mu(X_s, s)ds + \sigma(X_s, s)dB_s).$$

En appliquant le lemme d'Itô aux fonctions μ et σ , on obtient :

¹. On notera que ce critère ne fait pas référence à la convergence uniforme.

$$X_t = X_0 + \int_0^t \left\{ \mu(X_0, 0) + \int_0^s d\mu(X_s, s) \right\} dt + \int_0^t \left\{ \sigma(X_0, 0) + \int_0^s d\sigma(X_s, s) \right\} dB_t .$$

Les techniques de discrétisation approximative présentées *infra* reposent sur des développements d'Itô-Taylor.

2.2.3. Schéma d'Euler

Le procédé de discrétisation d'Euler consiste en l'approximation du processus continu X par le processus discret $\tilde{X}$ défini, avec les mêmes notations que précédemment, par :

$$\tilde{X}_{t+\delta} = \tilde{X}_t + \mu(\tilde{X}_t, t)\delta + \sigma(\tilde{X}_t, t)\sqrt{\delta}\varepsilon .$$

Ce procédé est obtenu, en négligeant le reste dans le développement d'Itô-Taylor au premier ordre :

$$X_t = X_0 + \mu(X_0, 0) \int_0^t ds + \sigma(X_0, 0) \int_0^t dB_s + R_1(0, t) ,$$

où le reste $R_1(0, t)$ est donné par :

$$\begin{aligned} R_1(0, t) &= \int_0^t \int_0^s L_0 \mu(X(u), u) du ds + \int_0^t \int_0^s L_1 \mu(X(u), u) dB_u ds \\ &\quad + \int_0^t \int_0^s L_0 \sigma(X(u), u) du dB_s + \int_0^t \int_0^s L_1 \sigma(X(u), u) dB_u dB_s \end{aligned}$$

Le schéma d'Euler revient donc à faire l'approximation naturelle suivante :

$$\int_0^t \mu(X_s, s) ds + \int_0^t \sigma(X_s, s) dB_s \approx \mu(X_0, 0)t + \sigma(X_0, 0)(B_t - B_0) .$$

Kloeden et Platen (1999) prouvent que sous certaines conditions de régularité², le schéma d'Euler présente un ordre de convergence forte de 0,5.

Par exemple dans le modèle de Cox, Ingersoll et Ross (CIR), le taux d'intérêt instantané est solution de l'EDS :

$$dr_t = a(b - r_t)dt + \sigma\sqrt{r_t} dB_t .$$

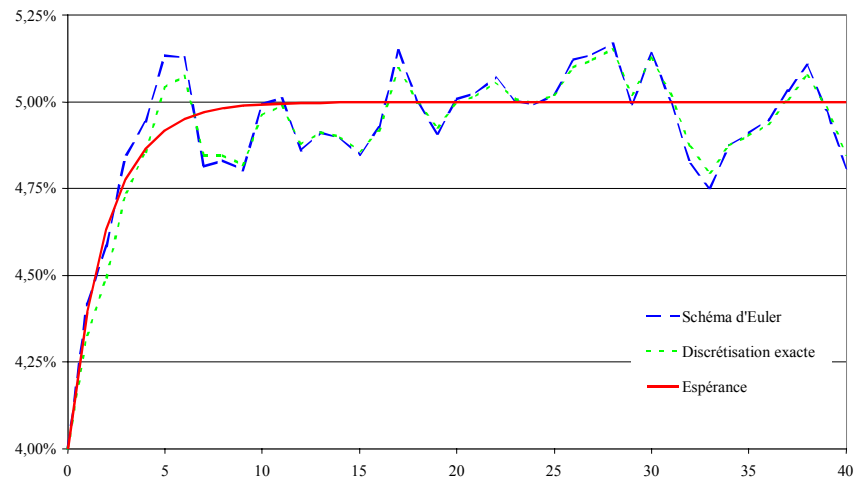
Aussi le processus discret $\tilde{r}$ déterminé par le schéma d'Euler peut s'écrire :

$$\tilde{r}_{t+\delta} = \tilde{r}_t + a(b - \tilde{r}_t)\delta + \sigma\sqrt{\tilde{r}_t} \times \delta \varepsilon .$$

². Il faut pour cela que μ et σ soient des fonctions C^4 avec des dérivées bornées jusqu'à l'ordre 4.

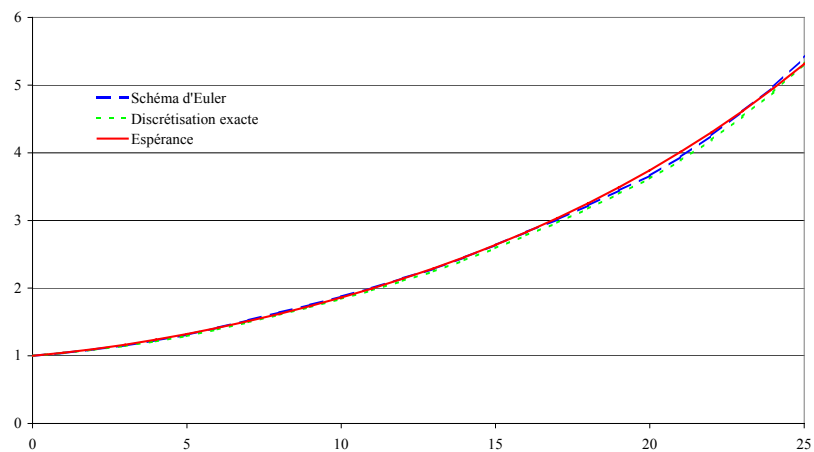
Le graphique suivant permet de comparer les évolutions moyennes de la diffusion du modèle de Vasicek selon que l'on se tourne vers la discrétisation exacte ou vers le schéma d'Euler.

Figure 31 - Évolution moyenne du taux modélisé par Vasicek selon le procédé de discrétisation retenu



Les discrétisations exacte et selon le schéma d'Euler sont relativement proches graphiquement. L'écart est en particulier négligeable lorsque l'on s'intéresse aux évolutions moyennes d'un bon de capitalisation qui évolue selon les taux simulés pour effectuer le graphique précédent.

Figure 32 - Modèle de Vasicek selon le procédé de discrétisation retenu



2.2.4. Schéma de Milstein

Le schéma de Milstein est obtenu en allant plus avant dans le développement d'Itô-Taylor. Le processus discret $\tilde{X}$ est alors défini par :

$$\tilde{X}_{t+\delta} = \tilde{X}_t + \mu(\tilde{X}_t, t)\delta + \sigma(\tilde{X}_t, t)\sqrt{\delta}\varepsilon + \frac{\sigma_x(\tilde{X}_t, t)\sigma(\tilde{X}_t, t)}{2}\delta(\varepsilon^2 - 1),$$

où $\sigma_x(\tilde{X}_t, t)$ désigne la dérivée par rapport au premier argument de la fonction $\sigma(.,.)$ évaluée en $(\tilde{X}_t, t)$.

Ce schéma est obtenu à partir d'un des développements d'Itô-Taylor au deuxième ordre :

$$\begin{aligned} X_t = & X_0 + \mu(X_0, 0) \int_0^t ds + \sigma(X_0, 0) \int_0^t dB_s \\ & + L_1 \sigma(0, X_0) \int_0^t \int_0^s dB_u dB_s + R_2(0, t) \end{aligned}$$

où le reste R_2 , qui sera négligé pour obtenir le schéma de Milstein, est donné par :

$$\begin{aligned} R_2(0, t) = & \int_0^t \int_0^s L_0 \mu(X(u), u) du ds + \int_0^t \int_0^s L_1 \mu(X(u), u) dB_u ds \\ & + \int_0^t \int_0^s L_0 \sigma(X(u), u) du dB_s \\ & + \int_0^t \int_0^s \int_0^u L_0 L_1 \sigma(X(v), v) dv dB_u dB_s \\ & + \int_0^t \int_0^s \int_0^u L_1 L_1 \sigma(X(v), v) dB_v dB_u dB_s \end{aligned}$$

Remarquons qu'il existe plusieurs développements d'Itô-Taylor à l'ordre 2 puisque plusieurs termes sont développés. En effet, dans le développement présenté ci-dessus, on a développé la quatrième intégrale double du reste R_1 , on aura pu choisir d'en développer d'autres et l'on aurait abouti à un autre développement du deuxième ordre. Toutefois le développement effectué a la « bonne idée » de nous faire aboutir à une expression qui ne demande la simulation que d'une unique variable aléatoire ε (que l'on utilisera deux fois). Cette expression permet donc, sans simuler davantage de réalisations de variables aléatoires, d'être plus précis que le schéma d'Euler.

Ce procédé de discrétisation présente, en général, un ordre de convergence forte de 1. Néanmoins la vitesse de convergence est identique à celle du schéma d'Euler pour des fonctions régulières (les plus fréquentes en pratique) ce qui

conduit à lui préférer alors le schéma d'Euler qui demande moins de temps de calcul (on simule moins de termes).

Par ailleurs, remarquons que si la volatilité $\sigma(\cdot)$ ne dépend que du temps, les procédés de discrétisation d'Euler et de Milstein conduiront à la même discrétisation, c'est le cas pour les modèles de Vasicek et de Hull et White. Alors on se tournera de préférence vers une discrétisation exacte. En revanche pour le modèle de CIR, il vient selon le schéma de Milstein :

$$\tilde{r}_{t+\delta} = \tilde{r}_t + a(b - \tilde{r}_t)\delta + \sigma\sqrt{\tilde{r}_t} \times \delta \varepsilon + \frac{\sigma^2}{4} \delta (\varepsilon^2 - 1).$$

Dans le cas du modèle CIR, ce schéma de discrétisation est intéressant dans la mesure où l'on obtient une meilleure approximation sans avoir à simuler davantage de réalisations de v.a.

En développant à des degrés supérieurs l'équation intégrale, il est possible d'obtenir des processus discrétisés d'ordre de convergence encore plus élevé. Toutefois ils nécessiteront des calculs plus nombreux et peuvent faire intervenir plus d'une variable aléatoire ce qui signifie des temps de simulation plus importants. De plus un résultat de Clark et Cameron (cf. Temam (2004)) prouve que, vis à vis de la norme L^2 , le schéma d'Euler est optimal dans la classe des schémas n'utilisant que les variables $B_{k\delta}$.

En outre, le schéma de Milstein peut poser des problèmes pratiques de mise en œuvre en dimension supérieure ou égale à 2.

3. Estimation des paramètres

L'estimation des paramètres est une étape délicate dans la simulation de trajectoires d'un processus continu car elle peut être l'origine d'un biais. En effet, le praticien peut se voir confronté à deux problèmes :

- le processus n'admet pas forcément de discrétisation exacte,
- la variable modélisée n'est pas toujours directement observable.

En effet si le processus considéré n'admet pas de discrétisation exacte, il sera impossible d'estimer les paramètres du modèle par la méthode du maximum de vraisemblance. Par ailleurs il est souvent impossible d'exprimer de manière exacte les densités transitoires entre deux observations. Il faudra alors se tourner vers des méthodes simulées telles que l'inférence indirecte pour estimer les paramètres.

3.1. Le biais associé à la procédure d'estimation

On reprend ici le cas simple de la modélisation du taux à court terme par un processus d'Ornstein-Uhlenbeck. On a vu précédemment que la discrétisation exacte d'un tel processus conduit à un processus auto-régressif d'ordre 1 (cf. le paragraphe 2.1.2).

L'estimation des paramètres du modèle s'effectue classiquement en régressant la série des taux courts sur la série décalée d'une période, que l'on écrit sous la forme usuelle :

$$Y = \alpha + \beta X + \sigma_1 \varepsilon,$$

avec

$$a = -\ln \beta, \quad b = \frac{\alpha}{1-\beta} \quad \text{et} \quad \sigma^2 = \sigma_1^2 \frac{2 \ln \beta}{\beta^2 - 1}.$$

L'estimateur du maximum de vraisemblance (EMV) coïncide avec l'estimateur des moindres carrés ; on obtient ainsi les estimateurs suivants des paramètres du modèle d'origine :

$$\hat{a}_{exact} = -\ln \hat{\beta}, \quad \hat{b}_{exact} = \frac{\hat{\alpha}}{1-\hat{\beta}} \quad \text{et} \quad \hat{\sigma}_{exact}^2 = \hat{\sigma}^2 \frac{2 \ln \hat{\beta}}{\hat{\beta}^2 - 1}.$$

On constate en particulier que ces estimateurs, s'ils sont EMV, sont biaisés. Ceci peut être gênant, notamment dans le cas où les paramètres du modèle ont une interprétation dans le modèle, comme par exemple b dans le cas présent, valeur limite du taux court. On voit sur cet exemple qu'on sera donc conduit à mettre en œuvre des procédés de réduction du biais.

3.2. Principe de l'estimation par inférence indirecte

Cette méthode est utilisée lorsque le processus n'admet pas de discrétisation exacte ou que sa vraisemblance, trop complexe, ne permet pas d'implémenter la méthode du maximum de vraisemblance. Elle consiste à choisir le paramètre $\hat{\theta}$ qui minimise (dans une métrique à définir) la distance entre l'estimation d'un modèle auxiliaire sur les observations et l'estimation de ce même modèle auxiliaire sur les données simulées à partir du modèle de base pour $\theta = \hat{\theta}$. Le modèle auxiliaire étant une discrétisation approximative de l'équation différentielle stochastique de base, le schéma d'Euler est souvent utilisé pour servir de modèle auxiliaire.

Cette méthode conduit à des estimations bien plus précises que celles obtenues à partir d'estimations « naïves » des paramètres du modèle auxiliaire sur les données observées.

Giet (2003) étudie l'impact du choix du procédé de discrétisation utilisé pour fournir le modèle auxiliaire sur l'estimation par inférence indirecte et montre qu'utiliser un procédé d'ordre de convergence plus élevé (comme la schéma de Milstein par rapport au schéma d'Euler) permet de réduire considérablement le biais sur l'estimation des paramètres.

Toutefois, que ce soit dans le cas d'une estimation par maximum de vraisemblance ou par inférence indirecte, on doit observer que dans le cas d'un modèle de taux tel que celui de Vasicek, l'estimation des paramètres de la dynamique d'évolution du taux court dans l'univers historique ne suffit pas et doit être complétée par l'estimation de la prime de risque, qui constitue un point délicat (cf. Lamberton et Lapeyre (1997) sur ce sujet). En pratique, on est amené à exploiter les informations contenues dans les prix observés (taux à terme, swaps, swaptions) qui reflètent les primes de risque et qui permettent ainsi de disposer directement du modèle dans l'univers risque-neutre.

3.3. Les méthodes d'estimation *ad hoc*

En pratique, dans les problèmes assurantiels, la grandeur modélisée par un processus continu n'est, le plus souvent, pas la grandeur d'intérêt : typiquement, on modélise le taux court, par exemple par un processus de diffusion, mais dans le cadre d'une problématique d'allocation d'actif pour un régime de rentiers, il sera déterminant que le modèle représente correctement les prix des obligations d'échéances longues. De manière équivalente, on est amené à estimer les paramètres du modèle pour représenter correctement la courbes des taux zéro-coupon, qui est dans une relation bijective avec la courbe du prix des obligations zéro-coupon :

$$P(0, T) = \exp\{-TR(0, T)\}.$$

L'idée est alors d'estimer les paramètres pour minimiser une distance (par exemple la distance quadratique) entre les prix prédits par le modèle et les prix observés sur le marché. Cette méthode est notamment utilisée dans Fargeon et Nissan (2003).

De plus, en pratique on pourra s'interroger sur les paramètres que l'on estime et les paramètres que l'on fixe arbitrairement compte tenu de la connaissance que l'on peut avoir du contexte par ailleurs. Ainsi, dans l'exemple évoqué, l'estimation simultanée des paramètres a , b et σ conduit à une courbe des taux quasi déterministe (σ est petit), ce qui peut apparaître irréaliste. On est ainsi conduit à fixer arbitrairement σ , puis à estimer les deux paramètres restants.

Dans cette approche il convient également d'être attentif au choix effectué quand à la paramétrisation du modèle : courbe des taux zéro-coupon ou courbe des prix des zéro-coupon ; en effet, la fonction de correspondance entre ces deux grandeurs accentue les écarts de courbure et le choix de l'une ou l'autre des paramétrisations peut conduire à des résultats sensiblement différents. Ce point est développé sur le plan théorique par Roncalli (1998) et est illustré sur un exemple dans la section suivante.

Au surplus, on notera que l'exploitation directe de données telles que le prix des zéro-coupon (ou les taux zéro-coupon) évite l'estimation de la prime de risque ; celle-ci est en effet incluse dans ces prix de marché et conduit ainsi à travailler naturellement dans l'univers corrigé du risque. L'estimation directe des paramètres sur des historiques de taux court nécessite, on l'a vu, l'estimation (délicate) de la prime de risque.

Dans les illustrations présentées ci-après, l'exploitation de données incluant les primes de risque est privilégiée.

3.4. Illustration dans le cas du modèle de Vasicek

A partir de la courbe des taux publiée par l'Institut des Actuariers, nous avons estimé les paramètres du modèle de Vasicek selon l'approche *ad hoc* décrite *supra*. L'estimation est menée d'une part directement à partir des valeurs des taux zéro-coupon de la courbe et d'autre part à partir des prix des obligations qui s'en déduisent. Cette courbe des taux est construite à partir des prix de marché et elle inclut les primes de risque.

Estimation *ad hoc* sur le prix des zéro-coupons : cette procédure consiste à déterminer par une méthode de type « moindres carrés » les paramètres du modèle de Vasicek qui permettent de représenter au mieux les prix des zéro-coupons. Cette méthode a été mise en œuvre en estimant les trois paramètres a , $\tilde{b}$ et σ puis en fixant σ et en n'estimant plus que les deux autres paramètres.

Comme on peut le voir dans le Tableau 8, l'estimation des paramètres du modèle de Vasicek donne des résultats très différents selon l'approche retenue :

Tableau 8 - Paramètres du modèle de Vasicek estimés selon les techniques d'estimation des paramètres

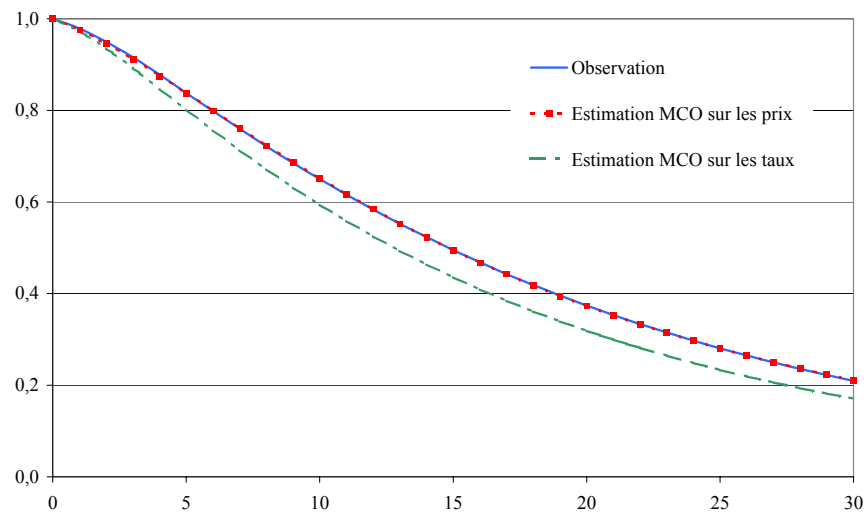
Paramètres estimés	EMC sur le prix des zéro-coupon			EMC sur les taux zéro-coupon		
	a, b et σ	a et b	a et b	a, b et σ	a et b	a et b
a =	0,2293	0,2353	0,5032	0,3951	0,3253	0,4232
b =	0,0572	0,0760	0,0701	0,0624	0,0745	0,0869
r_0 =	0,0207	0,0207	0,0207	0,0207	0,0207	0,0207
σ =	0,0000	0,0500	0,1000	0,0000	0,0500	0,1000

Le prix $P(t)$ en 0 d'un zéro-coupon d'échéance t est donné par la formule classique :

$$P(t) = \exp \left\{ -t * \left[\tilde{b} - \frac{\sigma^2}{2a} - \frac{1}{at} \left(\left(\tilde{b} - \frac{\sigma^2}{2a} - r_0 \right) (1 - e^{-at}) - \frac{\sigma^2}{4a} (1 - e^{-at})^2 \right) \right] \right\}.$$

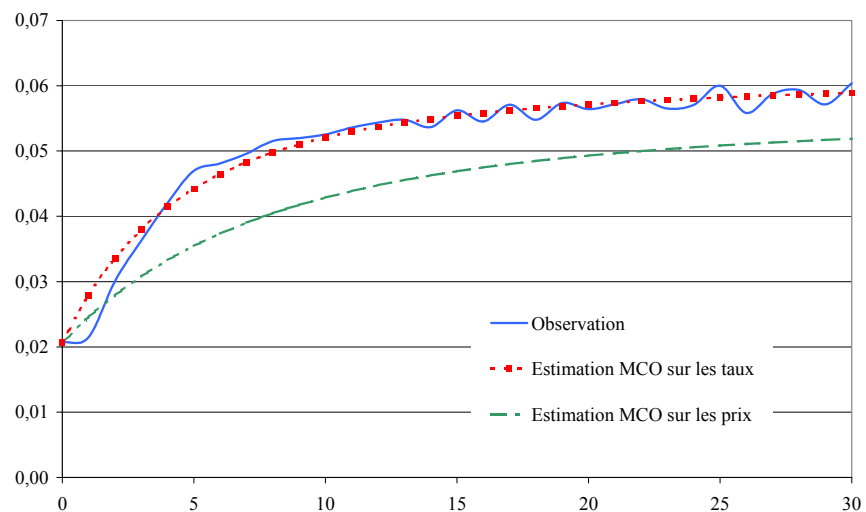
La Figure 33, qui indique le prix des zéro-coupons, en fonction de leur échéance, selon les différentes méthodes d'estimation des paramètres, nous permet d'observer que l'estimation des paramètres par le critère des moindres carrés sur les taux zéro-coupon s'avère moins performant que la méthode *ad hoc* si la variable d'intérêt est le prix des obligations zéro-coupon.

Figure 33 - Prix des zéro-coupons en fonction de leur échéance et de la méthode d'estimation des paramètres



En revanche, cette technique d'estimation des paramètres permet de générer des taux zéro-coupon nettement plus proches en espérance de la courbe des taux originelle que ceux obtenus par la méthode *ad hoc* comme on peut l'observer sur la Figure 34.

Figure 34 - Taux instantanés espérés en fonction de leur échéance et de la technique d'estimation des paramètres



Si les taux obtenus par la méthode *ad hoc* sur le prix des zéro-coupons sont, en espérance, éloignés des taux observés, ils permettent néanmoins d'approcher le prix des zéro-coupons avec une bonne précision.

Ainsi, en pratique, on sera conduit à privilégier pour une problématique d'allocation d'actifs d'un régime de rentes l'approche *ad hoc* sur le prix des zéro-coupons.

4. Génération des trajectoires

La génération des trajectoires passe nécessairement par la génération de nombres aléatoires. De manière pratique, il s'agit de générer des réalisations de variable aléatoire de loi uniforme sur le segment $[0, 1]$. En effet, si u est une telle réalisation, $F^{-1}(u)$ peut s'apparenter à une réalisation d'une variable aléatoire de fonction de répartition F . La technique d'inversion de la fonction de répartition permet ainsi à partir de réalisations de variables uniformes, d'obtenir de simuler des réalisations d'autres variables aléatoires. Lorsqu'on ne dispose pas de formule explicite pour F^{-1} , on utilisera des algorithmes d'approximation de cette fonction ou des algorithmes spécifiques à la loi que l'on souhaite traiter. Par exemple, pour simuler une réalisation d'une loi de Poisson de paramètre λ , on générera une suite (V_i) de réalisations de variable aléatoire exponentielle de paramètre 1. En effet $N = \inf \left\{ n \mid \sum_{i=1}^{n+1} V_i > \lambda \right\}$ suit une loi de Poisson de paramètre λ .

Les modélisations retenues en finance et en assurance faisant souvent intervenir des mouvements browniens, il est nécessaire de simuler des réalisations de variables aléatoires $\mathcal{N}(0;1)$. On ne dispose pas de formule exacte de l'inverse de la fonction de répartition inverse de la loi normale centrée réduite, mais l'algorithme de Box-Muller permet à partir de deux variables uniformes indépendantes sur $[0,1]$ de générer deux variables indépendantes $\mathcal{N}(0;1)$. En effet, si U_1 et U_2 sont des v.a. $\mathcal{U}(0;1)$ indépendantes alors en posant :

$$\begin{cases} X_1 = \sqrt{-2 \ln U_1} \cos(2\pi U_2) \\ X_2 = \sqrt{-2 \ln U_1} \sin(2\pi U_2) \end{cases}$$

Les variables aléatoires X_1 et X_2 sont indépendantes et suivent une loi gaussienne centrée réduite. Cette technique (exacte) est toutefois relativement longue à mettre en œuvre et nécessite l'indépendance des réalisations uniformes générées. Augros et Moréno (2002) présentent divers algorithmes pour approximer l'inverse de la fonction de répartition de la loi normale centrée réduite. Dans la suite, nous utiliserons l'algorithme de Moro (cf. Planchet, Thérond et al. (2005)) qui allie rapidité et précision.

4.1. Deux générateurs de nombres aléatoires : Rnd et le tore

4.1.1. Le générateur implémenté dans Excel / Visual Basic

Le générateur implémenté dans Excel (*Rnd*) est un générateur congruentiel, c'est à dire un générateur périodique issu d'une valeur initiale (on parle également de « graine » du générateur). Changer de valeur initiale permet de

changer de suite de nombres. Le lecteur se réfèrera à Planchet, Thérond et al. (2005) pour plus d'informations sur l'implémentation informatique de tels générateurs.

4.1.2. La translation irrationnelle du tore

Ce générateur multidimensionnel donne à la n -ème réalisation de la d -ème variable aléatoire uniforme à simuler la valeur u_n :

$$u_n = n\sqrt{p_d} - \left[n\sqrt{p_d} \right],$$

où p_d est le d -ème nombre premier et $[.]$ désigne l'opérateur partie entière.

4.1.3. Quelques éléments de comparaison

Les tests statistiques élémentaires (test d'adéquation du χ^2 , de Kolmogorov-Smirnov et d'Anderson-Darling) rappelés en annexe permettent d'apprécier la supériorité de répartition de la suite générée par le tore par rapport à celle générée par *Rnd*.

Figure 35 - Répartition et tests statistiques du générateur Rnd

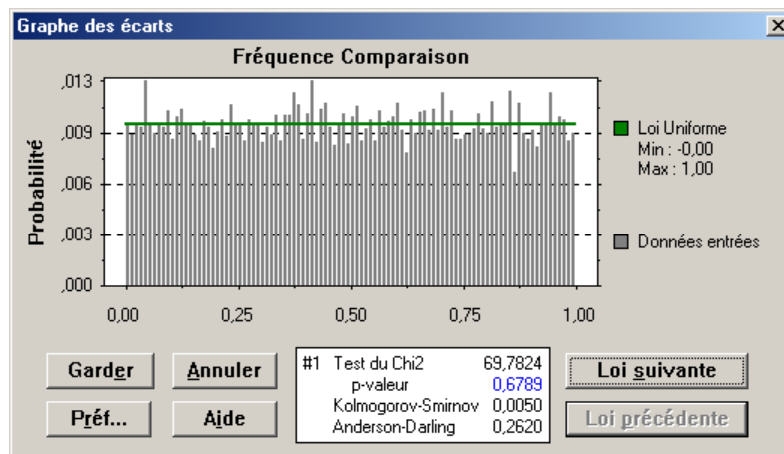
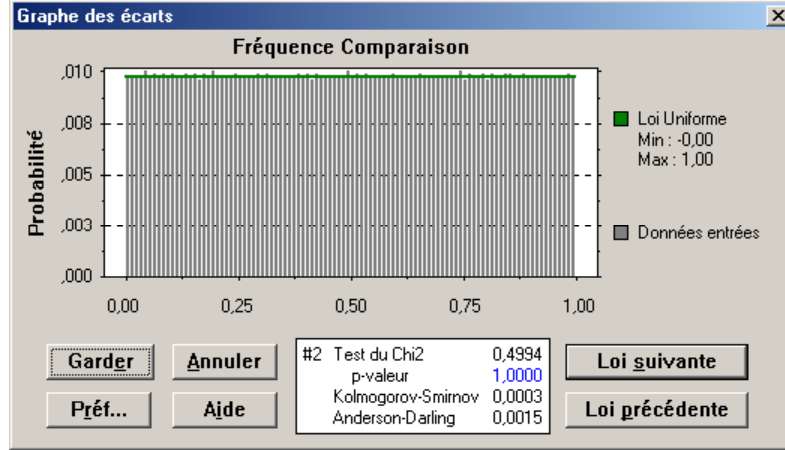


Figure 36 - Répartition et tests statistiques du générateur du tore



Par ailleurs, pour comparer les performances des deux générateurs, nous avons étudié la précision de l'estimation empirique du prix d'un call européen sur une action modélisée par un mouvement brownien géométrique sous la probabilité risque-neutre :

$$dS_t = rS_t dt + \sigma S_t d\tilde{B}_t.$$

où r est le taux sans risque et $\tilde{B}$ un mouvement brownien sous la mesure risque-neutre. Rappelons que, sous les hypothèses du modèle de Black et Scholes, la probabilité risque-neutre est l'unique probabilité sous laquelle les prix actualisés sont des martingales. La densité de Radon-Nikodym de cette probabilité par rapport à la probabilité historique est donnée par le théorème de Girsanov (cf. Augros et Moréno (2002)).

Ce processus dispose d'une discrétisation exacte :

$$S_{t+\delta} = S_t \exp \left\{ \left(r - \frac{\sigma^2}{2} \right) \delta + \sigma \sqrt{\delta} \varepsilon \right\},$$

où ε est une variable aléatoire de loi $\mathcal{N}(0;1)$.

Par ailleurs, nous disposons d'une formule fermée pour le prix en t d'un Call européen de prix d'exercice K et d'échéance T sur ce titre :

$$C_t(S, K, T) = S_t \Phi(d_1) - K e^{-r(T-t)} \Phi(d_2),$$

où :

$$\begin{aligned} - \quad d_1 &= \frac{\ln \frac{S_t}{K} + \left(r + \frac{\sigma^2}{2} \right) (T-t)}{\sigma \sqrt{T-t}}, \\ - \quad d_2 &= d_1 - \sigma \sqrt{T-t}, \end{aligned}$$

– Φ désigne la fonction de répartition de la loi normale centrée réduite.

Nous pouvons donc mesurer la performance des générateurs par le biais de l'erreur d'estimation du prix de l'option considérée. En effet, si s_T^i est le prix du titre considéré à la date T dans la i -ème simulation, l'estimateur naturel du prix de l'option en t est :

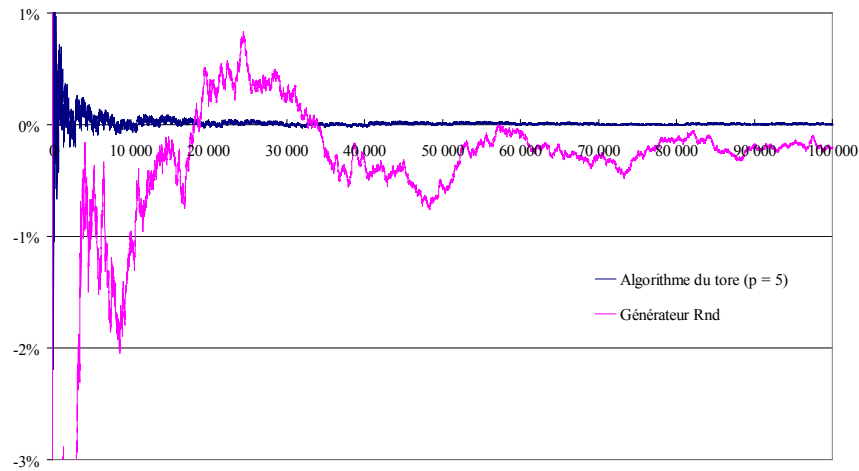
$$\hat{C}_t(S, K, T, n) = \frac{e^{-r(T-t)}}{n} \sum_{i=1}^n [s_T^i - K]^+.$$

L'erreur relative d'estimation peut donc s'écrire :

$$\rho = \frac{\hat{C}_t(S, K, T, n) - C_t(S, K, T)}{C_t(S, K, T)}.$$

Le graphique suivant présente l'évolution de ρ selon le nombre de simulations effectuées pour une action de volatilité 20 % et un Call d'échéance 6 mois.

Figure 37 - Erreur d'estimation en fonction du nombre de simulations



L'estimation à partir des valeurs générées par *Rnd* est biaisée de manière systématique de près de 0,2 % alors celle effectuée à l'aide du tore converge rapidement vers la valeur théorique : à partir de 7 000 simulations, l'erreur relative est inférieure à 0,1 %.

Il ressort de ces différentes comparaisons que le générateur du tore semble plus performant que *Rnd*. Son utilisation connaît toutefois un certain nombre de limites.

4.2. Limites de l'algorithme du tore

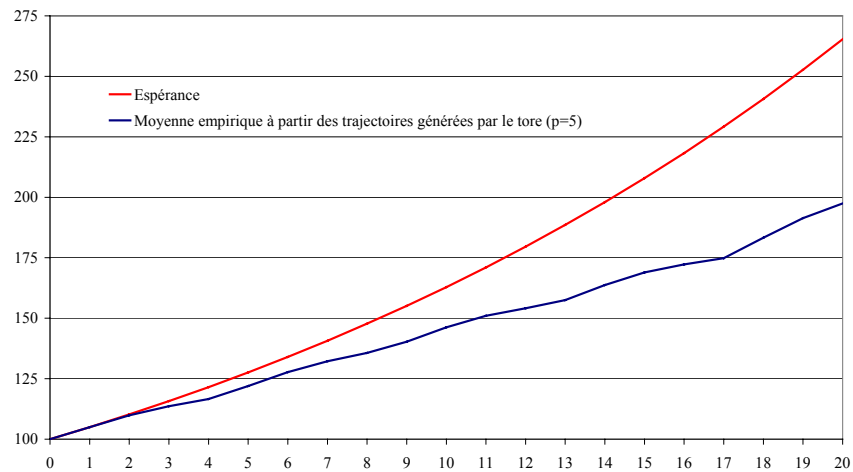
Les valeurs générées par le tore ne sont pas indépendantes terme à terme, ceci peut générer des erreurs non négligeables. Par exemple nous avons généré 10 000 trajectoires, avec un pas de discrétisation de 1, sur 20 ans d'un

mouvement brownien géométrique de volatilité 20 % et avons comparé l'évolution moyenne du cours estimée à partir des simulations et l'évolution moyenne théorique. En effet, si s_t^i est le cours du titre à la date t dans la i -ème simulation, l'estimateur empirique du cours moyen à la date t , $\bar{s}_t$ est donné par :

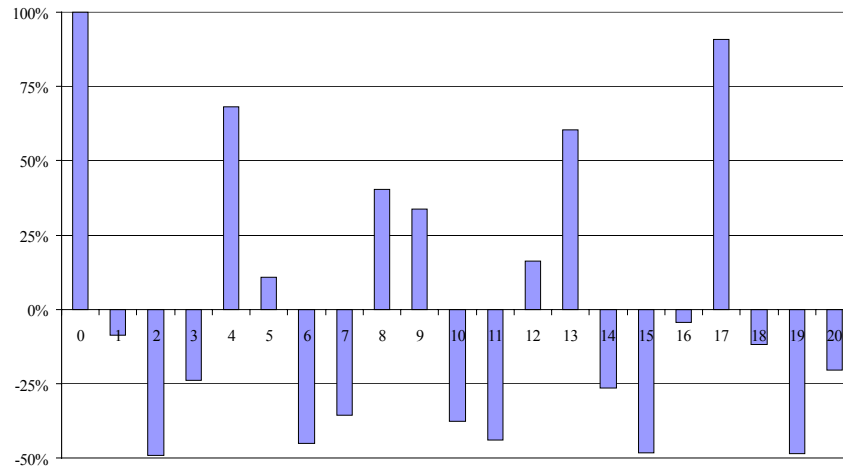
$$\bar{s}_t = \frac{1}{n} \sum_{i=1}^n s_t^i.$$

Le graphique suivant présente les deux trajectoires obtenues.

Figure 38 - Trajectoire moyenne simulée par le tore d'un mouvement brownien géométrique



A partir de $t = 2$, les trajectoires générées à partir de l'algorithme du tore sont en moyenne très en dessous de la trajectoire moyenne. Cet algorithme n'est donc pas utilisable en l'état pour générer des trajectoires. Ce biais s'explique par la dépendance terme à terme des valeurs générées par le tore que nous permet d'observer le corrélogramme de la suite générée par ce générateur.

Figure 39 - Corrélogramme de la suite générée par le tore ($p = 5$)

Rappelons que le h -ème terme du corrélogramme ρ_h s'écrit :

$$\rho_h = \frac{\sum_{k=1}^n (u_k - \bar{u})(u_{k+h} - \bar{u})}{\sum_{k=1}^n (u_k - \bar{u})^2},$$

où $\bar{u}$ désigne la moyenne empirique de la suite u .

Le « test du poker » permet également de mettre en évidence cette faille de l'algorithme du tore. L'idée de ce test est de comparer les fréquences théoriques des mains au poker avec les fréquences observées sur les simulations effectuées. De manière pratique, ce test consiste à prendre des listes de quatre chiffres tirés aléatoirement, de manière à observer les « combinaisons de valeurs » puis à effectuer un test d'adéquation du χ^2 pour voir si les fréquences observées correspondent aux fréquences théoriques. On distingue ainsi cinq cas :

- les quatre chiffres sont tous différents,
- la liste contient une et une seule paire,
- la liste contient deux paires,
- la liste contient un brelan (trois chiffres identiques),
- la liste contient un carré (quatre chiffres identiques).

En constituant les listes en prenant les chiffres dans l'ordre où ils sont simulés, le tableau 2 résume les fréquences suivantes.

Tableau 9 - Test du poker sur 4 000 réalisations générées par l'algorithme du tore

	Fréquences	Fréquences observées pour l'algorithme du tore									
	théorique	p = 2	p = 3	p = 5	p = 7	p = 11	p = 13	p = 17	p = 19	p = 23	p = 29
Carré	0,1%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
Brelan	3,6%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
Double paire	2,7%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%
Paire	43,2%	0%	0%	0%	37,2%	49,7%	0%	0%	23,3%	0%	0%
Tous différents	50,4%	100%	100%	100%	62,8%	50,3%	100%	100%	76,7%	100%	100%
p-valeur du χ^2	1,00	0,91	0,91	0,91	1,00	1,00	0,91	0,91	0,99	0,91	0,91

Les p-valeurs des tests d'adéquation du χ^2 sont proches de 1, toutefois quel que soit p , le tore ne conduit jamais à l'obtention d'un carré, ni d'un brelan, ni d'une double paire alors que ces trois situations représentent 6,4 % des cas.

Ce test, rudimentaire, permet donc de se rendre compte que les valeurs engendrées par l'algorithme ne sont pas indépendantes.

4.3. Générateur du tore mélangé

Pour contourner le problème posé par la dépendance terme à terme des valeurs générées par l'algorithme du tore, nous proposons une adaptation qui consiste à « mélanger » ces valeurs avant de les utiliser.

4.3.1. Descriptif de l'algorithme

Notons (u_n) la suite générée par le nombre premier p . Au lieu d'utiliser la valeur u_n lors du n -ième tirage de la loi uniforme sur $[0 ; 1]$, nous proposons d'utiliser u_m où m est choisi de manière aléatoire dans $\mathbb{N}$.

Le générateur ainsi obtenu présente les mêmes bonnes caractéristiques globales que l'algorithme du tore sans la dépendance terme à terme. Il nécessite toutefois davantage de temps de simulation du fait du tirage de l'indice m . Nous proposons l'algorithme suivant lorsque l'on souhaite générer N réalisations de variables uniformes :

$$u_m = u_{\varphi(n)},$$

avec :

$$\varphi(n) = [\alpha \times N \times \tilde{u} + 1],$$

où :

- $[.]$ désigne l'opérateur partie entière,
- $\alpha \geq 1$,
- $\tilde{u}$ est la réalisation d'une variable aléatoire de loi uniforme.

Le facteur α a pour vocation de réduire le nombre de tirages qui donneraient lieu au même indice et donc au même nombre aléatoire. En effet, plus α est grand plus la probabilité de tirer deux fois le même nombre aléatoire est faible. Dans la pratique, $\alpha = 10$ est satisfaisant.

4.3.2. Choix de la procédure de mélange

Pour la génération de $\tilde{u}$, nous avons retenu le générateur *Rnd* ou tout autre générateur congruentiel ayant une période importante. En effet, utiliser l'algorithme du tore pour effectuer le mélange ne réduit pas la corrélation observée comme le montre le corrélogramme suivant.

Figure 40 - Corrélogramme de la suite générée par le tore ($p_1 = 5$) mélangé par le tore ($p_2 = 19$)

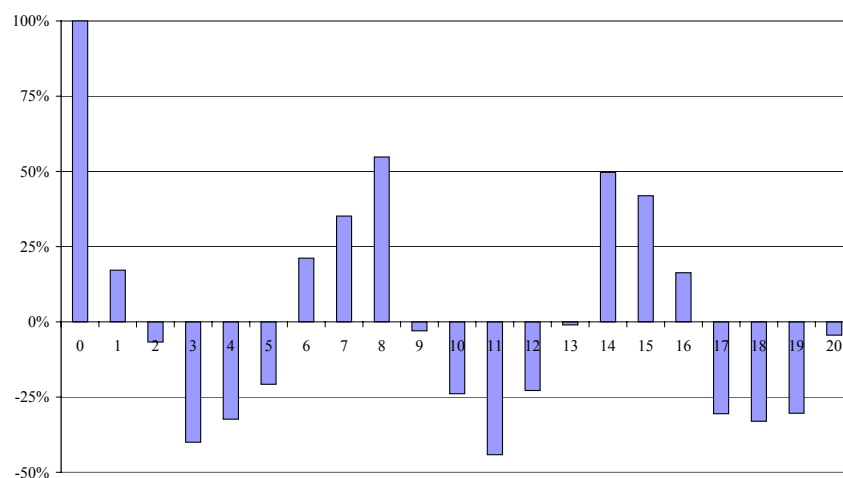
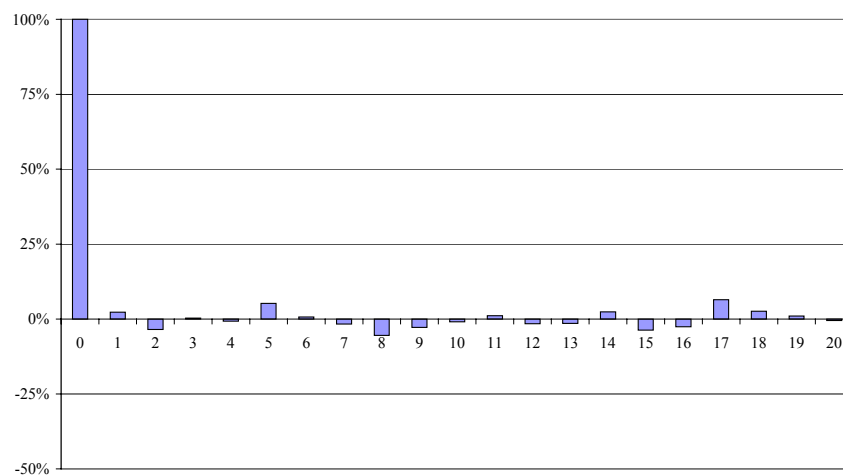


Figure 41 - Corrélogramme de la suite générée par le tore ($p_1 = 5$) mélangé par *Rnd*



En revanche l'utilisation de *Rnd* permet de réduire considérablement les corrélations. La Figure 41 nous permet de constater que le mélange a pratiquement fait disparaître la corrélation terme à terme, le générateur du *tore*

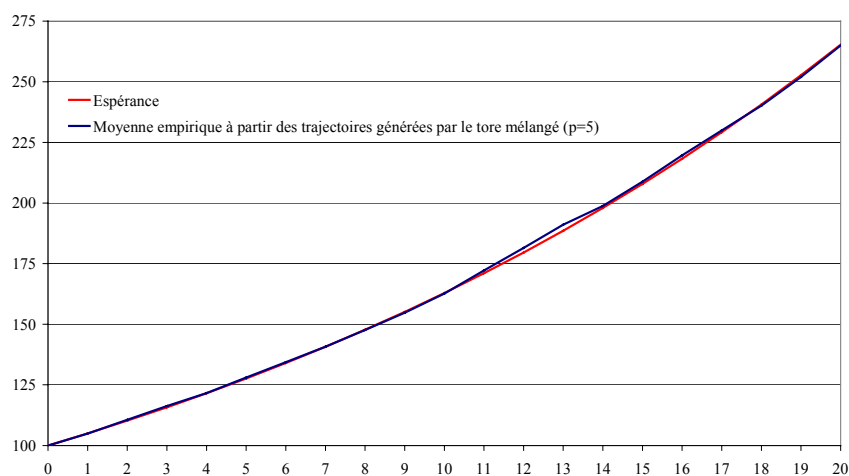
mélangé conservant néanmoins les propriétés de bonne répartition globale et de rapidité de convergence du tore.

De plus, comme le montre le Tableau 10, le tore mélangé par *Rnd* satisfait de manière satisfaisante le test du poker puisque toutes les « mains » sont représentées dans des proportions proches des fréquences théoriques.

Tableau 10 - Test du poker sur 40 000 réalisations générées par le tore mélangé par *Rnd*

	Fréquences théorique	Fréquences observées pour le tore mélangé par <i>Rnd</i> ($p=2$)
Carré	0,1%	0,1%
Brelan	3,6%	2,8%
Double paire	2,7%	2,5%
Paire	43,2%	37,2%
Tous différents	50,4%	57,4%
p-valeur du χ^2	1,00	1,00

Figure 42 - Trajectoires espérée et moyenne des 5 000 trajectoires simulées par le tore mélangé d'un mouvement brownien géométrique



La Figure 42 nous permet enfin de vérifier que les trajectoires générées à partir du tore mélangé sont satisfaisantes.

5. Simulation de la mortalité d'un portefeuille d'assurés

Plusieurs approches sont envisageables pour simuler la mortalité d'un ensemble d'assurés. Parmi elles, l'approche classique consiste à segmenter la période de temps qui sépare la date de simulation du terme du contrat en plusieurs sous intervalles puis à modéliser le passage de l'état « vivant » à l'état « décédé » successivement sur chaque intervalle de temps. En pratique, la segmentation retenue correspond le plus souvent à des périodes de référence

pour le versement de prestations (l'année, le trimestre ou le mois). Néanmoins cette approche s'avère être peu efficace en pratique, puisqu'elle conduit à utiliser un nombre conséquent de variables aléatoires pour simuler le parcours de l'assuré et par conséquent à simuler un grand nombre de réalisations de ces variables aléatoires.

Dans ce paragraphe, nous proposons une approche différente qui consiste à simuler, pour chaque assuré, directement sa date de décès. Dans la situation classique, où la loi de la mortalité est une table de mortalité, cette technique nécessite un travail préalable et l'utilisation d'une hypothèse de répartition des décès entre deux âges.

Le premier sous-paragraphe décrit l'algorithme permettant de simuler l'âge (entier) au décès, il est complété par le second qui permet d'obtenir l'âge précis au décès et donc la date du décès. En effet, les problématiques courantes en assurance nécessitent souvent de disposer de cette date. Par exemple, si l'on considère un contrat d'épargne tel que dans le Chapitre 1, un assuré peut mourir à 50 ans, le montant à reverser aux ayant-droits pourra être significativement différent selon que le décès survient avant ou après la date de revalorisation (31/12).

5.1. Simulation de l'âge au décès

Considérons un assuré d'âge réel y à la date de la simulation. L'objet de ce sous-paragraphe est de décrire un algorithme permettant de simuler l'âge au décès de cet assuré. Dans la suite, nous noterons t_s la date de début de la projection, t_{terme} la date de fin (éventuellement infinie) et τ la date aléatoire de décès. Enfin $[\cdot]$ désignera l'opérateur partie entière.

Partant d'une table de mortalité classique, il s'agit de déterminer les probabilités de décès de l'assuré sur les intervalles d'âges

$$\begin{cases} I_0 = [y, [y] + 1[, \\ I_i = [[y] + i, [y] + i + 1[, \end{cases}$$

sachant qu'il est en vie à la date de début de projection t_s . Notons $\theta(x, i)$ les probabilités :

$$\theta(x, i) = \Pr(\tau \in I_i \mid \tau > t_s).$$

Simuler un âge de décès revient à simuler une réalisation de variable aléatoire de loi uniforme sur l'intervalle $]0; 1[$ puis déterminer l'intervalle I_i correspondant.

Dans la suite on supposera que l'on dispose de la loi de mortalité des assurés et que celle-ci se présente sous la forme d'une table de mortalité classique où L_k représente le nombre d'assuré ayant atteint l'âge k . Avec les notations actuarielles classiques (cf. Planchet et Thérond (2006)), on détermine les probabilités conditionnelles q_k de décès à l'âge k sachant que l'on a atteint cet âge, de la manière suivante :

$$q_k = \frac{L_k - L_{k+1}}{L_k}.$$

En pratique, on observe généralement un âge non entier pour l'assuré. Pour passer d'une expression discrète à une expression non discrète, nous retenons l'approche de linéarisation de la fonction de survie, qui revient à supposer une répartition uniforme des sorties sur $[k, k+1[$ (hypothèse DUD, cf. Planchet et Thérond (2006) pour une revue des hypothèses de répartition des décès les plus classiques). Formellement, cela revient à déterminer :

$${}_{1-t}q_{[x]+t} = \frac{(1-t) \times q_{[x]}}{1-t \times q_{[x]}}.$$

De la même manière on détermine le L_x (où x n'est pas nécessairement entier) correspondant en utilisant la formule :

$$L_x = L_{[x]+t} = \frac{L_{[x]+1}}{1 - {}_{1-t}q_{[x]+t}}, \text{ où } t = x - [x],$$

et en remplaçant ${}_{1-t}q_{[x]+t}$ par son expression on obtient :

$$L_x = (1-t) \times L_{[x]} + t \times L_{[x]+1}.$$

On est maintenant en mesure de calculer les probabilités $\theta(x, i)$:

– pour la première période :

$$\theta(x, 0) = \frac{L_{[x]+1}}{L_{[x]+t}} \times {}_{1-t}q_{[x]+t},$$

– pour les périodes entières suivantes :

$$\theta(x, i) = \frac{L_{[x]+i}}{L_{[x]+t}} \times q_{[x]+i}.$$

– pour la dernière période non entière avant le terme :

$$\theta(x, n) = \frac{L_{[x]+n}}{L_{[x]+t}} \times {}_{\Lambda}q_{[x]+n},$$

où $\Lambda = t_{\text{terme}} - [x] + n$.

On construit alors la fonction $\Theta(x, \bullet)$ de la manière suivante :

$$\Theta(x, i) = \sum_{j=0}^i \theta(x, j), \text{ pour } 0 \leq i \leq n.$$

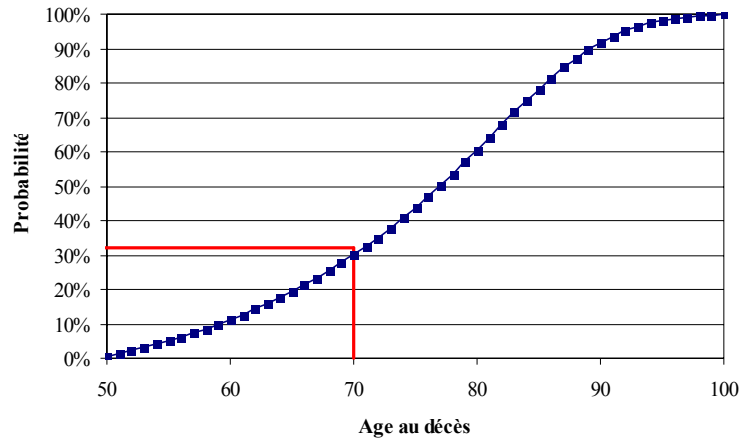
Par construction, $\Theta(x, i)$ représente la probabilité que l'assuré décède avant l'âge $[x] + i + 1$.

L'algorithme à suivre pour simuler la date de décès est le suivant :

- tirage d'une réalisation u de la loi uniforme sur $]0;1[$;
- si $u > \Theta(x, t_{\text{terme}})$ alors l'individu survit jusqu'au terme ;
- Si $\Theta(x, i) < u \leq \Theta(x, i+1)$ alors l'individu décède entre les âges entiers $[x]+i$ et $[x]+i+1$.

La Figure 43 reprend la fonction de répartition de l'âge au décès pour un assuré de 50 ans selon la TD 88-90.

Figure 43 - Distribution de l'âge au décès pour un assuré de 50 ans



Si le tirage de la v.a. uniforme est de l'ordre de 0,323, cela correspond à un décès à 70 ans.

5.2. Détermination de la date du décès

L'objet de ce sous-paragraphe est d'affiner l'algorithme présenté supra de manière à déterminer la date exacte du décès. Ce point est d'importance dans des situations où le versement de prestations (ou le montant de celles-ci) est conditionné par la présence de l'assuré à des dates prédéfinies. Par exemple, versement d'une rente si l'assuré est en vie en début de trimestre ou revalorisation de l'épargne selon des règles différentes selon que l'assuré décède avant ou après la date de revalorisation des contrats.

Considérons la réalisation u obtenue précédemment. On cherche le réel δ qui représente la partie non entière de l'âge réel au décès $[x]+i+\delta$. Il est solution de l'équation :

$$u - \Theta(x, i) = {}_{\delta}q_{[x]+i}.$$

Conformément à l'hypothèse DUD de répartition des décès, δ est donné par :

$$\delta = \frac{\delta q_{[x]+i}}{q_{[x]+i}} = \frac{u - \Theta(x, i)}{\Theta(x, i+1) - \Theta(x, i)}.$$

L'intérêt de la méthode décrite *supra* est évident en termes de temps de calcul.

Selon la TD 88-90, l'espérance de vie à 50 ans pour un homme est de l'ordre de 26 ans. L'approche pas à pas nécessitera donc en moyenne 26 simulations de variables aléatoires avant que l'assuré ne meurt. Un seul tirage est nécessaire dans notre méthode.

6. Conclusion

Les nouvelles normes comptables, prudentielles et économiques imposent le recours le plus fréquent aux méthodes de simulation. En particulier, elles sont à la base du modèle interne qui permettra de déroger à la formule standard de détermination du capital de solvabilité dans Solvabilité 2. Si les méthodes de Monte-Carlo s'avèrent des outils particulièrement intéressants puisqu'elles permettent d'obtenir des résultats numériques lorsque les processus étudiés sont trop complexes pour être étudiés analytiquement, il n'en demeure pas moins que leur utilisation nécessite des traitements préalables et des attentions nouvelles (optimisation des algorithmes de calcul, développements informatiques).

En particulier, la génération pratique de trajectoires de variables modélisées par des processus continus nécessite l'estimation des paramètres, la discrétisation du processus et la génération de nombres aléatoires.

L'étape de discrétisation conduit à arbitrer entre précision et temps de calcul à moins qu'une discrétisation exacte peu coûteuse en calculs soit disponible. Lorsque ce n'est pas le cas, on préférera le schéma de Milstein au schéma d'Euler car il demande le même nombre de tirages de nombres aléatoires pour une meilleure précision.

L'estimation des paramètres doit faire l'objet d'une attention particulière du fait des biais auxquels conduirait une estimation « naïve ». Il convient de préférer la méthode d'estimation par inférence indirecte à une estimation directe souvent porteuse de biais. Le principe de l'estimation *ad hoc* permettra également d'estimer des paramètres satisfaisants lorsque la variable d'intérêt n'est pas directement la variable simulée, ce qui est le plus souvent le cas dans les problématiques d'assurance.

Enfin nous recommandons l'utilisation du *tore mélangé* pour toute construction pas-à-pas de trajectoires. Ce générateur, simple à mettre en place, donne des résultats très satisfaisants, notamment comparé au *Rnd* d'Excel.

Par ailleurs, nous avons vu dans le cas de la simulation de la mortalité d'un portefeuille d'assurés que des méthodes alternatives aux méthodes les plus naturelles (simulation pas à pas qui « colle » à la forme des lois de mortalité) permet de gagner en temps de calcul.

Annexe : Tests d'adéquation à une loi

Le lecteur trouvera un descriptif plus détaillé de ces tests dans Saporta (1990) pour les deux premiers et dans Partrat et Besson (2005) pour le dernier.

1. Test du χ^2

Soit X une variable aléatoire à valeurs dans $\mathbf{E}$ de loi P_X inconnue et P_0 une loi connue sur $\mathbf{E}$. Soit $A_1, \dots, A_d$ une partition de $\mathbf{E}$ telle que $\pi_k = P_0[A_k] > 0$ pour tout k . Nous disposons d'un n -échantillon indépendant et identiquement distribué $(X_1, \dots, X_n)$ de X . Soit N_k le nombre de variables aléatoires X_i dans A_k . Si $P_X = P_0$ (H_0) considérons la statistique D^2 définie comme suit :

$$D^2 = \sum_{k=1}^d \frac{(N_k - n\pi_k)^2}{n\pi_k}.$$

La statistique D^2 est asymptotiquement distribuée comme une variable de χ_{d-1}^2 . La p-valeur de ce test est donc donnée par :

$$\hat{\alpha} = P\left[\chi_{d-1}^2 > D^2\right].$$

2. Test de Kolmogorov-Smirnov

Si F_n^* est la fonction de répartition empirique d'un n -échantillon d'une variable aléatoire de fonction de répartition F , alors la statistique $D_n = \sup \left| F_n^*(x) - F(x) \right|$ est asymptotiquement distribuée comme suit :

$$P\left[\sqrt{n} D_n < y\right] \rightarrow \sum_{k=-\infty}^{\infty} (-1)^k \exp\left\{-2k^2 y^2\right\}.$$

La statistique D_n est indépendante de F et le résultat précédent nous permet de tester :

$$\begin{cases} H_0 : F(x) = F_0(x) \\ H_1 : F(x) \neq F_0(x) \\ \alpha \end{cases}$$

3. Test d'Anderson-Darling

Ce test repose sur l'écart d'Anderson-Darling A_n^2 défini comme suit :

$$A_n^2 = n \int_0^1 \frac{\left(F_n^*(x) - F(x)\right)^2}{F(x)(1-F(x))} dF_0(x)$$

Cet écart permet de tester

$$\begin{cases} H_0 : F(x) = F_0(x) \\ H_1 : F(x) \neq F_0(x) , \\ \alpha \end{cases}$$

l'expression opérationnelle de A_n^2 étant :

$$A_n^2 = -n - \frac{1}{n} \sum_{i=1}^n (2i-1) \left\{ \ln F(x_{(i)}) + \ln \left[1 - F(x_{(n-i+1)}) \right] \right\} ,$$

où $x_{(i)}$ désigne la i -ème plus petite réalisation de X dans l'échantillon.

On rejettera H_0 si

$$-n - \frac{1}{n} \sum_{i=1}^n (2i-1) \left\{ \ln F_0(x_{(i)}) + \ln \left[1 - F_0(x_{(n-i+1)}) \right] \right\} ,$$

est supérieur à une valeur que la variable aléatoire A_n^2 a une probabilité α de dépasser.

Bibliographie

- Augros J.C., Moreno M. (2002) *Les dérivés financiers et d'assurance*, Paris : Economica.
- Black F., Scholes M. (1973) « The pricing of options and corporate liabilities », *Journal of Political Economy* 81 (3), 637-54.
- Cox J.C., Ingersoll J.E., Ross S.A. (1985) « A theory of the term structure of interest rates », *Econometrica* 53, 385-407.
- Fargeon L., Nissan K. (2003) *Recherche d'un modèle actuariel d'analyse dynamique de la solvabilité d'un portefeuille de rentes viagères*, Mémoire d'actuariat, ENSAE.
- Giet L. (2003) « Estimation par inférence indirecte des équations de diffusion : l'impact du choix du procédé de discrétisation », Document de travail n° 03A15 du GREQAM.
- Hami S. (2003) *Les modèles DFA : présentation, utilité et application*, Mémoire d'actuariat, ISFA.
- Hull J.C. (1999) *Options, futures and other derivatives*, 4th edition, Prentice-Hall.
- Kaufmann R., Gadmer A., Klett R. (2001) « Introduction to Dynamic Financial Analysis », *ASTIN Bulletin* 31 (1), 213-49.
- Kloeden P., Platen E. (1995) *Numerical Solution of Stochastic Differential Equations*, 2nd Edition, Springer-Verlag.
- Lamberton D., Lapeyre B. (1997) *Introduction au calcul stochastique appliqué à la finance*, 2^e édition, Paris : Ellipses.
- Partrat Ch., Besson J.L. (2005) *Assurance non-vie. Modélisation, simulation*, Paris : Economica.
- Planchet F., Thérond P.E. (2005c) « Simulation de trajectoires de processus continus », *Belgian Actuarial Bulletin* 5, 1-13.
- Planchet F., Thérond P.E. (2006) *Modèles de durée. Applications actuarielles*, Paris : Economica.
- Planchet F., Thérond P.E., Jacquemin J. (2005) *Modèles financiers en assurance. Analyses de risque dynamiques*, Paris : Economica.
- Roncalli T. (1998) *La structure par terme des taux zéro : modélisation et implémentation numérique*, Thèse de doctorat, Université Bordeaux IV.
- Saporta G. (1990) *Probabilités, analyse des données et statistique*, Paris : Technip.
- Vasicek O. (1977) « An equilibrium characterization of the term structure », *Journal of financial Economics* 5, 177-88.

CONCLUSION GÉNÉRALE

L'avènement des nouveaux référentiels comptable, prudentiel et de communication financière que constituent la phase II de la norme IFRS dédiée aux contrats d'assurance, la Directive européenne Solvabilité 2 et le cadre MCEV vont modifier profondément le rapport au risque des assureurs. En effet les assureurs vont devoir revoir leurs processus de valorisation des contrats d'assurance et, de manière plus générale, intégrer de nouvelles contraintes dans la gestion effective de la compagnie.

De nouvelles règles de valorisation

Concernant la valorisation des portefeuilles et des sociétés, les nouvelles règles se distinguent de l'actuel cadre prudentiel français et du Code des assurances.

En premier lieu, les nouveaux référentiels imposent un recours systématique aux hypothèses actuarielles les plus adaptées de manière à procéder à des évaluations dites *best estimate*. Au niveau des provisions techniques (comptables comme de solvabilité), ces évaluations *best estimate* seront complétées de marges pour risque en référence explicite au risque réellement supporté. En effet, ces marges pour risque auront pour vocation :

- de représenter la prime de risque sur laquelle l'engagement pourrait être échangé sur un marché organisé (conception IFRS) ;
- de permettre le transfert de l'engagement d'assurance d'un assureur vers un autre en finançant les coûts d'immobilisation du capital correspondants, en cas de faillite du premier assureur (conception Solvabilité 2).

Les deux conceptions se rejoignent partiellement en cela qu'elles font référence au transfert de l'engagement d'assurance. Néanmoins, elles se distinguent par les conditions dans lesquelles s'effectue ce transfert : la faillite dans le cas de Solvabilité 2 alors que ce sont des conditions standard de marché qu'évoquent les normes IFRS.

Le recours à la référence au marché est d'ailleurs omniprésent puisque, dans les trois référentiels, les risques financiers doivent être valorisés grâce aux méthodes économiques initialement développées pour donner un prix aux dérivés financiers. Ainsi des garanties de taux ou de participation aux bénéfices sont assimilées à des options financières. Conceptuellement, cette association amène deux commentaires :

- les actifs sur lesquels sont assises ces garanties ne sont généralement pas cotés sur des marchés (il s'agit généralement de l'actif général de l'assureur) et leur évolution n'est pas indépendante du comportement de l'assureur (ex : actif général d'un assureur) ;

- les agents économiques concernés (assurés et assureurs) n'ont pas nécessairement le comportement économique que celui que l'on prête à l'investisseur rationnel sur les marchés financiers.

En effet, la situation de l'assurance est complexe et la prise en compte uniquement des garanties financières minimales que procure l'assureur à ses assurés ne reflète pas la réalité de leurs relations. Par exemple, pour des contrats d'assurance vie de type épargne en euros, si l'assureur se contente de revaloriser l'épargne de ses assurés uniquement du minimum contractuel (éventuellement taux garanti et participation aux bénéfices), il risque d'être confronté à des vagues de rachat de contrats et (surtout) à un risque d'image qui réduira sa capacité à se développer en commercialisant de nouveaux contrats.

Le deuxième grand chantier est celui des exigences de capitaux de solvabilité.

Le critère annoncé par Solvabilité 2, i.e. le contrôle de la probabilité de ruine à 99,5 %, est très ambitieux. Il s'agit, en effet, pour les sociétés d'assurance, de disposer de ressources suffisantes pour se prémunir de la ruine, à horizon un an, dans 199 cas sur 200.

On l'a vu, les assureurs pourront déterminer l'exigence de fonds propres grâce à une formule standard commune à tous les assureurs ou par le biais d'un modèle interne qui leur sera propre et qui est censé mieux représenter le risque effectif auquel ils sont soumis.

La formule standard proposée dans la troisième étude d'impact quantitatif menée par le CEIOPS (QIS 3) repose sur une approche modulaire des risques : à chaque risque correspond un niveau de capital élémentaire, ceux-ci font l'objet d'une agrégation pour aboutir à l'exigence globale. Au final, le montant obtenu est censé approcher la Value-at-Risk à 99,5 % annoncée comme objectif. Le modèle repose cependant sur des hypothèses simplificatrices qui tendent à biaiser les résultats. En particulier, ce modèle est construit dans un univers gaussien alors même que les grandeurs en jeu reflètent des événements extrêmes éminemment non gaussiens.

Par ailleurs, la mise en place d'un modèle interne nécessite de modéliser l'ensemble des risques auquel est exposé l'assureur puis de simuler leur réalisation de manière à disposer « d'observations » pour estimer le capital de solvabilité. Ces modèles se heurtent principalement aux problèmes opérationnels du choix de modèles qui n'est pas immédiat, surtout lorsqu'il s'agit de représenter des queues de distribution et donc des phénomènes que l'on observe, par définition, rarement. De plus, la multiplicité des risques et leurs imbrications nécessite de modéliser leurs dépendances et le comportement des assurés et des assureurs, y compris dans des situations extrêmes, pour lesquelles on ne dispose pas d'historique (ex : quid du rachat de contrats d'épargne dans une économie avec des taux d'intérêt à 15 % ?). La complexité de ces modèles nécessite le recours aux techniques de simulation, dont on sait qu'elles sont gourmandes en ressource informatique et en temps de calcul et qu'elles peuvent conduire à un risque opérationnel important du fait de leur complexité et des pièges qu'elles comportent.

Entre formule standard et modèle interne devraient se développer des modèles partiels. Il s'agit en fait d'utiliser le format modulaire de l'approche

standard en utilisant comme capitaux élémentaires pour certains risques des résultats issus de modèles développés spécifiquement pour ces risques. En effet, les sociétés d'assurance de taille modeste et qui sont spécialisés dans un nombre restreint de branches, dont certaines mutuelles d'assurances, sont les plus affectés¹ en termes d'exigences de capitaux par la formule standard. Si ces sociétés n'ont pas les moyens financiers de développer ou d'acquérir un modèle interne global, elles auront certainement intérêt à développer un modèle partiel sur les risques qu'elles assurent et qu'elles maîtrisent particulièrement bien.

En parallèle le régulateur devra se donner des règles strictes d'appréciation des modèles internes (globaux ou partiels) en imposant, par exemple, un contrôle de leur adéquation sur les valeurs extrêmes.

De nouvelles règles de gestion

La mise en place des nouveaux standards prudentiels et de reporting financier n'est pas sans conséquence en matière de gestion de la compagnie.

La référence explicite au risque réellement supporté par l'assureur dans la détermination du capital de solvabilité a une grande incidence en terme de gestion technique de la société puisque toute mesure de gestion des contrats (revalorisation de contrats, abaissement de franchise, etc.), tout dispositif de réassurance et tout acte de gestion des actifs financiers a potentiellement un impact sur le niveau minimal des fonds propres dont doit disposer l'assureur. Les sociétés vont donc devoir poursuivre la normalisation des processus de gestion en vue de prévenir le risque opérationnel qui résulterait d'une erreur de gestion.

À terme, ces nouvelles règles devraient également participer à l'essor de la titrisation des risques d'assurance. Compte tenu des économies de capitaux qu'elle peuvent permettre de réaliser, ces opérations de titrisation vont tendre à se développer.

Le processus s'est d'ailleurs déjà accéléré avec les opérations d'AXA sur le risque auto² et de Swiss Re et Axa sur le risque systématique de mortalité. Concernant le risque de longévité, JP Morgan a mis en place un indice synthétique pour suivre l'évolution de la mortalité. Déjà disponibles pour les populations américaines, anglaises et galloises, cet indice devrait permettre le développement de dérivés permettant de se couvrir contre des variations adverses de l'évolution de la longévité. De telles initiatives devraient encore plus favoriser la titrisation des risques d'assurance.

Par ailleurs, les états financiers qui résultent de l'application des normes IFRS, de même que la *Market Consistent Embedded Value*, sont destinés aux marchés financiers. Compte tenu de l'importance de l'information financière sur la valorisation boursière des sociétés, les assureurs vont être amenés à considérer l'impact des décisions de gestion envisagées sur les états financiers

¹. Cf. les résultats de QIS 2 publiés par le CEIOPS au niveau global et par l'ACAM au niveau français.

². Cette opération était particulièrement novatrice dans la mesure où la titrisation n'avait jusqu'alors concernée que des risques à faible fréquence et à forte amplitude (risques catastrophiques essentiellement), l'objectif de cette opération était de se prémunir contre une dérive d'un risque de forte fréquence et de faible variabilité.

avant de trancher entre plusieurs alternatives. Dans le futur, on peut penser que les produits commercialisés auront été pensés en fonction de leur rentabilité attendue mais également de leur susceptibilité à produire de « bons » états financiers, par exemple à ne pas conduire à des jeux de comptes trop volatiles.

Valorisation versus gestion des risques

Si les nouveaux standards conduisent à mieux identifier et mesurer les risques supportés par les sociétés d'assurance, il convient néanmoins de distinguer valorisation et gestion des risques.

Lorsque l'on détermine, par exemple, un niveau de provision, on calcule une valeur destinée à alimenter des états financiers. Même si les règles préconisées exigent la prise en compte des risques dans la méthode d'évaluation, l'aboutissement de la démarche n'en demeure pas moins qu'une valeur, censée représenter l'engagement considéré.

La démarche de valorisation est en cela différente de la démarche de gestion effective du risque qui consiste à prendre des décisions sur la manière de se protéger (ou pas) contre le risque supporté. L'exemple caractéristique est celui des garanties plancher en cas de décès sur les contrats en unités de compte présenté dans le Chapitre 1 : les méthodes de valorisation nous donnent un prix mais l'assureur a le choix de gérer le risque en mettant en œuvre ou pas la stratégie de couverture qui lui permet de réduire considérablement le risque associé à l'engagement.

S'ils ne constituent pas des règles de gestion, les nouveaux référentiels devraient néanmoins permettre aux assureurs d'apprécier l'exhaustivité des risques auxquels ils sont soumis et d'encourager leur bonne gestion via des économies d'exigences de capitaux propres.

Annexe

Solvabilité 2 - les modèles proposés par QIS 3

Dans le cadre de l'élaboration du futur référentiel prudentiel commun aux assureurs européens, la Commission Européenne a demandé conseil au CEIOPS¹ pour élaborer de nouvelles normes de solvabilité et de contrôle. Ces travaux sont accompagnés d'études d'impact quantitatif (*Quantitative Impact Studies* - QIS) de manière à estimer les impacts quantitatifs du système de solvabilité en construction. Ces études d'impact reposent sur la participation des entreprises d'assurance européennes : le CEIOPS publie un cahier de spécifications techniques puis, sur ces bases, les assureurs peuvent réaliser les calculs demandés sur leur propre portefeuille et communiquer les résultats à leur autorité de tutelle (l'ACAM pour les assureurs français) qui les restitue au CEIOPS.

La première étude d'impact QIS 1 visait principalement à mesurer l'impact sur le niveau des provisions techniques qu'engendrerait la mise en place de nouvelles exigences de provisionnement : Value-at-Risk à 60 %, 75 % ou 90 %.

QIS 1 a rencontré un grand succès puisque 312 assureurs (dont 150 pratiquent des activités vie, 190 des activités non-vie et 4 sont des réassureurs) ont répondu à l'enquête. Ces entreprises ne furent néanmoins pas toutes capables de produire des chiffres avant la date limite de l'enquête. Cette étude a néanmoins permis de constater que, dans la plupart des cas, des provisions techniques calculées par une Value-at-Risk à 75 % ou 90 % sur des hypothèses *best estimate* aboutissaient à des montants inférieurs à ceux qui figurent aujourd'hui dans les comptes des assureurs.

La deuxième étude d'impact QIS 2 visait principalement à obtenir des informations pour élaborer la formule standard de calcul du capital de solvabilité qui doit approximer la Value-at-Risk à 99,5 % du risque global sur un an. Cette formule standard est construite sur la base d'agrégation de capitaux de solvabilité obtenus risque par risque. Un objectif annexe était d'approfondir les résultats de QIS 1 en faisant des tester des approches coût du capital (*Cost of Capital* - CoC) dans le calcul des provisions techniques.

Cette étude d'impact était donc particulièrement importante pour les assureurs dans la mesure où elle était destinée à élaborer la formule d'exigence

¹. *Committee of European Insurance and Occupational Pension Supervisors*. Le CEIOPS regroupe l'ensemble des autorités de contrôle des pays de l'Union européenne.

de capitaux propres d'une part et à indiquer quels éléments pouvaient venir en couverture de ce montant de capitaux.

Les assureurs européens ont massivement répondu présent puisque 514 d'entre eux ont participé à l'étude, ce qui représente 65 % du marché européen en vie et 56 % en non-vie. Il convient de remarquer que les sociétés françaises (76) et parmi elles les mutuelles (31) se sont particulièrement bien mobilisées. Néanmoins, les résultats de cette étude ne sauraient être sans biais dans la mesure où ce sont principalement les grandes entreprises des trois principaux marchés (Royaume-Uni, France et Allemagne) qui font le gros des réponses.

En moyenne, les exigences de capitaux qui résultent de la formule QIS 2 s'avèrent légèrement supérieures aux exigences Solvabilité 1 (marge de solvabilité). On remarque cependant que certaines catégories d'assureur sont particulièrement impactés par cette formule standard. C'est le cas des petites sociétés d'assurance non-vie qui, le plus souvent, ne pratiquent qu'une branche d'activité et ne bénéficient donc pas des facteurs de diversification qui leurs permettraient de réduire leur capital de solvabilité. À ce problème s'ajoute celui de leur taille, puisqu'ils ne bénéficient pas pleinement de l'effet de mutualisation du fait de leur taille.

Dans ce contexte, le CEIOPS a conçu une nouvelle étude d'impact quantitatif qui a pour but d'affiner la formule standard. QIS 3 a débuté le 1^{er} avril 2007 et les résultats doivent être retournés aux autorités de contrôle avant le 30 juin 2007. Cette date limite est particulièrement importante puisque le projet de directive cadre sur Solvabilité 2 doit être publié en juillet.

Ayant tiré les conclusions des QIS 1 et 2, QIS 3 propose une formule standard modifiée et donne la préférence à la méthode coût du capital dans le calcul des provisions techniques. Les grandes lignes de cette étude sont reprises ci-après.

1. Modèles d'évaluation : l'approche standard

L'objet de ce premier paragraphe est de décrire de manière synthétique les modèles retenus pour valoriser les provisions techniques d'une part et les actifs en représentation des engagements d'autre part.

1.1. Les actifs

Les différents actifs devront être évalués en valeur de marché. Ces valeurs devront être égales à des prix du marché fiables et observables que l'on trouve dans des marchés liquides. Pour des positions dites longues (positions à l'achat) sur l'actif, le prix approprié et coté du marché est le prix offert (*bid price*) pris à la date d'évaluation, alors que pour des positions dites courtes (positions à la vente), c'est le prix vendeur (*offer price*).

Dans les cas où il n'y a pas de valeur de marché disponible, une approche alternative peut être adoptée, mais celle-ci doit être encore cohérente avec les informations de marché.

Les actifs non-négociables ou peu liquides devront être évalués sur des bases prudentes, en prenant pleinement en compte la réduction de valeur due au crédit et à la liquidité des risques attachés :

- en l’absence d’indications suffisantes, la valeur de ces actifs ne devra pas être plus élevée que leur coût d’acquisition dont on aura déduit à la fois, la marge de profit estimée et faite par le vendeur, et la dépréciation due à l’utilisation ou la désuétude de l’actif en question ;
- en l’absence d’indications suffisantes, les actifs incorporels, les fournitures, les installations et actifs similaires avec un risque significatif de dépréciation devront être évalués comme nuls.

Si des avis fiables et indépendants d’experts sont disponibles, ils devront être pris en compte dans l’évaluation.

1.2. Les provisions techniques

L’évaluation des provisions techniques doit être effectuée en décomposant les passifs répliquables ou *hedgeable*, d’une part, et les passifs non-répliquables ou *non-hedgeable* d’autre part. Un passif est répliquable si les flux qu’ils engendrent peuvent être parfaitement répliqués ou couverts à l’aide d’instruments financiers se monnayant sur un marché liquide et transparent².

Lorsque la distinction entre passifs répliquables et non-répliquables est incertaine ou lorsque des valorisations *market-consistent* ne peuvent être obtenues, c’est l’approche non-répliquable (*best estimate* + marge pour risque) qui s’imposera.

Le niveau des provisions techniques ne pourra pas être réduit en raison de la solvabilité de l’assureur lui-même.

Les valeurs des passifs répliquables et non-répliquables devront être présentées distinctement. De plus, pour les passifs non-répliquables, le *best estimate* et la marge pour risque feront de même l’objet d’une présentation séparée.

Si un passif peut être parfaitement répliqué sur un marché liquide, profond et transparent, le portefeuille répliqué procure de façon immédiate un prix observable qui doit être utilisé pour valoriser le passif. On parle d’évaluation *marked to market*.

1.2.1. Le best estimate

Le CEIOPS définit le *best estimate* comme étant égal à la valeur actuelle probable des potentiels flux futurs de trésorerie estimés à partir d’informations courantes et fiables et à partir d’hypothèses spécifiques.

Une projection à horizon assez long saisissant tous les flux résultant de contrats ou de groupes de contrats devra être utilisée. Si l’horizon de la projection n’est pas étendu jusqu’au terme du dernier contrat ou de la dernière indemnisation, l’entreprise devra s’assurer que l’utilisation d’un horizon réduit n’affecte pas significativement les résultats.

Un certain nombre d’hypothèses est à prendre en compte :

². Le CEIOPS définit de tels marchés comme des marchés dans lesquels les participants peuvent rapidement exécuter un large volume de transactions avec de faibles impacts sur les prix.

- les flux financiers devront être actualisés au taux sans risque de marché observé pour des instruments de même durée³;
- la valeur actuelle des chargements sur contrats et la valeur actuelle des dépenses probables devront être examinées dans les projections de flux.

Pour ce qui est des passifs cédés en réassurance, le *best estimate* devra être estimé par un calcul brut et par un calcul net. Dans certaines compagnies de réassurance, l'existence d'un certain laps de temps entre l'enregistrement d'une indemnisation et le moment de l'indemnisation directe à l'assuré doit être prise en compte lors de l'évaluation du *best estimate* net. Lors du calcul du *best estimate* net, il est supposé que le réassureur ne fera pas défaut⁴.

Par ailleurs, les primes futures ne seront prises en compte que dans la mesure où :

- elles sont prévues contractuellement ;
- elles permettent à l'assuré de bénéficier de conditions plus favorables.

Dans le deuxième cas, il s'agira d'estimer la probabilité d'occurrence du versement de ces primes.

Il est précisé dans la cahier de spécification de QIS 3 que le *best estimate* devra être évalué en utilisant au moins deux différentes méthodes. Deux méthodes sont considérées différentes quand elles sont basées sur différentes techniques et sur différents groupes d'hypothèses. Au final, c'est le résultat le plus prudent qui sera retenu comme le *best estimate*.

1.2.2. La marge pour risque

C'est une méthode de coût du capital (CoC) qui doit être utilisée dans la détermination de la marge pour risque pour les passifs non-réplicables.

Le calcul de la marge CoC repose sur l'hypothèse d'insolvabilité de l'entreprise d'assurance qui, à la fin de l'année en cours ne dispose plus d'aucun capital. Les provisions techniques doivent être suffisamment importantes pour permettre le transfert des engagements de l'assureur ruiné vers un autre assureur. La marge pour risque est alors définie comme le coût de la valeur actuelle du futur SCR que devra immobiliser l'assureur qui recevra le passif jusqu'à sa liquidation.

Cette méthodologie est décrite plus précisément dans le sous-paragraphe suivant.

Concernant les actifs réplicables, la marge pour risque n'est pas à expliciter puisqu'elle est incluse dans leur prix de marché (cf. le Chapitre 1).

1.2.3. La méthode du coût du capital

Les étapes de calcul de la marge pour risque avec la méthode CoC se résument en trois étapes :

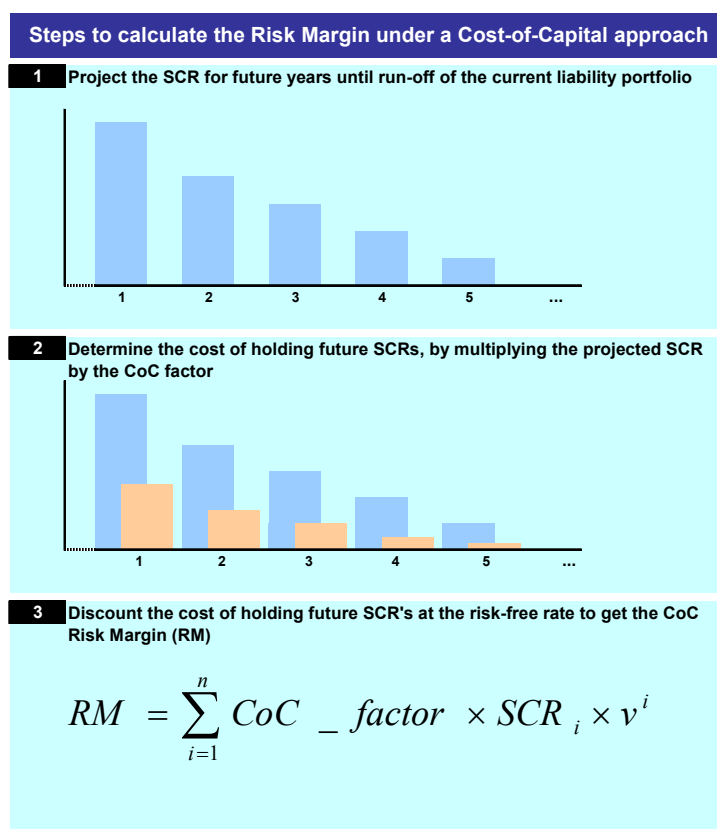
³. Lorsque les marchés ne recèlent pas une telle information, elle devra être extrapolée à partir des taux observés.

⁴. Ce risque est pris en compte à un autre niveau : celui du capital de solvabilité.

- Détermination du SCR pour chaque exercice futur jusqu'à la liquidation du portefeuille. Le SCR projeté sera estimée sur des bases simplifiées par rapport au SCR global.
- Multiplication de chacun des futurs SCR par le facteur de coût du capital (par exemple 6 % au dessus du taux sans risque) afin d'obtenir le coût de détention des futurs SCR.
- Actualisation des montants obtenus par la courbe des taux sans risque. La somme de ces montants actualisés correspond à la marge pour risque recherchée.

Cette démarche est résumée sur la Figure 44.

Figure 44 - Étapes de calcul de la marge pour risque (méthode CoC)



Tous les participants devront prendre un même facteur CoC de 6 % au-dessus de la courbe des taux sans risque pour l'évaluation de la marge pour risque.

2. Provisions techniques en assurance vie

Pour les activités en assurance vie, la segmentation générale suivante devra être opérée :

- contrats avec des clauses participatives de profit ;
- contrats pour lesquels le souscripteur supporte le risque d'investissement ;
- autres contrats sans clause participative ;
- réassurance.

Dans un deuxième temps, pour chaque premier niveau de segmentation, une deuxième subdivision consistera à distinguer :

- les contrats en cas de décès ;
- les contrats en cas de vie ;
- les contrats pour lesquels le principal risque est un risque de type arrêt de travail ;
- les contrats d'épargne.

L'évaluation doit être effectuée sur des données extraites police par police, mais des méthodes raisonnables et actuarielles peuvent être utilisées. En particulier la projection des flux futurs basée sur des polices types est autorisée : l'utilisation de données agrégées est donc possible.

Les projections de flux doivent aussi prendre en compte les éventuels options dont disposent les assurés. Ceci implique notamment que les règles de gestion de la société (gestion d'actifs, PB discrétionnaire, etc.) doivent être modélisées dans la mesure où elles ont un impact sur le comportement des assurés vis-à-vis de ces options (renonciation, rachat de contrats, etc.)

3. Provisions techniques en assurance non-vie

Les différentes valeurs à prendre en considération en assurance non vie doivent être indiquées pour chaque ligne d'activité définies à l'article 63 du *Council Directive* dans les comptes annuels et les comptes consolidés des entreprises d'assurance, à savoir selon les types d'activités suivants :

- accident et santé ;
- véhicules et responsabilité des tiers ;
- véhicules et autres classes ;
- marine, aviation, transport ;
- incendie et autres dégâts ;
- responsabilité des tiers ;
- crédit et caution ;
- dépenses légales ;
- assistance ;
- diverses assurance non vie.

4. Capital éligible

Le capital éligible joue le même rôle dans Solvabilité 2 que la marge de solvabilité dans l'actuel référentiel Solvabilité 1. Il s'agit donc des capitaux

éligibles à la représentation du SCR. Ce capital éligible se décompose en trois catégories en fonction de leur capacité à absorber les pertes. Ainsi les capitaux de première catégorie (*Tier 1*) comme par exemple les capitaux propres ou les réserves statutaires sont admis sans limite, alors que ceux de deuxième catégorie (*Tier 2*) sont admis avec des limites et que ceux de troisième catégorie, comme par exemple la capacité des mutuelles au rappel de cotisations, doivent faire l'objet d'une autorisation préalable du superviseur pour être admis.

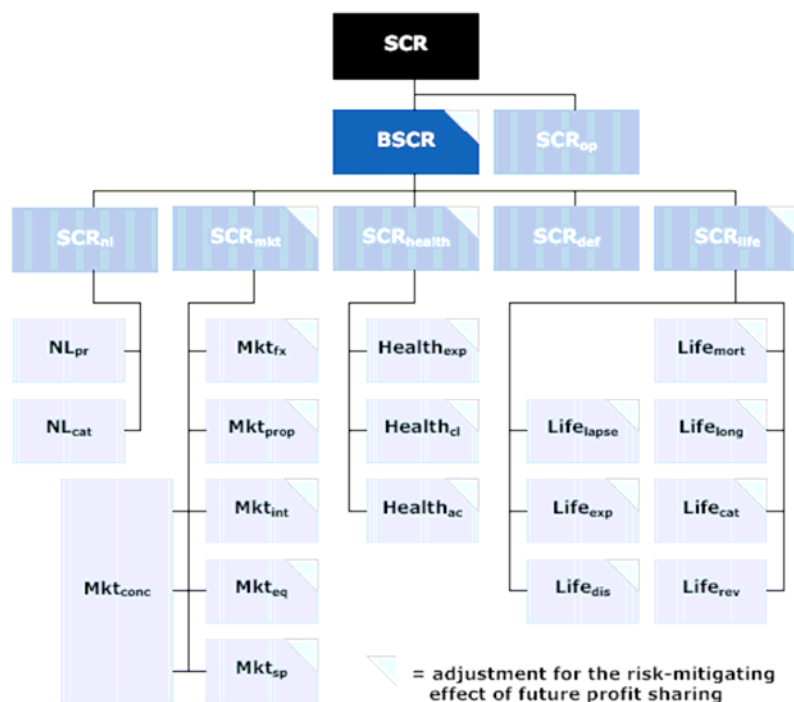
5. Capital de solvabilité (SCR) : formule standard

La formule standard telle que proposée dans QIS 3 repose sur une approche modulaire qui divise le SCR selon les différents risques identifiés. Après une présentation de cette approche, nous étudierons plus particulièrement la manière dont sont modélisés certains risques qui la composent.

5.1. Approche modulaire

L'approche modulaire retenue pour la formule standard est résumée par la Figure 45.

Figure 45 - Structure modulaire de la formule standard



Au premier niveau, le SCR provient de l'agrégation de capitaux obtenus pour deux composantes : le SCR de base (*basic SCR* ou BSCR) et le risque opérationnel (*SCR_{op}*).

Les paramètres et hypothèses retenues pour le calcul du SCR sont déterminés de manière à ce que le modèle estime une Value-at-Risk à 99,5 % sur le risque global de la compagnie à horizon un an.

En pratique, la démarche consiste, pour chaque risque, à déterminer un montant de capital (à partir d'un scénario défavorable) puis à agréger ensemble des risques de même niveau et ainsi de suite en remontant dans la hiérarchie de la formule modulaire.

Par exemple, au niveau le plus haut de la hiérarchie, le niveau final du SCR s'obtient par la formule d'agrégation suivante :

$$SCR = BSCR + SCR_{op},$$

où SCR_{op} désigne le capital associé au risque opérationnel et $BSCR$ le *basic* SCR.

Les deux termes sont directement additionnés pour donner le SCR, cela traduit le fait que les risques sous-jacents ne se mutualisent pas.

5.2. Basic SCR

Le *basic* SCR est obtenu grâce à l'agrégation des capitaux associés à cinq risques différents :

- le risque de marché SCR_{mkt} ;
- le risque de défaut SCR_{def} ;
- le risque de souscription en vie SCR_{life} ;
- le risque de souscription en non-vie SCR_{nl} ;
- le risque de santé SCR_{health} .

Pour représenter la capacité des assureurs à absorber des chocs grâce aux participations aux bénéfices futures, la formule fait également intervenir les grandeurs suivantes :

- le montant global des provisions techniques correspondant à la participation aux bénéfices discrétionnaire future FDB ;
- l'effet d'atténuation par la PB future du risque de souscription vie KC_{life} ;
- l'effet d'atténuation par la PB future du risque de marché KC_{mkt} ;
- l'effet d'atténuation par la PB future du risque santé KC_{health} .

La formule d'agrégation est la suivante :

$$BSCR = \sqrt{\sum_{r \times c} CorrSCR_{r,c} \times SCR_r \times SCR_c} - \min \left\{ \sqrt{\sum_{r \times c} CorrSCR_{r,c} \times KC_r \times KC_c}, FDB \right\}$$

où la matrice de corrélation $CorrSCR$ est donnée par le tableau suivant :

$CorrSCR =$	SCR_{mkt}	SCR_{def}	SCR_{life}	SCR_{health}	SCR_{nl}
SCR_{mkt}	1				
SCR_{def}	0.25	1			
SCR_{life}	0.25	0.25	1		
SCR_{health}	0.25	0.25	0.25	1	
SCR_{nl}	0.25	0.5	0	0	1

Il faut noter qu'il s'agit d'une matrice de corrélation entre les capitaux associés aux risques et non les risques eux-mêmes.

5.2.1. Risques de marché

Le risque de marché est une des principales composantes du basic SCR. Le capital qui lui est associé résulte également de l'agrégation des capitaux obtenus pour les sous-risques que sont les risques de taux, du marché action, de l'immobilier, d'inflation et de change. Ainsi le SCR_{mkt} s'obtient à partir de la formule d'agrégation suivante :

$$SCR_{mkt} = \sqrt{\sum_{r \times c} CorrMkt_{r,c} \times Mkt_r \times Mkt_c},$$

où la matrice de corrélation $CorrSCR$ est donnée par le tableau suivant :

$CorrMkt$	Mkt_{int}	Mkt_{eq}	Mkt_{prop}	Mkt_{sp}	Mkt_{conc}	Mkt_{fx}
Mkt_{int}	1					
Mkt_{eq}	0.5	1				
Mkt_{prop}	0.5	0.75	1			
Mkt_{sp}	0.25	0.25	0.25	1		
Mkt_{conc}	0	0	0	0	1	
Mkt_{fx}	0.25	0.25	0.25	0.25	0	1

Par ailleurs, pour chacun de ces risques, la propension des participations aux bénéfices futures à amortir des variations adverses des grandeurs financière a été quantifiée par les facteurs KC_{eq} , KC_{int} , KC_{prop} et KC_{fx} . Ces grandeurs sont agrégées en utilisant la même formule d'agrégation que précédemment :

$$KC_{mkt} = \sqrt{\sum_{r \times c} CorrMkt_{r,c} \times KC_r \times KC_c}.$$

5.2.2. Exemple : risque de taux

Le risque de taux est un des principaux risques auxquels sont soumis les assureurs du fait de la composition de leur actif essentiellement obligataire et des garanties de taux qu'ils peuvent proposer à leurs clients.

La formule standard de calcul du SCR prévoit que le capital élémentaire venant en contrepartie de ce risque est obtenu comme étant le niveau de capital dont doit disposer l'assureur pour ne pas être en ruine dans un scénario de forte hausse ou de forte baisse des taux d'intérêt.

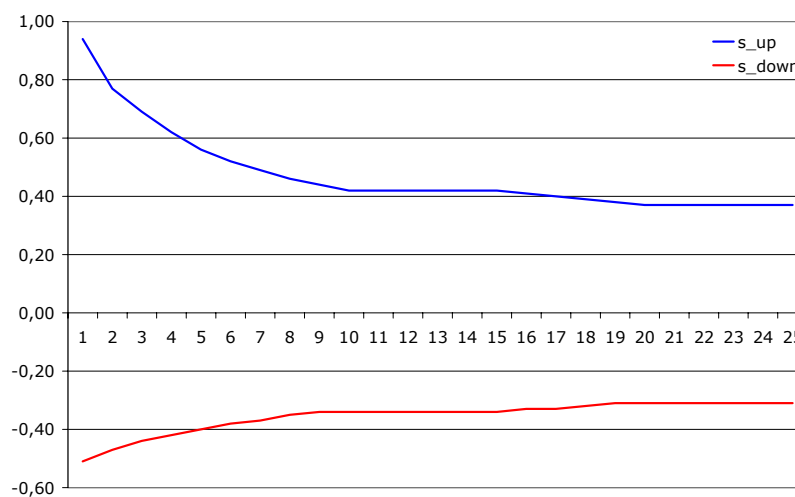
QIS 3 prévoit les déformations suivantes de la courbe des taux :

Maturity t (years)	1	2	3	4	5	6	7
relative change $s^{up}(t)$	0,94	0,77	0,69	0,62	0,56	0,52	0,49
relative change $s^{down}(t)$	-0,51	-0,47	-0,44	-0,42	-0,40	-0,38	-0,37

Maturity t (years)	8	9	10	11	12	13	14
relative change $s^{up}(t)$	0,46	0,44	0,42	0,42	0,42	0,42	0,42
relative change $s^{down}(t)$	-0,35	-0,34	-0,34	-0,34	-0,34	-0,34	-0,34

Maturity t (years)	15	16	17	18	19	20+
relative change $s^{up}(t)$	0,42	0,41	0,40	0,39	0,38	0,37
relative change $s^{down}(t)$	-0,34	-0,33	-0,33	-0,32	-0,31	-0,31

Figure 46 - Coefficients de déformation de la courbe des taux



Les courbes des taux déformées sont obtenues en multipliant la courbe des taux actuelle par les coefficients $(1 + s^{up})$ ou $(1 + s^{down})$ qui évoluent en fonction de la maturité considérée.

Sur la base de ces courbes de taux d'intérêt, il s'agit de projeter à nouveau les flux futurs de l'assureur et d'observer la différence entre les valeurs nettes de l'actif et du passif dans la situation de référence, d'une part, et dans les scénarios avec courbe des taux modifiée, d'autre part. Cette différence est

déterminée dans les scénarios de majoration et d'abattement de la courbe des taux et c'est le résultat le plus élevé qui est retenu comme Mkt_{int} .

Un deuxième calcul est ensuite effectué pour mesurer la capacité des clauses de participation aux bénéfices futures à absorber des chocs (cf. la formule du *basic SCR*). Il s'agit, en conservant, les revalorisations futures obtenues dans le scénario de référence de déterminer la valeur nette des actifs et des passifs dans les deux scénarios de chocs de la courbe des taux puis de comparer les résultats obtenus à ceux des projections dans lesquelles les revalorisations futures ont été adaptées aux modifications opérées sur la courbe des taux. La différence entre ces résultats permet d'obtenir KC_{int} qui représente une sorte de matelas de sécurité à la disposition de l'assureur pour amortir des variations adverses de la courbe des taux. C'est pour cette raison que cette grandeur vient, après agrégation, en diminution dans le calcul du BSCR.

BIBLIOGRAPHIE GÉNÉRALE

- AAI (2004) *A global framework for insurer solvency assessment*, <http://www.actuaires.org>.
- Artzner Ph., Delbaen F., Eber J.M., Heath D. (1999) « Coherent measures of risk », *Mathematical Finance* 9, 203-28.
- Augros J.C., Moreno M. (2002) *Les dérivés financiers et d'assurance*, Paris : Economica.
- Ballotta L. (2005) « A Lévy process-based framework for the fair valuation of participating life insurance contracts », *Insurance: Mathematics and Economics* 37 (2), 173-96.
- Bergonzat A., Cavals F. (2006) *Garanties en cas de vie sur contrats multisupports*, Mémoire d'actuariat, ENSAE.
- Black F., Scholes M. (1973) « The pricing of options and corporate liabilities », *Journal of Political Economy* 81 (3), 637-54.
- Blum K. A., Otto D. J. (1998) « Best estimate loss reserving : an actuarial perspective », *CAS Forum Fall* 1, 55-101.
- Bottard S. (1996) « Application de la méthode du Bootstrap pour l'estimation des valeurs extrêmes dans les distributions de l'intensité des séismes », *Revue de statistique appliquée* 44 (4), 5-17.
- Bühlmann H., Straub E. (1970) « Glaubwürdigkeit für Schadensätze », *Mitteilungen der Vereinigung Schweizerischer Versicherungsmathematiker* 70.
- Bühlmann H., Gisler A. (2005) *A course in credibility theory*, Berlin : Springer.
- Caperaa P., Fougères A.L., Genest C. (1997) « A nonparametric estimation procedure for bivariate extreme value copulas », *Biometrika* 84, 567-7.
- Christoffersen P., Hahn J., Inoue A. (2001) « Testing and comparing value-at-risk measures », *Journal of Empirical Finance* 8 (3), 325-42.
- Coles S., Powell E. (1996) « Bayesian methods in extreme value modelling: a review and new developments », *Internat. Statist. Rev.* 64, 119-36.
- Commission européenne (2003) « Conception d'un futur système de contrôle prudentiel applicable dans l'Union européenne – Recommandation des services de la Commission », MARKT/2509/03.
- Commission européenne (2004) « Solvency II – Organisation of work, discussion on pillar I work areas and suggestions of further work on pillar II for CEIOPS », MARKT/2543/03.

- Cox J.C., Ingersoll J.E., Ross S.A. (1985) « A theory of the term structure of interest rates », *Econometrica* 53, 385-407.
- Darsow W., Nguyen B., Olsen E. (1992) « Copulas and Markov Processes », *Illinois Journal of Mathematics* 36 (4), 600-42.
- Davidson R., MacKinnon J.G. (2004) « Bootstrap Methods in Econometrics », working paper.
- de Haan L., Peng L. (1998) « Comparison of tail index estimators », *Statistica Neerlandica* 52 (1), 60-70.
- Deelstra G., Janssen J. (1998) « Interaction between asset liability management and risk theory », *Applied Stochastic Models and Data Analysis* 14, 295-307.
- Deheuvels P. (1978) « Caractérisation complète des lois extrêmes multivariées et de la convergence des types extrêmes », *Publications de l'Institut de Statistique de l'Université de Paris* 23, 1-36
- Deheuvels P. (1979a) « Propriétés d'existence et propriétés topologiques des fonctions de dépendance avec applications à la convergence des types pour des lois multivariées », *C. R. Académie des Sciences de Paris, Série 1*, 288, 145-8
- Deheuvels P. (1979b) « La fonction de dépendance empirique et ses propriétés. Un test non paramétrique d'indépendance », *Académie Royale de Belgique – Bulletin de la Classe des Sciences – 5^e Série*, 65, 274-92
- Dekkers A., de Haan L. (1989) « On the estimation of the extreme-value index and large quantile estimation », *Annals of Statistics* 17, 1795-832.
- Dekkers A., Einmahl J., de Haan L. (1989) « A moment estimator for the index of an extreme-value distribution », *Annals of Statistics* 17, 1833-55.
- Denuit M., Charpentier A. (2004) *Mathématiques de l'assurance non-vie. Tome 1 : principes fondamentaux de théorie du risque*, Paris : Economica.
- Denuit M., Charpentier A. (2005) *Mathématiques de l'assurance non-vie. Tome 2 : tarification et provisionnement*, Paris : Economica.
- Dhaene J., Laeven R.J.A., Vanduffel S., Darkiewicz G., Goovaerts M.J. (2004) « Can a coherent risk measure be too subadditive? », Research Report OR 0431, Department of Applied Economics, Katholieke Universiteit Leuven.
- Dhaene J., Vanduffel S., Tang Q.H., Goovaerts M., Kaas R., Vyncke D. (2004) « Solvency capital, risk measures and comonotonicity: a review », Research Report OR 0416, Department of Applied Economics, K.U.Leuven.
- Dhaene J., Vanduffel S., Goovaerts M., Kaas R., Vyncke D. (2005) « Comonotonic approximations for optimal portfolio selection problems », *Journal of Risk and Insurance* 72 (2).

- Dhaene J., Wang S., Young V., Goovaerts M. (2000) « Comonotonicity and maximal Stop-Loss premiums », *Bulletin of the Swiss Association of Actuaries* 2, 99-113.
- Diebolt J., El-Aroui M., Garrido S., Girard S. (2005a) « Quasi-conjugate bayes estimates for gpd parameters and application to heavy tails modelling », *Extremes* 8, 57-78.
- Diebolt J., Guillou A., Rached I. (2005b) « Approximation of the distribution of excesses using a generalized probability weighted moment method », *C. R. Acad. Sci. Paris, Ser. I* 340 (5), 383-8.
- Djehiche B., Hörfelt P. (2005) « Standard approaches to asset & liability risk », *Scandinavian Actuarial Journal* 2005, n° 5, 377-400.
- Efron B. (1979) « Bootstrap methods: Another look at the Jackknife », *Ann. Statist.*, 1-26.
- Efron B. (1987) « Better bootstrap confidence intervals », *Journal of the American Statistical Association* 82, 171-200.
- Efron B., Tibshirani R.J. (1993) *An introduction to the bootstrap*, Chapman & Hall.
- Embrechts P. (2000) « Actuarial versus financial pricing of insurance », *Risk Finance* 1 (4), 17-26.
- Embrechts P., Furrer H., Kaufmann R. (2007) « Different Kinds of Risk » in *Handbook of Financial Time Series*, Eds. Andersen, Davis, Kreiss, and Mikosch.
- Embrechts P., Klüppelberg C., Mikosch T. (1997) *Modelling Extremal Events for Insurance and Finance*, Springer Verlag, Berlin.
- Embrechts P., McNeil A., Straumann D. (1999) « Correlation: Pitfalls and alternatives », *RISK Magazine*, May, 69-71.
- England P.D., Verall R.J. (1999) « Analytic and bootstrap estimates of prediction errors in claim reserving », *Insurance: mathematics and economics* 25, 281-93.
- Ewald F., Lorenzi J.H. (eds) (1998) *Encyclopédie de l'assurance*, Paris : Economica.
- Fargeon L., Nissan K. (2003) *Recherche d'un modèle actuariel d'analyse dynamique de la solvabilité d'un portefeuille de rentes viagères*, Mémoire d'actuariat, ENSAE.
- Fedor M., Morel J. (2006) « Value-at-risk en assurance : recherche d'une méthodologie à long terme », *Actes du 28e congrès international des actuaires*, Paris.
- Frantz C., Chenut X., Walhin J.F. (2003) « Pricing and capital allocation for unit-linked life insurance contracts with minimum death guarantee », *Proceedings of the 13th AFIR Colloquium*, Maastricht.

- Genest C., MacKay R. J. (1986) « The joy of copulas: Bivariate distributions with uniform marginals », *The American Statistician* 40, 280-3.
- Giet L. (2003) « Estimation par inférence indirecte des équations de diffusion : l'impact du choix du procédé de discrétisation », Document de travail n°03A15 du GREQAM.
- Goldie C., Smith R. (1987) « Slow variation with remainder : a survey of the theory and its applications », *Quarterly Journal of Mathematics Oxford* 38 (2), 45-71.
- Gouriéroux Ch. (1999) *Statistique de l'assurance*, Paris : Economica.
- Hami S. (2003) *Les modèles DFA : présentation, utilité et application*, Mémoire d'actuariat, ISFA.
- Hardy M.R., Panjer H.H. (1998) « A credibility approach to mortality risk », *ASTIN Bulletin* 28 (2), 269-83.
- Hill B. (1975) « A simple general approach to inference about the tail of a distribution », *Annals of Statistics* 3, 1163-74.
- Horowitz J.L. (1997) « Bootstrap methods in econometrics: theory and numerical performance », in *Advances in Economics and Econometrics: Theory and Application*, volume 3, 188-222, D. M. Kreps, K. F. Wallis (eds), Cambridge: Cambridge University Press.
- Hosking, J.R., Wallis, J.R. (1987) « Parameter and quantile estimation for the generalized pareto distribution », *Technometrics* 29, 339-49.
- Hull J.C. (1999) *Options, futures and other derivatives*, 4th edition, Prentice-Hall.
- Hürlimann W. (1999) « On risk and price : Stochastic orderings and measures », *Actes du 27^e colloque ASTIN*, Cancun, 22 mars 2002.
- Jal P., Partrat Ch. (2004) « Evaluation stochastique de la provision pour sinistres », Conférence scientifique de l'Institut des Actuaire, Paris, 20 janvier 2004.
- Joe H. (1997) « Multivariate models and dependence concepts », *Monographs on Statistics and Applied probability* 73, London : Chapman & Hall.
- Kass R., Dhaene J., Goovaerts M. (2000) « Upper and lower bounds for sums of random variables », *Insurance: Mathematics and Economics* 27, 151-68.
- Kaufmann R., Gadmer A., Klett R. (2001) « Introduction to Dynamic Financial Analysis », *ASTIN Bulletin* 31 (1), 213-49.
- Kloeden P., Platen E. (1995) *Numerical Solution of Stochastic Differential Equations*, 2nd Edition, Springer-Verlag.
- Lambert A. (1998) « Assurons l'avenir de l'assurance », *Rapport d'information du Sénat* n° 45 (98-99), tome II – commission des Finances.
- Lamberton D., Lapeyre B. (1997) *Introduction au calcul stochastique appliqué à la finance*, 2^e édition, Paris : Ellipses.

- Louis T., Shi J. (1995) « Inferences on the association parameter in copula models for bivariate survival data », *Biometrics* 51, 1384-99.
- Merlus S., Pequeux O. (2000) *Les garanties plancher des contrats d'assurance vie en UC: tarification et couverture*, Mémoire d'actuariat, ENSAE.
- Merton R.C. (1976) « Option pricing when underlying stock returns are discontinuous », *Journal of Financial Economics* 3, 125-44.
- Meyers G.G., Klinker F.L., Lalonde F.A. (2003) « The aggregation and correlation of insurance exposure », *CAS Forum*, 16-82.
- Nelsen R.B. (1999) *An introduction to Copulas*, Lecture Notes in Statistics 139, New-York: Springer Verlag.
- Partrat Ch., Besson J.L. (2005) *Assurance non-vie. Modélisation, simulation*. Paris : Economica.
- Pickands J. (1975) « Statistical inference using extreme orders statistics », *Annals of Statistics* 3, 119-31.
- Planchet F., Thérond P.E. (2003) « Évaluation de l'engagement de l'entreprise associé à un plan de stock-options », *Bulletin Français d'Actuariat* 6 (11), 149-66.
- Planchet F., Thérond P.E. (2004a) « Allocation d'actifs d'un régime de rentes en cours de service », *Proceedings of the 14th AFIR Colloquium*, Boston, 111-34.
- Planchet F., Thérond P.E. (2004b) « Les principes de valorisation des engagements sociaux », *Revue fiduciaire comptable*, n° 310, 18-25.
- Planchet F., Thérond P.E. (2005a) « Asset allocation : new constraints induced by the Solvency 2 project », *Proceedings of the 36th ASTIN Colloquium*, Zürich.
- Planchet F., Thérond P.E. (2005b) « L'impact de la prise en compte des sauts boursiers dans les problématiques d'assurance », *Proceedings of the 15th AFIR Colloquium*, Zürich.
- Planchet F., Thérond P.E. (2005c) « Simulation de trajectoires de processus continus », *Belgian Actuarial Bulletin* 5, 1-13.
- Planchet F., Thérond P.E. (2006) *Modèles de durée. Applications actuarielles*, Paris : Economica.
- Planchet F., Thérond P.E. (2007a) *Pilotage technique d'un régime de rentes viagères*, Paris : Economica.
- Planchet F., Thérond P.E. (2007b) « Allocation d'actifs selon le critère de maximisation des fonds propres économiques en assurance non-vie : présentation et mise en oeuvre dans la réglementation française et dans un référentiel de type Solvabilité 2 », *Bulletin Français d'Actuariat* 7 (13), 10-38.

- Planchet F., Thérond P.E., Jacquemin J. (2005) *Modèles financiers en assurance. Analyses de risque dynamiques*, Paris : Economica.
- Quenouille M. (1949) « Approximate tests of correlation in time series », *Journal of the Royal Statistical Society, Soc. Series B*, 11, 18-84.
- Ramezani C.A., Zeng Y. (1998) « Maximum likelihood estimation of asymmetric jump-diffusion processes: application to security prices », Working paper.
- Roncalli T. (1998) *La structure par terme des taux zéro : modélisation et implémentation numérique*, Thèse de l'Université Bordeaux IV.
- Saporta G. (1990) *Probabilités, analyse des données et statistique*, Paris : Technip.
- Schweizer B., Sklar A. (1958) « Espaces métriques aléatoires », *C. R. Acad. Sci. Paris* 247, 2092-4.
- Schweizer B. (1991) « Thirty years of copula, Advances in probability distributions with given marginals: beyond the copulas », ed. by G. Dall'Aglio, S. Kotz and G. Salinetti, *Mathematics and its applications* v. 67, Kluwer Academic Publishers, 13-50.
- Sklar A. (1959) « Fonctions de répartition à n dimensions et leurs marges », *Publications de l'Institut de Statistique de l'Université de Paris* 8, 229-31.
- Smith R.L. (1987) « Estimating tails of probability distributions », *Annals of Statistics* 15, 1174-207.
- Thérond P.E. (2003) *Impact des futures normes IFRS sur la tarification et le provisionnement des contrats d'assurance vie : mise en oeuvre de méthodes par simulation*, Mémoire d'actuariat, ISFA.
- Thérond P.E. (2005) « Contrôle de la solvabilité des compagnies d'assurance : évolutions récentes », Séminaire *Gestion des risques et assurance*, Hanoï, 28 février 2005.
- Thérond P.E., Bonche S. (2006) « Gestion du niveau de la franchise d'un contrat avec bonus-malus », *Actes du 28e congrès international des actuaires*, Paris.
- Thérond P.E., Girier G. (2005) « Solvabilité 2 en construction », *L'Argus de l'assurance*, Hors-série de décembre.
- Thérond P.E., Planchet F. (2007) « Provisions techniques et capital de solvabilité d'une compagnie d'assurance : méthodologie d'utilisation de Value-at-Risk », *Assurances et gestion des risques* 74 (4), 533-63 et in *Proceedings of the 37th ASTIN Colloquium*, Orlando.
- Vasicek O. (1977) « An equilibrium characterization of the term structure », *Journal of financial Economics* 5, 177-88.
- Wang S. (2002) « A risk measure that goes beyond coherence », *Actes du 12^e colloque AFIR*, Cancun, 18 mars 2002.

-
- Windcliff, H., Boyle, P. (2004) « The 1/n pension investment puzzle », *North American Actuarial Journal* 8 (3), 32-45.
- Zajdenweber D. (2000) *Économie des extrêmes*, Paris : Flammarion.
- Zeltermann D. (1993) « A semiparametric Bootstrap technique for simulating extreme order statistics », *Journal of American Statistical Association* 88 (422), 477-85.

Notations utilisées

Le tableau ci-dessous reprend l'essentiel des notations utilisées tout au long de la thèse.

Symbole	Description
$a(t) \propto b(t)$	Les fonctions a et b sont proportionnelles (leur rapport est indépendant de la variable t)
$a \wedge b$	Minimum de a et b
$a \vee b$	Maximum de a et b
$Y \prec_{st} X$	X domine stochastiquement Y
$Y \prec_{cx} X$	X domine Y selon l'ordre convexe
$Y \prec_{icx} X$	X domine Y selon l'ordre convexe croissant
$X \stackrel{loi}{=} Y$	X et Y ont la même distribution de probabilité
$E[X]$	Espérance de X
$F_X(x)$	Fonction de répartition de X prise au point x
$f(t)$	Densité prise au point t
$\mathbf{F}(F_1, F_2)$	Classe de Fréchet associée aux fonctions de répartition F_1 et F_2
$h(t)$	Fonction de hasard prise au point t
$H(t)$	Fonction de hasard cumulée prise au point t
$L(\theta, \mathbf{x})$	Vraisemblance de l'échantillon $\mathbf{x}$ et du paramètre θ
$L_X(t)$	Transformée de Laplace de X prise au point t
$\mathcal{LN}(\mu, \sigma)$	Loi log-normale issue d'une gaussienne $\mathcal{N}(\mu, \sigma)$
$\mathcal{N}(\mu, \sigma)$	Loi gaussienne de moyenne μ et d'écart-type σ

$\mathcal{P}(\lambda)$	Loi de Poisson de paramètre λ
P	Probabilité historique
Q	Probabilité risque-neutre
$S(t) = \bar{F}(t)$	Survie au delà de l'instant t
$S_u(t)$	Survie au delà de l'instant t , sachant que l'on est en vie en u
${}^t M$	Transposée de la matrice M
$\text{Var}[X]$	Variance de X
$\text{VaR}_q(X) = F_X^{-1}(q)$	Quantile (ou Value-at-Risk) d'ordre q de la loi de X
$X_{1,n} = X_{(n)}$	Maximum de $X_1, \dots, X_n$
$X_{k,n} = X_{(n-k+1)}$	k -ième plus grande valeur de $X_1, \dots, X_n$
$X_{n,n} = X_{(1)}$	Minimum de $X_1, \dots, X_n$
Φ	Fonction de répartition de la loi normale centrée réduite
$[x]$	Partie entière de x
$\rightarrow_d$	Convergence en loi
$\rightarrow_{\text{Pr}}$	Convergence en probabilité

Table des illustrations

Figure 1 - Value-at-Risk de la somme de deux v.a. de Pareto	17
Figure 2 - Quelques fonctions de distorsion	20
Figure 3 - Évolution des bornes de l'intervalle de confiance à 90 % en fonction du nombre de polices assurées.....	31
Figure 4 - Distribution des coûts futurs en absence de gestion active de la couverture	41
Figure 5 - Distribution des coûts futurs lorsque la gestion active de la couverture est mise en place	42
Figure 6 - Évolution du rendement moyen de l'actif	61
Figure 7 - Évolution de l'écart-type du rendement.....	61
Figure 8 - Distribution empirique de la revalorisation l'épargne au terme.....	62
Figure 9 - Évolution du taux de revalorisation moyen	63
Figure 10 - Évolution du coefficient de variation du taux de revalorisation.....	63
Figure 11 - Densité empirique de la valeur actuelle des prestations futures.....	68
Figure 12 - Densité empirique des prestations futures actualisées	69
Figure 13 - Densité de la copule de Franck de paramètre 1	79
Figure 14 - Distribution de la charge totale de sinistres	80
Figure 15 - Provision économique exprimée en pourcentage de la provision réglementaire en fonction de ω_1	84
Figure 16 - Probabilité de ruine en fonction de ω_1	85
Figure 17 - Capital cible en fonction de la part investie en actions ω_1	87
Figure 18 - Graphe de ϕ	88
Figure 19 - Capital cible en fonction de la part investie en actions ω_1	89
Figure 20 - Typologie des différents risques rencontrés.....	101
Figure 21 - Estimation d'un quantile extrême : erreur relative d'estimation ...	105
Figure 22 - Méthode bootstrap en provisionnement non-vie	108
Figure 23 - Intervalles de confiance (VaR à 99,5 %).....	111
Figure 24 - Intervalles de confiance (VaR à 75 %).....	111
Figure 25 - Modèle interne simplifié : identification des valeurs extrêmes....	114
Figure 26 - Modèle interne simplifié : identification des valeurs extrêmes (probabilités).....	114
Figure 27 - Rendement journalier du titre TOTAL : QQ-plot loi empirique vs loi normale	115
Figure 28 - Rendement journalier du titre TOTAL : QQ-plot loi empirique vs modèle de Merton	117
Figure 29 - Estimation de l'épaisseur de la queue d'une distribution de Pareto	128
Figure 30 - Fonction de répartition de la copule cubique	147
Figure 31 - Évolution moyenne du taux modélisé par Vasicek selon le procédé de discrétisation retenu.....	178
Figure 32 - Modèle de Vasicek selon le procédé de discrétisation retenu	178

Figure 33 - Prix des zéro-coupons en fonction de leur échéance et de la méthode d'estimation des paramètres	184
Figure 34 - Taux instantanés espérés en fonction de leur échéance et de la technique d'estimation des paramètres	184
Figure 35 - Répartition et tests statistiques du générateur Rnd.....	186
Figure 36 - Répartition et tests statistiques du générateur du tore	187
Figure 37 - Erreur d'estimation en fonction du nombre de simulations	188
Figure 38 - Trajectoire moyenne simulée par le tore d'un mouvement brownien géométrique	189
Figure 39 - Corrélogramme de la suite générée par le tore ($p = 5$).....	190
Figure 40 - Corrélogramme de la suite générée par le tore ($p_1 = 5$) mélangé par le tore ($p_2 = 19$)	192
Figure 41 - Corrélogramme de la suite générée par le tore ($p_1 = 5$) mélangé par Rnd	192
Figure 42 - Trajectoires espérée et moyenne des 5 000 trajectoires simulées par le tore mélangé d'un mouvement brownien géométrique	193
Figure 43 - Distribution de l'âge au décès pour un assuré de 50 ans.....	196
Figure 44 - Étapes de calcul de la marge pour risque (méthode CoC).....	209
Figure 45 - Structure modulaire de la formule standard	211
Figure 46 - Coefficients de déformation de la courbe des taux	214

Table des matières

Sommaire.....	1
INTRODUCTION GÉNÉRALE.....	3
PARTIE I NOUVELLES APPROCHES COMPTABLE, PRUDENTIELLE ET FINANCIÈRE DES RISQUES EN ASSURANCE.....	9
Chapitre 1 Traitement spécifique du risque : aspects théoriques	13
1. Les outils mathématiques de l'analyse des risques	13
1.1. Les mesures de risque.....	13
1.2. Comparaison des risques	26
2. Un traitement différencié du risque	29
2.1. L'approche assurantielle.....	29
2.2. L'approche économique	31
2.3. Confrontation des deux approches en assurance.....	34
3. Contrats en unités de compte : les garanties plancher.....	36
3.1. Contrats en unités de compte	36
3.2. Principaux risques.....	36
3.3. Valorisation et gestion des garanties plancher.....	38
4. Conclusion	42
Bibliographie	44
Chapitre 2 Traitement spécifique du risque : aspects pratiques	47
1. De nouveaux référentiels distincts	47
1.1. Contexte.....	47
1.2. Des finalités différentes	49
1.3. Un socle de principes communs	50
2. Des incidences opérationnelles.....	52
2.1. Hypothèses actuarielles	52
2.2. Futures décisions de gestion	57
3. Cas pratique : portefeuille d'assurance vie	64
3.1. Présentation du produit.....	64
3.2. Le portefeuille des assurés.....	67
3.3. Modélisation retenue	67
3.4. Analyse des résultats.....	68
3.5. Conclusion.....	69
Bibliographie	70
Chapitre 3 Incidence sur la gestion technique d'un assureur.....	71
1. Modélisation de la société d'assurance	72

2. Critère de maximisation des fonds propres économiques	73
2.1. Présentation générale	73
2.2. Mise en œuvre dans le cadre réglementaire français	75
2.3. Mise en œuvre dans le référentiel Solvabilité 2	76
3. Recherche de l'allocation optimale	77
3.1. Modélisation des risques	78
3.2. MFPE dans la réglementation française	82
3.3. MFPE dans un référentiel du type Solvabilité 2	85
4. Conclusion	88
Annexe A : Démonstration des résultats mathématiques	91
Annexe B : Simulation des réalisations de la charge de sinistres	93
Bibliographie	94

PARTIE II MODÉLISATIONS AVANCÉES EN ASSURANCE 95

Chapitre 4 Limites opérationnelles : la prise en compte des extrêmes 97

1. Calcul de VaR en assurance	99
2. Notations	101
3. Estimation de quantiles extrêmes	102
3.1. Estimation naturelle	102
3.2. Ajustement à une loi paramétrique	102
3.3. Approximation GPD	103
3.4. Estimateur de Hill	104
3.5. Illustration	105
4. Application du bootstrap	105
4.1. Présentation	106
4.2. Calcul d'un intervalle de confiance pour une VaR	108
4.3. Illustration numérique	110
5. Robustesse du SCR	111
5.1. Estimation des paramètres des variables de base	112
5.2. Simulation	112
5.3. Spécification du modèle	113
6. Conclusion	117
Annexe A : Loi de Pareto généralisée (GPD)	118
A.1 Définition	118
A.2 Quelques propriétés	118
A.3 Estimation des paramètres	119
Annexe B : Résultats probabilistes	121
B.1 Loi du maximum	121
B.2 Épaisseur des queues de distribution	122
B.3 Loi des excès au-delà d'un seuil	123
Annexe C : Estimation du paramètre de queue	125
C.1 Méthodes paramétriques	125
C.2 Méthodes semi-paramétriques	126
Bibliographie	129

Chapitre 5 Prise en compte de la dépendance.....	133
1. Analyse mathématique de la dépendance	134
1.1. Mesures de dépendance	134
1.2. Structures de dépendance.....	138
2. Rappels sur la dépendance linéaire	143
3. La théorie des copules.....	145
3.1. Copules, lois conjointes et dépendance	145
3.2. Inférence statistique	152
3.3. Copules et valeurs extrêmes	155
3.4. Copules et simulation	156
3.5. La méthode NORTA.....	159
3.6. Exemples d'application des copules	160
4. Rachat de contrats d'épargne : un modèle ad hoc.....	161
4.1. Notations.....	161
4.2. Hypothèses simplificatrices	162
4.3. Attribution de PB.....	162
4.4. Rachat.....	165
5. Capital de solvabilité : les méthodes d'agrégation.....	165
5.1. Un cadre gaussien	166
5.2. Les limites des modèles retenus.....	167
Bibliographie	168
Chapitre 6 Techniques de simulation.....	171
1. Introduction.....	171
2. Discrétisation de processus continus.....	172
2.1. Discrétisation exacte	173
2.2. Discrétisation approximative	175
3. Estimation des paramètres	180
3.1. Le biais associé à la procédure d'estimation.....	180
3.2. Principe de l'estimation par inférence indirecte	181
3.3. Les méthodes d'estimation ad hoc.....	182
3.4. Illustration dans le cas du modèle de Vasicek	182
4. Génération des trajectoires.....	185
4.1. Deux générateurs de nombres aléatoires : Rnd et le tore	185
4.2. Limites de l'algorithme du tore.....	188
4.3. Générateur du tore mélangé.....	191
5. Simulation de la mortalité d'un portefeuille d'assurés.....	193
5.1. Simulation de l'âge au décès	194
5.2. Détermination de la date du décès	196
6. Conclusion	197
Annexe : Tests d'adéquation à une loi	198
1. Test du χ^2	198
2. Test de Kolmogorov-Smirnov.....	198
3. Test d'Anderson-Darling	198
Bibliographie	200

CONCLUSION GÉNÉRALE	201
Annexe Solvabilité 2 - les modèles proposés par QIS 3	205
1. Modèles d'évaluation : l'approche standard	206
1.1. Les actifs	206
1.2. Les provisions techniques	207
2. Provisions techniques en assurance vie.....	209
3. Provisions techniques en assurance non-vie	210
4. Capital éligible.....	210
5. Capital de solvabilité (SCR) : formule standard	211
5.1. Approche modulaire.....	211
5.2. Basic SCR.....	212
BIBLIOGRAPHIE GÉNÉRALE.....	217
Notations utilisées	225
Table des illustrations	227
Table des matières	229

Mesure et gestion des risques d'assurance : analyse critique des futurs référentiels prudentiel et d'information financière

Résumé

L'avènement des nouveaux référentiels comptable (IFRS), prudentiel (Solvabilité 2) et de communication financière (EEV/MCEV) va modifier profondément le rapport au risque des assureurs au travers de nouvelles exigences de valorisation des contrats d'assurance et de calcul des exigences de capitaux propres. L'objectif de cette thèse est de présenter ces changements, d'identifier leurs limites et d'analyser leurs conséquences sur la gestion effective d'un assureur.

La première partie est consacrée à la présentation et à l'analyse des principes sur lesquels repose le traitement des différents risques requis par les différents référentiels. Une attention particulière est portée sur leur mise en œuvre pratique. Enfin, une illustration des incidences de ces nouveaux référentiels sur la gestion technique d'un assureur est proposée.

La seconde partie porte sur les aspects opérationnels de la mise en œuvre des référentiels. Une attention particulière est portée à la prise en compte des valeurs extrêmes sur laquelle repose le calcul des exigences de capitaux propres, sur la prise en compte de la dépendance entre les risques et revient sur les méthodes de simulations rendues incontournables par les approches retenues.

Mots-clés : Risque, assurance, solvabilité, *International Financial Reporting Standards* (IFRS), *Market Consistent Embedded Value* (MCEV), valeurs extrêmes, simulations de Monte Carlo, garanties de taux, participation aux bénéfices, option de rachat, allocation d'actifs, capital de solvabilité.

Risk measure and management in insurance : critical analysis of the future prudential and financial reporting frameworks

Abstract

The advent of the new accounting (IFRS), prudential (Solvency 2) and financial reporting (EEV/MCEV) frameworks will deeply change the risk treatment of the insurers. Indeed, they are facing new requirement for valuating contracts and for calculating solvency capital. This thesis objective is to present these changes, to identify their limits and to analyze their consequences on the effective management of the insurance company.

The first part is devoted to the presentation and the analysis of the principles of the new frameworks. A close attention is paid to their practical implementation. Lastly, an illustration of the incidences of these new frames on the technical management of an insurer is proposed.

The second part relates to the operational aspects of these new frameworks. A detailed attention is paid to the extremes on which is based the calculation of the solvency capital requirement, to the dependence between the risks and to the Monte Carlo methods which become uncountournable within the new standards.

Keywords : risk, insurance, solvency, International Financial Reporting Standards (IFRS), Market Consistent Embedded Value (MCEV), extremes, Monte Carlo simulations, guaranteed life insurance, participating contracts, surrender option, asset allocation, solvency capital.