



**Mémoire présenté devant le jury de l'EURIA en vue de l'obtention du
Diplôme d'Actuaire EURIA
et de l'admission à l'Institut des Actuaire**

le 2 septembre 2013

Par : Ludovic POIRAUD

Titre : Construction et validation d'une table de mortalité d'expérience par des méthodes non paramétriques

Confidentialité : non.

***Membre présent du jury de l'Institut
des Actuaire***

Signature :

Entreprise

Signature :

Membres présents du jury de l'EURIA

Directeur de mémoire en entreprise

Signature :

Invité

Signature :

***Autorisation de publication et de mise en ligne sur un site de diffusion
de documents actuariels***

Signature du responsable entreprise

Signature du candidat

Bibliothèque :

Secrétariat :

Remerciements

Je souhaite remercier tout particulièrement Frédéric PLANCHET, associé chez Prim'Act, et Aymric KAMEGA, directeur de mission chez Prim'Act, pour m'avoir encadré sur ce mémoire, pour leur disponibilité, leurs conseils avisés et le suivi régulier qu'ils ont effectué au cours de ces 6 mois de stage.

Je souhaite également remercier Gilles DEPOMMIER, actuaire associé chez Prim'Act, qui, avec sa qualité de certificateur, m'a permis de réaliser de nombreuses missions sur le thème de la certification de tables, en faisant partager son savoir et son expérience dans ce domaine.

Je n'oublierai pas Emmanuel AVRIL, actuaire consultant chez Prim'Act, qui était toujours disponible pour répondre à mes interrogations diverses et des échanges enrichissants que nous avons partagés, ainsi que toute l'équipe de Prim'Act de m'avoir intégrée au sein de la société, avec une pensée toute particulière à Geoffroy DELION.

Je remercie également l'ensemble de mon entourage, pour m'avoir soutenu en toutes circonstances, du début de mes études à la fin de ce stage.

Enfin, je souhaite remercier l'ensemble du corps professoral de l'Euro Institut d'Actuariat de Brest pour la qualité de leur enseignement dispensé durant ces trois dernières années.

Résumé

MOTS-CLÉS : Table de mortalité prospective, risque d'estimation, risque de tendance, modèle relationnel, méthodes non-paramétriques, outils de validation.
--

Dans l'optique de la mise en place de Solvabilité II, les assureurs sont tenus d'évaluer leurs risques et leurs engagements de la manière la plus précise possible. Dans le domaine de l'assurance-vie, la qualité des tables de mortalité est un paramètre incontournable et elle ne doit pas être remise en question. Pour cela et depuis 1993, les compagnies d'assurance ont la possibilité de construire leurs propres tables, adaptées à leurs portefeuilles d'assurés. La méthode de construction est laissée libre, ainsi nous proposons et développons au cours de ce mémoire une méthode relationnelle originale s'appuyant sur des modèles non-paramétriques pour le lissage des taux bruts. Cette démarche a pour but de conserver une bonne fiabilité dans le cadre de petits échantillons.

Cette méthode originale s'effectue en plusieurs étapes. Dans un premier temps, à partir des données, on calcule les taux bruts de mortalité, sans distinction des années. Ces taux sont ensuite mis en relation avec les taux d'une année d'une table de référence prospective choisie avec précaution. Contrairement aux démarches habituelles, c'est cette relation que nous proposons ensuite de lisser par différentes méthodes non-paramétriques et non directement les taux bruts. Cette étape a pour intérêt de profiter des informations transmises par la table de référence. L'étape qui suit consiste à prolonger la courbe lissée jusqu'à 130 ans à l'aide de méthodes de fermeture de tables classiques. Enfin, la construction se termine par la projection sur plusieurs

années des taux lissés en appliquant directement l'évolution utilisée par la table de référence prospective.

Cette démarche relativement simple présente le mérite de limiter au maximum le risque de modèle, tout en proposant des résultats très acceptables. Néanmoins, cette méthode a pour principale contrainte l'introduction de trois paramètres exogènes : un paramètre de lissage, une fonction de poids et une table de mortalité prospective de référence. La seconde partie de ce mémoire a pour objet la validation de la table ainsi construite. Il y sera décrit l'ensemble des risques sous-jacents à la construction d'une table de mortalité prospective, avec notamment le risque d'estimation et le risque de tendance. Les tests de validation que l'on propose par la suite analyse de façon complémentaire les différents aspects de la table. Enfin, le présent mémoire se termine par une application de l'ensemble des concepts développés au cours de la partie théorique.

Abstract

KEYWORDS : Prospective life table, estimation risk, trend risk, relational model, non-parametric methods, tools for validation.

In order to set up Solvency II, insurers must assess their risks and commitments as precisely as possible. In the life insurance field, life table quality is an essential parameter that can't be debated. Therefore, since 1993, insurance companies can build their own life tables, best suited to their insured portfolio. The building method is free, thus we propose and develop in this thesis an original relational approach based on non-parametric models for crude rate smoothing. This approach aim is to preserve a good reliability within the framework of small samples.

That innovative approach is done in several stages. At first, from the data, we calculate mortality crude rates, without distinction of years. After, this rates are related with reference rates of the optimal year, extract from the reference prospective life table, carefully chosen. Contrary to the usual initiatives, it's this relation which we suggest to smooth by various non-parametric methods and not immediatly on crude rates. This stage have for interest to take advantage of the informations submitted by the reference table. The Last stage consist in extending the curve smoothed until 120 years by using classic closing methods.

This fairly easy method enables to limit as much as possible the model risk, adding to giving very acceptable results. However, this approach imposes the introduction of three exogenous parameters : a smoothing parameter, a weight function and a reference prospective life table. The second

part of this thesis concerns the validation of the table built. It will describe all the underlying risks in the prospective life table, in particular the model risk and the trend risk. The validation tests we suggest afterwards analyze complementarily the various aspects of the table. Lastly, the hereby thesis ends with an application of all the concepts approached during the theoretical part.

Synthèse

De tout temps, il a été nécessaire aux assureurs d'estimer leurs risques pour adapter une bonne tarification ainsi que pour le calcul de leurs engagements. Sur ce point, les tables de mortalité sont particulièrement utiles. Depuis 1993, le code des Assurances, avec l'article A335-1 [1] autorise les compagnies d'assurance à développer leurs propres tables pour les adapter à leurs données. Cette approche est encouragée par la réforme européenne Solvabilité II qui impose des calculs *best-estimate* (c'est-à-dire de la façon la plus réaliste possible). Il est donc nécessaire pour les assureurs de construire des tables d'expérience prospectives de manière fiable.

1. Construction d'une table de mortalité prospective relationnel à l'aide d'outils non-paramétriques

La littérature actuarielle a proposé de multiples méthodes pour approcher la mortalité réelle d'un portefeuille d'assurés, en évoluant et se perfectionnant au fil du temps. Il est aujourd'hui aisé de calculer des taux de mortalité bruts, notamment avec les méthodes de Hoem et de Kaplan-Meier. La difficulté réside plutôt dans le lissage de ces taux pour obtenir une table réaliste, faiblement impactée par la volatilité des taux bruts issus des données. La présente étude propose une méthode originale mêlant l'utilisation d'un modèle relationnel avec un lissage non-paramétrique, ayant pour objectif de conserver une bonne fiabilité dans le cadre de faibles échantillons.

1.1 Les modèles relationnels classiques

Au cours de cette étude, il est donc présenté une démarche s'appuyant sur l'utilisation d'une table de mortalité de référence prospective, en l'occurrence la table "INSEE 2007-2060". Pour cela les taux bruts calculés sans distinction des années sont mis en relation avec une des années de la table de référence prospective choisie précautionneusement. C'est-à-dire que pour chaque âge on associe le taux brut avec le taux de référence correspondant. C'est cette fonction de passage $q_x^{brut} = f(q_x^{ref})$ que l'on lisse. Cela signifie que l'on applique le lissage non pas directement sur les taux bruts mais sur la fonction de passage entre les taux de référence et les taux bruts. L'intérêt de cette démarche est d'intégrer au modèle les variations d'une table existante et qui a fait ses preuves.

Concernant le lissage, de nombreuses méthodes semi-paramétriques ont été développées, plus ou moins réalistes. On peut notamment citer :

- le modèle à risque proportionnel de Cox (1972)
- le modèle de Brass (1971)
- le modèle de Hannerz (2001)

L'utilisation de l'un de ces modèles semi-paramétriques présente l'inconvénient de potentiellement s'écarter des données observées (risque de modèle). C'est pourquoi, d'autres approches ont été développées pour effectuer le lissage des taux bruts, notamment des approches non-paramétriques.

1.2 Application d'un lissage non-paramétrique

Dans l'optique d'éliminer tout risque de modèle issu de l'application d'un des lissages semi-paramétriques précédents, qui peut être très important, les mathématiciens ont mis en place de nombreuses méthodes non-paramétriques. Ceci a été rendu possible par le développement informatique de ce siècle dernier. Au cours de cette étude, on a choisi de s'attarder sur trois méthodes de lissage non-paramétriques :

- lissage par moyennes mobiles
- lissage Whittaker-Henderson
- lissage par méthode de Loess

Ces différents types de lissages non-paramétriques présentent l'avantage

d'être adaptables à tout type de données et de fournir des résultats très convenables. Néanmoins cette démarche intègre un autre risque : intitulé le risque d'avis d'expert. En effet, il y a nécessité de définir trois paramètres exogènes pour effectuer le lissage.

Une table de référence. C'est une table existante qui s'appuie généralement sur une population très large, telle que la population nationale. Ce paramètre est important car définit l'orientation que l'on va donner à nos lissages, de plus, c'est la tendance de cette table qui est retenue pour projeter les taux sur quelques années. Malheureusement, ce choix est très limité et donc imparfait. Pour cette étude, la table "INSEE 2007-2060" a été retenue pour deux principales raisons : c'est une table prospective très récente et qui se base sur la population française, comme notre portefeuille.

Un paramètre de lissage. Ce paramètre, noté λ , détermine la taille du voisinage utilisée pour le lissage de chaque point. Il peut être exprimé en distance ou en nombre de points (centré ou non sur le point à lisser). Delwarde et al. proposent, au cours d'un de leurs articles [16], d'utiliser un AIC modifié pour mesurer la qualité du lissage et ainsi d'optimiser la valeur de λ .

Une fonction de poids. Ce paramètre représente l'importance (le poids) pris par chaque point appartenant à la zone de lissage. Il est soumis à de nombreuses contraintes. Il n'existe pas de méthode intéressante pour déterminer la meilleure fonction de poids, ce choix est soumis à l'appréciation de l'expert.

Ces paramètres définis et le lissage réalisé, il reste à projeter nos taux sur plusieurs années, pour cela nous réutilisons la dérive de la table de référence à partir de l'année optimale. Le choix de l'année optimale servant de référence peut être important. C'est la raison pour laquelle quelques tests statistiques variés et complémentaires sont proposés afin de retenir l'année la plus pertinente tel que le test du χ^2 , la minimisation de la déviance ...

Cette première partie se termine donc sur l'obtention d'une table finale, lisse issue d'une table de référence et qui s'accorde à nos données du portefeuille. Reste néanmoins à certifier ce résultat, cela passe par de nombreux

tests présents dans cette seconde partie.

2. Validation de la table de mortalité prospective

Concernant la validation des tables de mortalité, la littérature actuarielle est peu fournie alors que cette étape est primordiale. Au cours de cette étude, il est proposé une série de tests permettant d'obtenir une réponse objective à la question : doit-on valider la table de mortalité prospective ? Pour cela des approches variées et complémentaires sont définies.

2.1 Une validation générale

Dans un premier temps, une approche qualifiée de générale est proposée. Elle consiste à vérifier que des caractéristiques simples et observables au sein de tout groupe de population sont présentes dans nos résultats. Par exemple, que la mortalité est décroissante en fonction de l'âge (en mettant de côté la mortalité infantile) et en fonction des années, et que la mortalité masculine est à tout âge supérieure à la mortalité féminine, sauf cas extrêmes où des facteurs extérieurs viendraient remettre en question ces règles. Un non-respect d'une de ces règles est généralement dû à une base de données trop peu fournie, laissant des âges avec des résultats aberrants ou à un lissage pas assez conséquent. Le dernier test que l'on qualifie de général est le calcul des SMR (*Standardized Mortality Ratio*) qui représentent les rapports entre les nombres de décès observés et les nombres de décès estimés : il doit être proche de 1.

2.2 Une validation des paramètres exogènes

Une seconde partie concerne l'appréciation de la qualité des paramètres exogènes introduits par le modèle, et principalement l'étude du lissage et du vecteur poids. Pour cela, il est nécessaire de se concentrer sur l'étude des résidus (résidus de la réponse, résidus de Pearson et résidus de la déviance) et de comprendre l'origine des résidus trop élevés, qui peuvent être dûs à un sous-lissage, et les séries de résidus consécutifs de même signe qui, au contraire, peuvent provenir d'un sur-lissage. D'autres tests sont proposés pour venir étayer cette analyse tel que le critère de régularité des taux lissés, le test des signes ... Concernant la table de référence, il n'est pas aisé de savoir si elle est bien adaptée, par contre il est intéressant de constater les

écarts qui apparaissent avec l'utilisation d'une autre table, pour ainsi évaluer l'impact d'un mauvais choix sur nos résultats.

2.3 Une validation des risques sous-jacents à la construction d'une table de mortalité

Enfin une troisième partie de la validation est mise en avant. Elle concerne l'évaluation des risques sous-jacents à la construction d'une table de mortalité. De manière indirecte, les tests précédents ont permis d'appréhender le risque d'avis d'expert. Deux autres risques non négligeables méritent d'être étudiés, le risque d'estimation et le risque de tendance, qui font l'objet de cette partie. Au sein de sa thèse [33], A. Kamega présente des démarches intéressantes pour l'évaluation du risque d'estimation sur les taux ajustés ou sur le calcul des provisions, à partir d'une resimulation des résidus, ou directement. Ces approches permettent le calcul de coefficients de variation apportant une mesure exploitable de ce risque. Concernant le risque de tendance, il a été choisi de regarder l'évolution de l'espérance de vie et de l'entropie à chaque âge, ce qui permet d'émettre un avis qualitatif après comparaison avec l'avis d'un expert. Toutefois, une réserve est posée sur ce test qui, contrairement aux autres, est particulièrement subjectif.

Synthesis

Insurers always needed to assess their risks to adapt the pricing and to calculate the commitments. On this specific point, life tables are very useful. Since 1993, the insurance code, and more precisely Article A335-1 [1] allows insurance companies to build their own tables to be suited to their data. This approach is encouraged by the European reform Solvency II which imposes best-estimate calculations (meaning as realistic as possible). Therefore, insurers must build prospective experience tables in a reliable way. This approach aim is to preserve a good reliability within the framework of small samples.

1. Building a relational prospective life table with non-parametric tools

1.1 Relational models

Actuarial literature suggested a various number of methods to handle the real mortality of the insured portfolio, changing and improving over time. Today, it's fairly easy to calculate crude mortality rates, for example with the Hoem and Kaplan-Meier's methods. The hardest point is actually to smooth these rates to get a realistic table, not significantly affected by crude rate volatility coming from the data.

This study is based on the use of relational models and more precisely on the prospective life insurance "INSEE 2007-2060". That means that we don't apply smoothing immediately on crude rates but on the function of tran-

sition between reference rates and crude rates. The point of this approach is to integrate to the model the variations of an existing table that was already successful. Concerning smoothing, several semi-parametrics methods have been developed, more or less realistic. We can quote for example :

- Cox proportional hazard model (1972)
- Brass model (1971)
- Hannerz model (2001)

The use of one of these models semi-parametrics the inconvenience to potentially move away from observed data (risk of model). That's why other methods have been developed to smooth crude rates.

1.2 Applying a non-parametric smoothing

In order to delete the model risk from the application of the previous smoothing, which can be very high, mathematicians set up a various number of non-parametric methods. This was made possible thanks to the development of computer science over the last century. This study will focus on three non-parametric smoothing methods :

- moving average smoothing
- Whittaker-Henderson smoothing
- Loess smoothing

These three kinds of non-parametric smoothing have the benefit of being adaptable to all sort of data and providing very satisfying results. Nevertheless, this approach has another risk called expert advice risk. Indeed, three exogenous parameters must be defined to apply smoothing.

A reference table. It's an existing table usually based on a very large population, such as the national population. This parameter is very important because it provides the direction we will give to the smoothing. Moreover, it's the trend of this table that is retained to project the rates over several years. Unfortunately, the choice is limited and therefore imperfect. In this study, the table "INSEE 2007-2060" was retained for two reasons : this is a very recent prospective table and it's based on the French population, as well as our portfolio.

A smoothing parameter. This parameter, called λ , determines the size of the vicinity used to smooth each point. It can be expressed as a distance or a number of points (centered or not on the point to be smoothed). Delward and al. [16] suggest, in one of their articles, to use a modified AIC to measure the smoothing quality and therefore to optimize λ .

A weight function. This parameter gives the weight of each point from the smoothing section, which is subject to many constraints. There is no interesting method to determine the best weight function. Then, it's a choice to be made by the expert according to his judgement.

Once the parameters are defined and the smoothing carried out, it only remains to project the rates over several years using the reference table again. The year taken as a baseline may be important, this is the reason why a few statistical tests, varied and complementary, are available to choose the most appropriate year. One of these tests is χ^2 test, deviance minimization. . .

At the end of the first part we get a smoothed resulting table coming from a reference table and appropriate to the portfolio data. Nevertheless, we must now certify the achieved result. This is possible thanks to many tests we will see in the second part.

2. Validation of the prospective life table

Although this step is essential, actuarial literature doesn't provide much information about prospective life table validation. This study gives a series of tests to get an objective answer to the following question : must we validate the prospective life table ? Varied and complementary approaches are defined to help us.

2.1 General validation

Firstly, an approach said "general" is suggested. This approach is to check if simple characteristics, observable within any sub-population, are present in our results. In instance, that mortality decreases with age (putting aside infant mortality) and over the years, or that male mortality is at

any age higher than female mortality, except for extraordinary cases where external factors would put into question these rules. A non-compliance of one of these rules is usually due to a too small database including outliers or a non-efficient smoothing. The last general test is the SMR calculation (Standardized Mortality Ratio). The SMR are the ratios between the number of deaths observed and the estimated rate : it must be around 1.

2.2 Exogenous parameters validation

This part deals with the validation of the exogenous parameters introduced by the model and mainly the smoothing and weight vector study. Consequently, we must focus on the residuals (response residuals, Pearson residuals, deviance residuals) and understand the origin of the highest residuals, which can be due to an "oversmoothing" and the series of consecutive residuals of the same sign that, in the opposite, can be caused by an "undersmoothing". Some other tests are given to support this analysis such as the test of regularity on smoothed rates, the sign test... It's not easy to know if the reference table is suited, but it's interesting to observe the differences appearing after the introduction of another table to assess the impact of a bad choice on our results.

2.3 Validation of the underlying risks of a life table building

This last part deals with the assessment of the underlying risks of a life table building. Indirectly, the tests seen previously enabled to approach the expert advice risk. But two other risks deserve to be studied as well, the estimation risk and the trend risk, which are the topic of this part. In his thesis [33], A. Kamega introduces interesting methods for estimation risk assessment on adjusted rates or on the liabilities calculation, either from a resimulation of the residuals either directly from the rates or from the provision calculation. These methods allow to calculate coefficients of variation and therefore an exploitable measurement of this risk. Concerning the trend risk, it was decided to observe life expectancy and entropy for each age. This enables a qualitative statement after the comparison to an expert advice. However, a reserve is put on this test which, contrary to the others, is especially subjective.

Table des matières

Remerciements	1
Résumé	1
Abstract	3
Synthèse	5
Syntheses	10
Introduction générale	18
 I	 21
1 Le contexte des tables de mortalité prospectives	22
1.1 Qu'est qu'une table de mortalité prospective?	22
1.2 Intérêts des tables de mortalité	24
1.3 Les tables de mortalité prospectives et la réglementation . . .	25
1.4 Comment construire une table de mortalité prospective? . . .	26
1.5 Conclusion et objectifs	28
 2 Construction d'une table de mortalité prospective relation-	
nelle par des modèles non-paramétriques	29
2.1 Introduction	30
2.2 Les modèles relationnels classiques	32
2.3 Résolution à l'aide d'outils non-paramétriques	34

2.3.1	Des choix communs à l'ensemble des méthodes non-paramétriques	34
2.3.1.1	Le choix de la table de mortalité de référence	35
2.3.1.2	Le paramètre de lissage	36
2.3.1.3	La fonction de poids	37
2.3.2	Les méthodes de construction des tables de mortalité non-paramétriques	38
2.3.2.1	La méthode des moyennes mobiles	38
2.3.2.2	La méthode de Whittaker-Henderson	39
2.3.2.3	La méthode de Loess	40
2.4	Fermeture de table	41
3	Projection et mesure des risques sous-jacents à la construction d'une table de mortalité	43
3.1	Introduction	43
3.2	Une cartographie des risques	44
3.3	Le risque d'estimation	46
3.3.1	Rééchantillonnage par simulation directe	47
3.3.2	Rééchantillonnage par simulation sur les résidus	48
3.3.3	Mesure du risque d'estimation sur les provisions	49
3.4	Incertitude sur la tendance : risque d'avis d'expert	50
3.4.1	Description du choix de la dérive	52
4	Les méthodes d'appréciation et de validation de la table prospective à certifier	55
4.1	Introduction	55
4.2	Validation générale	56
4.3	Appréciation et validation des paramètres exogènes	58
4.3.1	Étude du lissage et du vecteur poids	58
4.3.2	Table de référence	61
4.4	Intervalles de confiance	61
4.5	Appréciation du risque d'estimation et du risque de tendance	63
4.6	Conclusions, limites	64
II	Partie Application	66
5	Etude de la base de données	67

5.1	Introduction	67
5.2	Présentation du jeu de données	68
5.3	Caractéristiques de la base de données	69
5.4	Evolution des variables dans le temps	72
5.4.1	Évolution de la parité Homme/Femme au cours du temps	72
5.4.2	Moyenne des âges d'exposition et de décès au cours du temps	73
5.5	Conclusion	74
6	Construction d'une table de mortalité prospective	75
6.1	Introduction	75
6.2	Calcul des taux de mortalité bruts	76
6.3	Lissage des taux bruts	78
6.3.1	Choix des paramètres exogènes	78
6.3.2	Application des différentes méthodes de lissage non- paramétriques	80
6.3.2.1	Lissage par la méthode des moyennes mobiles	80
6.3.2.2	Lissage par la méthode de Whittaker-Henderson	81
6.3.2.3	Lissage par la méthode de Loess	81
6.3.3	Fermeture de la table	82
6.4	Projection des taux de la table du moment	83
6.5	Conclusion	84
7	Application des outils de validation de table de mortalité	86
7.1	Introduction	86
7.2	Le risque d'estimation	87
7.2.1	Ré-échantillonnage par simulation directe	87
7.2.2	Ré-échantillonnage par simulation des résidus	88
7.3	Le risque de tendance	89
7.4	Validation générale	91
7.5	Appréciation des paramètres exogènes	93
7.5.1	Analyse du paramètre de lissage	93
7.6	Conclusion	95
	Conclusion générale	97

Annexes	101
A Les principaux indicateurs d'analyse de tables de mortalité	103
B Les fonctions de poids classiques	106
C Outils de fermeture des tables de mortalité	108
C.1 Méthodes de Denuit et Goderniaux	108
C.2 Méthodes de Coale et Kisker	109
D Tableau d'exposition et de sinistralité de la base	110
E Tableau des résultats issus du lissage des taux bruts	112
F Représentation des courbes finales	114
F.1 Lissage par la méthode des moyennes mobiles	114
F.2 Lissage par la méthode de Whittaker-Henderson	115
F.3 Lissage par la méthode de Loess	115
Bibliographie	116

Introduction générale

Au vu de l'importance prise par les tables de mortalité prospectives dans l'actuariat et surtout de l'impact de l'utilisation d'une table mal adaptée, des réformes telles que Solvabilité II ou le SST (*Swiss Solvency Test*) encouragent les assureurs à évaluer leurs engagements de la façon la plus précise possible. Pour cela, les assureurs doivent être en mesure de déterminer les risques auxquels ils sont soumis à chaque instant. Dans ce but, il est donc important de pouvoir construire des tables en cohérence avec les données du portefeuille et dans un second temps, d'en vérifier la validité afin d'obtenir les tables les plus fiables possibles pour les projections futures. En effet, c'est le véritable objectif des tables de mortalité prospectives. En partant du principe qu'aucune table de mortalité n'est parfaite et que la littérature présentant des outils de contrôle ou de validation de ces tables est relativement peu fournie, l'objectif de ce mémoire est double. Le premier est de proposer une méthode de construction de table de mortalité originale tentant d'approcher au maximum la réalité et ainsi permettre aux assureurs d'améliorer la vision de leurs risques. Le second est de proposer des indicateurs pertinents, ainsi que des tests pour contrôler et pouvoir valider une table de mortalité prospective.

Pour parvenir à cet objectif, le présent mémoire se divise en deux grandes parties. Dans la première, la partie théorique, nous commençons par introduire les tables de mortalité de manière générale en précisant notamment leurs champs d'application dans le domaine de l'actuariat. Ceci permet de poser la problématique de ce mémoire. Ensuite, il est développé une démarche originale mêlant l'application de différentes méthodes non-

paramétriques pour réaliser une table de mortalité prospective avec l'utilisation d'une table de mortalité de référence. C'est une des méthodes préconisées par des travaux récents de l'Institut des Actuaires [46] pour construire une table de mortalité propre à un organisme, surtout dans le cadre de données réduites. L'emploi d'une de ces méthodes a pour vocation de limiter au maximum un des risques inhérents à la construction d'une table de mortalité : le risque de modèle et ainsi se rapprocher du risque réel du produit auquel est soumis l'assureur. Néanmoins, comme toutes les méthodes de construction de table de mortalité, elle possède d'autres risques. Ces risques, nous les détaillons puis nous les estimons au cours de ce mémoire, en proposant de multiples tests et indicateurs, variés et complémentaires, dans le but d'évaluer la pertinence de la table. La première partie se termine par un développement des méthodes retenues pour l'évaluation finale de la table et pouvoir ainsi donner un avis argumenté répondant à la question : doit-on, oui ou non, valider la table de mortalité considérée ?

La seconde partie, tournée vers la pratique, a pour but de mettre en application la théorie précédente sur un jeu de données issu d'un agglomérat de portefeuilles. Après une analyse approfondie de ces données, nous appliquons les démarches développées dans la première partie, à savoir : la construction de la table, l'évaluation des risques, la mise en place des indicateurs et des tests qui mènent à la validation ou au rejet de la table ainsi créée. Le mémoire se conclut sur une analyse de la procédure employée ainsi que sur les tests de validation réalisés en présentant les différents avantages, les défauts et les limites repérés.

Première partie

Partie théorique

Chapitre 1

Le contexte des tables de mortalité prospectives

Sommaire

1.1	Qu'est qu'une table de mortalité prospective ? .	22
1.2	Intérêts des tables de mortalité	24
1.3	Les tables de mortalité prospectives et la réglementation	25
1.4	Comment construire une table de mortalité prospective ?	26
1.5	Conclusion et objectifs	28

1.1 Qu'est qu'une table de mortalité prospective ?

Les tables de mortalité sont des éléments centraux pour l'ensemble des compagnies d'assurance exerçant dans la branche vie. Une table de mortalité est un outil présenté sous la forme d'un tableau, montrant le nombre de survivants par âge, et par année pour les tables prospectives, en fonction du nombre d'individus à l'âge initial. Le plus souvent, on retient 100 000 personnes à cet âge initial, comme pour l'exemple de la figure 1.1. Elle a pour fonction de représenter, et parfois d'anticiper, la mortalité d'une population. Cette modélisation de la mortalité permet notamment une évaluation de l'espérance de vie (à tout âge) d'une population donnée. Elle revêt également une importance capitale dans la tarification des produits d'assurance-vie,

TH00-02		TGH05				
Age	Lx	Age	ANNEE DE NAISSANCE			
			1900	1901	1902	1903
0	100000	92	-	-	-	-
1	99511	93	-	-	-	100 000
2	99473	94	-	-	100 000	79 024
3	99446	95	-	100 000	77 333	61 197
4	99424	96	100 000	75 327	58 321	46 205
5	99406	97	73 481	55 401	42 932	34 045
6	99390	98	52 620	39 697	30 782	24 426
7	99376	99	36 636	27 648	21 447	17 025
8	99363	100	24 738	18 670	14 483	11 498
9	99350	101	16 156	12 190	9 454	7 503
10	99338	102	10 175	7 672	5 946	4 718

FIGURE 1.1 – Extrait de la table du moment TH00-02 et de la table générationnelle TGH05

pour le calcul de rentes ou pour le calcul des provisions mathématiques ... Aujourd'hui, la littérature actuarielle répartit ces tables de mortalité en deux catégories : les tables du moment et les tables prospectives.

Les tables du moment représentent une photographie de la mortalité de la population à un moment donné. Elles sont mises en place suite à l'étude d'une population dans sa globalité, sans tenir compte de l'année de naissance. Toutefois, comme le montre un article de l'Institut National d'Etudes Démographiques (INED) [24], on constate sur la population française en général, une amélioration de l'espérance de vie à tous les âges. Ce phénomène peut être expliqué par les progrès de la médecine, des conditions sanitaires et de travail ... Cette remarque est valable sur la population française au global mais n'est pas directement observable sur un portefeuille d'assurés (à cause du manque de profondeur), ou pas forcément présente (population beaucoup plus restreinte que la population nationale, ce qui implique que le phénomène n'est peut-être pas représenté à l'identique).

Au contraire des tables du moment, les tables prospectives sont des tables bidimensionnelles qui tiennent compte de l'âge, mais aussi de l'année de naissance : elles permettent de détecter une tendance dans la mortalité de la population et ainsi projeter la mortalité future. Elles prétendent donc

refléter plus fidèlement la mortalité de la population actuelle et d'anticiper sur l'avenir. Comme nous le verrons par la suite, ces tables sont bien plus intéressantes pour l'assureur, mais également plus complexes à réaliser.

Ce présent mémoire s'appuie sur l'élaboration d'une table de mortalité prospective relationnelle, c'est-à-dire que l'on réalise une table de mortalité prospective qui se base sur une table de référence certifiée : en l'occurrence, les tables "INSEE 2007-2060". Il s'agit de la projection démographique nationale française sur la période 2007-2060 fournie par l'Institut National de la Statistique et des Études Économiques (INSEE) [25]. Ces tables sont préférables aux tables générationnelles TGH/F 05 car concernent directement les années futures (pas de nécessité de retraitements de la table pour l'adapter). De plus, il a été noté en conclusion de travaux sur les tables de mortalité de place, menés par l'Institut des Actuaire (IA), que les taux de mortalité des tables TGH/F 05 sont irréguliers en fonction de t à x fixé.

1.2 Intérêts des tables de mortalité

Pour l'ensemble des contrats d'assurance proposés par les organismes d'assurance, il y a nécessité de définir un tarif en relation avec les risques auxquels sont soumis les assureurs ainsi que les provisions correspondantes. Ces dernières viennent protéger les assureurs en cas de sinistralité.

Pour cela, l'assureur se doit de mesurer le plus précisément possible les risques auxquels il est soumis. La mortalité de l'assuré est l'un des principaux risques présents dans la grande majorité des contrats : que ce soit prépondérant pour une assurance-vie ou moins évident pour tout type de contrat IARD¹. Les tables de mortalité, prospectives ou non, ont donc pour but de définir sous quelle probabilité le contrat peut être clôturé pour cause de décès.

L'objectif des tables de mortalité est donc de représenter la mortalité d'une population qui est généralement segmentée par sexe. En effet même si dans l'objectif d'une égalité homme/femme, une tarification par sexe n'est plus permise depuis le 21 décembre 2012, les assureurs utilisent toujours des

1. IARD : Incendie, Accidents et Risques Divers

tables de mortalité par sexe pour évaluer au mieux les risques auxquels ils sont soumis et ainsi calcul au plus juste leurs provisions.

1.3 Les tables de mortalité prospectives et la réglementation

Les tables de mortalité utilisées par les assureurs sont soumises à une réglementation très précise, décrite par l'article A335-1 du code des Assurances [1]. Ils doivent soit utiliser une des tables homologuées par arrêté du ministre de l'Économie et des finances sur la base de données publiées par l'INSEE, soit une table certifiée par un actuaire agréé à cet effet et indépendant, qui est soumis à une charte de déontologie stricte rédigée par l'IA [26]. Les tables TGH05 et TGF05 introduites par l'arrêté du 1^{er} août 2006 portant sur l'homologation des tables de mortalité pour les rentes viagères [28], sont les tables de mortalités réglementaires utilisées pour les rentes viagères, respectivement pour les hommes et les femmes. Elles proposent une évolution de la mortalité pour les générations comprises entre 1900 et 2005.

De nos jours, nous assistons à une transition avec l'élaboration des réglementations telles que Solvabilité II, le *Swiss Solvency Test* (SST) ou encore les normes *International Financial Reporting Standards* (IFRS), qui ont pour optique que l'assureur doit toujours être en mesure de connaître le plus précisément possible son risque. Il y a nécessité d'évaluer les engagements d'assurance à partir d'hypothèses réalistes et fines. Cela signifie être dans un cadre d'évaluations *best estimate*, ce qui contraste avec les démarches actuelles de Solvabilité I. En effet, si ces deux cadres gardent une marge de prudence, celle-ci doit être explicite dans Solvabilité II. On pourra se reporter à l'ouvrage de Planchet et Thérond [40] pour plus de précisions. Il est alors recommandé de s'appuyer sur l'expérience du portefeuille et en toute rigueur d'utiliser des tables de mortalité prospectives quel que soit le risque : mortalité ou survie. L'utilisation de tables du moment en cas de décès était possible dans le cadre de Solvabilité I, car pour un risque en cas de décès, la prudence d'une table du moment croît avec le temps et il n'est pas toujours aisé d'avoir des données conséquentes sur plusieurs années. De plus, cette prudence évite la création d'un risque de modélisation (risque de tendance) qui peut être très important. Mais avec ces nouvelles régle-

mentations et d'une manière générale pour les contrats en cas de vie, cette prudence n'est plus souhaitable et l'utilisation d'un mécanisme tel que le recours à un décalage sur les tables ou le recours à des tables prospectives est alors indispensable. Néanmoins, on note qu'une certaine marge de prudence est toujours souhaitable dans Solvabilité II, lors de la réalisation d'une table de mortalité en sous-estimant les taux de mortalité en cas de vie et en les surestimant en cas de décès.

Si la réglementation est très précise concernant l'utilisation des tables de mortalité, elle ne donne aucune méthode quant à la réalisation de ces tables, ni approche pour en vérifier la qualité. Ceci offre une très grande liberté à l'actuaire pour choisir la méthode à adopter et la validation est laissée à la discrétion de l'actuaire agréé chargé de la certifier.

1.4 Comment construire une table de mortalité prospective ?

Comme nous le rappelle Booth et Tickle [4], les recherches sur la prévision de la mortalité, qui se sont fortement développées ces dernières décennies, ont révélé trois approches pour réaliser une table de mortalité prospective. Ces trois approches sont très différentes mais ne sont pas indépendantes les unes des autres, généralement la construction d'une table de mortalité fait appel à l'ensemble de ces méthodes.

L'approche la plus utilisée par les actuaires et également la plus naturelle consiste à se référer au passé, en effet nous disposons uniquement des informations sur la mortalité passée. On remarque que son évolution au cours du temps présente des régularités. Il est possible d'appliquer un modèle sur les taux de mortalité que l'on calibre sur les observations disponibles, puis que l'on projette pour l'obtention des taux futurs et ainsi créer une table prospective. Cette approche qui est également la méthode la plus développée, présente une importante difficulté pour les études sur de petits échantillons. En effet, comme le montrent Alho et Spencer [2], la période d'observation doit être deux à trois fois plus grande que la période de prédiction. Pour cette approche, nous pouvons citer quelques modèles "classiques" couramment utilisés : modèle de Lee-Carter (1992), log-Poisson, log-linéaires, lo-

gistique décalée ... qui sont détaillés au sein de l'ouvrage de Planchet et Thérond [41].

Une seconde approche, plus subjective, consiste à se référer à un avis d'expert. Cet avis consiste en l'apport d'une ou plusieurs contraintes exogènes qui peuvent se porter sur une certaine évolution de la mortalité, l'espérance de vie à un certain âge, le choix d'une table de référence ... Généralement, cette approche est utilisée en complément de l'approche précédente en apportant des informations supplémentaires dans le cas de petits échantillons, malgré le biais possible en cas d'erreur de cet avis. Par exemple, il n'est pas aisé d'anticiper une amélioration de la mortalité due à des découvertes médicales. Dans ce domaine, nous pouvons citer les travaux réalisés par le *Continuous Mortality Investigation* (CMI), un organisme de recherche de l'IA au Royaume-Uni et un lecteur intéressé pourra se référer aux documents suivants [10] [11] [12] qui en sont issus.

Enfin une dernière approche, beaucoup moins utilisée par les actuaires, est envisageable. On peut adopter une vision plus structurelle en décomposant la mortalité d'une population en facteurs de mortalité, tels que le sexe, la situation professionnelle, le secteur géographique, fumeur ou non ... Il est également possible d'utiliser une autre segmentation en tenant compte des différentes causes de mortalité, ceci dans le but de s'approcher au plus près de la réalité et notamment proposer des modèles épidémiologiques. Ces types de segmentation que l'on peut qualifier d'extrêmes nécessitent un important jeu de données, très bien renseigné, ce qui est complexe à réaliser, en particulier si l'on veut tenir compte des causes multiples de décès. Cet élément semble indispensable pour obtenir des résultats cohérents. De plus, on note que l'actuaire, contrairement au démographe, n'a pas la nécessité de comprendre précisément les raisons des décès mais la tendance générale pour tarifier un produit ou calculer des engagements. On peut néanmoins citer l'ouvrage sur le sujet de Burgio et Frova [5] ou plus récemment les travaux de Kamega et Planchet [30] et de Guibert et Planchet [20], ou encore le mémoire de Guérin [21].

Pour conclure, il n'existe pas de bonnes ou de mauvaises méthodes, chacune ayant ses avantages et ses inconvénients. Elles peuvent généralement se

compléter. Il convient plutôt de choisir en fonction des objectifs (par exemple si la connaissance précise des causes de mortalité est nécessaire) ainsi que de la taille des données : plus le nombre de données est faible, plus la présence d'un ou plusieurs avis d'un expert est nécessaire. Au cours de cette étude, les trois approches décrites sont combinées, comme cela est généralement fait. En effet, l'étude s'appuie sur la mortalité passée qui est complétée par certains paramètres exogènes tels qu'une table de référence par sexe. Ceci permet de réaliser une table de mortalité prospective différenciée selon le sexe des individus.

1.5 Conclusion et objectifs

Cette première partie nous a donc éclairés sur le concept des tables de mortalité, c'est un outil simple dans sa présentation ainsi que dans son utilisation, et sur leur importance pour les assureurs. Le tout est dicté par une réglementation plus ou moins précise, car une grande liberté est laissée pour la procédure de construction de ces tables. Ceci a l'avantage d'offrir le choix de la méthode la plus adaptée aux données et aux objectifs. Mais comme le montre différents mémoires réalisés sur le sujet tels que Lelieur [35] ou Ducros [19], on constate que l'application de deux méthodes différentes, même si elles peuvent sembler proches, peuvent donner des résultats sensiblement différents.

Il existe donc un réel besoin d'offrir des outils pour contrôler les tables de mortalité, question que la littérature aborde peu. Ce mémoire ambitionne de proposer quelques indicateurs et quelques tests pour palier à ce manque. On note que l'évolution des risques est étroitement liée à la procédure de construction choisie, il en va de même pour les éléments de validation que nous proposons. Ils sont en partie issus de travaux menés dans le cadre d'un groupe de travail de l'IA (cf. Tomas et Planchet [45]).

Pour faciliter la lecture, le lecteur trouvera en annexe un descriptif exhaustif des notations actuarielles employées tout au long de ce mémoire.

Chapitre 2

Construction d'une table de mortalité prospective relationnelle par des modèles non-paramétriques

Sommaire

2.1	Introduction	30
2.2	Les modèles relationnels classiques	32
2.3	Résolution à l'aide d'outils non-paramétriques .	34
2.3.1	Des choix communs à l'ensemble des méthodes non-paramétriques	34
2.3.1.1	Le choix de la table de mortalité de référence	35
2.3.1.2	Le paramètre de lissage	36
2.3.1.3	La fonction de poids	37
2.3.2	Les méthodes de construction des tables de mortalité non-paramétriques	38
2.3.2.1	La méthode des moyennes mobiles	38
2.3.2.2	La méthode de Whittaker-Henderson	39
2.3.2.3	La méthode de Loess	40
2.4	Fermeture de table	41

2.1 Introduction

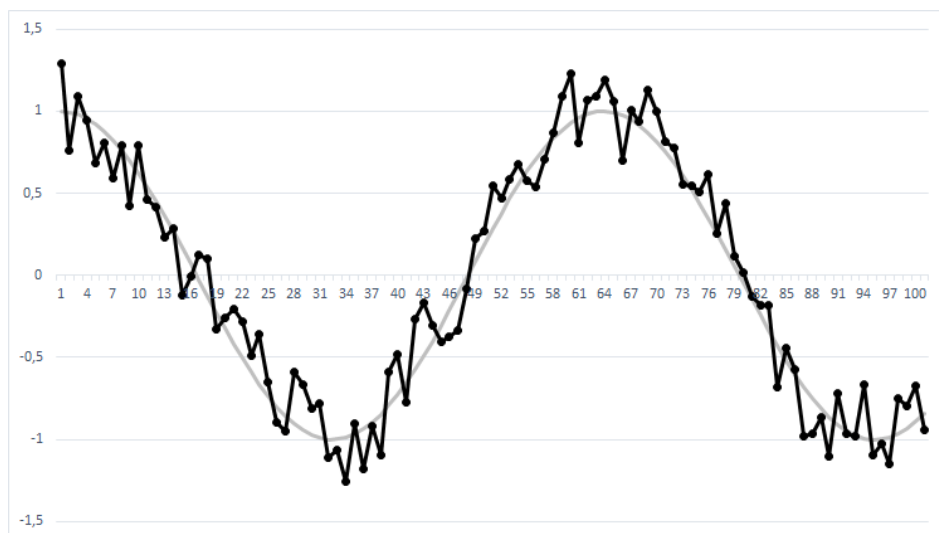


FIGURE 2.1 – La courbe grise, plus régulière, est beaucoup plus intéressante que la courbe noire

Au cours de ce chapitre, nous allons présenter différentes méthodes de construction de tables de mortalité qui auront comme point commun d'être non-paramétriques. Elles ont pour objectif de limiter au maximum les risques associés à la création de ces tables. Néanmoins, il est important de savoir que, quelle que soit la méthode utilisée, toute création de table de mortalité (prospective ou non) est soumise à de multiples risques. La procédure qui est employée au cours de ce mémoire ne fait pas exception. Les risques étant présents, cela implique qu'il n'existe pas de table de mortalité parfaite et on note que la "meilleure" des tables de mortalité n'est pas forcément celle qui correspond le mieux aux données, mais celle qui les utilise le mieux pour les extrapoler. Ainsi la courbe la moins lissée (noire) de la figure 2.1 est considérée comme moins bonne que la courbe grise car elle ne donne pas une représentation de la tendance des données, malgré le fait qu'elle passe par l'ensemble des observations. En effet, l'objectif est avant tout de prévoir l'évolution la plus probable de la mortalité future à l'aide des données passées, mais aussi d'informations extérieures que nous regroupons sous la

notion d'avis d'expert.

Au cours de cette partie, nous allons détailler la méthode et les choix retenus pour la construction de notre table de mortalité prospective. Elle a pour objectif constant de limiter au maximum les risques inhérents à la réalisation de ce type de tables.

Cette construction se déroule en plusieurs étapes. Tout d'abord, on calcule classiquement les taux bruts à partir de notre jeu de données par la méthode de Hoem ou la méthode de Kaplan-Meier, sans distinction des années. L'originalité de la démarche proposée est de mettre en relation pour tous les âges, les taux bruts obtenus avec ceux d'une certaine année d'une table de mortalité prospective, telle que cela est présenté sur la figure 2.2.

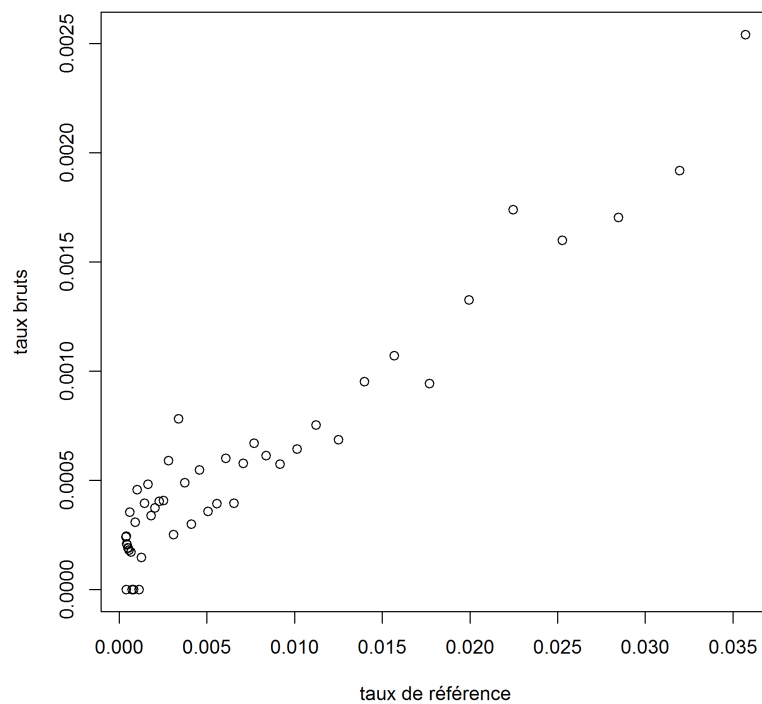


FIGURE 2.2 – Mise en relation des taux de sortie bruts avec les taux de sortie de référence pour tous les âges

Par la suite, c'est cette fonction de passage, entre les taux de référence et les taux bruts, que l'on lisse. Une partie importante de ce chapitre est consacrée au développement de différentes méthodes de lissage non-paramétriques

qui permettent d'obtenir une table de mortalité du moment complète. Le présent chapitre se termine par le développement de différentes méthodes de fermeture de table. Le passage de cette table du moment à une table prospective sur quelques années fait l'objet d'une section du chapitre suivant.

2.2 Les modèles relationnels classiques

En général, les assureurs sont soumis à une forte contrainte concernant la mise en place de tables de mortalité ajustées à leurs données. Cette contrainte est la taille de ces données, tout aussi bien au niveau de la quantité que de la profondeur. Celles utilisées au cours de ce mémoire sont constituées de 500 000 individus sur une période de 6 années. Il n'est donc pas toujours aisé d'en extraire des résultats fiables. Pour contrer ceci, une idée simple est d'utiliser une table de référence externe qui a été construite sur une population beaucoup plus étendue, et qui se rapproche de notre échantillon d'assurés.

Ainsi, cette démarche respecte notre objectif de limiter au maximum les risques en se basant sur une table, qui sans être parfaite, présente des garanties suffisantes. On peut s'y fier si l'évolution des indicateurs de notre portefeuille est considérée comparable à celle de la table de référence. L'introduction de ce paramètre exogène peut sembler déroutante (apparition d'un risque d'avis d'expert) mais nous verrons que cela apporte de nombreux autres atouts.

Cette approche consiste donc à mettre en relation entre notre jeu de données et la table de référence des indicateurs démographiques dépendant de l'âge x . Il n'est généralement pas tenu compte de la dépendance du temps t , car les jeux de données sont, le plus souvent, faibles en volume, ce qui ne permet pas une segmentation en année. Les indicateurs sont donc déterminés par âge et mis en relation avec l'année de la table de référence la plus adaptée. On peut citer comme indicateur :

- la probabilité de survie (p_x) qui représente à chaque âge la probabilité d'atteindre vivant l'âge supérieur ;
- la probabilité de décès (q_x) qui représente à chaque âge la probabilité de décéder avant d'atteindre l'âge supérieur. Tout naturellement $q_x =$

- $1 - p_x$;
 - le taux de mortalité instantané (μ_x) égal à $-\ln(1 - q_x)$;
 - l'espérance de vie à l'âge x (e_x) qui représente le nombre d'année moyen restant à vivre pour un individu d'âge x .
- Ces différents indicateurs sont détaillés en annexe.

Pour chaque âge, il est ainsi créé un point ayant en abscisse le taux de référence et en ordonnée le taux brut issu de la méthode de Hoem ou de Kaplan-Meier (détaillées en annexe) correspondant, voir figure 2.2 à la page 31. La table de référence est une table qui a déjà été certifiée, donc possédant une certaine fiabilité. Plus de précisions sur les indicateurs économiques sont disponibles en annexe.

Historiquement les principaux modèles relationnels (modèle reliant un indicateur démographique brut avec un indicateur démographique de référence) qui ont été proposés sont des modèles semi-paramétriques dépendant de l'âge x et du temps t :

- Le modèle à risques proportionnels de Cox (1972) :

$$\mu_{x,t} = \theta \mu_{x,t}^{ref}$$

- Le modèle de Brass (1971) :

$$\text{logit}(q_{x,t}) = \theta_1 + \theta_2 \text{logit}(q_{x,t}^{ref})$$

- Le modèle de Hannerz (2001) :

$$\text{logit}(q_{x,t}) = \text{logit}(q_{x,t}^{ref}) + \theta_0 - \frac{\theta_1}{x} + \theta_2 \frac{x^2}{2} + \theta_3 \frac{\exp(cx)}{c}$$

L'avantage de l'utilisation de ces modèles, notamment pour les démographes, est de pouvoir donner une interprétation de chaque paramètre. Par exemple, pour le modèle de Brass, le premier paramètre θ_1 peut être vu comme le niveau de la mortalité et θ_2 représente la variation de la mortalité par rapport à la mortalité de la table de référence. Mais pour l'actuaire, dont le but est avant tout de pouvoir tarifier des contrats ou d'évaluer des montants de provisions, ceci a beaucoup moins d'importance. C'est à ce titre que l'utilisation d'un modèle relationnel non-paramétrique est envisageable.

2.3 Résolution à l'aide d'outils non-paramétriques

Le souci des méthodes relationnelles classiques est d'imposer un modèle comme ceux décrit précédemment et qui peut ne pas correspondre à la dynamique de notre jeu de données. Il y a donc l'insertion d'un fort risque de modèle, risque que nous tentons de limiter au maximum tout au long de ce mémoire.

Considérer un problème non-paramétrique répond à cette problématique car il ne pose pas une relation précise entre la table de référence et le portefeuille de l'assureur, mais vise à l'estimation directe de la fonction de passage (et non pas des paramètres) sur la base des statistiques de mortalité. Cette fonction notée f a pour seules contraintes le fait d'être régulière (c'est-à-dire dérivable en tout point) et que

$$\forall x, \forall t \quad q_{x,t} = f(q_{x,t}^{ref})$$

Un second avantage privilégiant le recours à une méthode de lissage et non à un modèle paramétrique réside dans la quantité d'informations qui est utilisée. En effet, dans le cas des modèles comme ceux présentés précédemment, on résume l'ensemble de l'information fournie par le portefeuille en quelques paramètres. Au contraire, une méthode de lissage non-paramétrique peut être considérée comme représentant un modèle possédant une infinité de paramètres. Leur principal défaut, par rapport à une méthode dite "classique", est l'importante augmentation du temps de calcul. Mais ce n'est plus aujourd'hui un facteur discriminant grâce au développement important des possibilités de calculs informatiques de ces dernières années. Sur ce point, on peut évoquer l'article de Nathan Keyfitz de 1964 [34] qui montre l'impact de l'informatique, lors de son apparition, pour aider les démographes.

2.3.1 Des choix communs à l'ensemble des méthodes non-paramétriques

En plus des données sur les deux indicateurs que l'on veut mettre en relation, les différents algorithmes que nous allons présenter pour la réalisation du lissage de la fonction de passage nécessitent trois paramètres exogènes :

une table de mortalité de référence, un paramètre de lissage et une fonction de poids.

2.3.1.1 Le choix de la table de mortalité de référence

Lors de la réalisation d'une table de mortalité prospective relationnelle, une des principales difficultés réside dans le choix de la table de référence (et de l'année de référence dans le cas de table de mortalité prospective), car cela peut avoir un impact significatif sur les résultats. Sur ce point, nous pouvons citer le mémoire de Mindelson [37]. Cette table de référence est utilisée à deux reprises. La première fois lors de l'élaboration de la fonction de passage f , il faut une table sur laquelle la variation de l'indicateur démographique retenu soit "proche" des valeurs observées sur le jeu de données. Par exemple, si l'on désire trouver une relation sur les taux de mortalité :

$$\forall x, \forall t, \frac{\partial q_{x,t}}{\partial x} \simeq \frac{\partial q_{x,t}^{ref}}{\partial x}$$

Pour la seconde utilisation de la table de référence, lors de la projection de la tendance dans le futur, il faut une table dont l'évolution de la tendance soit comparable à notre jeu de données. C'est-à-dire que dans l'idéal :

$$\forall x, \forall t, \frac{\partial q_{x,t}}{\partial t} \simeq \frac{\partial q_{x,t}^{ref}}{\partial t}$$

Par mesure de simplicité et pour conserver une certaine cohérence, on utilise la même table pour les deux situations.

Le détail du choix de l'année de référence dans le cas où l'on retient une table de mortalité de référence prospective est décrit lors de l'étape de la description du choix de la dérive.

Malheureusement, il n'existe pas un grand nombre de tables de référence et pour satisfaire ces deux contraintes, on se borne généralement aux critères classiques et naturels de la nationalité. Cela signifie que l'on retient une table de référence réalisée sur un échantillon d'individus de la même nationalité que notre jeu de données, et dont la date de construction est la plus récente possible. Ainsi pour la France, les tables proposées par l'INSEE : "INSEE 2007-2060" [25] semblent adaptées.

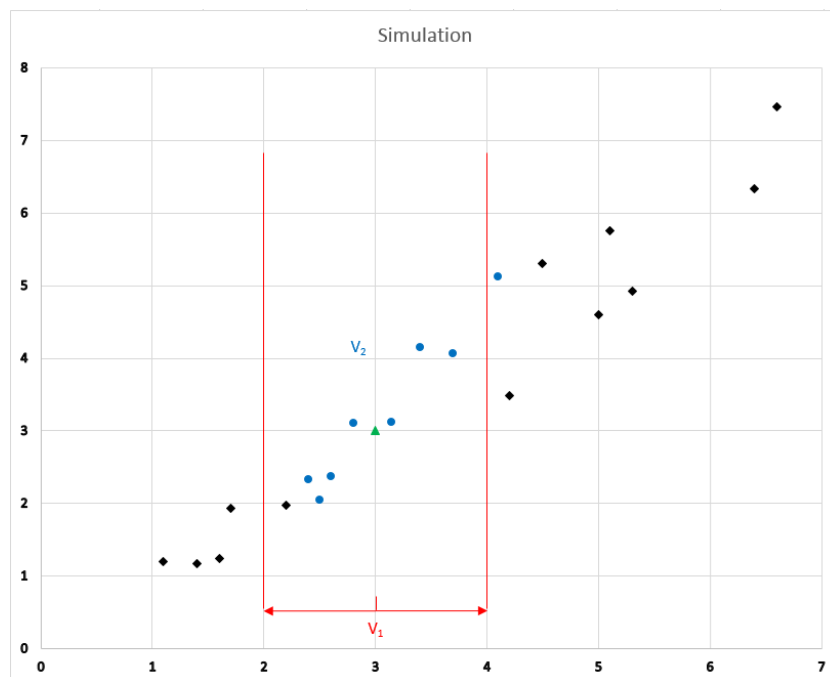


FIGURE 2.3 – Les différents types de voisinage

Par la suite, nous omettrons l'indice t des notations. En effet, le temps ne sera pas une variable dans l'obtention des taux bruts à cause de la trop faible profondeur de nos données. Cela permet également un allègement des notations.

2.3.1.2 Le paramètre de lissage

Ce second paramètre que nous présentons, appelé paramètre de lissage, représente la quantité de données utilisées pour chaque ajustement. Il est noté λ . Plusieurs approches sont possibles pour définir ce paramètre exogène. Nous pouvons considérer λ comme une distance, et ainsi prendre l'ensemble des données x_i ayant une distance au point de calcul inférieure à λ . Ce premier mécanisme est représenté par V_1 sur la figure 2.3, c'est un premier type de voisinage autour du point en forme de triangle (coordonnée (3,3)). Le paramètre λ peut aussi être un pourcentage du nombre de données x_i en prenant les λ données les plus proches (quel que soit leur distance) ou en conservant une symétrie autour du point de calcul. Ce second mécanisme

est représenté par V_2 sur la figure 2.3 qui représente ce second type de voisinage. L'ensemble des points ainsi sélectionné est appelé le voisinage, est noté $V(x_i)$. Plus ce voisinage est important et plus la courbe qui en résulte est lisse, mais au contraire s'éloigne des données.

Dans leur article, pour sélectionner le degré de lissage optimal, Delwarde et al. [16] proposent de se référer à un AIC modifié défini comme suit

$$AICC1 = n \ln(\hat{\sigma}^2) + n \frac{\frac{\delta_1}{\delta_2} (n + \text{trace}(\hat{H}))}{\frac{\delta_1}{\delta_2} - 2}$$

avec $\hat{H}$ la matrice de passage entre q_x^{ref} et q_x , $\frac{\delta_1}{\delta_2}$ représente le degré de liberté et n le nombre d'observations, ces notations étant décrites de manière plus précise en annexe. Le premier terme représente la qualité de l'ajustement tandis que le second représente la complexité du modèle : on cherche à sélectionner le modèle minimisant le $AICC1$.

2.3.1.3 La fonction de poids

En plus de ce paramètre de lissage, ces méthodes nécessitent la définition d'une fonction de poids, qui dépend fortement du paramètre de lissage retenu. En effet, il semble naturel d'émettre l'hypothèse que deux points "proches" sont plus liés que deux points "éloignés". La fonction de poids, notée W , a pour effet de favoriser les points u les plus proches de la zone d'ajustement x_i avec comme contraintes :

$$\begin{cases} \forall u, W(u) \geq 0 \\ \forall u \notin V(x_i), W(u) = 0 \\ W \text{ doit être symétrique et décroissant en fonction de la distance à } x_i \end{cases}$$

La littérature propose une multitude de fonctions de poids que l'on peut retrouver en annexe.

Une fois ces paramètres choisis ainsi que l'indicateur démographique servant aux futurs calculs d'ajustements sélectionnés, il reste à retenir une méthode de lissage. La littérature mathématique est abondante dans ce domaine et nous allons en présenter trois différentes.

2.3.2 Les méthodes de construction des tables de mortalité non-paramétriques

Il existe de nombreuses méthodes de lissage, nous en avons retenu trois (moyennes mobiles, Whittaker-Henderson et Loess) que nous présentons et appliquons sur un exemple simple. Ces méthodes sont comparées lors de la partie pratique. Nous pouvons citer la thèse de Julien Tomas qui est très complète sur ces approches [47]. Lors de cette partie, la description des méthodes de lissage se fait dans un cas général, pour l'appliquer à notre étude on remplacera $x_i = q_x$. Les taux de sortie de la table de référence, q_x^{ref} intervenant dans le calcul des distances pour la fonction de poids.

2.3.2.1 La méthode des moyennes mobiles

Cette méthode est l'une des premières méthodes de lissage à avoir été développée et présente le mérite d'être la plus simple à comprendre et à utiliser. Elle consiste à prendre, pour chaque point à évaluer, la moyenne des points appartenant à son voisinage, pondérés par une fonction de poids qui dépend de la distance.

La formule générale pour ce lissage nécessite une fonction de poids qui dépend de l'écart de rang entre x_i et x_j :

$$f(x_i) = \sum_{j \in V(x_i)} W_i(x_j) x_j$$

Par exemple, à partir d'un vecteur y simulé ainsi :

$$y = x + \epsilon \text{ avec } x \in \{1; \dots; 10\} \text{ et } \epsilon \in \mathcal{N}(0, 1)$$

si on applique la méthode de lissage par moyennes mobiles avec comme vecteur de poids autour de chaque point $(\frac{1}{3}; \frac{1}{3}; \frac{1}{3})$ et $(\frac{1}{2}; \frac{1}{2})$ pour les points extrêmes, alors nous obtenons les résultats suivants :

x	1	2	3	4	5	6	7	8
y	1,30	2,55	2,50	4,20	4,54	5,64	6,84	7,23
Moyennes mobiles	1,93	1,91	2,91	3,61	4,51	5,45	6,83	7,04

TABLE 2.1 – Exemple d'un lissage moyennes mobiles

2.3.2.2 La méthode de Whittaker-Henderson

La méthode de Whittaker-Henderson est un modèle non-paramétrique simple utilisé entre autres par le Bureau Commun des Assurances Collectives (BCAC) pour le lissage des barèmes de provisions mathématiques en incapacité et invalidité. La méthode s'appuie sur deux paramètres dont on cherche à minimiser la combinaison linéaire :

- un paramètre de fidélité F :

$$F = \sum_{i=1}^n W_i (x_i - \hat{x}_i)^2$$

- un paramètre de régularité S :

$$S = \sum_{i=1}^{n-z} (\Delta^z x_i)^2$$

avec w_i représentant les poids et z le degré du polynôme utilisé pour le critère de régularité.

On cherche donc à minimiser la fonction

$$M = F + hS$$

h représente l'importance que l'on veut donner au paramètre de régularité par rapport au paramètre de fidélité, ainsi plus il est élevé et plus la courbe sera lisse. Cette dernière équation nous amène à résoudre l'ensemble d'équations suivant :

$$\forall i \quad \frac{\partial M}{\partial x_i} = 0$$

Comme le rappelle Planchet et Thérond au sein de leur ouvrage [41], on montre que la minimisation de ce critère mène à l'équation matricielle suivante :

$$\hat{y} = (W + hK_z^T K_z)^{-1} W y$$

avec W la matrice de poids et K_z une matrice de dimension $(n - z)z$ où les termes sont les coefficients binomiaux d'ordre z avec des signes alternés en commençant par un positif si z est pair.

Si l'on reprend l'exemple précédent avec $h = 1$, $z = 2$ et en appliquant des poids identiques à tous les points (fonction de poids uniforme) :

x	1	2	3	4	5	6	7	8
y	1,30	2,55	2,50	4,20	4,54	5,64	6,84	7,23
Whittaker-Henderson	1,23	1,85	2,80	3,69	4,40	5,52	6,93	7,98

TABLE 2.2 – Exemple d'un lissage Whittaker-Henderson

2.3.2.3 La méthode de Loess

Cette méthode pour évaluer la fonction f est aussi appelée régression polynomiale pondérée localement. Elle a été imaginée par Cleveland [8], puis mise en place par Cleveland et Devlin [7]. Elle combine la simplicité des régressions linéaires des moindres carrés avec la flexibilité des régressions non-linéaires en construisant une fonction qui décrit la partie déterministe de la variation des données, point par point.

La méthode part de la remarque qu'au niveau de chaque point de l'ensemble de données, on peut appliquer une fonction polynomiale de faible degré (1 ou 2) ajustée par la méthode des moindres carrés en ne tenant compte que des points situés au voisinage. On cherche donc à minimiser la fonction :

$$g(\beta_0; \dots; \beta_1) = \sum_{k=1}^n W_k(x_i) [y_k - \beta_0 - \beta_1 x_k - \dots - \beta_d x_k^d]^2$$

avec d le degré de la fonction retenue. Nous notons la solution de cette minimisation $\hat{\beta}_j(x_i)$ et donc :

$$\hat{y}_i = \sum_{j=0}^d \hat{\beta}_j(x_i) x_i^j$$

La théorie donne pour solution de l'équation précédente avec $\hat{\beta} = (\hat{\beta}_0; \dots; \hat{\beta}_p)$

$$\begin{cases} \hat{\beta}_0(x_i) = \hat{f}(x_i) = e_i^T (X^T W X)^{-1} X^T W y \\ \forall k \in 1, \dots, p \quad \hat{\beta}_k(x_i) = (X^T W X)^{-1} X^T W y \end{cases}$$

avec e_i , un vecteur de dimension p rempli de 0 et un 1 en $i^{\text{ème}}$ position.

Si l'on reprend l'exemple précédent, en appliquant la méthode de *Loess* à l'aide de polynômes de degré 2 et de la fonction de poids *Triweight* :

x	1	2	3	4	5	6	7	8
y	1,30	2,55	2,50	4,20	4,54	5,64	6,84	7,23
Loess	1,18	2,02	2,87	3,84	4,12	5,43	6,98	8,31

TABLE 2.3 – Exemple d'un lissage par Loess

L'application de l'une de ces trois approches nous permet de poursuivre notre objectif de limiter au maximum les risques, en effet, le risque de modèle est réduit à son minimum car ce sont les données qui directement construisent le modèle. Il est à noter que ceci est rendu possible uniquement parce que nous avons retenu le choix d'utiliser une table de référence, ces méthodes n'étant pas applicables en l'état sur un faible jeu de données.

2.4 Fermeture de table

Du fait du faible nombre de personnes d'âge élevé (supérieur à 80 ou 90 ans) présent au sein des portefeuilles d'assurance, il existe un fort risque d'estimation entraînant une forte volatilité des résultats. Par conséquent, il est généralement choisi de ne pas intégrer dans l'étape précédente les individus

d'âge élevé pour la fabrication de la table. Néanmoins, il est nécessaire d'utiliser ces grands âges pour compléter notre modèle. De nombreuses études ont été réalisées sur le sujet tels que Vaupel [48] ou Horiuchi et Wilmoth [23] et de multiples méthodes ont été proposées : nous pouvons citer par exemple celle de Denuit et Goderniaux, celle de Coale et Kisker, le modèle de Kannistö ou encore le modèle logistique ...

L'étude approfondie réalisée par Quashie et Denuit [43] a comparé ces différentes méthodes de fermeture de table, car comme ils nous le rappellent le choix d'une méthode de fermeture est loin d'être innocent. En concordance avec leur conclusion, et pour la simplicité de la méthode, nous allons appliquer la méthode récente proposée par Denuit et Goderniaux [18]. Par souci de comparaison, nous testerons également une seconde méthode plus ancienne, celle de Coale et Kisker [9]. Les méthodes sont détaillées en annexe, et nous présentons ci-dessous uniquement les résultats que nous appliquerons (on note "a" l'âge limite inférieur des âges élevés) :

Pour Denuit et Goderniaux :

$$\forall x \geq a, \ln(\hat{q}_x) = c (130^2 - 260x + x^2) + \epsilon_x$$

La valeur de c est choisie afin de raccorder les deux courbes obtenues : les âges "faibles" par une des méthodes de lissage décrites précédemment et les âges "élevés" par cette méthode de fermeture.

Pour Coale et Kisker :

$$\forall x \geq 80, \mu_x = \mu_{x-1} \exp\left(\frac{\ln(\frac{\mu_{80}}{\mu_{65}})}{15} + s(x - 80)\right)$$

Avec

$$s = -\frac{\ln(\frac{\mu_{79}}{\mu_{110}}) + 31 k_{80}}{465}$$

De plus, pour les deux méthodes, afin d'éviter une cassure entre les courbes d'âges faibles et d'âges élevés, il est possible d'appliquer un simple lissage par moyennes mobiles comme décrit précédemment en prenant en compte la valeur des points aux alentours de 80 ans. Ceci permet de respecter notre hypothèse de départ, à savoir que la fonction relationnelle est régulière.

Chapitre 3

Projection et mesure des risques sous-jacents à la construction d'une table de mortalité

Sommaire

3.1	Introduction	43
3.2	Une cartographie des risques	44
3.3	Le risque d'estimation	46
3.3.1	Rééchantillonnage par simulation directe	47
3.3.2	Rééchantillonnage par simulation sur les résidus	48
3.3.3	Mesure du risque d'estimation sur les provisions	49
3.4	Incertitude sur la tendance : risque d'avis d'expert	50
3.4.1	Description du choix de la dérive	52

3.1 Introduction

Au cours de sa thèse [33], A. Kamega a montré que les méthodes de construction de tables de mortalité à l'aide de méthodes classiques, c'est-à-dire paramétriques, sont soumises à de nombreux risques, Figure 3.1. Au cours de cette partie, nous relevons les risques présents dans le cadre d'une construction selon un modèle non-paramétrique et les identifions en les comparant avec le cas classique des méthodes paramétriques. Certains de ces risques sont ensuite analysés.

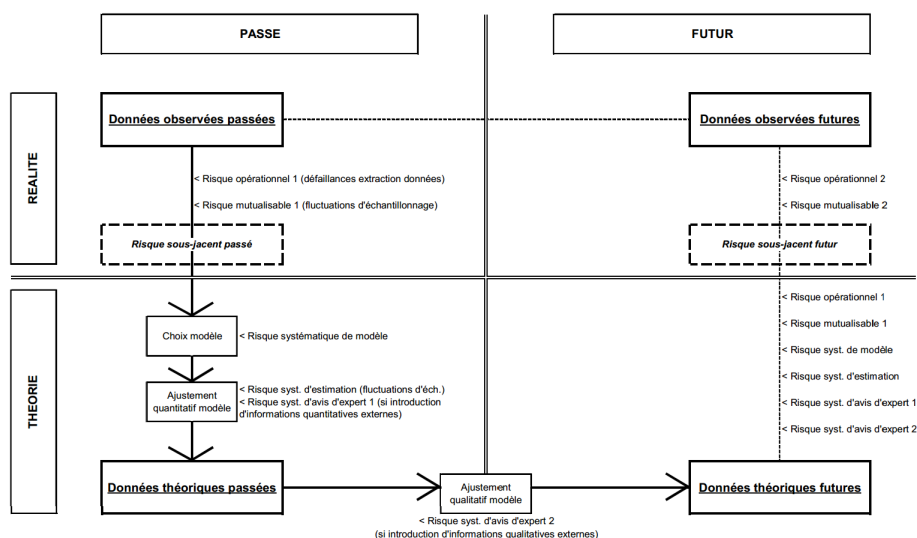


FIGURE 3.1 – Les risques inhérents à la réalisation d’une table de mortalité prospective extrait de la thèse de Kamega [33]

3.2 Une cartographie des risques

L’étude en détail des différentes étapes de la méthode originale, que nous proposons pour construire une table de mortalité prospective, révèle que la démarche est soumise à 3 risques majeurs :

- le risque opérationnel
- le risque d’estimation
- le risque d’avis d’expert

Ainsi la démarche que nous proposons réduit le nombre de risque par rapport à une étude classique : le risque de modèle est éliminé. Néanmoins, cet effet est compensé par l’augmentation du risque d’avis d’expert, notamment avec l’intégration d’une table de référence.

Comme tout travail se basant sur un jeu de données, celles sur lesquelles s’appuie la construction de notre table de mortalité sont soumises à deux risques : le risque opérationnel et le risque d’échantillonnage. Le risque opérationnel symbolise l’ensemble des erreurs commises lors de l’élaboration de la base de données. Il comprend ainsi les fautes de saisies, la dégradation

de la table survenant lors d'un transfert ou une extraction de données ... Ce risque, qui peut se révéler important, est généralement très complexe à évaluer et à maîtriser (cf. Decroocq [14]).

Le risque d'estimation est un risque mutualisable, il représente les erreurs provenant de la taille de l'échantillon. En effet, ayant une base de données limitée, nous pouvons uniquement obtenir des résultats approchés : l'erreur est d'autant plus grande que la base de données est petite. Au cours de ce mémoire, nous ne sommes pas acteurs de la création de la base de données. Nous nous contenterons d'étudier les données en proposant différentes méthodes d'évaluation de ce second risque et de mettre en place une procédure afin d'émettre un avis sur la qualité des données. Une fois la base de données sélectionnée et les taux bruts obtenus, nous devons lisser les taux bruts et mettre en place une dérive pour permettre la projection des taux sur quelques années. Pour cela, on applique à notre table la tendance de la table de référence à partir de l'année jugée optimale.

En fonction de la méthode de construction retenue, deux risques très dépendants peuvent apparaître. Le premier, le risque de modèle, est un risque systématique qui peut être plus ou moins important, surtout si on applique aucun avis d'expert en complément. En effet, deux modèles complètement différents, s'ils sont bien calibrés, peuvent être cohérents avec les données observées, mais peuvent fournir des résultats de projections sensiblement différents. Au sein de ce mémoire, nous voulons réduire au maximum le risque de modèle grâce à un ajustement par une méthode non-paramétrique. Cette méthode présente l'avantage de ne pas appliquer un modèle général à l'ensemble des données. De plus pour limiter ce risque, nous avons recours à l'introduction d'informations extérieures telles que le choix d'une table prospective de référence.

Mais cet apport extérieur implique l'apparition d'un second risque : le risque d'avis d'expert (pourquoi cette table de référence et pas une autre?). Nous détaillons et proposons des méthodes d'estimation de ce risque dans la suite de ce chapitre et pour cela, nous nous appuyons principalement sur les résultats développés au sein de l'ouvrage de Kamega et Planchet [32].

Ces risques sont donc bien connus des assureurs et détaillés au sein de la thèse de Kamega [33]. Pour mémoire, ils ont été rappelés au niveau de

la figure 3.1 à la page 44. Le plus important et le plus difficile est donc de limiter ces risques et de les évaluer afin de définir la marge d'erreur de la table finale. C'est encore un domaine où la recherche actuarielle a sa place. Au cours de ce chapitre, nous nous attacherons à définir et mesurer le risque d'estimation ainsi que le risque de tendance, c'est-à-dire à obtenir une valeur chiffrée de ces risques en faisant l'hypothèse forte qu'ils sont indépendants des autres risques. Nous verrons dans la partie suivante, dont une partie est consacrée à l'étude de ces résultats, comment interpréter ces résultats.

3.3 Le risque d'estimation

Si l'on se place dans le contexte de la nouvelle réglementation qui impose des calculs *best estimate*, il convient de prendre en compte l'hétérogénéité du portefeuille en le segmentant. De manière générale, une table de mortalité se contente d'une division homme/femme, en plus de la segmentation par âge, car les données sont le plus souvent restreintes. Une plus grande segmentation, telle que la prise en compte des caractères médicaux (par exemple fumeur/non fumeur) montrerait ses limites. Pour plus d'informations sur le sujet on peut se référer à l'article de Planchet et Leroy [39]. Dans ce cadre de données concernant de petits échantillons, ce qui est généralement le cas lors de la réalisation d'une table de mortalité basée sur des portefeuilles d'assurance, le risque d'estimation prend toute son importance et est amplifié par cette segmentation. En effet, c'est un risque systématique (c'est-à-dire non mutualisable) donc potentiellement dangereux.

Au cours de cette section, nous présentons le risque d'estimation au titre des fluctuations d'échantillonnage à travers différentes méthodes. Les deux premières consistent à resimuler la distribution des taux de mortalité ou des résidus issus des taux de mortalité suivant une loi normale. Nous utilisons pour cela la méthode de Monte-Carlo. Ces deux méthodes ont été développées respectivement par Kamega et Planchet [29] et par Delwarde [17]. La mesure suivante qui est utilisée, présentée dans la thèse de Kamega [33], a pour but de mesurer ce risque d'estimation en fonction de son impact sur les provisions.

3.3.1 Rééchantillonnage par simulation directe

Préliminaire : au cours des deux étapes de rééchantillonnage par simulation, quelques notations spécifiques sont employées :

- q_x/d_x représentent les variables brutes à l'âge x , le taux de mortalité est calculé par la méthode de Kaplan-Meier ;
- q_x^{Hoem} représente le taux de mortalité calculé par la méthode de Hoem à l'âge x ;
- $\hat{q}_x^{(i)}/\hat{d}_x^{(i)}/\hat{r}_x$ représentent les variables simulés pour l'âge x , i représentant le numéro d'ordre de la variable ;
- q_x^{adj}/d_x^{adj} représentent les variables lissés (ajustés) ;
- ER_x représente l'exposition au risque, c'est-à-dire le nombre de personnes exposés au risque à chaque âge. Elle est calculée au prorata-temporis¹.

La première méthode que nous développons consiste à resimuler les taux de mortalité q_x . Il existe deux approches pour y parvenir :

- soit par une simulation directe avec une loi normale, les taux de mortalité sont déterminés par la méthode de Hoem (voir annexe) :

$$\hat{q}_x \sim \mathcal{N}(q_x^{hoem}; \sqrt{\frac{q_x^{hoem} (1 - q_x^{hoem})}{ER_x}})$$

- soit en passant par l'estimateur du nombre de décès :

$$\hat{d}_x \sim Bin(\lfloor ER_x \rfloor; q_x^{brut})$$

puis

$$\hat{q}_x = \frac{\hat{d}_x}{ER_x}$$

Ensuite nous ajustons les taux simulés $\hat{q}_x$ obtenus pour obtenir les $q_{x,t}$ ajustés (lissage et projection) que l'on note $q_{x,t}^{adj}$. Le risque d'estimation peut

1. Par exemple, une personne présente seulement 6 mois compte pour 0,5 d'exposition.

alors être mesuré par le coefficient de variation empirique qui est une mesure de dispersion des taux de décès simulés autour des taux de décès ajustés. Si nous avons effectué k simulations des taux de mortalité :

$$cv_{x,t} = \frac{\sqrt{Var(q_{x,t}^{adj,k})}}{E[q_{x,t}^{adj,k}]}$$

Cet indicateur, qui représente le rapport entre l'écart-type de la variable et sa moyenne et est ainsi assimilable au coefficient de variation. Il semble plus pertinent que la simple variance (ou écart-type) car il est sans unité et offre ainsi la possibilité de comparer deux séries de données.

3.3.2 Rééchantillonnage par simulation sur les résidus

Dans le but de mesurer le risque d'estimation, cette seconde méthode se base sur une simulation des résidus des taux de mortalité pour rééchantillonner les taux bruts. Après avoir vérifié l'adéquation des résidus issus du jeu de données avec la distribution d'une loi normale, on en détermine la moyenne et la variance empirique pour générer une nouvelle série de résidus, puis une nouvelle série de q_x , noté $\hat{q}_x$.

Comme pour la première méthode, nous pouvons passer par la variable concernant le nombre de décès pour calculer les résidus. Sur l'exemple de Delwarde avec l'utilisation des résidus de Pearson :

$$\hat{r}_x = \frac{\hat{d}_x - d_x}{\sqrt{d_x}}$$

que l'on centre :

$$\tilde{r}_x = \hat{r}_x - r$$

avec r la moyenne arithmétique des résidus de Pearson. Si nous avons effectué k simulations des résidus :

$$r = \frac{1}{k} \sum_{i=1}^k \hat{r}_x^{(i)}$$

Nous obtenons ainsi un nouveau jeu de résidus que l'on peut appliquer pour obtenir une nouvelle série :

$$\forall i \in \llbracket 1; k \rrbracket, \quad \hat{d}_x^{(i)} = d_x + \sqrt{d_x} \tilde{r}_x^{(i)}$$

avec $\tilde{r}_x^{(i)}$ tirés aléatoirement dans le nouveau jeu de résidus. De la même manière que pour la première méthode présentée, les taux de mortalité $\hat{q}_x$ sont obtenus par la relation :

$$\forall i \in \llbracket 1; k \rrbracket, \quad \hat{q}_x^{(i)} = \frac{\hat{d}_x^{(i)}}{ER_x}$$

et les taux de mortalité ajustés $q_{x,t}^{adj}$ par une des méthodes de lissage décrites précédemment.

Le risque d'estimation est à nouveau mesuré par le coefficient de variation empirique, avec pour k simulations des taux mortalité :

$$cv_{x,t} = \frac{\sqrt{Var(q_{x,t}^{adj,k})}}{E[q_{x,t}^{adj,k}]}$$

3.3.3 Mesure du risque d'estimation sur les provisions

Une autre approche proposée pour mesurer le risque d'estimation est proposée par Kamega dans sa thèse [33]. Elle consiste à mesurer directement l'impact pour l'assureur (ou pour l'assuré), on peut regarder les conséquences sur le niveau des provisions (ou de la prime pure). Dans le cas des provisions, que nous calculons de manière déterministe, les flux probables des prestations à payer en j pour un individu d'âge x lors de l'année t (pour k simulations des taux de mortalité réalisés à partir de l'une des deux méthodes précédentes) sont :

$$\forall i \in \llbracket 1; k \rrbracket, \quad \begin{cases} F_{x,t}^{(i)}(j) = C \hat{q}_{x,t}^{(i)} \text{ si } j = 0 \\ F_{x,t}^{(i)}(j) = C \hat{q}_{x+j,t}^{(i)} \prod_{i=0}^{j-1} (1 - \hat{q}_{x+i,t}^{(i)}) \text{ sinon} \end{cases}$$

Avec C le montant de la prestation.

Nous obtenons donc

$$L_0^{(i)} = \sum_{t=0}^{d-1} F_{x,t}^{(i)}(j)(1 + r_{j+1})^{-t-\frac{1}{2}}$$

en supposant les décès en milieu d'année. Le paramètre r_j représente le taux d'actualisation des flux à échéance j . Enfin, nous pouvons déterminer

un coefficient de variation empirique :

$$cv = \frac{\sqrt{E[(L_0^{(i)} - L_0)^2]}}{\bar{L}_0}$$

avec

$$\bar{L}_0 = \frac{1}{k} \sum_{j=1}^k L_0^{(j)}$$

3.4 Incertitude sur la tendance : risque d'avis d'expert

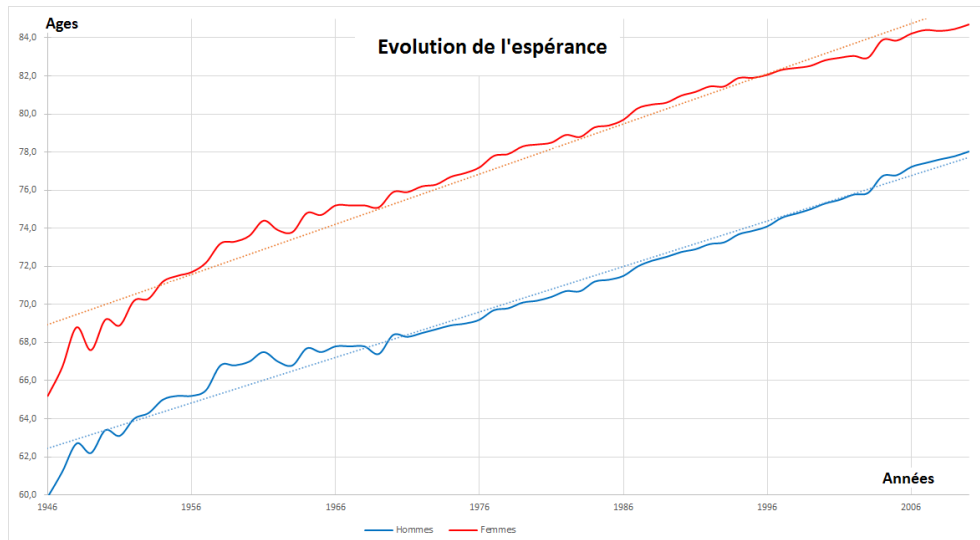


FIGURE 3.2 – Évolution de l'espérance de vie en France (Source INSEE)

Nous vivons plus longtemps que nos grands-parents. Ce constat est confirmé par les statistiques de l'INSEE qui montrent que depuis une trentaine d'années, l'espérance de vie à la naissance pour la population française présente une croissance régulière avec le gain d'environ un trimestre tous les ans, figure 3.2. Cette variation de l'espérance de vie au cours du temps est le principal argument à l'origine de la mise en place de tables de mortalité prospective pour aider l'assureur à évaluer au mieux son risque. Mais cette amélioration de l'espérance de vie ne peut être directement transposable à

un portefeuille d'assurés, car la représentation statistique par âge est différente. Par exemple, si la mortalité infantile venait à disparaître (mortalité avant 1 an) il y aurait d'après la table TH00-02 un gain de l'espérance de vie à la naissance d'environ 4 mois et demi (de 76,00 ans à 76,37 ans) alors que cela n'aura aucune incidence sur notre portefeuille car aucun adhérent n'a moins de 1 an. De plus, certaines études telles que Currie et al. [13] ont montré que le taux instantané de mortalité présente, aux différents âges, des variations erratiques autour de la tendance, variations non expliquées par les fluctuations d'échantillonnage (Cairns et al. [6]), voir Figure 3.3.

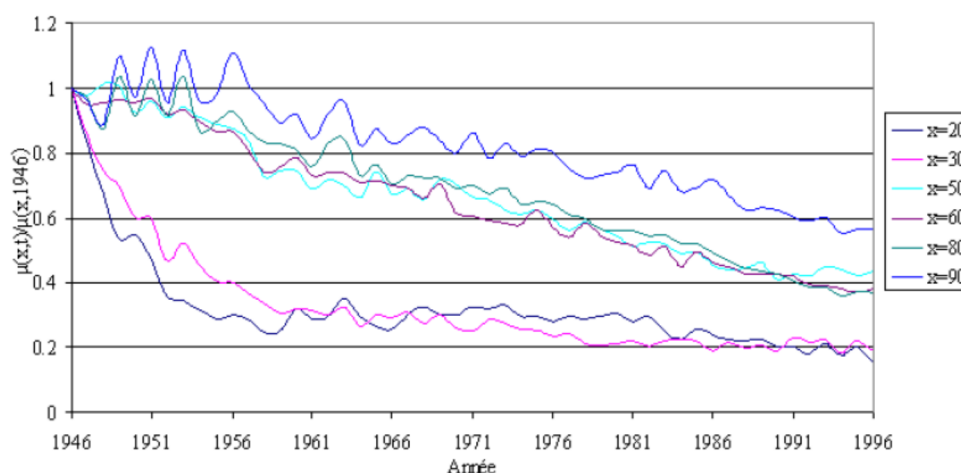


FIGURE 3.3 – Évolution des taux de mortalité instantanée, extrait de la thèse de Planchet [42]

Notons que ces variations sont non mutualisables, donc dangereuses pour l'assureur. Une méthode classique est de se référer à l'évolution des taux passés de mortalité par âge et de prolonger cette évolution dans le futur. Comme cela est décrit dans la thèse de Planchet [42], cette méthode possède plusieurs sources d'incertitude tels que le choix de la période d'observation, les fluctuations stochastiques des taux de mortalité ... et l'on constate que les différents modèles prospectifs, qui ont été réalisés sur des rentiers par cette méthode, ont tendance à sous-estimer l'augmentation de l'espérance de vie. De plus, étant confronté à une base de données peu fournie, cette méthode devient totalement inapplicable. Nous allons donc, au cours de ce chapitre,

nous baser une nouvelle fois sur une table prospective de référence, dont la tendance a pu être réalisée sur un échantillon de population beaucoup plus important.

3.4.1 Description du choix de la dérive

Comme nous l'avons vu dans l'introduction de cette section, le recours à une table externe est indispensable pour l'intégration d'une tendance à notre table du moment. La dérive, ou la tendance, est une quantité facilement calculable :

$$\forall x, \forall t \quad Tendance_{x,t} = T_{x,t} = \frac{q_{x,t+1}}{q_{x,t}}$$

La difficulté réside dans le choix de l'année de référence par rapport à laquelle la projection est effectuée. En effet prendre les mêmes années entre notre table d'expérience et celle de référence n'est pas forcément le choix le plus judicieux. Notre base de données peut-être comparée à un échantillon des individus ayant servi à la construction de la table de référence, on peut donc supposer que l'évolution de leur mortalité en est proche mais il peut néanmoins y avoir un décalage dans le temps. Nous allons donc chercher au sein de la table de référence, quelle année calendaire peut être considérée comme la plus "proche" de nos données afin d'appliquer la même tendance.

Au cours d'une note de travail, Tomas et Planchet [44] ont proposé différents outils permettant de retenir le meilleur ajustement. En partant de l'estimateur du maximum de vraisemblance :

$$\forall x, \forall t \quad q_{x,t}^{ref} = \frac{d_{x,t}^{ref}}{ER_{x,t}^{ref}}$$

- **La déviance.** C'est un indicateur de la qualité de l'ajustement. Pour un modèle de Poisson, elle se formule de la manière suivante :

$$Deviance = \sum_{(x,t)} 2 \left(d_{x,t} \ln\left(\frac{d_{x,t}}{ER_{x,t}^{ref} q_{x,t}^{ref}}\right) - (d_{x,t} - ER_{x,t}^{ref} q_{x,t}^{ref}) \right)$$

Cet indicateur doit être le plus faible possible.

- **Le test du R^2 .** Ce coefficient détermine s'il y a une corrélation linéaire entre deux séries de données. Il représente la part de variance expliquée par rapport à la variance globale :

$$R^2 = 1 - \frac{SumSquare_{err}}{SumSquare_{tot}} = 1 - \frac{\sum_{(x)} (q_{x,t}^{adj} - q_{x,t}^{ref})^2}{\sum_{(x)} (q_{x,t}^{adj} - \frac{1}{n} \sum_{(x)} q_{x,t}^{adj})^2}$$

avec n le nombre d'observations. Plus le coefficient est proche de 1 (sachant qu'il est compris entre 0 et 1) et plus la corrélation entre les deux séries de données est effective.

- **Le test du χ^2 .** C'est un test statistique permettant de tester l'adéquation entre une série de données observées et une série de données attendues.

$$\chi^2 = \sum_{(x)} \frac{(q_{x,t}^{adj} - q_{x,t}^{ref})^2}{q_{x,t}^{ref}}$$

plus le résultat de cet indicateur est faible et plus on peut considérer que les deux séries sont "proches".

- ***Mean Absolute Percentage Error (MAPE)*.** Cet indicateur mesure l'exactitude de l'ajustement par rapport aux observations, il exprime un pourcentage d'erreur de précision entre deux séries de données.

$$MAPE = \frac{1}{n} \sum_{(x)} \left| \frac{q_{x,t}^{adj} - q_{x,t}^{ref}}{q_{x,t}^{adj}} \right|$$

Ce coefficient est également à minimiser pour considérer les deux séries comme "proche".

- ***Standardized Mortality Ratio (SMR)*.** Ce dernier indicateur est un ratio spécifique à la mortalité, il représente le rapport entre le nombre de décès observés, et le nombre de décès prévus, nous allons l'adapter à notre cas de figure :

$$SMR = \frac{\sum_{(x,t)} d_{x,t}}{\sum_{(x,t)} ER_{x,t} q_{x,t}^{ref}}$$

Il doit tout naturellement être proche de 1 pour confirmer le lien entre les observations et les prédictions.

Ainsi, l'analyse de ces différents résultats nous permet d'obtenir un point de référence, correspondant à l'année qui renvoie les meilleurs résultats pour les indicateurs présentés. Ceci nous permettra d'appliquer les facteurs de dérives de la table de référence avec notre table de mortalité. Cette démarche permet donc de retenir la meilleure année de référence pour projeter les taux, ce qui permet de passer d'une table du moment à une table prospective, mais rien ne nous garantit si la projection est satisfaisante. C'est un élément qui est déterminé dans le chapitre suivant consacré à la validation de la table de mortalité.

Chapitre 4

Les méthodes d'appréciation et de validation de la table prospective à certifier

Sommaire

4.1	Introduction	55
4.2	Validation générale	56
4.3	Appréciation et validation des paramètres exo- gènes	58
4.3.1	Étude du lissage et du vecteur poids	58
4.3.2	Table de référence	61
4.4	Intervalles de confiance	61
4.5	Appréciation du risque d'estimation et du risque de tendance	63
4.6	Conclusions, limites ...	64

4.1 Introduction

Les chapitres précédents ont permis d'établir une méthode de construction de tables de mortalité prospectives, de référencer les risques qui y sont associés et d'en mesurer certains. Mais il nous reste une étape délicate qui consiste à évaluer la table construite au global, car nous avons implicitement émis l'hypothèse forte que les risques décrits sont indépendants entre eux,

ce qui n'est pas le cas. Cette évaluation se fait en plusieurs parties. Dans un premier temps nous effectuons quelques tests simples pour vérifier que la table créée suit les standards, suivis d'un rebouclage (ou *back-testing*). Dans un second temps, nous allons donner un avis sur les risques qui ont été mesurés dans le chapitre précédent, c'est-à-dire savoir s'ils sont acceptables, ou trop importants ? Nous jugeons ensuite la qualité des paramètres exogènes que l'on a choisis avec de multiples tests : cette partie est importante car elle est au cœur de la méthode de construction qui a été proposée. Enfin, nous terminons ce chapitre par une prise de recul sur la méthode en mettant en exergue les avantages ainsi que les limites de la démarche employée pour la construction de la table de mortalité et sa validation.

4.2 Validation générale

La première étape de notre procédure de validation consiste à conserver un point de vue que l'on qualifie de général, c'est-à-dire que nous testons des éléments indépendants de la méthode de construction choisie et qui se doivent d'être vérifiés sur toutes les tables de mortalité.

Les premiers éléments à regarder, et que l'on peut considérer comme problématiques s'ils ne sont pas respectés, sont :

- sur toutes les populations, on observe que les probabilités de décès augmentent avec l'âge (sauf pour les tout premiers âges où la mortalité peut être importante), nous devons donc vérifier que l'évolution des $q_{x,t}$ est croissante à t fixé :

$$\forall t, \forall x \geq 5, \quad q_{x+1,t} > q_{x,t}$$

- de la même façon, on constate une baisse des taux de mortalité, pour un âge donné, au cours du temps :

$$\forall t, \forall x, \quad q_{x,t} > q_{x,t+1}$$

D'autres éléments peuvent dépendre de la segmentation retenue, par exemple :

- nous observons également que la mortalité masculine est à tout âge supérieure à la mortalité féminine. Cet aspect doit se retrouver dans notre table :

$$\forall t, \forall x, q_{x,t}^h > q_{x,t}^f$$

Deux raisons principales peuvent expliquer le fait qu'une ou plusieurs de ces trois règles ne soient pas respectées (en supposant que la population étudiée est classique, c'est-à-dire qu'il n'y a pas de causes explicites de surmortalité d'une tranche du portefeuille).

La première raison du non-respect de l'une de ces règles peut être la présence de taux de mortalité bruts aberrants qui auraient été conservés avant d'appliquer le lissage : ceci est très fréquent lors de l'étude de petits échantillons où certains âges peuvent avoir des taux de mortalité bruts nuls ou alors très élevés. Dans cette situation, nous pouvons effectuer quelques modifications des taux bruts avant d'effectuer le lissage : par exemple, il est possible de corriger les taux aberrants (trop élevés) en les réduisant.

Une seconde explication, si nous n'avons pas de points considérés comme aberrants, est dû à un sous-lissage, c'est-à-dire que la courbe lissée est trop fidèle aux données en subissant la volatilité des taux de mortalité et donc ne capte pas la tendance. Pour contrer ce phénomène, il est possible de jouer sur le paramètre λ en faisant augmenter la taille du voisinage ou de modifier la fonction de poids. Il est également possible de modifier les paramètres propres aux méthodes de lissage telles que les valeurs de h et z pour Whittaker-Henderson.

Nous terminons cette première partie de validation par un test capital pour valider ou refuser la table construite. Il consiste en la réalisation d'un rebouclage (ou *Back-testing*) sur des données n'ayant pas servies à la construction de la table. Ce test porte sur la comparaison entre le nombre de décès observés et le nombre de décès estimés par la table de mortalité prospective : c'est-à-dire le *Standardized Mortality Ratio* vu précédemment :

$$SMR = \frac{\sum_{(x,t)} d_{x,t}}{\sum_{(x,t)} ER_{x,t} \hat{q}_{x,t}}$$

plus le résultat est proche de 1 et plus la valeur est considérée comme pertinente et donc que le modèle est fiable.

Ces premiers tests que nous présentons apportent des informations sur la qualité de la table que nous avons construite, c'est-à-dire si elle correspond aux attentes que l'on a d'une table de mortalité et si elle correspond aux données. Les différents tests qui suivent permettent d'apporter des éléments complémentaires à l'aide de mesures de cohérence et de sensibilité, pour déterminer si la table de mortalité qui a été construite est fiable ou non.

4.3 Appréciation et validation des paramètres exogènes

La méthode de construction de la table de mortalité que l'on a retenue au cours de ce mémoire fait appel à trois paramètres exogènes. Il est indispensable de mesurer l'introduction de ce risque d'avis d'expert

4.3.1 Étude du lissage et du vecteur poids

Lors de la phase de lissage, il a été retenu une certaine valeur pour λ . Il est intéressant de voir si un décalage de cette valeur modifie, et avec quelle intensité, nos résultats. Cela permet de porter un avis sur la qualité du lissage qui sera d'abord visuel. Ensuite pour confirmer ce choix de paramètres, il est réalisé une étude de trois types de résidus.

- Résidus de la réponse :

$$r_{x,t} = q_{x,t} - \hat{q}_{x,t}$$

c'est la première approche qui représente directement l'écart entre l'ajustement effectué et les observations. Ces résidus ne sont pas très pertinents si les observations présentent de l'hétéroscédasticité.

- Résidus de Pearson :

$$r_{x,t} = \frac{d_{x,t} - l_{x,t}\hat{q}_{x,t}}{\sqrt{\text{Var}(l_{x,t}\hat{q}_{x,t})}}$$

ou

$$r_{x,t} = \frac{q_{x,t} - \hat{q}_{x,t}}{\sqrt{\text{Var}(\hat{q}_{x,t})(1 - h_{i,j})}}$$

ces résidus représentent une normalisation des résidus précédents dans le but d'éliminer le phénomène d'hétéroscédasticité.

- Résidus de la déviance :

$$r_{x,t} = \text{signe}(d_{x,t} - l_{x,t}\hat{q}_{x,t})\sqrt{\text{Deviance}}$$

tout comme les résidus de Pearson, ceux de la déviance sont une mesure globale résumant la distance entre l'observé et l'estimé.

L'étude de ces trois types de résidus consiste plus particulièrement à la compréhension des résidus élevés et la succession des résidus de mêmes signes qui peuvent expliquer un surlissage. Ils doivent suivre une loi normale $\mathcal{N}(0; \sigma^2)$. D'autres tests statistiques complémentaires sont réalisés pour contrôler la qualité du lissage.

- **Le critère de régularité des taux lissés.** Ce critère mesure l'écart de valeur entre une série de points consécutifs

$$CR = \sum_{Inf(x)}^{Sup(x)-i} (\Delta^i \hat{q}_{x,t})^2$$

par exemple,

$$\left\{ \begin{array}{l} \text{si } i = 1, \quad \sum_{Inf(x)}^{Sup(x)-i} (\hat{q}_{x+1,t} - \hat{q}_{x,t})^2 \\ \text{si } i = 2, \quad \sum_{Inf(x)}^{Sup(x)-i} (\hat{q}_{x+2,t} - 2\hat{q}_{x+1,t} + \hat{q}_{x,t})^2 \end{array} \right.$$

Si l'indicateur est proche de 0, cela signifie que les écarts entre les taux sont faibles, donc plus l'indicateur est faible et plus la courbe est lissée.

- **Le test des signes.** Les individus de notre portefeuille sont supposés indépendants entre eux, cette hypothèse forte suppose une indépendance entre les $sig_{x,t} = \hat{q}_{x,t} - q_{x,t}$. Cette indépendance stipule que le nombre de signes négatifs et positifs suit une loi binomiale de paramètres (nombre de signes; $\frac{1}{2}$). Nous pouvons donc aisément vérifier l'adéquation d'une loi binomiale sur nos résultats. Si cette adéquation n'est pas vérifiée, cela provient d'une mauvaise estimation des taux de mortalité. En effet si les $sig_{x,t}$ sont plus souvent positifs (respectivement négatifs) que négatifs (respectivement positifs) cela signifie que les taux de mortalité estimés surestiment (respectivement sous-estiment) la mortalité du jeu de données.

Nous mesurons cette adéquation entre la loi théorique et le jeu de données à l'aide la *p-value*. La *p-value* est un test statistique qui représente la probabilité de rejeter à tort l'hypothèse nulle :

$$H_0 : \text{nombre de signes positifs} = \frac{\text{nombre de signes}}{2}$$

- **Le test de Wilcoxon.** Ce test consiste à intégrer au test des signes précédents la grandeur des différences. On calcule pour chaque observation l'écart entre la valeur brute et la valeur ajustée. On range ces valeurs dans un ordre croissant en valeur absolue en supprimant les éventuelles valeurs nulles. Ensuite, on associe à chaque observation son rang (ou moyenne des rangs en cas d'égalité). On note ω_+ la somme des rangs positifs, ω_- la somme des rangs négatifs et $\omega = \max(\omega_+; \omega_-)$. Nous savons que si les variables aléatoires ω_+ et ω_- suivent la même loi alors

$$\zeta = \frac{\frac{\omega_+ - 1}{2} - \frac{n(n+1)}{4}}{\sqrt{n(n+1) \frac{2n+1}{24}}}$$

suit une loi normale $\mathcal{N}(0; 1)$. Pour les mêmes raisons que pour le test précédent, l'inadéquation de la loi normale sur la variable ζ provient d'une surestimation ou d'une sous-estimation des taux de mortalité. Nous mesurerons cette adéquation entre la loi théorique et le jeu de données à l'aide la *p-value* et de l'hypothèse nulle :

$$H_0 : \zeta = 0$$

- **Le test de Wald-Wolfowitz (ou test des *runs*).** Le principe de ce test est de contrôler si les éléments d'une séquence binaire sont mutuellement indépendants. On appelle un *run* une suite d'éléments identiques. En notant $nb(0)$ le nombre de 0 observé et $nb(1)$ le nombre de 1 observé, le nombre de *runs* d'une suite uniquement composée de 0 et de 1 indépendant les uns des autres suit une loi normale de moyenne :

$$\mu = \frac{2 \, nb(0) \, nb(1)}{nb(0) + nb(1)} - 1$$

et de variance :

$$\sigma^2 = \frac{2 \, nb(1) \, nb(0)(2 \, nb(1) \, nb(0) - (nb(0) + nb(1)))}{(nb(0) + nb(1))^2(nb(0) + nb(1) - 1)}$$

Contrairement aux deux tests précédents, ce test contrôle la présence de sur ou sous lissage localement et non pas au global.

Nous mesurerons cette adéquation entre la loi théorique et le jeu de données à l'aide la *p-value* et de l'hypothèse nulle :

$$H_0 : \text{nombre de runs} = \mu$$

4.3.2 Table de référence

Le dernier paramètre exogène que l'on a introduit est le recours à une table de référence. C'est le principal pilier de la méthode de construction employée au cours de ce mémoire. Le choix auquel on fait face est limité et donc imparfait. Il est donc envisageable de reproduire la démarche de construction en utilisant d'autres tables de mortalité prospectives qui semblent moins pertinentes car plus anciennes ou issues d'une autre population, afin de visualiser l'impact d'un mauvais choix sur nos résultats. Les tests concernant la table de référence ne portent pas de réels jugements mais plus un éclairage sur la validité des méthodes relationnelles.

4.4 Intervalles de confiance

Comme cela a été présenté tout au long de ce mémoire, le recours à une méthode relationnelle et non-paramétrique présente de nombreux avantages. En plus de limiter les risques, ces méthodes présentent également l'avantage

d'obtenir des intervalles de confiance des ajustements réalisés. En effet si l'on note $\hat{y} = (\hat{f}(x_1); \dots; \hat{f}(x_n))'$ le vecteur des valeurs ajustées et $\hat{\epsilon} = (\hat{\epsilon}_1; \dots; \hat{\epsilon}_n)'$ le vecteur des résidus alors la théorie montre qu'il existe une matrice $\hat{H}$, dépendant uniquement de x et du degré de lissage λ de sorte que :

$$\hat{y} = \hat{H} y \text{ et } \hat{\epsilon} = (1 - \hat{H}) y$$

Ces deux éléments obéissent aux lois normales multivariées de matrice de variance-covariance $\sigma^2 H H'$ et $\sigma^2 (1 - H)(1 - H)'$.

Si on note $h_{i,j}$ les éléments constitutifs de $(1 - H)(1 - H)'$ et

$$\delta_k = \sum_i h_{ii}^k$$

avec $k \in \{1; 2\}$. Alors on montre facilement que

$$E[\sum_i \hat{\epsilon}_i^2] = \sigma^2 \delta_1 \Rightarrow \hat{\sigma}^2 = \frac{1}{\delta_1} \sum_i \hat{\epsilon}_i^2$$

Nous avons donc un estimateur sans biais de la variance et

$$\frac{\delta_1^2}{\delta_2} \times \frac{\hat{\sigma}}{\sigma}$$

suit approximativement une loi du χ^2 avec $\frac{\delta_1^2}{\delta_2}$ degrés de liberté.

Et nous obtenons que

$$\frac{\hat{f}(x_i) - f(x_i)}{\hat{\sigma} \sqrt{\sum_{j=1}^n h_{ij}^2(x)}}$$

suit approximativement une loi de Student à $\frac{\delta_1^2}{\delta_2}$ degrés de liberté. Nous pouvons ainsi obtenir les intervalles de confiance de la fonction f .

Enfin, on obtient

$$\hat{f}(x_i) = \sum_{i=0}^n \hat{\beta}_i(x_i) x_i^i$$

en minimisant la fonction :

$$\sum_{k \in V(x_i)} w_k(x_i) (y_k - \sum_{j=0}^n \hat{\beta}_j(x_i) x_k^j)$$

4.5 Appréciation du risque d'estimation et du risque de tendance

Comme nous l'avons vu lors de la partie précédente, l'application des différentes méthodes d'échantillonnage nous permet de calculer le coefficient de variation. Sa valeur, plus ou moins importante apporte un premier élément sur l'importance de ce risque. La distribution des taux de mortalité qui ont été simulés pour l'obtention de ce coefficient de mortalité suit une loi normale, ce qui nous permet d'obtenir les intervalles de confiance des taux de mortalité et donc du risque d'estimation :

$$IC(\alpha) = [E[q_{x,t}^{(k)}] (1 - u \frac{cv_{x,t}}{\sqrt{n_x}}) ; E[q_{x,t}^{(k)}] (1 + u \frac{cv_{x,t}}{\sqrt{n_x}})]$$

avec u le quantile d'ordre $1 - \frac{\alpha}{2}$ de la loi $\mathcal{N}(0; 1)$ et n_x l'exposition à l'âge x .

Ce résultat théorique est à comparer avec les résultats empiriques que l'on peut obtenir sur les taux de décès ainsi que sur les provisions et qui tiennent uniquement compte de ce risque d'échantillonnage. En effet, empiriquement nous avons :

$$IC(\alpha) = [Inf\{q_{x,t}^i \text{ simulés} / P(q_{x,t}^{(k)} \leq q_{x,t}^i) \geq \alpha\} ; Inf\{q_{x,t}^i \text{ simulés} / P(q_{x,t}^{(k)} \leq q_{x,t}^i) \geq 1 - \alpha\}]$$

Que faire dans le cas où la taille des IC est considérée comme trop importante et donc que les taux de mortalité q_x ne sont pas suffisamment fiables ? Différents éléments peuvent expliquer l'apparition d'intervalles de confiance non-satisfaisants. Le principal est le fait que la base de données est trop réduite et entraîne une importante volatilité des taux de mortalité. Une seconde explication peut provenir de la présence de données de mauvaises qualités, qui faussent les résultats. C'est une conséquence du risque opérationnel, une meilleure étude des données est à envisager afin d'éliminer des données erronées.

Concernant le risque de tendance, Planchet a montré au cours de sa thèse [42] qu'une mauvaise évaluation de celle-ci peut avoir d'importante répercussion sur les engagements de l'assureur. L'étape de projection propose

la meilleure tendance en fonction des données et de la table de référence, mais en aucun cas ne permet de valider cette projection. Pour cela, nous proposons de regarder les conséquences de la projection sur deux indicateurs démographiques : l'espérance de vie à chaque âge, et l'entropie à chaque âge. Les résultats obtenus sont comparés aux prévisions attendues par des experts, ces tests restent donc très subjectifs.

- Le calcul de l'espérance de vie à chaque âge, c'est-à-dire la durée de vie moyenne restante à chaque âge et pour chaque année, que l'on peut comparer avec ce qui est attendu par les experts :

$$e_{x,t} = \frac{\sum_{k=1}^{+\infty} l_{x+k,t+k}}{l_{x,t}}$$

- L'approche est identique pour l'entropie, qui peut s'interpréter comme un gain relatif d'espérance de vie si l'on a survécu à l'âge précédent, voir Delwarde et Denuit [15] :

$$entropie_{x,t} = - \frac{\sum_{k=1}^{+\infty} l_{x+k,t+k} \ln(l_{x+k,t+k})}{\sum_{k=1}^{+\infty} l_{x+k,t+k}}$$

4.6 Conclusions, limites ...

En conclusion de cette partie théorique, il est important de rappeler que cette démarche, qui consiste à lisser notre fonction de passage entre les taux de référence et les taux bruts, a le mérite de proposer un angle d'approche nouveau et adapté à la réalisation d'une table de mortalité prospective dans le cadre de petits échantillons. Elle n'est applicable que depuis l'émergence des outils informatiques. Comme cela a été décrit au cours de ce mémoire, les trois méthodes développées (moyennes mobiles, Whittaker-Henderson et Loess) ne posent pas de modèle a priori, contrairement aux méthodes relationnelles classiques. L'ensemble des méthodes relationnelles s'appuie au maximum sur la table de référence retenue, qui est l'élément primordial. Cette table de référence a le mérite d'exister et d'avoir fait ces preuves, elle

semble donc être une excellente base à la construction de nouvelles tables adaptées à des données plus précises. Quelques légères réserves peuvent néanmoins être émises sur la démarche proposée, notamment concernant le choix des paramètres de lissage et de la fonction de poids. En effet, les possibilités sont infinies et peu d'éléments peuvent venir dicter ce choix qui est laissé à l'appréciation de l'expert. De plus, la méthode de fermeture de la table vient contraster avec la démarche de lissage que l'on propose car il y a nécessité de poser un modèle, chose que l'on a cherché à éviter en appliquant des modèles relationnels non-paramétriques.

Concernant les tests qui ont été proposés dans ce quatrième chapitre sur la validation de la table, ils sont variés dans leur approche et s'attaquent à des angles de validation différents et complémentaires. On note qu'ils n'ont pas tous la même importance dans le choix final de validation ou non de la table, mais ils apportent tous des informations qui bout à bout donnent une assez bonne appréciation de la qualité de la table de mortalité construite. La partie suivante, la partie pratique va permettre de mettre en application l'ensemble de la procédure décrite précédemment, et surtout de mettre en place les tests détaillés au cours de ce chapitre, et de voir dans une situation concrète quelles conséquences on peut en tirer.

Deuxième partie

Partie Application

Chapitre 5

Etude de la base de données

Sommaire

5.1	Introduction	67
5.2	Présentation du jeu de données	68
5.3	Caractéristiques de la base de données	69
5.4	Evolution des variables dans le temps	72
5.4.1	Évolution de la parité Homme/Femme au cours du temps	72
5.4.2	Moyenne des âges d'exposition et de décès au cours du temps	73
5.5	Conclusion	74

5.1 Introduction

En assurance vie, les tables de mortalité sont utilisées pour de multiples travaux et les objectifs d'usages sont bien différents. Qu'il s'agisse de tarification, de provisionnement, de modélisations ... une meilleure connaissance des particularités de la mortalité des assurés par rapport à la population générale et de l'impact de la nature des garanties sur le niveau de la mortalité est essentielle. Pour cela et comme pour tous les travaux ayant recours à l'utilisation d'une base de données, la construction d'une table de mortalité nécessite que celle-ci soit de très bonne qualité.

Son étude est donc primordiale pour la suite. Au cours de ce chapitre, le lecteur est mis à la place d'un actuaire certificateur, cela impose une

description exhaustive de la base de données utilisées : son origine, sa période d’observation, ses différentes caractéristiques ... Ensuite, il est décrit l’ensemble des traitements appliqué en vue de l’étude qui va suivre : suppression des données en double, questionnement sur les données incomplètes ou inutiles ...

5.2 Présentation du jeu de données

La base de données qui est utilisée au cours de cette étude est une agrégation de différents portefeuilles d’assurance concernant de multiples types de contrats tels que des contrats d’épargne, d’emprunts, de retraites, de temporaires décès ... Une remarque importante est que l’intégration des différentes couvertures dans le portefeuille se fait à des dates différentes. Cet élément peut être gênant en pratique pour suivre l’évolution du portefeuille dans le temps car cela a pour conséquence de rendre la base inhomogène, mais dans notre cas cela n’a pas d’incidence, ce que nous expliquons par la suite.

L’ensemble de ces données est issu d’observations comprises entre le 01/01/2004 et le 31/12/2009. On note la présence d’une troncature (à gauche), par exemple pour les contrats ayant commencé avant le 01/01/2004, ainsi que d’une censure (à droite), notamment pour les contrats toujours en cours le 31/12/2009. La période d’observation est relativement courte (6 ans), on peut donc utiliser l’ensemble des données pour construire une table du moment car la dérive due à l’allongement de l’espérance de vie est considérée comme négligeable sur ce court laps de temps.

La base de données est issue de travaux réalisés pour un projet de statistiques en 2011 [36] dont le rapport nous sert d’appui pour la compréhension des données. Lors de cette étude, de nombreux retraitements ont été menés pour obtenir la base finale sur laquelle nous allons nous appuyer, dont notamment :

- suppression des individus sans date de naissance et ceux dont la relation d’ordre entre la date de naissance, la date d’entrée et la date de sortie ne sont pas correctes ;
- suppression des individus qui sont sortis ou sinistrés avant la période d’observation (du 01/01/2004 au 31/12/2009) ;

- uniformisation des formats de variables tels que le sexe, le statut ...
- calcul des variables "AgeEntree" et "AgeSortie".

Cette étude initiale avait pour objet la comparaison de la mortalité entre deux populations très différentes, en l'occurrence la population française avec la population de différents pays africains et notamment en mesurant l'hétérogénéité via différentes approches. Pour notre étude, il a été conservé uniquement les contrats représentant la population française.

Chaque ligne de la base de données est renseignée de plusieurs informations, la suite de l'étude utilise uniquement les variables suivantes :

- l'identifiant du contrat, il représente le numéro de dossier du contrat qui permettra la détection de doublon,
- le sexe de l'assuré, féminin ou masculin,
- la date de naissance de l'assuré,
- la date de début d'observation du contrat,
- la date de fin d'observation du contrat,
- le statut, c'est-à-dire si la sortie est due à un décès ou à une censure.

5.3 Caractéristiques de la base de données

Sur les informations fournies par ces données, différents retraitements ont été effectués tels que la suppression des doublons (même identifiant, même sexe et même date de naissance) ou de lignes dont l'ensemble des informations ne sont pas renseignées. On rappelle que notre base de données est issue de l'agrégation de plusieurs portefeuilles, il est donc possible qu'une même personne ayant souscrit à différents types de contrats apparaisse plusieurs fois sous des numéros de contrats différents. Selon notre sélection, elle ne sera pas considérée comme un doublon et viendra biaiser nos résultats. Ce facteur du risque opérationnel n'a pas d'incidence sur l'objet de cette étude. En effet, elle n'a pas vocation à réaliser une table parfaite, mais seulement l'utilisation de données réalistes pour la comparaison de méthode de lissage et de tests de validation, quelle que soit la qualité de la base de données. Suite à ce retraitement, nous obtenons une base constituée d'environ 1,8 million de lignes.

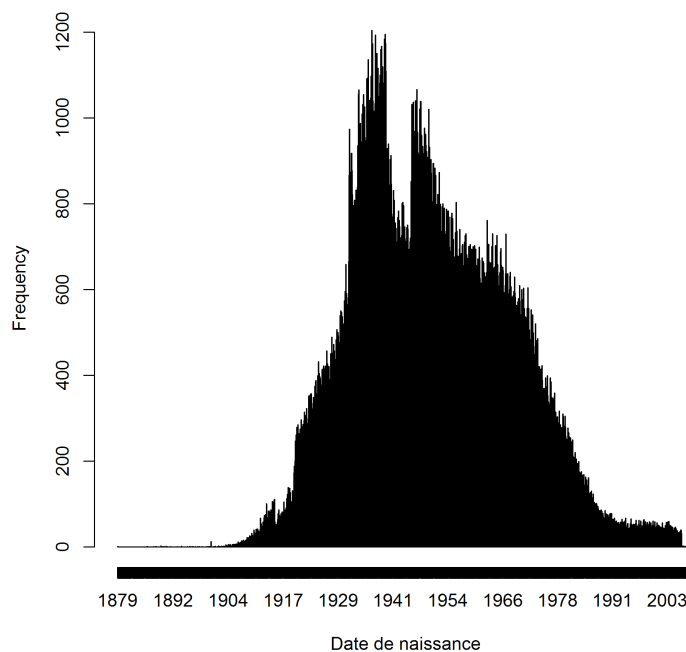


FIGURE 5.1 – Représentation de la répartition des dates de naissance du portefeuille

Afin d'anticiper sur l'étape du *back-testing*, seulement une partie de ces données doivent être utilisée pour la construction de la table de mortalité. Généralement, on ne conserve que un tiers des données pour la construction de la table. Les données servant directement à la construction de la table de mortalité sont sélectionnées aléatoirement. La base de données finales est ainsi composée de 500 000 lignes tirées aléatoirement dans la base de données (fonction "sample" du logiciel R) avec les caractéristiques suivantes :

- la base de données est constituée de 56,3 % de femmes (soit 43,7 % d'hommes),
- les dates de naissance des assurés sont comprises entre le 26/01/1880 et le 22/10/2009, et ne semblent pas présenter de date préférentielle comme le montre la figure mensuelle 5.1. Néanmoins, on observe un léger "trou" entre les dates de 1941 à 1945 que l'on explique par la

baisse de la natalité lors de la Seconde Guerre mondiale.

- les dates d'entrée et de sortie sont comprises entre le 01/01/2004 et le 31/12/2009 soit la taille de la période d'observation, la figure 5.2 présente la répartition des âges d'entrée et de sortie. Elles n'impliquent pas de remarque particulière.
- la base est constituée de 15 790 lignes avec un statut décès, soit 3,2 % des lignes.

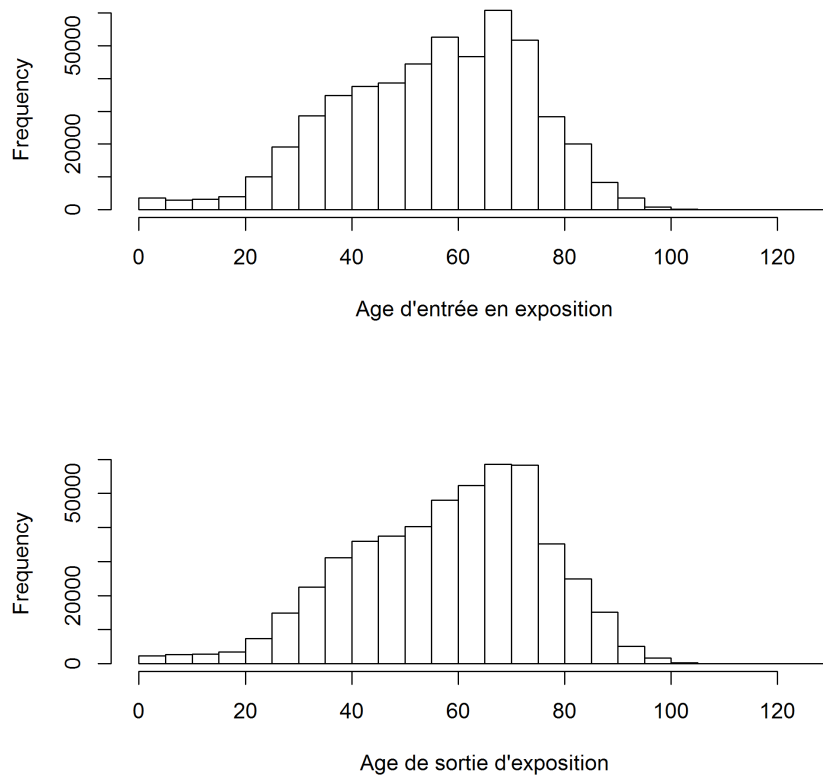


FIGURE 5.2 – Représentation des âges d'entrée et de sortie du portefeuille

Pour avoir une représentation visuelle des données fournies par la base de données, l'exposition et le nombre de décès par âge ont été représentés sur la figure 5.3. L'exposition est calculée au prorata-temporis, c'est-à-dire qu'une personne présente du 01/07/2005 au 31/03/2006 comptabilisera une

exposition de 0,5 pour l'année 2005 et une exposition de 0,25 pour l'année 2006.

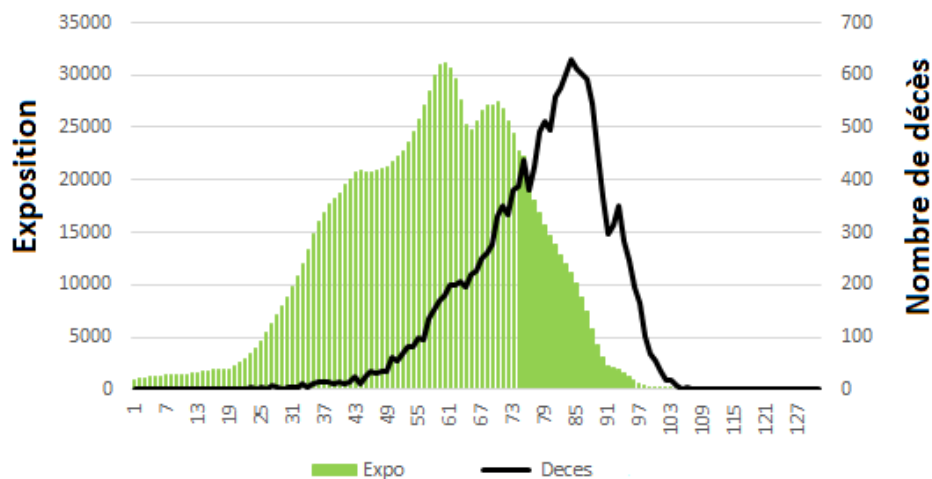


FIGURE 5.3 – Représentation de l'exposition et du nombre de décès par âge

5.4 Evolution des variables dans le temps

Un point important pour juger la qualité d'une base de données est d'observer l'évolution des différents attributs au cours du temps. En effet, pour obtenir des résultats fiables, il y a nécessité de vérifier que la table est stable au cours du temps.

5.4.1 Évolution de la parité Homme/Femme au cours du temps

Lors de la construction d'une table de mortalité, le caractère homme ou femme est très important, car tout au long de leur vie, la mortalité des femmes est plus faible que celle des hommes. La réglementation européenne interdit aux assureurs d'effectuer une différenciation de tarif en fonction du sexe, néanmoins cette segmentation est encore utilisée, par exemple pour le calcul des provisions. Ainsi, il est toujours intéressant de savoir si la proportion homme/femme est stable au cours du temps, ce qui pourrait fausser les résultats dans le cas contraire. Pour notre cas pratique on ne se base pas sur un unique portefeuille, mais sur un agglomérat de portefeuilles, on

constate ainsi que la proportion homme/femme possède une année d'observation atypique (année 2006) avec une surexposition de 13 % de la proportion de femmes (voir figure 5.4), principalement à cause de l'entrée au sein du portefeuille des contrats de type retraite au cours de l'année 2006, où la gent féminine est fortement représentée.

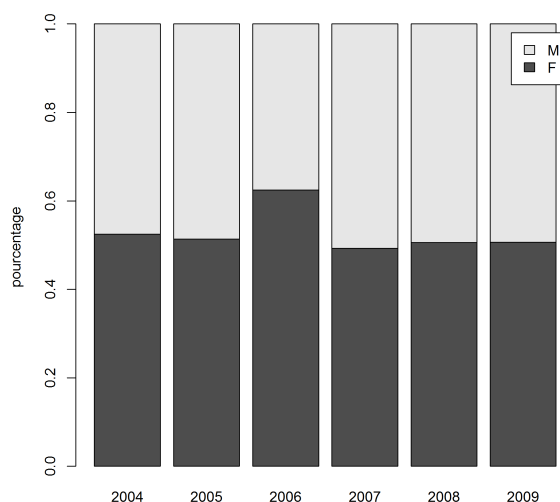


FIGURE 5.4 – Evolution annuelle de la proportion homme/femme du portefeuille

5.4.2 Moyenne des âges d'exposition et de décès au cours du temps

Pour les mêmes raisons que pour la parité homme/femme, il est intéressant de connaître l'âge moyen d'exposition de notre base de données. Le haut de la figure 5.5 présente cette donnée sur l'ensemble des individus de la base, sans distinction de sexe, au début de chaque mois. Cet âge moyen se situe entre 52 et 59 ans selon les années, la variation est expliquée par l'entrée massive de contrats à certaines dates, comme cela a été décrit en début de chapitre. Par souci d'exhaustivité de l'étude, nous avons également regardé l'évolution de l'âge de moyen de décès qui est assez régulier dans le temps : compris entre 74 et 77 ans, voir le bas de la figure 5.5.

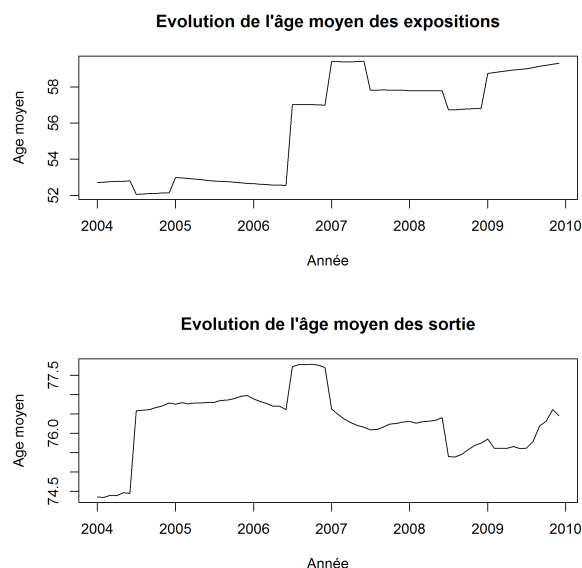


FIGURE 5.5 – Evolution de l'âge moyen d'entrée et de sortie du portefeuille

5.5 Conclusion

Nous avons pu remarquer au cours de ce chapitre, consacré à l'étude de la base de données, que nous sommes en possession d'une base de données finale très riche (1,8 million de lignes comportant l'ensemble des informations nécessaires à l'étude de la mortalité). Elle couvre l'ensemble des âges de 1 à 120 ans, même si l'exposition ne semble exploitable qu'entre les âges 25 à 90 ans où elle dépasse les 3 000 contrats. Les différentes vérifications que l'on a présentées n'ont pas montré de pics particuliers sur la répartition des dates où sur celles des âges d'entrée et de sortie d'exposition qui auraient mérité une étude approfondie. Par contre, un point très important est que l'évolution de la base d'étude n'est pas uniforme dans le temps, que ce soit au niveau de la proportion homme/femme ou de l'âge moyen d'entrée et de sortie d'exposition. Mais cet élément, qui est très gênant dans une situation réelle où la table construite a vocation à être appliquée, a beaucoup moins d'importance dans notre situation. En effet, on cherche à appliquer des méthodes et non pas obtenir des résultats chiffrés.

Chapitre 6

Construction d'une table de mortalité prospective

Sommaire

6.1	Introduction	75
6.2	Calcul des taux de mortalité bruts	76
6.3	Lissage des taux bruts	78
6.3.1	Choix des paramètres exogènes	78
6.3.2	Application des différentes méthodes de lissage non-paramétriques	80
6.3.2.1	Lissage par la méthode des moyennes mobiles	80
6.3.2.2	Lissage par la méthode de Whittaker-Henderson	81
6.3.2.3	Lissage par la méthode de Loess	81
6.3.3	Fermeture de la table	82
6.4	Projection des taux de la table du moment . . .	83
6.5	Conclusion	84

6.1 Introduction

Ce chapitre a pour finalité de passer de la base de données que l'on a étudiée au cours du chapitre précédent à une table de mortalité prospective segmentée par sexe. Pour cela, on rappelle que la démarche se décompose

en plusieurs étapes. Dans un premier temps, on extrait des données les taux de mortalité bruts, par la méthode de Hoem ou de Kaplan-Meier, sans distinction d'années, décrit en annexe. Ces taux sont ensuite mis en relation avec leur homologue de la table de référence, la table "INSEE 2007-2060", par rapport à l'âge, après avoir déterminé l'année de référence la plus adaptée. C'est la fonction de passage entre les taux de référence et les taux bruts que l'on est amenée à lisser. Les méthodes de lissage sont non-paramétriques et donc nécessitent le recours à deux paramètres exogènes supplémentaires (le paramètre de lissage et la fonction de poids). À l'issue de ce lissage, on obtient une table de mortalité du moment que l'on projette sur quelques années en appliquant directement la dérive observée sur la table de mortalité de référence à partir de l'année que l'on a considérée d'optimale.

6.2 Calcul des taux de mortalité bruts

Pour l'élaboration de la table de mortalité, nous avons calculé les taux bruts à l'aide de deux méthodes différentes, la méthode de Hoem et la méthode de Kaplan-Meier, qui sont décrites en annexe. Ces calculs se sont fait indépendamment du temps, uniquement en fonction de l'âge. Les figures 6.1 et 6.2 présentent ces résultats par âges, respectivement pour les hommes et pour les femmes. Les deux méthodes montrent quelques différences à certains âges. Ces écarts sont dus la répartition de l'exposition au cours de l'année. En effet, nous avons pour certaines années une importante intégration de contrats en milieu d'année, par exemple les contrats temporaires décès ont une date d'observation qui commence le 01/07/2006 et surtout une mortalité différente des contrats précédents. Ce phénomène biaise beaucoup plus les résultats fournis par la méthode de Hoem, ainsi pour la suite de l'étude, il a été conservé les résultats de la méthode de Kaplan-Meier.

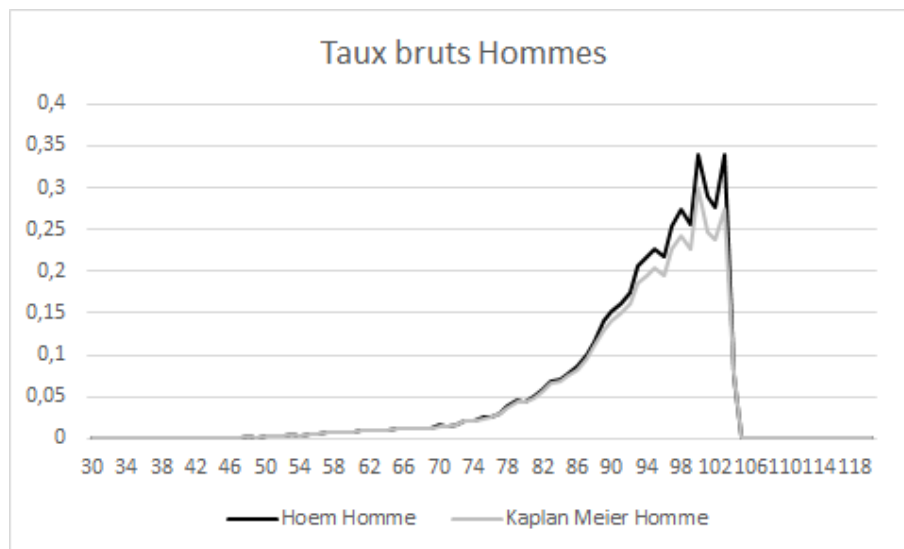


FIGURE 6.1 – Représentation des taux bruts masculins par âges, méthode Kaplan-Meier et méthode Hoem

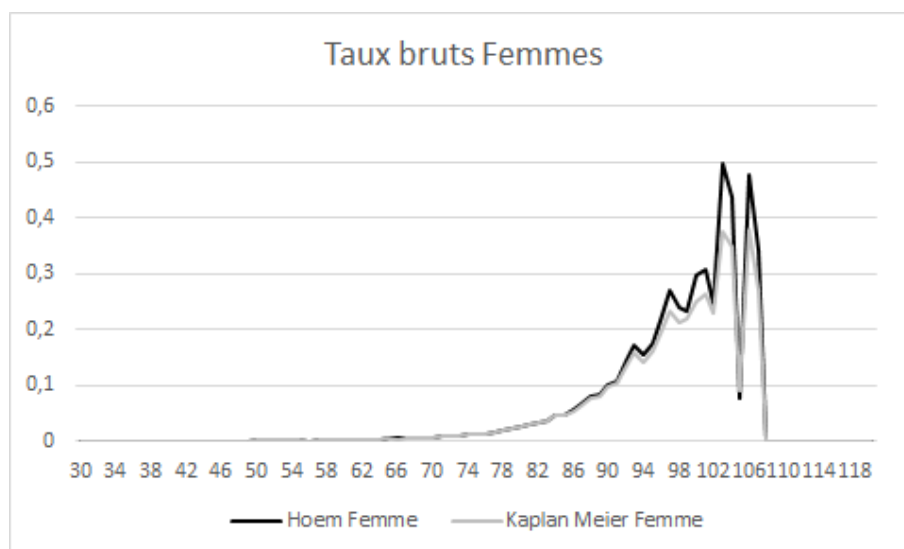


FIGURE 6.2 – Représentation des taux bruts féminins par âges, méthode Kaplan-Meier et méthode Hoem

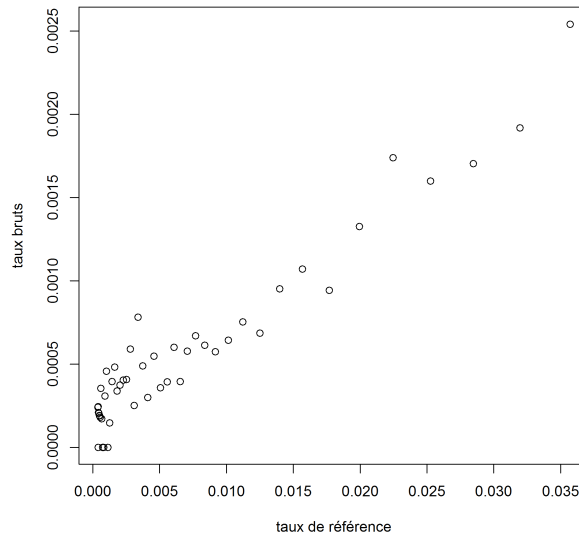


FIGURE 6.3 – Mise en relation des taux de sortie bruts avec les taux de sortie de référence pour les âges compris entre 35 et 80 ans, sur la population masculine

6.3 Lissage des taux bruts

Les taux bruts obtenus par la méthode de Kaplan-Meier montrent une certaine volatilité, ce qui nécessite un lissage. Les taux bruts calculés présentent peu ou pas de décès au niveau des âges faibles, de plus la table de référence que l'on a retenu ("INSEE 2007-2060") commence à partir de l'âge de 35 ans : il a donc été décidé de construire les tables de mortalité qu'à partir de cet âge (35 ans). On rappelle que le lissage est appliqué sur la fonction de passage entre les taux de mortalité bruts et les taux de mortalité de l'année retenue de la table de référence prospective et non pas directement des taux de mortalité bruts, voir figure 6.3.

6.3.1 Choix des paramètres exogènes

Comme cela a été décrit au cours de la partie théorique, la démarche employée pour le lissage nécessite trois paramètres exogènes.

Le premier paramètre que l'on sélectionne est la table de référence prospective et de l'année de référence. Le choix s'est rapidement tourné vers la

table "INSEE 2007-2060" pour deux raisons principales : il s'agit d'une table de mortalité prospective très récente et basée sur la population française. Le couple de tables TGH05 et TGF05 aurait également pu prétendre à être choisie, mais de récents travaux menés par Julien Tomas et Frédéric Planchet ont montré que l'utilisation de cette table est possible dans le cas de rente mais pas pour des risques en cas de décès car elle est erratique [44]. La table de référence générationnelle "INSEE 2007-2060", est composée de 56 tables du moment. Pour appliquer les lissages, il a été choisi de se référer à la même année calendaire (année 2007) que pour la projection de la tendance. Les tests appliqués pour le choix de l'année calendaire optimale sont détaillés au sein du chapitre concernant la sélection de la tendance page 83.

Le second paramètre que l'on sélectionne est la fonction de poids. N'ayant pas de raisons évidentes et objectives de préférer une fonction à une autre, il a été choisi de prendre une des fonctions les plus simples qui à la particularité d'offrir un poids proportionnel à la distance. C'est la fonction de poids triangulaire :

$$W = (1 - u) I_{\{u \leq 1\}}$$

Le troisième paramètre à définir est le paramètre de lissage λ . Nous avons appliqué les différentes méthodes de lissage (moyennes mobiles, Whittaker-Henderson et Loess) en faisant varier le paramètre de lissage afin de calculer les différents AIC modifiés définis page 37.

Par exemple pour le lissage par moyennes mobiles appliqué sur les taux bruts hommes et femmes, un extrait des AIC modifiés obtenus est présenté dans le tableau ci-dessous :

λ	7	8	9	10	11	12	13
AICC1 Hommes	-22,11	-29,53	-30,11	-31,27	-30,34	-29,38	-29,02
AICC1 Femmes	-52,28	-59,71	-60,28	-61,43	-60,50	-59,54	-59,19

TABLE 6.1 – Extrait du calcul des AIC modifiés

La valeur de λ correspond au nombre de points les plus proches que l'on sélectionne.

Comme nous le montre cet extrait du tableau, le AIC modifié minimum est atteint pour un paramètre de lissage égal à 10. Il est intéressant de noter que, quelle que soit la méthode de lissage appliquée et quelle que soit la courbe (représentée par les deux sexes), le paramètre de lissage est le même. Ceci nous laisse supposer que ce paramètre ne dépend pas ou peu de ces éléments, mais son optimisation se situe plutôt dans la proportion de points qu'il utilise : dans notre cas nous avons 45 points et $\lambda = 10$ soit environ 22 % des points. Des études complémentaires sont nécessaires pour confirmer ou infirmer cette remarque.

Une fois ces trois paramètres exogènes sélectionnés, nous pouvons appliquer les méthodes de lissage non-paramétriques définies au cours de la partie théorique.

6.3.2 Application des différentes méthodes de lissage non-paramétriques

Les trois méthodes de lissage non-paramétriques présentées ci-dessous ont été codées sous le logiciel *R*. On présente les résultats graphiques obtenus en utilisant les paramètres exogènes définis lors de la section précédente et les résultats chiffrés sont disponibles en annexe. On rappelle que le lissage est réalisé sur la fonction de passage entre les taux de référence et les taux bruts et non pas directement sur les taux bruts, ceci pour répercuter sur nos taux les variations de la table de référence.

6.3.2.1 Lissage par la méthode des moyennes mobiles

Ce premier lissage est effectué par la méthode des moyennes mobiles, pour les hommes et pour les femmes. Les graphiques comparent les taux bruts issus de Kaplan-Meier, les taux lissés et les taux issus de la table de référence pour les âges compris entre 35 et 80 ans. La figure 6.4 présente le résultat de ce lissage.

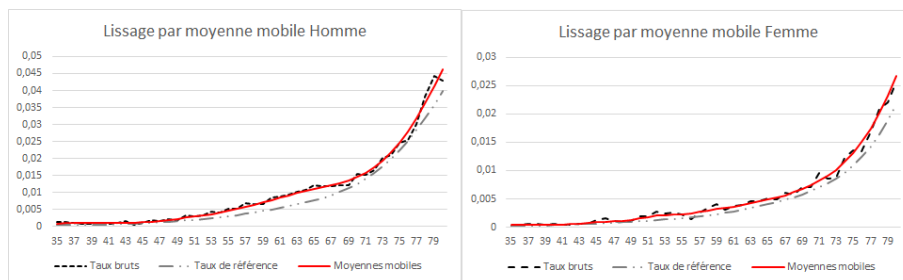


FIGURE 6.4 – Courbes issues du lissage par moyennes mobiles

6.3.2.2 Lissage par la méthode de Whittaker-Henderson

Ce second lissage est effectué par la méthode de Whittaker-Henderson (avec les paramètres propres classiques $h = 30$ et $z = 2$), pour les hommes et pour les femmes. Les graphiques comparent les taux bruts issus de Kaplan-Meier, les taux lissés et les taux issus de la table de référence pour les âges compris entre 35 et 80 ans. La figure 6.5 présente le résultat de ce lissage.

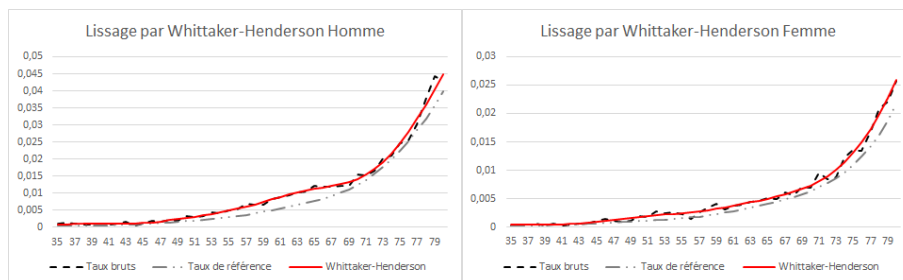


FIGURE 6.5 – Courbes issues du lissage Whittaker-Henderson

6.3.2.3 Lissage par la méthode de Loess

Le dernier lissage est effectué par la méthode de Loess (s'appuyant sur des polynômes de degré 2), pour les hommes et pour les femmes. Les graphiques comparent les taux bruts issus de Kaplan-Meier, les taux lissés et les taux issus de la table de référence pour les âges compris entre 35 et 80 ans. La figure 6.6 présente le résultat de ce lissage.

L'étude de ces courbes ne semble pas montrer de gros écarts entre les différentes méthodes de lissage que ce soit pour les hommes ou les femmes (les valeurs chiffrées de ces résultats sont disponibles en annexe). Cela est

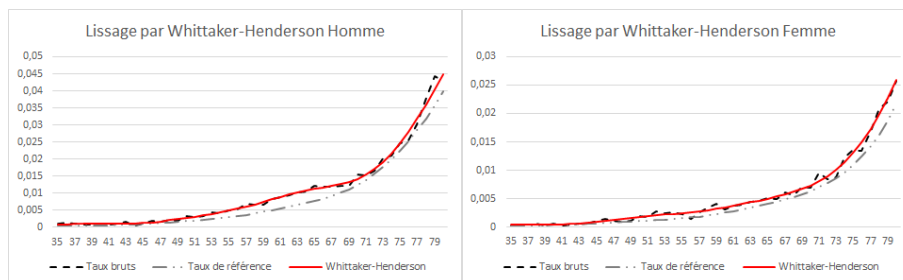


FIGURE 6.6 – Courbes issues du lissage par Loess

principalement dû à deux effets différents. Tout d’abord, les taux bruts possédaient déjà une certaine régularité (faible volatilité) car l’exposition sur la tranche d’âge 35-80 ans est importante (plus de 15 000 contrats). Le fait d’appliquer le lissage sur la fonction de passage et non pas directement sur les taux bruts vient également atténuer les différences entre les méthodes de lissage non-paramétriques.

6.3.3 Fermeture de la table

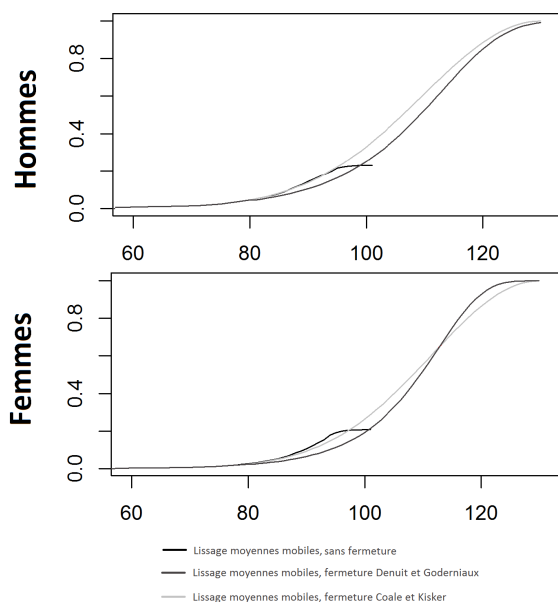


FIGURE 6.7 – Fermeture des taux lissés

Les tables lissées que l’on a obtenues ne pouvaient pas être fiables pour

les âges élevés (au-delà de 80 ans) à cause de la trop faible exposition dans notre base de données. Comme cela a été détaillé au cours de la partie théorique, il a été mis en place les méthodes de Denuit et Goderniaux et de Coale et Kisker. On présente ci-dessous les résultats graphiques pour chaque méthode de lissage et par sexe.

Les résultats graphiques de la figure 6.7, qui est représentative des trois méthodes de lissage, montrent que sur les âges 80-90 ans, la méthode de Denuit et Goderniaux est beaucoup plus proche des données que la méthode Coale et Kisker. C'est donc cette méthode que l'on retient pour la suite de la construction et notamment pour la projection des taux à venir.

6.4 Projection des taux de la table du moment

La projection des taux nécessite à nouveau le recours à la table de référence. Les tables 6.2 et 6.3 présente les résultats des différents tests statistiques détaillés lors de la partie théorique afin de sélectionner l'année optimale pour la projection. Le premier tableau concerne les femmes, le second concerne les hommes. Les années étudiées ont été restreintes aux années disponibles (à partir de 2007) et proches de la période d'observation.

Femmes	2007	2008	2009	2010	2011
Déviance	416	391	366	342	318
R^2	0,8547	0,8445	0,8338	0,8225	0,8107
χ^2	358	385	415	446	478
MAPE	0,4822	0,4948	0,5075	0,5213	0,5361
SMR	1,386	1,4057	1,4245	1,4444	1,4652

TABLE 6.2 – Résultats des tests de sélection de l'année de référence (1/2)

Pour la population masculine, le choix de l'année calendaire 2007 est clair. Ce choix produit la déviance, le χ^2 et le MAPE les plus faibles, le R^2 le plus élevé et le SMR le plus proche de 1. Concernant la population féminine, le choix a été également de retenir l'année calendaire 2007, malgré le fait que l'on ne minimise pas la déviance, l'optimisation des autres paramètres

Hommes	2007	2008	2009	2010	2011
Déviance	664	712	763	816	872
R^2	0,8808	0,8716	0,8618	0,8515	0,5406
χ^2	798	862	930	1002	1077
MAPE	0,2582	0,2652	0,2722	0,2792	0,2861
SMR	1,386	1,404	1,423	1,443	1,464

TABLE 6.3 – Résultats des tests de sélection de l'année de référence (2/2)

est respecté. Un exemple d'application des tendances issues des tables de référence à partir de l'année 2007 est présenté sur la figure 6.8 pour le lissage par moyennes mobiles, les résultats pour les autres lissages sont présentés en annexe.

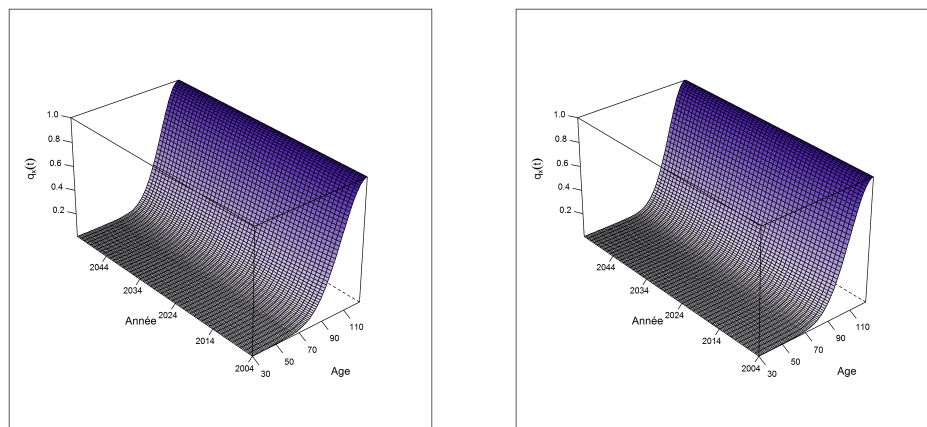


FIGURE 6.8 – Projection des taux lissés par la méthode des moyennes mobiles sur la population masculine et féminine

6.5 Conclusion

Dans un premier temps, au cours de ce chapitre, nous avons calculé les taux bruts déterminés par la méthode de Kaplan-Meier (voir annexes). Le fait d'avoir un nombre limité d'observations, nous oblige à appliquer un lissage pour améliorer la qualité des taux de mortalité en réduisant la volatilité. Cette étape a donc été réalisée suivant trois méthodes de lissage différentes,

par moyennes mobiles, Whittaker-Henderson et Loess, et l'utilisation d'une table de référence, en faisant une distinction par sexe. Les résultats obtenus sont visuellement satisfaisants sur les âges présentés, entre 35 et 80 ans. L'étape suivante a consisté dans le prolongement des taux au-delà de 80 ans. C'est la méthode de Denuit et Goderniaux qui apporte les meilleurs résultats (par rapport à la méthode de Coale et Kisker) en venant mieux se superposer à la courbe des taux lissés sur les âges 80 à 90 ans, pour chaque méthode. Enfin, nous avons terminé ce chapitre par la projection des taux sur 54 années grâce à la table de référence "INSEE 2007-2060" et l'année considérée comme optimale : l'année 2007. Toutefois, on pose une réserve sur cette année car c'est l'année minimale de la table "INSEE 2007-2060" et qui se situe au milieu de notre période d'observation, et il aurait été intéressant de tester des années plus anciennes. À l'issue de cette procédure, nous obtenons une table de mortalité prospective de 2007 (milieu de la période d'observation) à 2060, par sexe et de 35 à 120 ans.

Chapitre 7

Application des outils de validation de table de mortalité

Sommaire

7.1	Introduction	86
7.2	Le risque d'estimation	87
7.2.1	Ré-échantillonnage par simulation directe	87
7.2.2	Ré-échantillonnage par simulation des résidus	88
7.3	Le risque de tendance	89
7.4	Validation générale	91
7.5	Appréciation des paramètres exogènes	93
7.5.1	Analyse du paramètre de lissage	93
7.6	Conclusion	95

7.1 Introduction

Ce dernier chapitre est consacré à la validation des tables du chapitre précédent. La majorité des tests se veulent indépendants de la procédure de construction employée. Dans un premier temps, il est analysé deux des principaux risques sous-jacents à la construction d'une table de mortalité prospective : le risque d'estimation et le risque de tendance. Dans un deuxième temps, on se focalise sur les tests que l'on a qualifiés de généraux au cours du chapitre 4 de ce mémoire, et notamment le test principal de validation des tables de mortalité à savoir un "rebouclage" (ou *back-testing*) avec le calcul

des SMR. Enfin, la dernière partie porte sur l'appréciation des paramètres exogènes.

L'ensemble des tests qui sont proposés se veulent au minimum des outils de mesure de la qualité de la table de mortalité construite, en effet, seul le test de "rebouclage" qui contrôle la qualité de prédiction de la table permet de valider (ou refuser la table de mortalité construite).

7.2 Le risque d'estimation

Lors de la partie théorique, trois méthodes différentes ont été présentées dans le but de chiffrer le risque d'estimation par le biais du calcul de coefficients de variations. L'ensemble des résultats sont comparables grâce au calcul d'un indicateur sans unité : le coefficient de variation, déterminé par la relation suivante :

$$cv_{x,t} = \frac{\sqrt{Var(q_{x,t}^{adj,k})}}{E[q_{x,t}^{adj,k}]}$$

7.2.1 Ré-échantillonnage par simulation directe

On rappelle que cette première méthode consiste à resimuler les taux bruts de mortalité à partir des taux bruts de Hoem selon la loi suivante :

$$\hat{q}_{x,t} \sim \mathcal{N}(q_{x,t}^{hoem}, \sqrt{\frac{q_{x,t}^{hoem} (1 - q_{x,t}^{hoem})}{ER_{x,t}}})$$

Il a été retenu d'effectuer 5 000 simulations de taux de mortalité pour chaque âge. Des tests ont été effectués avec plus de simulations, mais ils n'amélioreraient pas significativement la précision des résultats. Ensuite, les simulations ont été lissées selon les trois méthodes de lissage de cette étude : Moyennes Mobiles (MM), Whittaker-Henderson (WH) et Loess (LO). Ainsi nous obtenons pour les 6 cas différenciés (3 méthodes par sexe), 5 000 tables de mortalité lissées, ce qui nous permet de calculer les coefficients de variation. Les résultats suivants ont été obtenus pour les âges 35 à 80 ans :

Méthode	Sexe	cv méthode 1
MM	Hommes	5,78 %
MM	Femmes	7,89 %
WH	Hommes	6,84 %
WH	Femmes	7,85 %
LO	Hommes	6,79 %
LO	Femmes	7,91 %

TABLE 7.1 – Coefficients de variation obtenus par la première approche

7.2.2 Ré-échantillonnage par simulation des résidus

Cette seconde méthode est similaire à la première avec l'idée de calculer les coefficients de variation à partir d'une resimulation des résidus de Pearson :

Méthode	Sexe	cv méthode 2
MM	Hommes	7,86 %
MM	Femmes	10,03 %
WH	Hommes	8,85 %
WH	Femmes	10,11 %
LO	Hommes	8,79 %
LO	Femmes	10,03 %

TABLE 7.2 – Coefficients de variation obtenus par la seconde approche

Une première remarque que l'on peut faire sur ces résultats est que la moyenne des coefficients de variation est plus élevée, avec cette méthode de simulation des résidus de Pearson, que lors de la simulation directe des taux de sortie (de l'ordre de 30 %). La raison provient du fait que les résidus de Pearson comparent le nombre de décès observés avec le nombre de décès estimés, il y a l'intégration d'un risque supplémentaire lors de l'estimation du nombre de décès. Cet élément permet d'expliquer la différence de valeurs des coefficients de variation.

La seconde remarque concerne la différence entre les hommes et les

femmes. En effet, les résultats des coefficients de variation des femmes sont 15 % supérieurs pour les méthodes Whittaker-Henderson et Loess et 30 % pour la méthode moyennes mobiles. Ceci s'explique par la différence d'exposition observée entre les deux sexes (16 % plus faible pour les femmes).

Enfin, concernant la comparaison entre les méthodes, nous n'observons pas de nettes différences à part le taux masculin de la méthode moyennes mobiles qui est bien plus faible que les autres.

La troisième approche qui a été proposée lors de la partie théorique n'a pas été mise en application.

7.3 Le risque de tendance

Mesurer le risque de tendance n'est pas chose aisée. Au sein de cette étude pratique, on propose de regarder l'évolution de l'espérance de vie à l'âge de 35 ans (âge départ de nos tables). L'avantage de cet estimateur (par rapport à l'entropie) est qu'on peut facilement le comparer avec des estimations effectuées par des organismes statistiques tels que l'INSEE. Cet estimateur va nous renseigner sur deux éléments : sa valeur est une première information sur la table de mortalité, son évolution en est une seconde, c'est celle-ci qui va déterminer la qualité de la tendance. Les figures 7.1 et 7.2 sont des représentations de l'espérance de vie homme et femme entre 2007 et 2026.

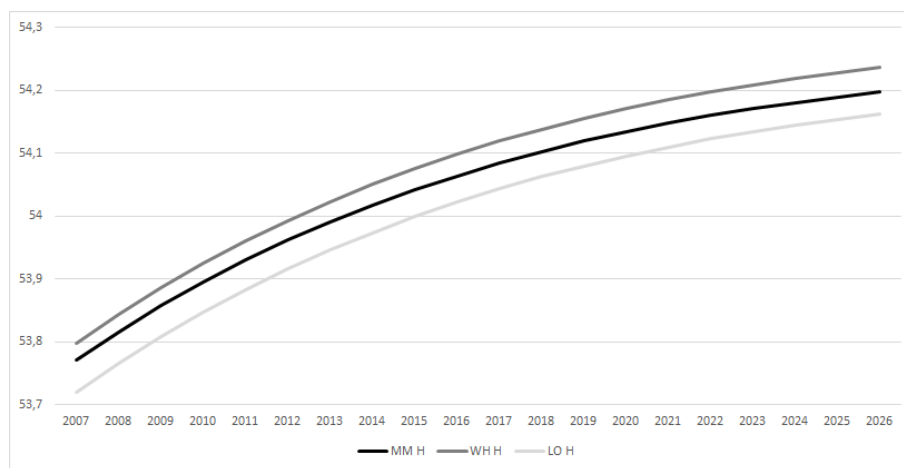


FIGURE 7.1 – Evolution de l'espérance de vie à 35 ans des hommes selon les différents lissages

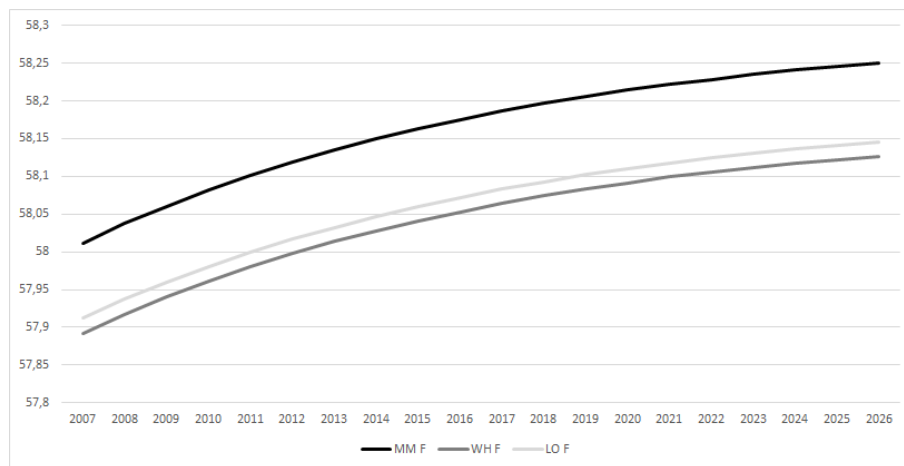


FIGURE 7.2 – Evolution de l'espérance de vie à 35 ans des femmes selon les différents lissages

Selon les données distribuées par l'INSEE¹, l'espérance de vie à 35 ans est de 42,8 ans pour les hommes et 49,4 ans pour les femmes sur la période 2000-2008. On constate donc que nos lissages sont très supérieurs aux prévisions, tous donnent environ 10 ans de plus d'espérance de vie, (niveau se rapprochant de celui du groupe de populations regroupant les cadres). Deux raisons peuvent expliquer cela : la principale est une sélection particulière d'individu dans notre base de données qui possède une mortalité beaucoup plus faible au-delà de 35 ans ou une méthode de fermeture de table mal adaptée. La réponse ne peut être tranchée avec les informations dont nous disposons.

Le second constat que l'on peut faire sur ces résultats est l'évolution des courbes, nos données fournissent une augmentation de l'espérance de vie d'une demi-année sur 20 ans, alors que lors des 20 dernières années l'espérance de vie en France à 35 ans a progressé d'environ 5 ans pour la population générale. Cette baisse de l'amélioration de l'espérance de vie, pour devenir presque nulle est une trajectoire envisagée par l'INSEE. En effet, ces dernières années, nous avons assisté à l'effet combiné de nombreuses avancées médicales et d'une nette amélioration des conditions de vie, ce que

1. Source : INSEE, Echantillon démographique permanent

l'INSEE ne pense pas reproductible dans les années à venir. Néanmoins, comme nous l'avons dit lors de la partie théorique, ce test est très subjectif car comment estimer l'espérance de vie future ?

7.4 Validation générale

Après avoir analysé le risque d'estimation et le risque de tendance, l'étude se poursuit avec l'utilisation de différentes approches dans le but d'émettre un avis sur la qualité de la table de mortalité. Dans un premier temps, il est vérifié que les caractéristiques classiques des tables de mortalité sont vérifiées.

Le premier aspect à contrôler concerne la pente de la courbe suivant les âges. En effet, l'expérience montre que la mortalité augmente avec l'âge de l'assuré. Sur l'ensemble de la table de mortalité prospective, nous avons :

- pour la table construite à partir d'un lissage par moyennes mobiles, 2,02 % des données masculines et 1,02 % des données féminines ne respectent pas cette règle,
- pour la table construite à partir d'un lissage par Whittaker-Henderson, 3,38 % des données masculines et 1,94 % des données féminines ne respectent pas cette règle,
- pour la table construite à partir d'un lissage par Loess, 1,02 % des données masculines et 1,02 % des données féminines ne respectent pas cette règle.

Ce premier test montre déjà des résultats très différents, selon les méthodes, dont se dégage la méthode de Loess.

Le second aspect que l'on propose de regarder est de vérifier si le respect de la baisse de la mortalité est respectée au cours du temps. Dans notre cas, la vérification est superflue, en effet, la tendance que l'on a appliquée est directement issue de la table de référence "INSEE 2007-2060" qui vérifie cette règle. Il est ainsi intéressant de noter que suite à l'emploi de notre procédure, cette règle ne dépend en aucun cas de la méthode de lissage choisie, mais uniquement de la table de référence sélectionnée.

Le troisième élément que l'on vérifie provient de la segmentation qui a

été retenue. C'est la différence de mortalité entre la population masculine et féminine. Dans la population en général, on observe une surmortalité à tout âge de la population masculine par rapport à la population féminine. Notre table de données ne présente pas de caractéristiques particulières sur ce point, car quel que soit la méthode de lissage et le sexe, cette troisième règle est vérifiée à tous les âges.

Enfin, nous terminons cette section par la réalisation d'un *back-testing*. Pour effectuer ce test, on calcule l'écart relatif entre le nombre de décès observés et le nombre de décès estimés pour chaque âge. Pour ne pas fausser les résultats, il est important de prendre des données qui n'ont pas servi à l'étude : c'est la raison pour laquelle nous nous sommes limités à 500 000 données initiales alors que la base de données contient environ 1,8 million de lignes. Ce test est le principal indicateur pour étudier la qualité de la table à prévoir la mortalité pour de nouveaux individus, et c'est actuellement le principal test effectué par les actuaires certificateurs pour contrôler la qualité de la base. Dans le cadre de cette étude, on a appliqué un rebouclage sur nos différentes tables de mortalité et nous avons obtenu les résultats moyens suivants :

Méthode	Sexe	SMR	Volatilité
MM	Hommes	1,042	0,186
MM	Femmes	0,961	0,181
WH	Hommes	1,019	0,175
WH	Femmes	0,981	0,192
LO	Hommes	1,000	0,157
LO	Femmes	0,991	0,191

TABLE 7.3 – Résultats du test SMR

Une représentation graphique a également été réalisée pour représenter cet indicateur avec les intervalles de confiance correspondants, disponibles en annexe. D'après les coefficients moyens du SMR, ce test est concluant pour les trois méthodes, cela signifie que les tables de mortalité prédisent le bon nombre de décès global.

À l'issue de ces premiers tests, si différentes tables de mortalité dont on est en train de procéder à la validation peuvent prétendre avoir une certaine fiabilité, on remarque que la méthode de Loess présente les meilleurs résultats pour les deux sexes, notamment pour le test SMR qui le plus important de cette partie.

7.5 Appréciation des paramètres exogènes

7.5.1 Analyse du paramètre de lissage

L'objectif de cette partie est de mettre en place différents tests pour émettre un avis sur la qualité du lissage, c'est-à-dire savoir s'il y a eu un phénomène de sur ou sous lissage. Il peut être la conséquence du choix du paramètre λ ou de la fonction de poids W . Notre première analyse concerne l'étude des résidus. Pour l'ensemble des tables, il a été calculé la moyenne des trois types de résidus : résidus de la réponse, résidus de Pearson et résidus de la déviance.

Méthode	Sexe	Réponse	Pearson	Déviance
MM	Hommes	4,37 e-06	-0,007	-2,50
MM	Femmes	3,92 e-05	0,049	11,24
WH	Hommes	-2,73 e-05	-0,011	-1,12
WH	Femmes	1,34 e-05	0,043	10,11
LO	Hommes	3,83 e-06	-0,007	-3,75
LO	Femmes	1,29 e-05	0,042	6,54

TABLE 7.4 – Moyenne des différents types de résidus

Même si ces résultats sont proches les uns des autres, la méthode de Loess semblent une nouvelle fois donner les meilleurs résultats. Pour pousser la réflexion sur les résidus, des représentations graphiques ont été réalisées identiques à celle de la figure 7.3. Peu d'informations supplémentaires ont été apportées. On remarque comme pour l'ensemble des cas, que la volatilité des résidus vient augmenter avec l'âge. Ceci est dû à la répartition de l'exposition, décroissante, que l'on voit sur la figure 5.3 à la page 72.

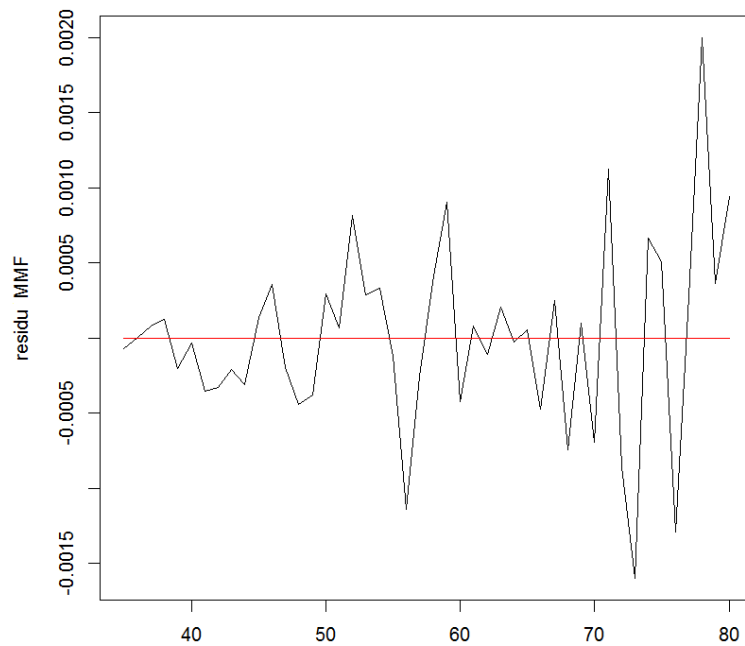


FIGURE 7.3 – Evolution des résidus de la réponse en fonction de l'âge pour le cas moyennes mobiles et sur la population féminine

Pour venir compléter l'étude du lissage, trois autres tests viennent apporter des approches nouvelles à cette validation. Elles se basent sur l'analyse des résidus de la réponse. Ces trois tests (test des signes, test de Wilcoxon et test de Wald-Wolfowitz) ont pour approches de comparer des caractéristiques de la suite des résidus de la réponse en les comparant à une loi sous-jacente (respectivement, une loi binomiale et deux lois normales), nous effectuons une analyse de ces résultats à l'aide de la p-value, c'est-à-dire la probabilité de rejeter à tort l'hypothèse nulle.

Les résultats sont présentés dans le tableau suivant avec les symboles (+) et (-) servant d'indicateurs pour renseigner les tests où la p-value est à maximiser ou à minimiser :

Méthodes	Sexe	Signes (+)	Wilcoxon (-)	Wald-Wolfowitz (-)
MM	Hommes	0,461	0,024	0,345
MM	Femmes	0,883	0,886	0,892
WH	Hommes	0,461	0,018	0,238
WH	Femmes	0,461	0,917	0,138
LO	Hommes	0,461	0,014	0,074
LO	Femmes	0,883	0,960	0,175

TABLE 7.5 – Résultats des différents tests sur les résidus de la réponse

L'analyse de ces résultats, nous renseigne sur plusieurs points. Dans un premier temps, on constate qu'il n'y a pas de grandes différences entre chaque méthode de lissage, mais une très grande variation entre les résultats de la population féminine et de la population masculine, notamment le lissage est considéré comme satisfaisant quelle que soit la méthode pour les femmes, mais n'est pas considérée satisfaisante pour les hommes.

Néanmoins, il est à noter que dans le cadre de cette étude, avec à notre disposition 36 résidus (correspondant aux âges entre 35 et 80 ans), ces tests ne sont pas robustes. Par exemple, le test des signes nécessite 77 résidus afin d'obtenir un intervalle de confiance de marge d'erreur de 5 % avec une probabilité de 95 %. On note que l'on peut aisément obtenir ce nombre de résidus sur une étude de population plus étendue au niveau des âges et ainsi porter un avis plus objectif.

7.6 Conclusion

La série de tests qui a été mise en place pour cette étude avait plusieurs objectifs. Tout d'abord elle devait permettre de sélectionner le modèle le plus pertinent parmi les méthodes de lissage (moyennes mobiles, Whittaker-Henderson et Loess), sur ce point, il semble que la méthode de Loess soit la plus efficace. En effet propose les meilleurs résultats au niveau des tests généraux et principalement au *Back-testing*. La méthode Loess est également celle qui semble la plus appropriée au regard des résultats de l'étude des résidus, même si cela est moins visible. Il n'est pas étonnant que ce lissage

soit préférable à la simple méthode des moyennes mobiles, car son optimisation du lissage de la fonction de passage entre les taux de mortalité de référence et les taux de mortalité bruts par des polynômes de faibles degrés est beaucoup plus efficace qu'une simple moyenne pondérée. La présente étude n'apporte pas non plus un avis négatif sur le lissage par la méthode des moyennes mobiles qui par son approche simple et facilement adaptable offre des résultats tout à fait satisfaisants. Par contre, la méthode de lissage par Whittaker-Henderson est à mettre un peu à part car en plus des trois paramètres exogènes que nous avons définis, la méthode nécessite en plus deux autres paramètres (notés dans ce mémoire h et z). Il a été posé des valeurs pour ces paramètres généralement reconnus par la littérature, néanmoins un test avec d'autres valeurs pour ces paramètres aurait mérité d'être réalisé pour tenter d'améliorer l'interprétation des résultats.

Le second intérêt de cette partie est de se faire un avis sur les différents outils de validation. Tous n'ont pas été appliqués par manque de temps, néanmoins un large tour d'horizon a été réalisé. L'importante quantité de données initiales a permis la construction de table que l'on peut considérer fiable comme le confirment l'étude des risques et le test SMR. Ce dernier est un test essentiel car son application fournit une information précieuse sur la justesse de la table de mortalité prospective. On peut citer les quelques tests qualifiés de généraux, qui ont l'avantage d'être visuels et donc aisément réalisables.

Par ailleurs, on note que le calcul de coefficients de variation dans le cadre du calcul du risque d'estimation permet la mise en place d'intervalles de confiance. On peut le relier à l'étude des règles classiques que l'on a définies. Le fait que la mortalité augmente selon les âges et diminue selon les années et que la mortalité masculine est à tout instant supérieure à la mortalité féminine. En effet, on peut être amené à réaliser des réajustements "manuels" des taux ne respectant pas ces règles élémentaires dans la limite des intervalles de confiance. Néanmoins, ces tests permettent plus une appréciation de la qualité de la base qu'une validation de celle-ci.

Enfin, parmi les tests qui ont été testés, certains tests n'ont pas été concluants pour plusieurs raisons. Tout d'abord, l'étude de la tendance a

été menée par des tests particulièrement subjectifs et ne peuvent être chiffrés, en effet cela nécessite de s'appuyer sur l'avis d'un expert pour évaluer l'évolution de l'espérance de vie à venir, chose qui n'est pas aisée. Les tests s'appuyant sur l'étude des résidus de la réponse n'ont pas été satisfaisants du fait du trop faible nombre d'éléments à analyser ce qui donne des résultats trop peu fiables.

Conclusion

Les assureurs ont depuis longtemps eu recours à des tables de mortalité, que ce soit pour tarifier leurs produits ou pour les provisionner. Au cours de ce mémoire, il a été proposé une démarche différente des approches habituelles qui posent un modèle prédéfini et tente de le calibrer pour l'adapter à leurs données. Cette démarche est généralement satisfaisante. Toutefois, il est utopique de penser à l'existence d'un tel modèle général et adaptable à l'ensemble des données. C'est dans cet esprit que se sont développés les modèles non-paramétriques, rendu possible grâce aux développements informatiques de ces dernières années.

Ainsi, le présent mémoire détaille la mise en place d'une démarche innovante mêlant l'utilisation d'un modèle relationnel avec l'application d'un lissage non-paramétrique. Les trois méthodes de lissage non-paramétrique présentées au cours de ce mémoire sont plus ou moins complexes : lissage par moyennes mobiles, par Whittaker-Henderson et par la méthode de Loess. Ces méthodes ont nécessité le recours de trois paramètres exogènes :

Une table de référence. L'étude se place dans le cadre d'une construction de table de mortalité prospective menée par une compagnie d'assurance. Le nombre de données disponibles est généralement limité et une méthode relationnelle est tout à fait indiquée en présence de petits échantillons. En effet, elle permet de s'appuyer sur une table existante ayant déjà fait ses preuves. Lors de la réalisation pratique, il a été retenu d'utiliser la table de référence "INSEE 2007-2060" pour de multiples raisons. C'est une table de mortalité prospective, très récente et basée sur la population française. L'autre table de référence envisageable, le couple TGH/F 05, présente le

désavantage d'être erratique comme l'a montré de récents travaux menés par Tomas et Planchet [44].

Un paramètre de lissage. Ce paramètre noté λ définit le nombre de points (voisinage) utilisé pour le lissage de chaque point. Cet élément est important, car il distingue une courbe "accidentée" ou "droite" de la bonne courbe. Au cours de ce mémoire, nous avons utilisé un indicateur, le AIC modifié, pour déterminer la valeur de λ optimale. Il s'est avéré que cette valeur ne dépend pas de la méthode de lissage retenue mais semble plutôt dépendre de la quantité de points de l'étude. Dans notre cas $\lambda = 10$.

Une fonction de poids. Le dernier paramètre définit l'importance (le poids) pris par les points appartenant au voisinage sélectionné par le paramètre précédent. Aucun élément objectif n'a fait pencher notre choix vers une fonction plus qu'une autre. Nous nous sommes donc arrêtés sur une fonction simple et proposant des poids proportionnels à la distance : c'est la fonction triangulaire.

Ces paramètres choisis, il a été possible de construire des tables de mortalité prospectives unisexes. La méthode de fermeture retenue a été celle de Denuit et Goderniaux approchant plus les taux lissés entre 80 et 90 ans que la méthode de Coale et Kisker.

La seconde partie a consisté à développer et à mesurer des risques sous-jacents à la construction d'une table de mortalité prospective. Un des objectifs du mémoire était de limiter au maximum le risque de modèle, cela a été en grande partie fait par la mise en place de lissages non-paramétriques, c'est-à-dire sans poser de modèle a priori. En pratique, les deux principaux risques détaillés au cours de ce mémoire sont le risque d'estimation et le risque de tendance. Le risque d'estimation, lié aux fluctuations d'échantillonnage a été évalué grâce à la resimulation des taux de sortie par deux approches (directe ou par les résidus). Les valeurs issues du calcul des coefficients de variation donnent des résultats acceptables (entre 5 et 10 % pour les deux méthodes présentées). Le risque de tendance est un risque important, c'est le risque d'une mauvaise projection. Pour notre exemple, nous nous sommes appuyés sur la tendance de la table de référence, ce qui offre une amélioration de l'espérance de vie de 0,5 an à 35 ans pour les 20 prochaines années (augmentation de l'ordre de 5 ans sur les 20 dernières années avec la table de référence) pour les hommes comme pour les femmes,

ceci est principalement dû à une combinaison importante de facteurs améliorant l'espérance de vie (médecine, condition de vie ...) apparue lors de ces dernières années que l'on ne pense pas retrouvé par la suite. Néanmoins, l'évaluation est basée sur un avis d'expert donc très subjective.

Cette démarche est intéressante sous plusieurs points de vue. En premier lieu, le recours à un lissage sur la fonction de passage entre les taux bruts et les taux de référence permet l'intégration "artificielle" de données supplémentaires en s'appuyant sur une table qui a fait ces preuves : ainsi la démarche peut sembler plus complexe que les démarches dites "classiques", mais apporte une réelle plus-value pour une application sur de petits échantillons de données. Dans un second temps, l'application de l'évolution de la mortalité issue d'une table de référence prospective est une démarche très simple et applicable quelle que soit la profondeur de la base de données.

Cette démarche originale est donc particulièrement adaptée pour les assureurs sur deux points importants :

- le premier point est réglementaire, en effet la méthode proposée s'appuie sur une table de référence validée par arrêté ministériel. Ceci est un élément important dans l'homologation de la méthode qui est détaillée au cours de ce mémoire ;
- le second point concerne l'efficacité de la démarche dans le cadre de petits échantillons. En effet, l'apport d'une table de référence externe permet l'augmentation artificielle du nombre de données.

Le dernier chapitre du mémoire a consisté à appliquer différents tests, qui se veulent universels (c'est-à-dire indépendants de la méthode de construction choisi) en variant les approches. Tous n'ont pas la même importance. Le test à privilégier est le *back-testing*, qui compare le nombre de décès observés, sur une population n'ayant pas servi à la construction de la table, au nombre de décès estimés. Ce test est actuellement le principal test effectué lors du contrôle et de la certification d'une table de mortalité par les actuaires certificateurs. Les tests que l'on place à un second niveau sont les tests généraux : la mortalité croît avec l'âge et diminue avec le temps et la mortalité masculine est supérieure à tout âge à la mortalité féminine. Ils fournissent une première bonne approche visuelle sur la fiabilité des tables.

D'autres tests ont été testés et présentés dans cette étude avec plus ou moins de résultats. La conclusion de ces outils de validation est que la méthode de lissage par Loess sort gagnante sur la majorité totalité des tests. Cela n'est pas étonnant, en effet, la méthode de lissage par moyennes mobiles est très simple mais offre néanmoins un très bon rendement simplicité/efficacité sur notre exemple et la méthode par Whittaker-Henderson propose des résultats encourageant mais n'a pas forcément été optimisée car cette fonction à la particularité de posséder deux paramètres endogènes supplémentaires. Ainsi, les trois méthodes que l'on a comparées sont utilisables et sont utilisées par les assureurs, mais des nouvelles applications notamment sur des échantillons plus réduits sont envisageables pour approfondir cette étude.

Annexes

Annexe A

Les principaux indicateurs d'analyse de tables de mortalité

De manière générale, pour représenter l'évolution de l'effectif de chaque cohorte aux âges successifs d'un portefeuille, l'actuaire utilise le diagramme de Lexis semblable à la figure A.1.

Un évènement concernant un individu d'une population peut être repéré par trois caractéristiques : l'âge de l'individu, l'année de naissance de l'individu et la date de l'évènement. Ces trois caractéristiques sont redondantes, seule l'utilisation de deux d'entre elles apporte toute l'information disponible. Au cours de ce mémoire, nous utilisons pour décrire chaque observation y , l'âge de l'individu noté x et la date de l'évènement, qui pour les tables de mortalité est le décès, notée t .

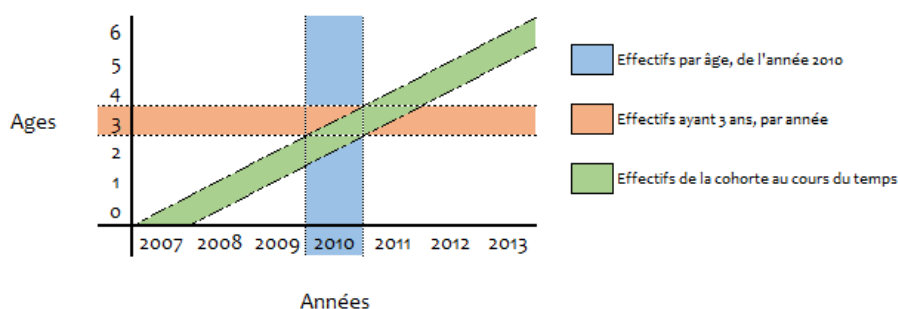


FIGURE A.1 – Diagramme de Lexis

A partir d'un jeu de données classique, il est plus ou moins aisé d'obtenir le nombre de décès à chaque âge en jour, noté $d_{x,t}^j$, et l'exposition au jour le jour, noté $ER_{x,t}^j$. L'exposition représente le nombre de personne exposé au risque. Ces deux premiers éléments sont indispensables pour obtenir les taux de mortalité. La littérature a proposé une multitude de méthodes pour le calcul des taux de mortalité annuel, nous présentons les deux principaux qui ont chacun leurs défauts (notons a le nombre de jours dans l'année d'observation)

- Méthode de Hoem :

$$q_{x,t}^H = \frac{\sum_{i=1}^a d_{x,t}^j(i)}{\frac{1}{a} \sum_{i=1}^a ER_{x,t}^j(i)}$$

Le principal souci de cette méthode est d'autoriser l'apparition de taux de mortalité supérieur à 1. En effet, si nous avons uniquement deux individus en début d'année et qui décèdent en milieu d'année, nous obtenons (si $a = 365$)

$$q_{x,t}^H = \frac{1 + 1}{\frac{1}{365}(332,5 + 332,5)} = 2$$

- Méthode de Kaplan-Meier :

$$q_{x,t}^{KM} = 1 - \prod_{i=1}^a \left(1 - \frac{d_{x,t}^j(i)}{ER_{x,t}^j(i)}\right)$$

Contrairement à la méthode précédente, les taux de mortalité qui résultent de cette méthode seront compris entre 0 et 1, mais le défaut de cette méthode est de ne considérer que les dates où il y a une sortie par décès. En effet, si durant l'année, 1 seul décès est observé alors il ne sera tenu compte que de l'exposition le jour du décès : cela est d'autant plus gênant s'il y a 100% de décès à une certaine date ce qui entraîne une probabilité de décès de 1 pour l'année correspondante.

Néanmoins, de manière générale, si l'exposition est suffisante, la méthode de Kaplan-Meier est à privilégier car se rapproche plus la mortalité réelle.

Ces taux de mortalité permettent l'obtention de deux nouveaux indicateurs démographiques, les taux de survie notés $p_{x,t}$:

$$p_{x,t} = 1 - q_{x,t}$$

et le nombre de survivants $l_{x,t}$ en fonction du nombre de vivants à la date initiale :

$$l_{x,t} = l_{x-1,t} (1 - q_{x-1,t})$$

ce dernier indicateur est celui que l'on retrouve dans la présentation des tables de mortalité comme par exemple pour la figure 1.1.

Le taux instantané de mortalité est noté $\mu_{x,t}$ et nous posons l'hypothèse classique qu'il est constant sur chaque carré du diagramme de Lexis, ainsi :

$$\mu_{x,t} = -\ln(1 - q_{x,t})$$

A l'ensemble de ces indicateurs démographiques, nous allons inclure un indicateur actuarielle proposée par Delwarde et Denuit [15] : la prime pure d'un contrat de rente viagère. Elle peut se voir comme une version pondérée de l'espérance de vie, donnant plus d'importance aux premières années par le jeu de l'actualisation. Nous supposons un seul versement par an d'une valeur de 1€, alors

$$PrimePure_{x,t} = \sum_{k \geq 1} {}_k p_{x,t} v^k = \sum_{k \geq 1} \prod_{j=1}^{k-1} p_{x+j,t} v^k$$

Les fonctions de poids classiques

La littérature actuarielle propose une multitude de fonctions de poids, notée W , qui a pour effet de favoriser les points u les plus proches de la zone d'ajustement x_i en respectant les contraintes suivantes :

$$\left\{ \begin{array}{l} \forall u, W(u) \geq 0 \\ \forall u \notin V(x_i), W(u) = 0 \\ W \text{ doit \AA tre sym\AA trique et d\AA croissant en fonction de la distance \AA } x_i \end{array} \right.$$

Nous pr\AA sentons ici une liste non-exhaustive. Le choix est subjectif, il d\AA pend de la variation d'importance, que veut donner la personne construisant la table, aux points appartenant au voisinage du point d'\AA valuation.

- La fonction rectangulaire :

$$W(u) = \frac{1}{2}I(|u| \leq 1)$$

- La fonction triangulaire :

$$W(u) = (1 - |u|)I(|u| \leq 1)$$

- La fonction Epanechnikov :

$$W(u) = \frac{3}{4}(1 - u^2)I(|u| \leq 1)$$

- La fonction *Biweight* :

$$W(u) = \frac{15}{16}(1 - u^2)^2I(|u| \leq 1)$$

- La fonction *Triweight* :

$$W(u) = \frac{35}{32}(1 - u^2)^3 I(|u| \leq 1)$$

- La fonction tricube :

$$W(u) = (1 - u^3)^3 I(|u| \leq 1)$$

Outils de fermeture des tables de mortalité

Le développement des méthodes de fermeture des tables présentées est repris de l'ouvrage de Delwarde et Denuit [2006] [15]

C.1 Méthodes de Denuit et Goderniaux

Cette méthode émet différentes hypothèses sur les taux de mortalité en ajustant par les moindres carrés, le modèle log-quadratique :

$$\ln(\hat{q}_{x,t}) = a + bx + cx^2 + \epsilon_{x,t} \text{ avec } \epsilon_{x,t} \sim \mathcal{N}(0, \sigma^2)$$

avec les deux contraintes suivantes :

$$\hat{q}_{130,t} = 1 \tag{C.1}$$

$$\frac{\partial \hat{q}_{x,t}}{\partial x} = 0 \tag{C.2}$$

La première équation nous donne

$$a + 130b + 130^2c = 0$$

Et la seconde

$$b + 2c = 0$$

Et donc

$$\forall x, \ln(\hat{q}_{x,t}) = c(130^2 - 260x + x^2) + \epsilon_{x,t}$$

C.2 Méthodes de Coale et Kisker

Cette seconde méthode de fermeture des tables travaille à partir des taux de mortalité instantanés. La méthode consiste à extrapoler les taux aux grands âges à partir de l'hypothèse de départ :

$$\forall x \geq 64, \mu_{x,t} = \mu_{65,t} \exp(k_{x,t}(x - 65))$$

La fonction $k_{x,t}$ représente l'évolution moyenne des taux.

Après des constatations empiriques, Coale et Kisker ont montré que les $k_{x,t}$ présentaient un pic aux alentours de 80 ans, ce qui a conduit à ajouter l'hypothèse suivante :

$$k_{x,t} = k_{80} + s(x - 80)$$

Le paramètre s représente le coefficient de la pente. Pour la suite des calculs avec la résolution de ces équations, il y a nécessité de poser les deux contraintes exogènes suivantes :

$$\begin{cases} \mu_{110,t} = 1 \text{ pour les hommes} \\ \mu_{110,t} = 0,8 \text{ pour les femmes} \end{cases}$$

La combinaison de ces hypothèses implique que :

$$\mu_{110,t} = \mu_{79,t} \exp\left(\sum_{x=80}^{110} k_{x,t}\right) = \mu_{79,t} \exp\left(\sum_{x=80}^{110} k_{80,t} + s(x - 80)\right)$$

Ce qui implique

$$s = -\frac{\ln\left(\frac{\mu_{79}}{\mu_{110}}\right) + 31k_{80}}{465}$$

et

$$\mu_{x,t} = \mu_{79,t} \exp\left(\sum_{y=80}^x (k_{80} + s(y - 80))\right)$$

ce qui se réécrit

$$\mu_{x,t} = \mu_{x-1,t} \exp(k_{80} + s(x - 80))$$

avec

$$k_{80} = \frac{\ln\left(\frac{\mu_{80}}{\mu_{65}}\right)}{15}$$

Annexe D

Tableau d'exposition et de sinistralité de la base

Le tableau suivant présente l'exposition et le nombre de décès par sexe et par âge de la base utilisée pour les calculs de Kaplan-Meier. On rappelle que l'exposition est calculée au prorata-temporis. Une personne présente du 01/07/2005 au 31/03/2006 comptabilisera une exposition de 0,5 pour l'année 2005 et une exposition de 0,25 pour l'année 2006.

Age	Exposition féminine	Nombre de décès féminin	Exposition masculine	Nombre de décès masculin	Age	Exposition féminine	Nombre de décès féminin	Exposition masculine	Nombre de décès masculin	Age	Exposition féminine	Nombre de décès féminin	Exposition masculine	Nombre de décès masculin
1	0	0	0	0	45	9119	11	11627	14	89	2507	209	1743	245
2	529	0	565	0	46	9139	14	11616	21	90	1808	183	1224	185
3	519	0	563	0	47	9128	10	11696	20	91	1414	155	884	143
4	564	0	585	0	48	9323	9	11726	27	92	1270	183	751	131
5	576	0	633	0	49	9476	11	11776	23	93	1172	203	706	146
6	594	0	660	0	50	9711	19	11994	41	94	1003	155	590	128
7	635	0	697	0	51	10117	19	12114	37	95	828	145	443	100
8	672	0	727	0	52	10436	29	12367	40	96	608	131	294	64
9	690	0	750	0	53	10889	26	12727	57	97	437	118	186	47
10	677	0	768	0	54	11487	30	13110	53	98	295	71	117	32
11	658	0	800	0	55	12104	28	13690	70	99	205	48	78	20
12	690	0	836	0	56	12833	19	14203	75	100	121	36	53	18
13	709	0	864	0	57	13623	35	14830	101	101	88	27	31	9
14	738	0	932	0	58	14465	50	15470	102	102	53	13	18	5
15	797	0	966	0	59	15127	63	15841	106	103	30	15	9	3
16	810	0	1012	0	60	15305	47	15808	133	104	18	8	13	1
17	796	0	1025	0	61	15132	58	15523	140	105	13	1	11	0
18	790	0	1031	0	62	14542	57	15017	142	106	8	4	10	0
19	826	0	1106	0	63	13469	61	14075	145	107	6	2	8	0
20	937	0	1254	1	64	12465	58	12808	137	108	4	0	7	0
21	1097	0	1415	2	65	11964	61	12783	157	109	2	0	9	0
22	1288	1	1612	0	66	11816	59	13716	166	110	0	0	3	0
23	1586	2	1799	1	67	11489	71	15024	180	111	0	0	4	0
24	1903	0	2016	1	68	11227	64	15875	194	112	0	0	2	0
25	2287	1	2305	3	69	11190	80	15868	197	113	0	0	3	0
26	2696	0	2650	2	70	11522	81	15967	248	114	0	0	1	0
27	3128	1	3129	6	71	11323	109	15521	240	115	0	0	1	0
28	3497	0	3571	3	72	10906	93	14601	242	116	0	0	2	0
29	3846	0	4093	1	73	10656	94	13820	285	117	1	0	4	0
30	4209	0	4575	4	74	10546	130	12129	259	118	1	0	3	0
31	4622	3	5111	0	75	10288	140	11998	297	119	1	0	2	0
32	5071	2	5707	3	76	10036	135	9538	247	120	0	0	0	0
33	5617	4	6410	8	77	9584	163	8474	260	121	0	0	2	0
34	6140	3	7257	2	78	9136	193	7681	298	122	0	0	0	0
35	6658	2	8120	10	79	8647	192	7024	319	123	0	0	0	0
36	7183	3	8768	12	80	8205	214	6515	281	124	0	0	0	0
37	7530	4	9329	11	81	7740	248	6072	311	125	0	0	0	0
38	7859	5	9767	8	82	7198	238	5611	336	126	0	0	0	0
39	8127	3	10139	9	83	6835	258	5127	346	127	0	0	1	0
40	8213	5	10533	10	84	6418	303	4704	328	128	0	0	0	0
41	8444	3	11062	8	85	5900	287	4160	327	129	0	0	0	0
42	8717	4	11365	12	86	5162	295	3539	306	130	0	0	0	0
43	8964	6	11668	18	87	4393	295	2969	296	131	0	0	0	0
44	9103	6	11740	6	88	3443	281	2283	265	Total	140983	58	169915	141

FIGURE D.1 – Exposition et nombre de décès par sexe et par âge

Annexe **E**

Tableau des résultats issus du lissage des taux bruts

Le tableau suivant présente l'ensemble des résultats obtenus par âge et par sexe selon chaque méthode de lissage non-paramétrique présenté.

Age	Lissage par moyennes mobiles		Lissage par Whittaker-Henderson		Lissage par Loess	
	Hommes	Femmes	Hommes	Femmes	Hommes	Femmes
35	0,091%	0,046%	0,088%	0,049%	0,091%	0,035%
36	0,095%	0,048%	0,094%	0,050%	0,093%	0,037%
37	0,098%	0,049%	0,098%	0,050%	0,095%	0,040%
38	0,098%	0,049%	0,100%	0,049%	0,097%	0,044%
39	0,096%	0,048%	0,099%	0,048%	0,100%	0,048%
40	0,096%	0,049%	0,098%	0,049%	0,105%	0,054%
41	0,098%	0,050%	0,097%	0,053%	0,110%	0,059%
42	0,100%	0,055%	0,099%	0,059%	0,117%	0,066%
43	0,104%	0,064%	0,104%	0,067%	0,125%	0,074%
44	0,110%	0,076%	0,115%	0,078%	0,135%	0,083%
45	0,127%	0,089%	0,130%	0,091%	0,147%	0,094%
46	0,146%	0,099%	0,152%	0,106%	0,162%	0,107%
47	0,170%	0,105%	0,178%	0,122%	0,180%	0,121%
48	0,199%	0,116%	0,208%	0,140%	0,206%	0,140%
49	0,226%	0,133%	0,241%	0,158%	0,235%	0,159%
50	0,262%	0,158%	0,277%	0,177%	0,269%	0,177%
51	0,295%	0,180%	0,315%	0,194%	0,311%	0,195%
52	0,330%	0,203%	0,356%	0,210%	0,354%	0,210%
53	0,370%	0,216%	0,399%	0,223%	0,400%	0,220%
54	0,414%	0,220%	0,445%	0,236%	0,452%	0,230%
55	0,470%	0,226%	0,496%	0,249%	0,501%	0,244%
56	0,521%	0,243%	0,552%	0,263%	0,549%	0,264%
57	0,580%	0,271%	0,612%	0,281%	0,603%	0,284%
58	0,641%	0,297%	0,677%	0,301%	0,664%	0,306%
59	0,704%	0,321%	0,747%	0,325%	0,735%	0,330%
60	0,777%	0,341%	0,818%	0,351%	0,810%	0,355%
61	0,847%	0,366%	0,889%	0,378%	0,887%	0,382%
62	0,919%	0,394%	0,958%	0,407%	0,961%	0,410%
63	0,989%	0,424%	1,020%	0,437%	1,021%	0,440%
64	1,054%	0,454%	1,076%	0,469%	1,074%	0,469%
65	1,112%	0,485%	1,124%	0,502%	1,123%	0,497%
66	1,157%	0,520%	1,167%	0,537%	1,176%	0,533%
67	1,208%	0,566%	1,211%	0,577%	1,233%	0,577%
68	1,268%	0,617%	1,262%	0,621%	1,290%	0,629%
69	1,347%	0,673%	1,329%	0,672%	1,361%	0,683%
70	1,457%	0,737%	1,419%	0,732%	1,461%	0,742%
71	1,585%	0,819%	1,542%	0,805%	1,580%	0,820%
72	1,743%	0,903%	1,703%	0,895%	1,727%	0,903%
73	1,937%	1,009%	1,910%	1,005%	1,904%	1,003%
74	2,172%	1,150%	2,164%	1,141%	2,107%	1,123%
75	2,468%	1,310%	2,465%	1,305%	2,372%	1,273%
76	2,797%	1,497%	2,808%	1,500%	2,714%	1,463%
77	3,184%	1,738%	3,189%	1,726%	3,128%	1,688%
78	3,634%	2,015%	3,599%	1,983%	3,573%	1,948%
79	4,124%	2,321%	4,035%	2,270%	4,046%	2,244%
80	4,622%	2,674%	4,497%	2,589%	4,578%	2,584%

FIGURE E.1 – Résultats des taux après lissage

Représentation des courbes finales

Voici ci-dessous, la représentation de l'ensemble de nos résultats lissés, fermés et projetés par sexe :

F.1 Lissage par la méthode des moyennes mobiles

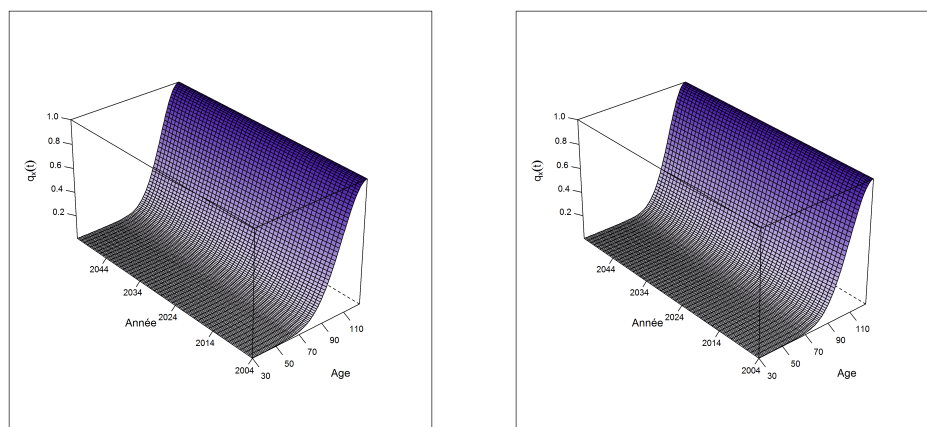


FIGURE F.1 – Projection des taux lissés par la méthode des moyennes mobiles sur la population masculine et féminine

F.2 Lissage par la méthode de Whittaker-Henderson

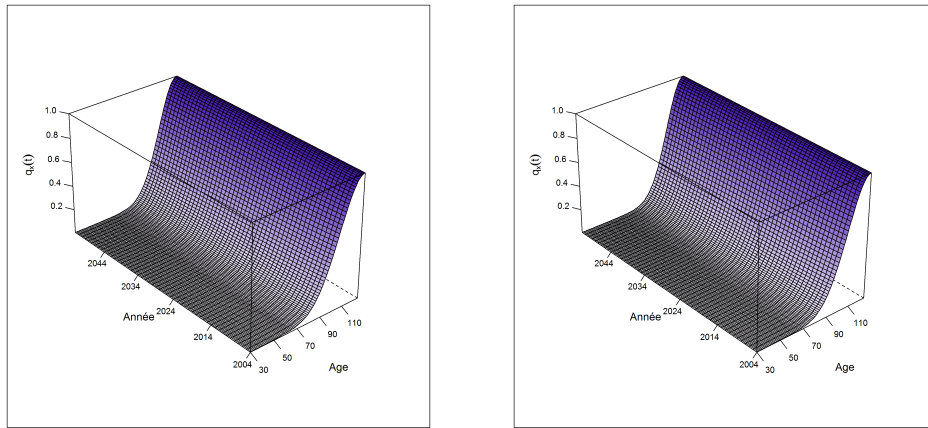


FIGURE F.2 – Projection des taux lissés par la méthode de Whittaker-Henderson sur la population masculine et féminine

F.3 Lissage par la méthode de Loess

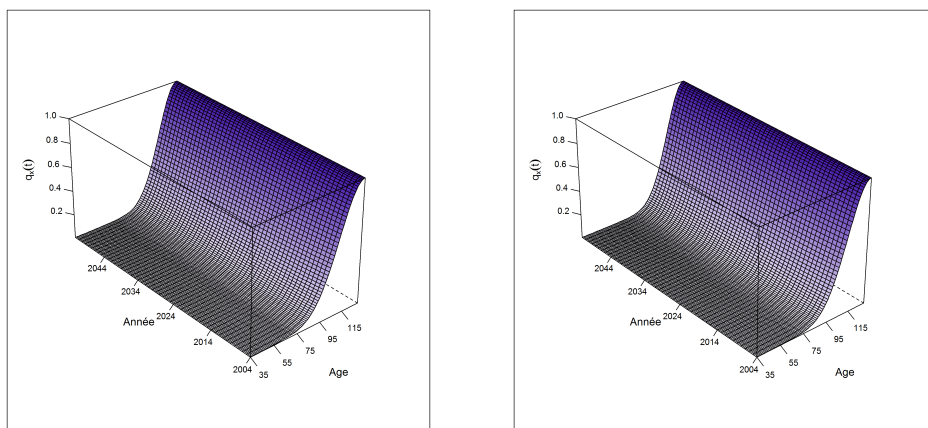


FIGURE F.3 – Projection des taux lissés par la méthode de Loess sur la population masculine et féminine

Bibliographie

- [1] (Décembre, 2012). Article A335-1. *Code des Assurances*.
- [2] Alho J. and Spencer B. (2005). *Statistical Demography and forecasting*. Springer.
- [3] Avril E. (Novembre, 2012). Risque systématique de longévité dans le cadre de données limitées. Application à un portefeuille français de rentes viagères. *Mémoire d'actuariat - Université Paris Dauphine*.
- [4] Booth H. and Tickle L. (Septembre, 2008). Mortality modeling and forecasting : A review of methods. *Annals of Actuarial Science - Volume 3, Issue 1-2, pp 3-43*.
- [5] Brugio A. and Frova L. (1995). *Projection de la mortalité par cause de décès : Extrapolation tendancielle ou modèle âge-période-cohorte*. Population, 50e année, n°4-5, pp 1031-1051.
- [6] Cairns D., Blake A.J.G., and Dowd K. (November, 2004). Pricing frameworks for securitization of mortality risk. *Proceedings of the International AFIR Colloquium, Boston, pp 509-540*.
- [7] Cleveland W.S. and Devlin S.J. (Septembre, 1988). Locally weighted regression : An approach to regression analysis by local fitting. *Journal of the American Statistical Association, Vol. 83, pp. 596-610*.
- [8] Cleveland W.S. (Décembre, 1979). Robust locally-weighted regression and smoothing scatterplots. *Journal of the American Statistical Association, Vol. 74, pp. 829-836*.
- [9] Coale A.J. and Kisker E.E. (1990). Defect in data on old age mortality in the United States : New procedure for calculating approximatively

- accurate mortality schedules and life tables at the highest ages. *Asian and Pacific Population Forum Vol. 4 no. 1.*
- [10] Continuous Mortality Investigation. (Juin, 2009). A prototype mortality projections model : Part one - an outline of the proposed approach. *Working Paper 38 - Institute of Actuaries and Faculty of Actuaries.*
 - [11] Continuous Mortality Investigation. (Juillet, 2009). A prototype mortality projections model : Part two - detailed analysis. *Working Paper 39, Institute of Actuaries and Faculty of Actuaries.*
 - [12] Continuous Mortality Investigation. (Septembre, 2011). User guide for the CMI mortality projections model, 'CMI_2011'. *User guide, Institute of Actuaries and Faculty of Actuaries.*
 - [13] Currie I., Durban M., and Eilers P. (Décembre, 2004). Smoothing and forecasting mortality rates. *Statistical Modelling - Vol. 4 no. 4 279-298.*
 - [14] Decroocq J.F. (Mars, 2011). La qualité des données. <http://www.varm.fr/content/La%20qualite%20des%20donnees%2003%202011.pdf>.
 - [15] Delwarde A. and Denuit M. (2006). *Construction de tables de mortalité périodiques et prospectives.* ECONOMICA.
 - [16] Delwarde A., Kachachidze D., Olie L., and Denuit M. (2004). Modèles linéaires et additifs généralisés, maximum de vraisemblance local et méthodes relationnelles en assurance vie. *Bulletin Français d'Actuariat, Vol. 6, no. 12.*
 - [17] Delwarde A. (2005). Modèle log-bilinéaire pour l'élaboration de tables de mortalité prospectives. *Mémoire d'actuariat - UCL.*
 - [18] Denuit M. and Godeniaux A.-C. (2005). Closing and projecting life tables using log-linear models. *Université catholique de Louvain - WP04-04.*
 - [19] Ducros J.-C. (Mars, 2013). Construction d'une table de mortalité prospective pour des contrats de retraite en points. *Mémoire d'actuariat - ISFA.*
 - [20] Guibert Q. and Planchet F. (Mars, 2013). Construction de lois d'expérience en présence d'événements concurrents – Application à l'estimation des lois d'incidence d'un contrat dépendance. *Les documents de travail de l'ISFA 2013.6.*

- [21] Guérin L. (2009). Construction de tables de mortalité féminines par génération à partir d'une approche par grandes causes médicales de décès - modélisation lee-carter. *Mémoire du Centre d'Études Actuarielles*.
- [22] Hannerz H. (2001). An extension of relational methods in mortality estimation.
- [23] Horiuchi S. and Wilmoth J.R. (Novembre, 1998). Deceleration in the age pattern of mortality at older ages. *Demography Vol. 35 no. 4 pp 391-412*.
- [24] Institut National d'Études Démographiques. (Décembre, 2010). Espérance de vie : peut-on gagner trois mois par an indéfiniment ? *Population & Société, n°473*.
- [25] Institut National d'Études Démographiques. (2010). Projections de population 2007-2060 pour la france métropolitaine : méthode et principaux résultats. *Série des Documents de Travail de la direction des statistiques Démographiques et Sociales F1008*.
- [26] Institut des Actuaires. (2002). Charte de la certification et de suivi des tables de mortalite et des lois de maintien en incapacite de travail et en invalidite. [http://www.ressources-actuarielles.net/EXT/ISFA/fp-isfa.nsf/0/1430AD6748CE3AFFC1256F130067B88E/\\$FILE/ChartePublique.pdf?OpenElement](http://www.ressources-actuarielles.net/EXT/ISFA/fp-isfa.nsf/0/1430AD6748CE3AFFC1256F130067B88E/$FILE/ChartePublique.pdf?OpenElement).
- [27] Institut des Actuaires. (Décembre, 2005). Notice d'utilisation des tables de mortalité TH00-02 et TF00-02, arrêté du 29 décembre 2005.
- [28] JORF n°197 page 12577 texte n°11. (Août, 2006). Arrêté du 1er août 2006 portant homologation des tables de mortalité pour les rentes viagères et modifiant certaines dispositions du code des assurances en matière d'assurance sur la vie et de capitalisation. http://www.cbanque.com/r/rente_viagere_2007.pdf.
- [29] Kamega A. and Planchet F. (Novembre, 2010). Mesure du risque d'estimation associé à une table d'expérience. *Cahiers de recherche de l'ISFA, WP2136*.
- [30] Kamega A. and Planchet F. (Février, 2011). Analyse et comparaison des populations générale et assurée en Afrique subsaharienne francophone pour anticiper la mortalité future. *Cahiers de recherche de l'ISFA*.

- [31] Kamega A. and Planchet F. (Juin, 2011). Hétérogénéité : mesure du risque d'estimation dans le cas d'une modélisation intégrant des facteurs observables. *Demographic research, Vol. 4 Art. 10.*
- [32] Kamega A. and Planchet F. (2012). *Actuariat et assurance vie en Afrique subsaharienne francophone.* Seddita.
- [33] Kamega A. (Décembre, 2011). Outils théoriques et opérationnels adaptés au contexte de l'assurance vie en Afrique subsaharienne francophone - Analyse et mesure des risques liés à la mortalité. *Thèse d'actuariat, université Lyon I.*
- [34] Keyfitz N. (1964). Utilisation de machines électroniques pour les calculs démographiques. *Population, Vol. 19 no. 4 pp. 673-682.*
- [35] Lelieur V. (2005). Utilisation des méthodes de lee-carter et log-poisson pour l'ajustement de tables de mortalité dans le cas de petits échantillons. *Mémoire d'actuariat - ISFA.*
- [36] Lu Y., Kacem R., and De Myttenaere A. (2011). Constitution et analyse d'une base de données sur la mortalité dans les portefeuilles d'assurance. *Synthèse de projet de statistique appliquée.*
- [37] Mindelsohn E. (2007). Les tables de mortalité recommandées par les normes dans le cadre des évaluations des engagements de retraite et les best estimates de la mortalité. *Mémoire d'actuariat - Université PARIS Dauphine.*
- [38] Planchet F., Faucillon L., and Juillard M. (Octobre, 2006). Quantification du risque systématique de mortalité pour un régime de rentes en cours de services. *assurances et gestion du risque Vol. 74(3).*
- [39] Planchet F. and Leroy G. (Décembre, 2009). Quel niveau de segmentation pertinent. *La Tribune de l'Assurance, n°142.*
- [40] Planchet F. and Therond P. (2007). *Pilotage technique d'un régime de rentes viagères : identification et mesure des risques, allocation d'actif, suivi actuariel.* ECONOMICA.
- [41] Planchet F. and Therond. P. (2011). *Modélisation statistique des phénomènes de durée : Applications actuarielles.* ECONOMICA.
- [42] Planchet F. (Novembre, 2006). Pilotage technique d'un régime de rentes viagères : identification et mesure des risques, allocation d'actif, suivi actuariel. *Thèse d'actuariat - ISFA.*

- [43] Quashie A. and Denuit M. (Février, 2005). Modèles d'extrapolation de la mortalité aux grands âges. *Université catholique de Louvain - WP04-13*.
- [44] Tomas J. and Planchet F. (2013). Construction et Validation des Tables de Mortalité Best Estimate. *Institut des Actuaire - Note de travail III291-11 v1.3*.
- [45] Tomas J. and Planchet F. (2013). Critère de Validation : Aspects méthodologiques. *Institut des Actuaire - Note de travail II291-14 v1.1*.
- [46] Tomas J. and Planchet F. (2013). Méthodes de positionnement : Aspects méthodologiques. *Institut des Actuaire - Note de travail III291-12 v1.1*.
- [47] Tomas J. (Janvier, 2013). Quantifying biometric life insurance risks with non-parametric smoothing methods. *Thèse, Faculty of Economics and Business, Amsterdam School of Economics Research Institute*.
- [48] Vaupel J.W. (1997). *Between Zeus and Salmon : The Biodemography of longevity*. National Research Council.
- [49] Viville M.-B. (2008). Comparaison de méthodes d'ajustement de la mortalité des rentiers dans un but prospectif. *Mémoire d'actuariat - ISFA*.

Table des figures

1.1	Extrait de la table du moment TH00-02 et de la table g�n�rationnelle TGH05	23
2.1	La courbe grise, plus r�guli�re, est beaucoup plus int�ressante que la courbe noire	30
2.2	Mise en relation des taux de sortie bruts avec les taux de sortie de r�f�rence pour tous les �ges	31
2.3	Les diff�rents types de voisinage	36
3.1	Les risques inh�rents � la r�alisation d'une table de mortalit� prospective extrait de la th�se de Kamega [33]	44
3.2	�volution de l'esp�rance de vie en France (Source INSEE) . .	50
3.3	�volution des taux de mortalit� instantan�e, extrait de la th�se de Planchet [42]	51
5.1	Repr�sentation de la r�partition des dates de naissance du portefeuille	70
5.2	Repr�sentation des �ges d'entr�e et de sortie du portefeuille .	71
5.3	Repr�sentation de l'exposition et du nombre de d�c�s par �ge	72
5.4	Evolution annuelle de la proportion homme/femme du portefeuille	73
5.5	Evolution de l'�ge moyen d'entr�e et de sortie du portefeuille	74
6.1	Repr�sentation des taux bruts masculins par �ges, m�thode Kaplan-Meier et m�thode Hoem	77

6.2	Représentation des taux bruts féminins par âges, méthode Kaplan-Meier et méthode Hoem	77
6.3	Mise en relation des taux de sortie bruts avec les taux de sortie de référence pour les âges compris entre 35 et 80 ans, sur la population masculine	78
6.4	Courbes issues du lissage par moyennes mobiles	81
6.5	Courbes issues du lissage Whittaker-Henderson	81
6.6	Courbes issues du lissage par Loess	82
6.7	Fermeture des taux lissés	82
6.8	Projection des taux lissés par la méthode des moyennes mobiles sur la population masculine et féminine	84
7.1	Evolution de l'espérance de vie à 35 ans des hommes selon les différents lissages	89
7.2	Evolution de l'espérance de vie à 35 ans des femmes selon les différents lissages	90
7.3	Evolution des résidus de la réponse en fonction de l'âge pour le cas moyennes mobiles et sur la population féminine	94
A.1	Diagramme de Lexis	103
D.1	Exposition et nombre de décès par sexe et par âge	111
E.1	Résultats des taux après lissage	113
F.1	Projection des taux lissés par la méthode des moyennes mobiles sur la population masculine et féminine	114
F.2	Projection des taux lissés par la méthode de Whittaker-Henderson sur la population masculine et féminine	115
F.3	Projection des taux lissés par la méthode de Loess sur la population masculine et féminine	115

Liste des tableaux

2.1	Exemple d'un lissage moyennes mobiles	39
2.2	Exemple d'un lissage Whittaker-Henderson	40
2.3	Exemple d'un lissage par Loess	41
6.1	Extrait du calcul des AIC modifiés	79
6.2	Résultats des tests de sélection de l'année de référence (1/2) .	83
6.3	Résultats des tests de sélection de l'année de référence (2/2) .	84
7.1	Coefficients de variation obtenus par la première approche . .	88
7.2	Coefficients de variation obtenus par la seconde approche . .	88
7.3	Résultats du test SMR	92
7.4	Moyenne des différents types de résidus	93
7.5	Résultats des différents tests sur les résidus de la réponse . .	95