
**Elaboration d'une table d'expérience :
comparaison de méthodes de lissage analytique
et d'ajustement statistique**

—
*Application à la population
d'un régime de retraite marocain*

Mémoire rédigé par

Olivier CLEMENT

EURIA Promotion 2003

Remerciements

D'abord je tiens à remercier l'équipe JWA-International du bureau de Bruxelles pour son accueil chaleureux, et chacun de ses membres pour ses précieux conseils tout au long de mon stage.

J'adresse ma gratitude toute particulière à M.PLOUVIER pour son dévouement et son soutien pédagogique au cours des différentes tâches qui m'ont été confiées ainsi que pour le suivi de ce mémoire.

D'autre part je veux remercier M.ACQUAVIVA pour son encadrement et ses conseils dans le domaine des statistiques.

Enfin je souhaite saluer Melle MISAK et les membres du i10 de Telecom-Brest pour leur soutien moral et logistique.

Résumé

Résumé

Ce mémoire se fonde sur une mission que j'ai effectuée au sein du cabinet de conseil en actuariat Joël Winter & Associés (JWA-International), pour le compte d'une caisse de retraite marocaine. Du fait de la modernisation des textes du Code des Assurances marocain et du souci des dirigeants de la caisse de pouvoir piloter au mieux ce régime de retraite à prestations définies, une série d'études actuarielles visant à maîtriser les paramètres du régime a été mise en place par le cabinet.

Dans ce cadre, et en vue de mettre à jour la table de décès utilisée dans un premier temps pour le calcul des rentes, j'ai réalisé pour le cabinet une étude portant sur la mortalité observée dans la population adhérente au cours de la décennie 1990. Le regroupement des données traitées a permis d'obtenir la fonction de répartition des taux de mortalité bruts par âge, sur laquelle nous avons effectué un lissage analytique de type Makeham. La comparaison des taux lissés ainsi construits à ceux de la table certifiée PF 60-64 a conforté le choix de cette dernière comme table de référence du régime.

Ce mémoire, au cours duquel nous présenterons les bases théoriques mathématiques de la méthode de lissage par Makeham ainsi que son application, a pour objectif de confronter cette méthode à une méthode statistique utilisée pour la construction de tables d'expérience, la méthode de Whittaker-Henderson, basée sur un lissage par série chronologique. Les résultats obtenus seront détaillés au moyen de plusieurs indicateurs, et les différentes tables d'expérience appliquées à un cas particulier de modélisation du régime afin de mettre en avant l'impact du choix de la table sur le passif de la caisse, et donc sur le pilotage du régime. Finalement, nous ferons un point sur la place de ces modèles internes dans le cadre de la réglementation marocaine, et envisagerons la construction d'une table prospective d'expérience.

Mots clés par domaine

- **Vocabulaire mathématique et statistique** : traitement de données, lois analytiques, Makeham, méthode du maximum de vraisemblance, lissage statistique, series chronologiques, Whittaker-Henderson, tests statistiques, test du Khi2, test de Kolmogorov-Smirnov, test d'Anderson-Darling ;
- **Vocabulaire d'assurance** : régime de retraite à prestations définies, table d'expérience, taux de mortalité (bruts, lissés), espérance de vie, PF 60-64.

Abstract

Abstract

This dissertation is based on a mission I have realized within JWA-International actuaries consulting firm, on behalf of a Moroccan pension fund. As Moroccan insurance law is slightly shifting and bringing up to date, and for the purpose of supplying all assets to the management in its piloting of the scheme, JWA consulting group has set up an actuarial studies battery to monitor all scheme's variables.

Within this context, I have updated the deceases table used in a first step to reckon annuity amounts, on the grounds of historical data linked to the death rate observed in pensioners population over the last decade. When gathering data, we have been able to carry out the distribution law of gross mortality rate by age, from which we have achieved an analytic smoothing, Makeham likely. Last, we have compared this smoothing to the certified table PF 60-64, and chosen the latter as the scheme's recommended table.

The aim of this dissertation is to collate the smoothing process by Makeham and the statistic process called Whittaker-Henderson's mean, ordinarily used to set up mortality tables by a smoothing with chronological series. We will describe the theoretical basis of both methods, and then apply them to our data. The results will be detailed and compared with different indicators. We will also explain the impact of the choice between both methods, applied to a specific pattern of the pension scheme. Finally, we will take a bearing to the situation of these inner methods inside the Moroccan insurance law, and pave the way for a prospective study of mortality in this particular population.

Key-words by field

- **Mathematical and statistical vocabulary** : data processing, analytic laws, Makeham, maximum likelihood method, statistic smoothing, chronological series, Whittaker-Henderson, statistical tests, Khi2 test, Kolmogorov-Smirnov test, Anderson-Darling test ;
- **Insurance vocabulary** : defined benefit pension scheme, table of experience, mortality rates (gross, smoothed), life expectancy, PF 60-64.

Table des matières

I	Rapport de stage	11
II	Mémoire	16
1	Principes théoriques de construction d'une table d'expérience	18
1.1	Rappels sur les tables de mortalité	18
1.2	Indicateurs utilisés	19
1.3	Exemple préliminaire	20
1.3.1	Calcul des taux et quotients de mortalité	20
1.3.2	Relations entre taux et quotients de mortalité par âge . . .	23
1.4	Méthode d'étude d'une population	25
1.4.1	Traitement des données	25
1.4.2	Estimation dans le cas de données incomplètes	26
1.4.3	Estimation dans le cas de données complètes	27
1.4.3.1	Modèle de survie	27
1.4.3.2	Fonction de vraisemblance dans le cas général	29
1.4.3.3	Estimateur dans l'hypothèse de distribution uni- forme des décès	29
1.4.3.4	Estimateur dans l'hypothèse de force constante de mortalité	30
1.4.4	Intervalle de confiance des estimations	31
2	Méthodes de lissage	33
2.1	Ajustement par loi analytique	33
2.1.1	Caractéristiques de la loi de Makeham	33
2.1.2	Le double intérêt de la loi de Makeham	34
2.1.3	Lissage par Makeham	35
2.1.3.1	Première estimation : méthode de King & Hardy	36
2.1.3.2	Seconde estimation : le développement de Taylor	37
2.1.3.3	Convergence du procédé et solution	38
2.1.4	Intervalle de confiance des $\tilde{q}_x$ lissés	38
2.2	Ajustement par lissage statistique	40
2.2.1	Rappels	40

2.2.1.1	Différences avant	40
2.2.1.2	Différences arrières	41
2.2.1.3	Différences centrales	41
2.2.1.4	Notations utilisées	41
2.2.2	La moyenne de Whittaker-Henderson	41
2.2.2.1	Critère de fidélité	41
2.2.2.2	Critère de régularité	42
2.2.2.3	Moyenne de Whittaker-Henderson	42
2.2.3	Lissage par Whittaker-Henderson	43
3	Adéquation des ajustements obtenus	46
3.1	Tests d'adéquation de l'ajustement	46
3.1.1	Principe statistique	47
3.1.2	Test du Khi-deux	48
3.1.2.1	Application du test du Khi-deux	48
3.1.2.2	Limites du test du Khi-deux	48
3.1.2.3	Intérêt du test du Khi-deux	50
3.1.3	Tests d'adéquation de Kolmogorov-Smirnov	50
3.1.3.1	Application du test de Kolmogorov-Smirnov	50
3.1.3.2	Intérêt et limites du test de Kolmogorov-Smirnov	51
3.2	Méthodes de comparaison des ajustements obtenus	51
3.2.1	critère statistique	51
3.2.2	critère de l'espérance de vie	52
3.2.3	Comparaison à la table PF60/64	53
4	Présentation de l'étude	55
4.1	Présentation du régime	55
4.1.1	Ressources du régime	55
4.1.2	Prestations du régime	55
4.1.3	Evaluation des engagements	56
4.1.4	Premières conclusions	59
4.2	Présentation de l'étude	60
4.2.1	Enjeux	60
4.2.2	Présentation des données	61
4.2.3	Traitement des données	61
5	Estimation des taux bruts de mortalité	64
5.1	Taux bruts de mortalité	64
5.2	Intervalles de confiance - avec la variance estimée	65
5.3	Intervalles de confiance - avec la variance exacte	65

6	Ajustement des taux bruts de mortalité par la loi de Makeham	68
6.1	Ajustement par Makeham	68
6.2	Intervalles de confiance des taux lissés	70
6.3	Tests d'adéquation statistiques	71
7	Ajustement des taux bruts de mortalité par la moyenne de Whittaker-Henderson	75
7.1	Ajustement par Whittaker-Henderson	75
7.2	Choix de la courbe la plus adaptée	79
7.3	Tests d'adéquation statistiques	79
8	Comparaison des deux méthodes retenues	82
8.1	Limite des tests statistiques	82
8.2	Comparaison sur les espérances de vie	83
8.3	Impact sur le régime de retraite	84
8.4	Vers une table d'expérience prospective	87
	Bibliographie	90
	Annexe 1 : Taux instantané de mortalité	92
	Annexe 2 : Tests statistiques	93
	Annexe 3 : Récapitulatif des taux de mortalité	96
	Annexe 4 : Code Matlab utilisé pour la méthode de Whittaker-Henderson	99

Table des figures

1.1	Allure d'une courbe de mortalité	20
1.2	Diagramme de Lexis	21
4.1	Engagement, cas d'un actif	57
4.2	Engagement, cas d'une pension de réversion	58
5.1	Taux bruts de mortalité par âge	65
5.2	Intervalles de confiance, variance estimée	66
5.3	Intervalles de confiance, variance exacte	66
6.1	Intervalle d'ajustement	68
6.2	Lissage des taux bruts par Makeham	69
6.3	Intervalle de confiance du lissage par Makeham	70
6.4	Résidus du lissage par Makeham - sur les $\hat{q}_x$	71
6.5	Fonctions de répartition brutes et lissées	72
6.6	Résidus du lissage par Makeham - sur les fonctions de répartition	73
6.7	Test de cohérence du lissage par Makeham	74
7.1	Lissage par Whittaker-Henderson (poids 0/1 ; h = 100 ; z = 3)	76
7.2	Lissage par Whittaker-Henderson (poids 0/1 ; h = 100 ; z = 5)	76
7.3	Lissage par Whittaker-Henderson (poids 0/1 ; h = 1000 ; z = 3)	77
7.4	Lissage par Whittaker-Henderson (poids 0/1 ; h = 1000 ; z = 5)	77
7.5	Lissage par Whittaker-Henderson (poids 0/1 ; h = 10000 ; z = 3)	78
7.6	Lissage par Whittaker-Henderson (poids 0/1 ; h = 10000 ; z = 5)	78
7.7	Résidus de répartition - whit(3)	80
7.8	Test de cohérence - whit(3)	81
8.1	Comparaisons des espérances de vie	84
8.2	Espérances de vie à 60 ans selon l'année d'observation	87

Liste des tableaux

6.1	Lissage par Makeham - test du Khi2	71
6.2	Lissage par Makeham - test de Kolmogorov-Smirnov	72
7.1	Lissage par Whittaker-Henderson, whit(3) - test du Khi2	79
7.2	Lissage par Whittaker-Henderson, whit(5) - test du Khi2	79
7.3	Lissage par Whittaker-Henderson, whit(3) - test de Kolmogorov-Smirnov	80
8.1	Comparaison des lissages obtenus - test du Khi2	83
8.2	Comparaison des lissages obtenus - espérances de vie	83
8.4	Espérances de vie selon l'année d'observation	87
8.5	Dérive de l'espérance de vie	88

Première partie

Rapport de stage

Rapport de stage

Cette partie a pour objectif d'illustrer les divers travaux que j'ai pu effectuer au cours de cette première immersion dans le milieu professionnel du conseil en actuariat. Tout d'abord j'y relaterai les modalités d'obtention et de vie de ce stage, avec une présentation du cabinet JWA-International et de ses activités, puis nous exposerons succinctement le déroulement des missions qui ont occupé la majeure partie de mon temps.

Introduction

La formation d'actuaire à l'EURIA comporte plusieurs stages : nous effectuons un stage ouvrier en fin de première année, et notre stage de fin d'études se déroule entre la deuxième et troisième année. Ce dernier stage doit se dérouler, autant que faire se peut, à l'étranger, et doit permettre aux étudiants d'appliquer les connaissances théoriques acquises au cours des deux premières années d'études.

Lors de la recherche de ce stage, ma priorité est restée de trouver une véritable expérience professionnelle hors frontières, tâche qui s'est avérée assez périlleuse vue la conjoncture du moment (quelques deux-cent cinquante envois de demande de stage à travers l'Europe anglophone, hispanophone, hellénique et francophone). De plus, après une première expérience au sein du milieu de la banque de proximité, mon désir était d'aller à la rencontre du métier d'actuaire comme il est exercé, depuis près d'un siècle, dans les sociétés d'assurances et, plus récemment, dans les cabinets de conseil en actuariat.

C'est dans cet esprit et au terme de cette recherche que j'obtenais un stage de six mois au sein du cabinet de conseil en actuariat JWA-International, à Bruxelles, stage effectué de juillet à décembre 2001 sous la direction de M. PLOUVIER Patrice, actuaire I.A.

Présentation du cabinet Joël Winter & Associés

Le cabinet Joël Winter & Associés, créé en 1983, est un cabinet indépendant d'actuaire-conseils, situé en tête du marché français de l'actuariat. Il est présent à Paris, Lyon, Nice, Rabat, Bruxelles. Il poursuit son développement dans plusieurs pays étrangers tels que la Belgique, le Maroc et la Chine.

Son activité consiste à proposer des solutions aux problèmes qui lui sont soumis en matière de retraite, prévoyance et assurance de personne. Sa clientèle est composée de compagnies d'assurances, d'institutions de prévoyance, de régimes privés de retraite, et d'une façon générale de sociétés industrielles et commerciales.

De manière plus concrète, les missions classiquement réalisées sont les suivantes :

- les appels d'offres prévoyance ou retraite,
- la mise en place de procédures et d'outils pour piloter les régimes de prévoyance ou de retraite,
- l'évaluation de passifs sociaux,
- l'audit des provisions techniques d'une mutuelle ou d'une compagnie d'assurances.

Outre les missions qu'il effectue pour les clients, le cabinet développe un certain nombre de logiciels de grande qualité permettant d'automatiser les calculs actuariels courants, et ce grâce à un pôle informatique performant. Ces logiciels sont de véritables outils pour les consultants lors des missions.

Equipe JWA à Bruxelles

L'annexe du cabinet JWA à Bruxelles est une petite structure. L'équipe est formée de M.Joël WINTER, qui partage son temps entre les différents sites où son entreprise est implantée, de M.Marc PATIGNY, actuaire belge et associé du cabinet, de M.Patrice PLOUVIER, actuaire I.A., et de Mme Valérie VANSCHPDAEL, secrétaire.

Travailler dans une telle formation présente un avantage certain de tranquillité et de connaissance de tous les membres de l'équipe. De plus, étant donné que la quasi-totalité des missions que j'ai pu traiter étaient supervisées par M.Patrice PLOUVIER, mon maître de stage, celui-ci est resté mon interlocuteur privilégié. J'ai pu bénéficier de sa pédagogie et me familiariser avec le travail du conseil en actuariat. Vu que j'ai traité uniquement des missions françaises ou marocaines, le revers de la médaille a été une faiblesse de contacts avec les clients.

Présentation d'une mission type

Mis à part les quelques missions auxquelles j'ai participé pour des études d'engagements de régimes de retraite, d'élaboration de tables d'expérience de moment, ou d'audit de compagnies d'assurances, la très grande majorité des missions auxquelles j'ai été confronté reposent sur l'évaluation d'engagements sociaux : IFC et médailles du travail. Nous présentons dans cette section le déroulement d'une telle mission type :

- réception des données 'client'
- analyse des données avec confrontation éventuelle aux données de l'année précédente
- rédaction du rapport d'analyse de données, demande de validation du client
- après validation, calcul des engagements sociaux
- tests individuels, globaux et tests de cohérence, ainsi qu'une confrontation éventuelle aux montants d'engagements de l'année précédente
- Rédaction du rapport final pour validation du client

Au cours de ces journées certes un peu monotones, j'ai pu apprécier la robustesse voire le confort des outils mis en oeuvre par le cabinet, notamment pour l'évaluation des engagements sociaux. Le logiciel LASER, outil très polyvalent puisqu'il permet d'effectuer à peu près toutes sortes de projections d'engagements, fait gagner énormément de temps pour ces missions routinières.

Cadre de vie

Partir travailler, même dans le cadre d'un stage, dans un pays étranger tel que la Belgique est loin d'être désagréable. Du moins c'est mon impression avérée après ce séjour prolongé dans le plat pays.

D'une part, du point de vue de la vie quotidienne, j'ai trouvé que le partage entre journée de travail et journée 'après le travail' constituait un bon équilibre, équilibre qui n'existe pas vraiment pendant les études (et en particulier pendant la période de rédaction du présent document).

D'autre part j'ai pleinement utilisé mes week-end à découvrir les villes belges, qui sont loin d'être dénuées de charme.

Conclusion

Pour conclure sur cette période de ma formation, mais également de ma vie, je dirai que cette expérience s'est révélée profitable à bien des égards. D'abord pour mes débuts dans la vie professionnelle d'actuaire, avec l'apprentissage du métier. Ensuite pour les personnes que j'ai pu rencontrer et qui ont été autant de chances.

Enfin, pour corriger l'idée un peu noire de la vie active, ou du monde du travail, que l'on peut se faire quand on est étudiant.

Deuxième partie

Mémoire

Introduction

L'objet de ce mémoire est l'élaboration d'une table d'expérience, ses fondements théoriques, sa mise en application. Les données utilisées en application sont issues d'une mission réalisée pour le compte d'un régime de retraite marocain.

Nous traiterons dans un premier temps des méthodes d'estimation des taux bruts de mortalité, qui passent par la construction d'intervalles d'observation, une méthode d'étude de la population, et un traitement de données adéquat.

Ensuite nous aborderons deux méthodes de lissage pour l'ajustement des taux bruts estimés, qui permettent d'obtenir une révision des estimations et d'aboutir à une courbe de mortalité lissée, en nous intéressant aux mesures permettant de juger statistiquement de la qualité du lissage produit.

Une fois ce cadre théorique posé, nous nous intéresserons au cas du régime de retraite, en essayant de dégager l'impact de la mortalité sur les engagements de la caisse et les spécificités de ce modèle à prestations définies. Nous appliquerons les méthodes décrites de façon théorique à la mortalité de la population observée, en testant leur adéquation de différentes manières.

Finalement, nous essaierons de dépasser les considérations purement techniques pour appréhender au mieux les phénomènes intervenant dans l'évaluation des engagements du régime, et pour se faire une idée de ce que peut être une table de mortalité prudente, prévisionnelle.

Chapitre 1

Principes théoriques de construction d'une table d'expérience

Ce chapitre a pour objectif de fixer un cadre théorique à l'étude qui suit. Nous y aborderons successivement les méthodes d'observation d'une population et les mesures de mortalité permettant de définir des indicateurs adéquats sur celle-ci. Ce travail préliminaire nous permettra d'aboutir à une série d'observations de taux bruts de mortalité, série exploitable pour les lissages qui constituent l'objet principal de ce mémoire.

1.1 Rappels sur les tables de mortalité

Comme les autres phénomènes démographiques, la mortalité peut être étudiée selon deux approches. Dans une perspective transversale (table de *moment*), le démographe mesure la mortalité au cours d'une période, l'année par exemple ; pour cela il considère les décès qui se produisent, durant cette année, à tous les âges. A l'opposé, s'il adopte un point de vue longitudinal (table de *génération*), il mesure la mortalité d'un groupe d'individus nés une même année tout au long de leur existence, ou durant une même période.

Le premier point de vue, celui qui nous intéresse, est le plus souvent adopté car la mortalité dépend étroitement des conditions du moment : développement économique et social, niveau sanitaire et médical... En outre les données nécessaires à une analyse transversale sont plus facilement disponibles. C'est en particulier ce type de table qui est élaboré et utilisé par les compagnies d'assurances ou les caisses de retraites pour caractériser les décès de la population affiliée (on parle alors de table d'*expérience*). Exemples (tables certifiées) : TV/TD 88-90, PF 60-64 ...

Le point de vue longitudinal, qui demande que la collecte des données se poursuive

jusqu'à l'extinction du groupe d'individus considéré, se justifie en particulier pour l'étude de questions spécifiques : évolution de la mortalité après atténuation d'une cause de décès, ou mesure de la mortalité selon les milieux sociaux. Exemple (table certifiée) : TPG 93.

Nous allons nous intéresser aux étapes de l'élaboration d'une table d'expérience et nous exposerons les mesures en analyse transversale, fondées sur le calcul des taux de mortalité.

1.2 Indicateurs utilisés

La présentation traditionnelle des tables de mortalité est basée sur plusieurs variables. Si l'on considère une population fermée¹, deux indicateurs sont étudiés :

- l_x : nombre théorique de survivants à l'âge entier x
- $d_x = l_{x+1} - l_x$: nombre théorique de décès entre x et $x+1$

A partir de ces indicateurs, l'élaboration d'une table d'expérience fait appel à différentes fonctions :

- Fonction de survie :

$$S_x = \frac{l_x}{l_0}$$

- Probabilité de survie à l'âge x :

$$p_x = \frac{l_{x+1}}{l_x}$$

- Quotient de mortalité à l'âge x :

$$q_x = \frac{d_x}{l_x} = 1 - p_x$$

- Espérance abrégée de vie :

$$e_x = \sum_{j=1}^{\omega} \frac{l_{x+j}}{l_x}$$

où ω est l'âge limite de la vie biologiquement envisageable

Par ailleurs nous utiliserons les notations classiques en la matière : minuscules pour les valeurs exactes (q_x), majuscules circonflexes pour les valeurs estimées ($\hat{Q}_x$), et "tildé" pour les valeurs lissées ($\tilde{q}_x$).

La série des q_x sert de base à la construction de la table de mortalité de moment.

¹si on suppose qu'il n'y a pas d'entrants

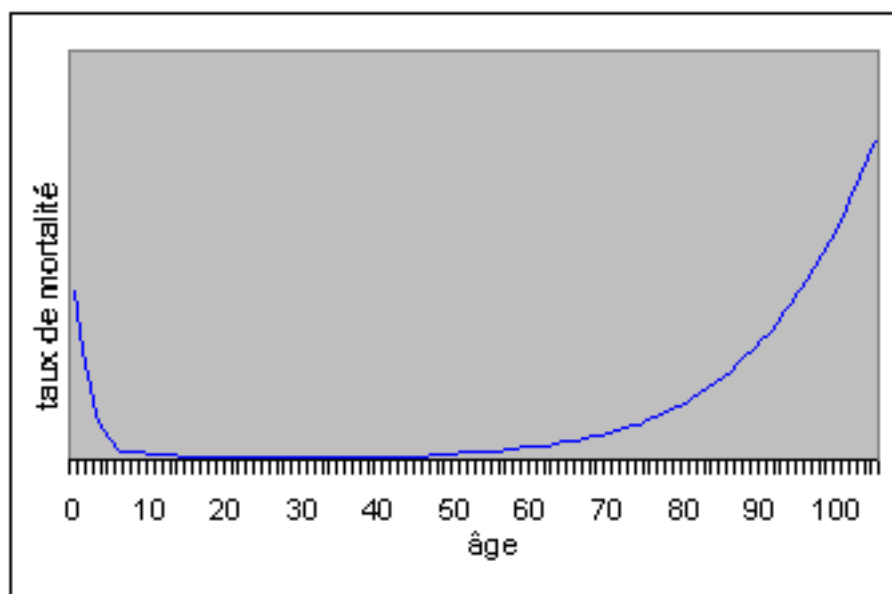


FIG. 1.1 – Allure d'une courbe de mortalité

Celle-ci se fonde sur l'artifice de la cohorte fictive. En termes généraux cet artifice consiste à faire parcourir tous les âges de la vie à un effectif arbitraire de nouveaux-nés, en lui faisant subir à ces divers âges les risques de mortalité qui auront été observés durant une période de temps (une ou plusieurs années) et dans les diverses générations présentes durant cette période. La table du moment ainsi obtenue constitue une mesure de la mortalité pour la période considérée. L'allure générale de la courbe des q_x bruts est représentée sur la figure 1.1.

Mais la détermination des q_x bruts (basés directement sur les données), série que nous devrons lisser par la suite, n'est déjà pas immédiate. En effet, selon la base de données observée, l'estimation de la mortalité exacte n'est pas directement accessible. Il faut alors procéder à une étape préliminaire, qui passe par le calcul des taux bruts de mortalité.

1.3 Exemple préliminaire

1.3.1 Calcul des taux et quotients de mortalité

La première étape pour construire une table de mortalité consiste à estimer à partir des données disponibles la population soumise au risque et le nombre de décès par âge. Cette population sera supposée homogène, c'est à dire présentant des caractéristiques de risque identiques, comme le sexe, la profession, le pays, le type de garantie d'assurances, etc., à l'exception de l'âge. De même nous n'aborderons

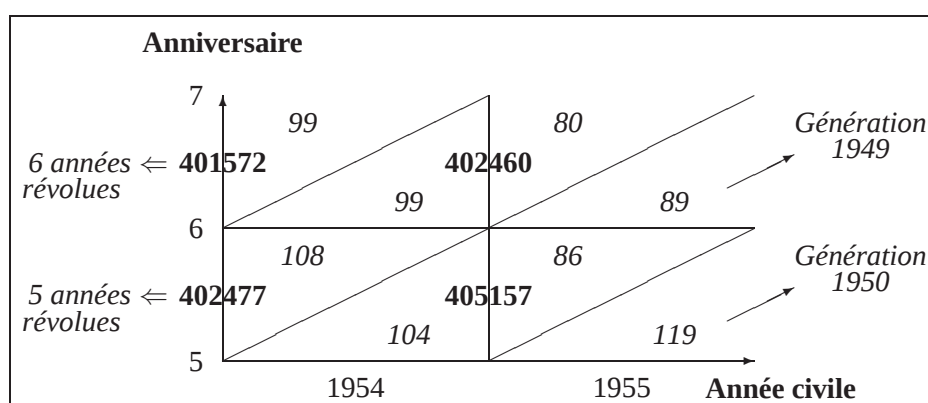


FIG. 1.2 – Diagramme de Lexis

pas les mécanismes d'autosélection et d'antisélection, largement décrits dans la littérature, qui peuvent venir biaiser la mesure de la mortalité de la population.

Le nombre de décès observés dans une population dépend du niveau de la mortalité, mais aussi de l'effectif de cette population et de sa structure par âge. Les mesures de mortalité visent donc à éliminer l'influence de ces deux derniers facteurs, afin d'isoler ce qui est proprement l'effet de la mortalité.

Le calcul le plus simple consiste à rapporter l'ensemble des décès survenus dans une population considérée, au cours d'une période donnée, à la moyenne arithmétique des effectifs de la population en début et en fin de période, pondérée par la durée de la période, c'est le taux brut de mortalité : sur un an,

$$m = \frac{\text{Deces}}{\frac{Pop_{dbut} + Pop_{fin}}{2}}$$

Lorsque les données disponibles en permettent le calcul, les taux bruts de mortalité par âge offrent une vision plus fine de ce phénomène étroitement lié à l'âge, et ne souffrent pas des distorsions introduites par les structures d'âge. Il est égal au rapport des décès survenus dans un ensemble d'âge donné à l'effectif de la population de même âge, affectée du coefficient représentant la durée d'observation exprimée en années : sur un an,

$${}_n m_x = \frac{d(x, x+n)}{Pop_{moyenne}(x, x+n)}$$

((x, x + n) représentant la longueur de l'intervalle d'âge décrit).

Exemple de calcul d'un taux brut de mortalité par âge à partir du diagramme de Lexis d'une population :

Le diagramme présenté figure 1.2 est un outil permettant de présenter, de façon ordonnée, les statistiques de base. De manière classique, les dates de calendrier sont portées en abscisses, les durées écoulées depuis la naissance, soit l'âge, figurent en ordonnées, l'échelle du temps étant la même. Dans ce système d'axes, le déroulement de la vie d'un individu est figuré par une ligne diagonale, ou ligne de vie, interrompue lors du décès en un point (*point mortuaire*). Les diagonales limitent donc les lignes de vie des personnes nées au cours d'une même année (ensemble appelé génération). On a donc indiqué sur ce diagramme :

- l'effectif de la génération, aux intersections avec les axes verticaux (effectif d'individus du même âge en années révolues). Ex : l'effectif de la population âgée de 5 ans en 1954 est égal à 402477.
- le nombre de décès localisés dans les surfaces délimitées par le système. Ex : le nombre de décès dans la population née en 1948, âgée de 5 ans, en 1954, est de 108.

Si l'on s'intéresse au taux de mortalité brut à 5 ans sur l'année 1954, on s'aperçoit que ce taux de mortalité met en jeu les décès :

- sur une année civile,
- dans deux générations,
- dans l'intervalle d'âge allant de 5 ans exacts à 6 ans exacts,

On a donc pour l'année 1954 :

$${}_1M_5 = M_5 = \frac{108 + 104}{\frac{402477 + 405157}{2}} = 0,524 \text{ pour mille}$$

Le calcul direct du quotient de mortalité quant à lui mobilise théoriquement, pour sa détermination, les données de deux années civiles consécutives. Si l'on s'intéresse au quotient de mortalité à 5 ans, on doit utiliser les décès :

- sur deux années civiles,
- dans une génération,
- dans l'intervalle d'âge allant de 5 ans exacts à 6 ans exacts,

La génération en cause a un effectif de 405157 au 1^{er} janvier 1955, soit $405157 + 104 = 405261$ à 5 ans exacts, et ce en l'absence de migrations ; donc :

$$Q_5 = \frac{104 + 86}{405261} = 0,469 \text{ pour mille}$$

Or, même si ce calcul reste valable s'il existe un courant migratoire régulier au

cours de l'année, il suppose de disposer de données sur deux années consécutives et surtout de pouvoir distinguer les décès par génération.

Dans la pratique, la nature des données, souvent présentées simplement par âge et année civile, exclut ce calcul direct. Dans notre exemple, sur l'année civile 1954, on aura l'effectif global des enfants âgés de 5 ans : 212, mais pas la décomposition par générations (108 + 104). On en est donc ramené à trouver une relation entre le quotient q_x et le taux m_x , dont le calcul est accessible.

1.3.2 Relations entre taux et quotients de mortalité par âge

Les taux de mortalité par âge (par année d'âge ou groupe de plusieurs années) sont des indices dont le calcul est très souvent possible. Comme par ailleurs ils ont une parenté très étroite avec les quotients de mortalité s'appliquant à la même étendue d'âge, c'est très fréquemment par leur intermédiaire que l'on construit une table de mortalité.

On peut établir de façon rigoureuse la relation qui existe entre taux et quotient de mortalité dans l'hypothèse d'une population stationnaire. Plaçons-nous dans l'intervalle d'âge $(x, x + n)$ et notons ${}_ne_x$ la durée moyenne de vie après l'âge x des décédés de l'intervalle $(x, x + n)$; l_x étant les survivants à l'âge x , et $d(x, x + n)$ les décès entre x et $x + n$, on doit établir une relation entre :

– le quotient

$${}_nq_x = \frac{d(x, x + n)}{l_x}$$

– le taux

$${}_nm_x = \frac{d(x, x + n)}{n[l_x - d(x, x + n)] + {}_ne_x \cdot d(x, x + n)}$$

Le dénominateur de ce taux représentant bien le nombre de personnes-années durant l'année du calcul. A partir de ces deux expressions on établit la relation :

$${}_nq_x = \frac{n \cdot {}_nm_x}{1 + (n - {}_ne_x) \cdot {}_nm_x}$$

On peut se rapprocher davantage des conditions réelles dans lesquelles se trouvent les populations en faisant l'hypothèse d'une population stable de taux d'accroissement r ; ${}_nm'_x$ étant le taux de mortalité dans une telle population, on établit cette fois que l'on a :

$${}_nq_x = \frac{K \cdot n \cdot {}_nm'_x}{1 + (n - {}_ne_x) \cdot K \cdot {}_nm'_x}$$

avec

$$K = \frac{1 - r \cdot {}_nA_x}{1 - r \cdot {}_ne_x}$$

${}_nA_x$ étant l'ancienneté dans le groupe d'âges $(x, x + n)$ des personnes de ce groupe dans la population stationnaire (ou encore l'âge moyen diminué de x).

La mise en oeuvre des formules 1.3 et 1.4 suppose la connaissance des quantités ${}_ne_x$ et ${}_nA_x$. Dans la pratique on utilisait le plus souvent la formule 1.3 avec l'hypothèse que ${}_ne_x = n/2$ (c'est-à-dire que les personnes mourant dans l'intervalle d'âge $(x, x + n)$ décèdent, en moyenne, au milieu de l'intervalle, ce qui revient à admettre pratiquement que la fonction de survie $S(x)$ est linéaire dans l'intervalle $(x, x + n)$. On a ainsi :

$${}_nq_x = \frac{2 \cdot n \cdot {}_nm_x}{2 + n \cdot {}_nm_x}$$

et, dans les deux cas les plus courants d'utilisation de la formule,

– quotient annuel

$${}_1q_x = q_x = \frac{2 \times m_x}{2 + m_x}$$

– quotient quinquennal

$${}_5q_x = \frac{10 \times {}_5m_x}{2 + 5 \times {}_5m_x}$$

A priori, l'hypothèse d'une population stationnaire limite un peu la portée de ces formules. Si cette limite est plus théorique que pratique pour la première (quotient annuel), elle peut paraître plus sérieuse pour la seconde (quotient quinquennal). De plus l'hypothèse de linéarité de $S(x)$, également acceptable pour un quotient annuel, apparaît plus que douteuse pour des calculs effectués sur cinq ans.

Ceci a motivé la recherche de liaisons entre taux et quotients sur d'autres bases. Notamment, dès le début du siècle, des tables empiriques sont apparues, étudiant par régression la correspondance existant entre taux et quotient de mortalité. Les tables de Reed et Merrell, par exemple, aboutissent à la formule :

$${}_nq_x = 1 - e^{-n \cdot {}_nm_x - 0,008 \cdot n^3 \cdot {}_nm_x^2}$$

Cette fonction empirique, croissante selon n , ressemble aux équations que l'on peut obtenir en supposant l'hypothèse de force constante de mortalité que nous verrons par la suite. Pour identifier le phénomène de mortalité et choisir ce type d'hypothèses en connaissance de cause, il faut établir des bases mathématiques rigoureuses pour l'estimation des données.

1.4 Méthode d'étude d'une population

1.4.1 Traitement des données

Depuis l'avènement des bases de données et des méthodes informatiques de traitement de données, le champ d'étude s'est considérablement amélioré. Le fait de pouvoir bénéficier de base de données individuelles en très grand nombre est certainement l'amélioration la plus importante apportée à ce jour dans le domaine qui nous concerne ici. En effet, nous allons pouvoir utiliser des méthodes bien plus efficaces en utilisant le traitement des données individuelles. Pour mener à bien une étude de mortalité, nous avons besoin des données suivantes :

- dates de début et de fin de la période d'observation
- pour chaque individu :
 - date de naissance
 - date d'entrée dans l'effectif (naissance, affiliation, entrée dans l'entreprise...)
 - date de sortie ou date de décès

Selon que nous disposerons de manière précise ou non de cette dernière clé, nous parlerons de données complètes ou incomplètes. L'estimation préalable des données brutes dépend essentiellement de cette différence dans la qualité des données.

Cette estimation repose sur le principe de contribution d'un individu à l'intervalle d'observation, c'est à dire sa proportion représentative de l'intervalle qui doit être retenue pour l'estimation. Plus précisément, si on note (i) le ième individu, on définit :

- âge à l'entrée dans l'étude : $x_i^{\text{début}}$
- âge prévu en fin d'observation : x_i^{fin}

Et sur l'intervalle d'observation d'âge $(x, x + 1]$ (dont les décès sont comptés à l'âge $x + 1$), on peut avoir trois types de contribution :

- pas de contribution à $(x, x + 1]$

$$\left\{ \begin{array}{l} x_i^{\text{fin}} \leq x \\ \text{ou} \\ x_i^{\text{début}} > x + 1 \end{array} \right.$$

- contribution pleine à $(x, x + 1]$

$$\left\{ \begin{array}{l} x_i^{\text{début}} \leq x \\ x_i^{\text{fin}} > x + 1 \end{array} \right.$$

– contribution partielle à $(x, x + 1]$

$$\begin{cases} x < x_i^{\text{début}} \leq x + 1 \\ \text{ou/et} \\ x < x_i^{\text{fin}} \leq x + 1 \end{cases}$$

Par la prise en compte de ces différentes contributions, on s'affranchit des problèmes de "bords" de l'intervalle d'observation. Le traitement des données pour le cas pratique est exposé au paragraphe 4.2.

1.4.2 Estimation dans le cas de données incomplètes

On suppose ici ne disposer que de l'âge de décès des individus dans la population observée. Nous sommes donc ramenés à étudier une population d'effectif N_x dans l'intervalle $(x, x + 1]$ où l'on observe D_x décès et $N_x - D_x$ survivants à l'âge $x + 1$.

On attache à chaque individu (i) de la population étudiée un indicateur X_i prenant à la fin de l'observation la valeur 1 si la personne décède et 0 si la personne survit :

$$X_i = \begin{cases} X_i = 1 & \text{si décès} \\ X_i = 0 & \text{sinon} \end{cases}$$

X_i suit une loi de Bernoulli de paramètre q_x , paramètre que l'on cherche à estimer. On fait l'hypothèse plausible d'indépendance des X_i (décès ou survie d'une personne indépendants de ceux d'une autre), et ainsi $D_x = \sum_{i=1}^{N_x} X_i$ suit une loi binomiale (N_x, q_x) .

La fonction de vraisemblance est simplement la probabilité binomiale de l'échantillon² soit :

$$L = \frac{N_x!}{D_x!(N_x - D_x)} (q_x)^{D_x} (1 - q_x)^{N_x - D_x} \propto (q_x)^{D_x} (1 - q_x)^{N_x - D_x}$$

L'estimateur du maximum de vraisemblance $\hat{Q}_x$ de q_x est tel que $L(\hat{Q}_x) \geq L(q_x)$, $\forall q_x$, ce qui revient à maximiser L , donc résoudre l'équation $\frac{dL}{dq_x} = 0$ (restant à vérifier par la suite que $\frac{d^2L}{dq_x^2} > 0$)

Ceci peut s'écrire en passant au logarithme népérien :

$$l = \ln L = D_x \cdot \ln q_x + (N_x - D_x) \cdot \ln(1 - q_x)$$

D'où :

$$\frac{dl}{dq_x} = \frac{D_x}{q_x} - \frac{N_x - D_x}{1 - q_x} = 0$$

On en déduit l'estimateur

$$\hat{Q}_x = \frac{D_x}{N_x}$$

² $\propto$ signifie "est proportionnel à". Comme la méthode ne dépend pas des coefficients multiplicateurs, on omettra de les écrire dans la suite

De plus, nous avons :

$$E[\hat{Q}_x] = \frac{E[D_x]}{N_x} = q_x$$

Donc $\hat{Q}_x$ est un estimateur sans biais de q_x .

On a, en outre,

$$Var[\hat{Q}_x] = \frac{Var[D_x]}{N_x^2} = \frac{q_x \cdot (1 - q_x)}{N_x}$$

Si D_x et N_x sont suffisamment grands, on peut approximer par la loi normale en utilisant le théorème de la limite centrale :

$$\frac{D_x - E[D_x]}{\sigma_{D_x}} \sim \mathcal{N}(0, 1) \implies \frac{\hat{Q}_x \cdot N_x - q_x \cdot N_x}{\sqrt{N_x \cdot q_x \cdot (1 - q_x)}} \sim \mathcal{N}(0, 1)$$

Si on ne voit pas bien ici l'utilité d'une telle théorie, vu que l'on retrouve un résultat somme toute évident, il n'en va pas de même quand on peut bénéficier de données complètes, c'est à dire des dates de décès exactes.

1.4.3 Estimation dans le cas de données complètes

Lorsque l'âge précis du décès est connu, nous pouvons utiliser cette information dans la procédure d'estimation. En notant X_i l'âge exact du décès d'un individu dans l'intervalle d'observation $(x, x + 1]$, nous prenons en compte chaque décès individuellement et la fonction de vraisemblance se construira comme le produit de chaque contribution aux décès pour les personnes décédées et de la contribution des survivants :

$$L = \underbrace{(1 - q_x)^{N_x - D_x}}_{\text{contribution des survivants}} \cdot \underbrace{\prod_{i=1}^{D_x} L_i}_{\text{contribution aux décès}}$$

Toutefois, pour obtenir la vraisemblance pour le ℓ^{me} décès, il faut connaître la fonction de répartition de la mortalité a priori, pour avoir la probabilité de décéder à l'âge X_i sachant que l'individu était en vie à l'âge x .

1.4.3.1 Modèle de survie

Pour cela, on construit un modèle de survie en temps continu, à partir des variables aléatoires suivantes :

- X âge au décès de l'individu
- $T(x) = X - x$ temps de vie futur d'une personne d'âge x

La fonction de répartition de la mortalité s'écrit alors :

$$F(x) = P[X \leq x] = 1 - S(x)$$

où $S(x)$ est la fonction de survie correspondante, définie au paragraphe (1.2) :

$$S(x) = P[X > x] = \frac{l_x}{l_0}$$

On peut alors définir les probabilités citées au paragraphe (1.2) en temps continu :

- Probabilité de décès entre x et $x + t$ pour une personne d'âge x :

$${}_tq_x = P[X < x + t | X \geq x] = 1 - \frac{S(x+t)}{S(x)}$$

- Probabilité de survie jusqu'à l'âge $x + t$ pour une personne d'âge x :

$${}_tp_x = 1 - {}_tq_x = P[X \geq x + t | X \geq x] = \frac{S(x+t)}{S(x)}$$

Le taux instantané de décès est la probabilité conditionnelle que l'individu décède entre l'âge x et l'âge $x + dx$, $dx \rightarrow 0$, sachant qu'il a survécu jusqu'à l'âge x :

$$P[x < X < x + dx | X > x] = \frac{F(x+dx) - F(x)}{1 - F(x)} \simeq \frac{f(x) \cdot dx}{1 - F(x)} = \mu_x \cdot dx$$

Où μ_x est la densité de probabilité conditionnelle, couramment appelée taux instantané de mortalité ou intensité de mortalité, et qui peut se simplifier par :

$$\mu_x = \frac{f(x)}{1 - F(x)} = \frac{\frac{d}{dx}F(x)}{S(x)} = \frac{-\frac{d}{dx}S(x)}{S(x)} = \frac{-\frac{d}{dx}l_x}{l_x}$$

Ou encore :

$$\mu_x = -\frac{d}{dx} \ln S(x) = -\frac{d}{dx} \ln l_x$$

Une autre façon de décrire le taux instantané de mortalité est présentée en annexe (1.1).

Si on connaît la forme de la fonction μ_x , on a alors, réciproquement et par intégration, la forme de la fonction de vie ${}_tp_x$:

$${}_tp_x = e^{-\int_x^{x+t} \mu_y dy}$$

Et on peut obtenir une forme analytique des quotients de mortalité ${}_tq_x$:

$${}_tq_x = 1 - {}_tp_x = 1 - e^{-\int_x^{x+t} \mu_y dy}$$

1.4.3.2 Fonction de vraisemblance dans le cas général

On peut maintenant expliciter la forme de la fonction de vraisemblance pour le i^{eme} décès par :

$$L_i = f(X_i | X > x) = \frac{f(x_i)}{S(x)} = \frac{S(X_i) \cdot \mu_{X_i}}{S(x)}$$

Ou encore, si l'on note $T_i = X_i - x$ le moment du i^{me} décès dans l'intervalle $(x, x + 1]$, alors nous avons :

$$L_i = \frac{S(x + T_i) \cdot \mu_{x+T_i}}{S(x)} = T_i p_x \cdot \mu_{x+T_i}$$

Et la contribution à L de la combinaison de tous les décès donne :

$$L = (1 - q_x)^{N_x - D_x} \cdot \prod_{i=1}^{D_x} T_i p_x \cdot \mu_{x+T_i}$$

Pour résoudre l'équation en $\hat{Q}_x$, il est nécessaire de faire des hypothèses sur la forme du taux instantané de mortalité μ_x pour pouvoir exprimer $T_i p_x \cdot \mu_{x+T_i}$ en termes de q_x . Nous allons considérer deux hypothèses de ce type.

1.4.3.3 Estimateur dans l'hypothèse de distribution uniforme des décès

Cette hypothèse correspond à ce qui était fait ³ et ce qui est encore réalisé de manière classique pour les périodes d'observation annuelles. On a alors, de manière analogue au résultat obtenu dans l'exemple préliminaire, pour un intervalle $(x, x+t]$:

$$q_x = \frac{\mu_{x+t}}{1 + t \cdot \mu_{x+t}}$$

Ou, autrement dit, $t q_x = t \cdot q_x$. On en déduit que :

$$T_i p_x \cdot \mu_{x+T_i} = q_x \quad 0 < T_i \leq 1$$

Et on retrouve la même fonction de vraisemblance qu'avec des données incomplètes, donc le même estimateur $\hat{Q}_x = \frac{D_x}{N_x}$. Cette hypothèse est le plus souvent retenue quand on veut utiliser a posteriori un lissage de type Makeham (sur une série observée annuelle) qui utilise déjà une hypothèse de mortalité exponentielle applicable à la série des $\hat{Q}_x$ ⁴.

³cf exemple, paragraphe 1.3.2

⁴cf paragraphe 2.1.1

1.4.3.4 Estimateur dans l'hypothèse de force constante de mortalité

Sous l'hypothèse exponentielle, on considère que la densité μ_x est une constante, $\mu_x = \mu$, et donc que l'on a :

$${}_tq_x = 1 - e^{-\mu \cdot t}$$

On en déduit que :

$${}_T p_x \cdot \mu_{x+T} = \mu \cdot e^{-\mu \cdot T} \quad 0 < T \leq 1$$

On en tire donc une fonction de vraisemblance exprimable en termes de μ :

$$L = (e^{-\mu})^{N_x - D_x} \cdot \prod_{i=1}^{D_x} \mu \cdot e^{-\mu \cdot T_i}$$

Que l'on réécrit sous la forme :

$$L = \mu^{D_x} \cdot e^{-\mu((N_x - D_x) - \sum_{i=1}^{D_x} T_i)}$$

On passe à la log-vraisemblance :

$$l = \ln L = D_x \cdot \ln \mu - \mu \left((N_x - D_x) + \sum_{i=1}^{D_x} T_i \right)$$

D'où :

$$\frac{dl}{d\mu} = \frac{D_x}{\mu} - \left((N_x - D_x) + \sum_{i=1}^{D_x} T_i \right)$$

Et l'on trouve comme estimateur du maximum de vraisemblance de μ :

$$\hat{\mu} = \frac{D_x}{(N_x - D_x) + \sum_{i=1}^{D_x} T_i}$$

Le dénominateur de cette expression est l'exposition exacte de la population (on voit bien l'analogie avec m_x tel qu'il avait été décrit dans l'exemple en (3.2)).

De plus nous obtenons l'estimateur du maximum de vraisemblance de q_x :

$$\hat{Q}_x = 1 - e^{-\hat{\mu}}$$

Cette méthode permet d'introduire l'expérience dans l'estimation même des données. En effet on peut, à partir des observations passées, choisir des hypothèses de mortalité particulières pour des plages d'âge préalablement définies.

Plus largement que l'étude à laquelle nous nous sommes restreints ici, il est bien sûr possible de tester d'autres hypothèses, mais surtout de prendre en compte les entrants et les censures⁵ dans les intervalles d'âge $(x, x + 1]$. Nous n'irons pas plus loin dans ce domaine vu le traitement des données que nous avons effectué pour l'étude pratique, et conseillons donc au lecteur de se reporter à la littérature existante⁶

⁵sorties volontaires ou non, sans décès

⁶voir Robert LANGMEIER intitulé *Etudes de différentes méthodes d'ajustement de tables de mortalité : application aux données d'une compagnie d'assurance* (présent dans la bibliographie), où sont approfondies les notions d'estimateurs tenant compte des censures, pour plus de développements.

1.4.4 Intervalle de confiance des estimations

L'analyse des résultats de cette étude passe ensuite par la définition d'*intervalles de confiance*.

L'intervalle de confiance permet de situer avec un niveau de confiance fixé la plage dans laquelle les taux de décès doivent se situer au regard des observations effectuées et du modèle d'estimation retenu.

Pour calculer l'intervalle de confiance, on peut soit :

- Approximer la quantité $q_x \cdot (1 - q_x)$ par la quantité empirique $\hat{Q}_x \cdot (1 - \hat{Q}_x)$;
- Utiliser la variance exacte.

Dans un premier temps, intéressons nous à l'**approximation de la quantité** $q_x \cdot (1 - q_x)$ **par la quantité empirique** $\hat{Q}_x \cdot (1 - \hat{Q}_x)$.

On définit α comme le niveau de confiance souhaité de l'estimation et C le seuil critique correspondant⁷. La loi normale vérifie :

$$\exists C \ P[-C < \mathcal{N}(0, 1) < C] = 1 - \alpha$$

et, en notant ϕ la fonction de répartition de la loi normale :

$$\phi(C) - \phi(-C) = 1 - \alpha$$

D'où :

$$C = \phi^{-1}\left(1 - \frac{\alpha}{2}\right) \text{ car } \phi(-C) = 1 - \phi(C)$$

En approximant la variance par $\frac{\hat{Q}_x \cdot (1 - \hat{Q}_x)}{N_x}$, les résidus de l'estimation doivent se comporter de façon normale :

$$P\left[-\phi^{-1}\left(1 - \frac{\alpha}{2}\right) < \frac{\hat{Q}_x - q_x}{\sqrt{\frac{\hat{Q}_x \cdot (1 - \hat{Q}_x)}{N_x}}} < \phi^{-1}\left(1 - \frac{\alpha}{2}\right)\right] = 1 - \alpha$$

On en déduit l'intervalle de confiance :

$$\left[\hat{Q}_x - \sqrt{\frac{\hat{Q}_x \cdot (1 - \hat{Q}_x)}{N_x}} \cdot \phi^{-1}\left(1 - \frac{\alpha}{2}\right) ; \hat{Q}_x + \sqrt{\frac{\hat{Q}_x \cdot (1 - \hat{Q}_x)}{N_x}} \cdot \phi^{-1}\left(1 - \frac{\alpha}{2}\right) \right]$$

Et si l'estimation est correcte, la valeur observée doit se trouver dans cet intervalle, à un niveau de confiance α .

⁷cf paragraphe 3.1.2

Pour utiliser **la méthode de la variance exacte**, on utilise de même :

$$P \left[\frac{|\hat{Q}_x - q_x|}{\sqrt{\frac{q_x \cdot (1-q_x)}{N_x}}} < \phi^{-1} \left(1 - \frac{\alpha}{2} \right) \right] = 1 - \alpha$$

Il faut donc résoudre l'équation de second degré suivante :

$$q_x^2 \cdot \left(1 + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} \right) - q_x \cdot \left(2\hat{Q}_x + \frac{1}{N_x} \cdot (\phi^{-1}(1 - \frac{\alpha}{2}))^2 \right) + (\hat{Q}_x^2) = 0$$

Par conséquent, on obtient les bornes de l'intervalle de confiance suivantes :

$$q_x^1 = \frac{2\hat{Q}_x + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} - \frac{\phi^{-1}(1 - \frac{\alpha}{2}) \cdot \sqrt{\left(4\hat{Q}_x \cdot (1 - \hat{Q}_x) + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} \right)}}{\sqrt{N_x}}}{2 \left(1 + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} \right)}$$

et

$$q_x^2 = \frac{2\hat{Q}_x + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} + \frac{\phi^{-1}(1 - \frac{\alpha}{2}) \cdot \sqrt{\left(4\hat{Q}_x \cdot (1 - \hat{Q}_x) + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} \right)}}{\sqrt{N_x}}}{2 \left(1 + \frac{(\phi^{-1}(1 - \frac{\alpha}{2}))^2}{N_x} \right)}$$

De manière usuelle, on acceptera l'estimation pour un degré de confiance à $\alpha = 95\%$, pour lequel on a $C = \phi^{-1}(1 - \frac{\alpha}{2}) \approx 1,96$.

Au terme de cette première étape, et espérant ne pas avoir déjà fatigué le lecteur, nous aurons donc abouti à une série de quotients de mortalité estimés $\hat{Q}_x$, valables selon un certain degré de confiance. Cette série sera désormais exploitable pour les méthodes de lissage global qui font l'objet du chapitre à venir.

Chapitre 2

Méthodes de lissage

Nous abordons ici la deuxième étape théorique de notre étude, au coeur de l'élaboration d'une table de mortalité. Les deux lissages que nous exposons ici débouchent, si tout se passe bien, sur la série des quotients de mortalité ajustés, notée $\tilde{q}_x$. Cette série constitue une table de mortalité, alors utilisée pour les évaluations actuarielles de l'instance concernée.

Nous allons présenter successivement deux méthodes de lissage, par loi analytique et par ajustement statistique. D'abord le lissage par la loi de Makeham, le plus ancien et encore le plus usité pour des raisons que nous exposerons ; puis l'ajustement par la moyenne de Whittaker-Henderson, méthode se rapprochant d'une moyenne mobile et ne permettant pas de donner de loi analytique de mortalité, mais qui s'est répandue du fait de la plus grande facilité d'accès aux données individuelles.

2.1 Ajustement par loi analytique

Dans cette section, nous allons introduire la loi de mortalité dite de Makeham, étudier ses caractéristiques, détailler son intérêt, puis présenter le lissage, réalisé une nouvelle fois au moyen de la méthode du maximum de vraisemblance.

2.1.1 Caractéristiques de la loi de Makeham

L'hypothèse fondamentale de Makeham est que le taux instantané de mortalité à partir d'un certain âge s'exprime par :

$$\mu_x = \alpha + \beta \cdot c^x$$

avec $\alpha > 0$; $\beta > 0$; $c > 1$.

L'interprétation de cette formule est la décomposition de la mortalité en :

- Un élément indépendant de l'âge et qui est en relation avec la mortalité accidentelle : α
- Un élément croissant exponentiellement avec l'âge : $\beta \cdot c^x$

Remarquons qu'après intégration¹, cette formule peut se mettre sous la forme :

$$l_x = k \cdot s^x \cdot g^{c^x}$$

avec $k = l_{x_0} \cdot e^{\alpha x_0 + \frac{\beta}{\ln c} \cdot c^{x_0}}$, constante positive arbitraire ;
et $s = e^{-\alpha} < 1$; $g = e^{-\frac{\beta}{\ln c}} < 1$.

De plus, si une loi de mortalité suit une loi de Makeham, le graphique représentant la fonction $\ln(|q_{x+1} - q_x|)$ en fonction de l'âge doit être ajustable par une fonction linéaire.

C'est donc à partir de ce graphique que l'on déterminera l'intervalle sur lequel sera appliqué l'ajustement makehamien.

D'expérience, on sait qu'il faudra supposer x suffisamment grand (≥ 30 ans).

2.1.2 Le double intérêt de la loi de Makeham

La loi de Makeham présente deux principaux intérêts.

Tout d'abord, notons qu'une conséquence de la formule initiale est :

$$\ln(p_x) = \ln(1 - q_x) = -\alpha - \beta \cdot \frac{(c - 1)}{\ln(c)} \cdot c^x$$

Sous l'aspect de cette formule, le premier intérêt de la loi de de Makeham est d'être un procédé d'ajustement (ou de lissage) applicable à des observations brutes de taux annuels de mortalité. Un tel lissage - nécessaire pour faire disparaître des irrégularités aléatoires - peut être réalisé par d'autres moyens, et notamment par une rectification de données brutes pour donner une certaine régularité à la série estimée des $\hat{Q}_x$.

Cette formule a eu pendant très longtemps une faveur particulière auprès des assureurs. Cela tenait aux propriétés simples qu'elle présente dans les calculs d'assurances sur plusieurs têtes. Ainsi, le second intérêt de cette loi réside dans la propriété de l'âge moyen et de son vieillissement uniforme pour un groupe disparaissant au premier décès : si l'on note $G = (x_1, x_2, x_3, \dots, x_n)$ un groupe de n individus, on a

$$\mu_G = \mu_{x_1, x_2, x_3, \dots, x_n} = n \cdot \alpha + \beta \cdot (c^{x_1} + c^{x_2} + \dots + c^{x_n})$$

¹Le lecteur pourra se reporter à la section portant sur la méthode de King & Hardy pour une démonstration plus détaillée

. Si la loi de survie de toutes les têtes composant le groupe est une loi de Makeham, on pourra remplacer dans tous les calculs les âges effectifs par un seul âge moyen ξ évaluable à partir de la constante c de la formule.

En effet, posons $n \cdot c^\xi = c^{x_1} + c^{x_2} + \dots c^{x_n}$; on a alors :

$$\mu_{x_1, x_2, x_3, \dots, x_n} = n \cdot (\alpha + \beta \cdot c^\xi) = \underbrace{\mu_{\xi, \xi, \xi, \dots, \xi}}_{n \text{ fois}}$$

Cela permet en particulier d'établir des barèmes pour des assurances sur deux têtes avec une seule entrée pour l'âge.

2.1.3 Lissage par Makeham

On utilise la méthode du maximum de vraisemblance pour estimer les paramètres de la loi d'ajustement, en l'occurrence la loi de Makeham.

Supposons que l'on dispose d'un ensemble d'observations de taux annuels de mortalité entre les âges entiers x_0 et x_M sous la forme de nombres observés de vivants d'âge x . On reprend les notations précédentes : N_x l'effectif en début de période, et D_x le nombre de décès entre les âges x et $x + 1$, $\forall x$ tel que $x_0 \leq x \leq x_M$.

La forme générale de la fonction de vraisemblance associée, au coefficient multiplicateur près, est alors, compte tenu de l'hypothèse d'indépendance des têtes :

$$L = \prod_{x=x_0}^{x_M} (q_x^{D_x}) \cdot (p_x)^{(N_x - D_x)}$$

On écrira l'hypothèse de Makeham sous la forme :

$$\ln(p_x) = -\alpha - b \cdot e^{\gamma x}$$

Avec :

$$b = \beta \cdot \frac{(c - 1)}{(\ln(c))} \text{ et } \gamma = \ln(c)$$

De la sorte, la fonction suivante :

$$\Phi = \ln(L) = \sum_{x=x_0}^{x_M} (D_x \cdot \ln(q_x) + (N_x - D_x) \cdot \ln(p_x))$$

est une fonction de α , b , et γ . Les meilleurs coefficients de la formule seront ceux qui maximisent la fonction de vraisemblance donc la fonction Φ pour les valeurs estimées $\hat{Q}_x$ (avec bien sûr $\hat{P}_x = 1 - \hat{Q}_x$). Une condition nécessaire pour cela est l'annulation des dérivées partielles de Φ par rapport à α , b , et γ , restant à vérifier par la suite que l'optimum trouvé est bien un maximum.

On a :

$$\begin{cases} d = \frac{\partial \Phi}{\partial \alpha} = -\sum N_x + \sum \frac{D_x}{1-p_x} & = 0 \\ e = \frac{\partial \Phi}{\partial b} = -\sum e^{\gamma x} \cdot N_x + \sum \frac{e^{\gamma x}}{1-p_x} \cdot D_x & = 0 \\ f = \frac{\partial \Phi}{\partial \gamma} = -b \cdot \sum x \cdot e^{\gamma x} N_x + b \cdot \sum \frac{x e^{\gamma x}}{1-p_x} \cdot D_x & = 0 \end{cases}$$

Afin de résoudre ce système de 3 équations non linéaires en α , b , et γ , nous utiliserons les deux estimations détaillées ci-après.

2.1.3.1 Première estimation : méthode de King & Hardy

On part d'une première valeur estimée des coefficients α_0 , b_0 , et γ_0 en utilisant la méthode de King & Hardy, valeurs pour lesquelles :

$$\begin{cases} e = e_0 \\ d = d_0 \\ f = f_0 \end{cases}$$

On a vu que le taux instantané de mortalité s'exprimait par : $\mu_x = \alpha + \beta \cdot c^x$

On en déduit :

$$l_x = l_{x_0} \cdot e^{-\int_{x_0}^x \alpha + \beta \cdot c^t dt}$$

D'où :

$$l_x = k \cdot s^x \cdot g^{c^x}$$

avec : $k = l_{x_0} \cdot e^{\alpha x_0 + \frac{\beta}{\ln c} \cdot c^{x_0}}$; $s = e^{-\alpha} < 1$; $g = e^{-\frac{\beta}{\ln c}} < 1$.

Si l'on remplace x par $x + l$ et x_0 par x , nous aboutissons à :

$$p_x = \frac{l_{x+1}}{l_x} = s \cdot g^{c^x \cdot (c-1)}$$

D'où :

$$\ln(p_x) = \ln(s) + c^x \cdot (c-1) \cdot \ln(g)$$

Dans la pratique, sur la plage d'âge $(x, x + n]$ que l'on cherche à lisser, on commence par calculer la quantité

$$A_x = \sum_x^{x+n-1} \ln(\hat{P}_x) = n \cdot \ln(s) + (c-1) \cdot \ln(g) \cdot c^x \cdot \frac{c^n - 1}{c - 1}$$

Puis on calcule

$$\frac{A_{x+n} - A_{x+2n}}{A_x - A_{x+n}} = \frac{\ln(g) \cdot (c^n - 1) \cdot (c^{x+n} - c^{x+2n})}{\ln(g) \cdot (c^n - 1)(c^x - c^{x+n})} = c^n$$

Et donc :

$$c = \left(\frac{A_{x+n} - A_{x+2n}}{A_x - A_{x+n}} \right)^{\frac{1}{n}}$$

La valeur de g est déduite de l'expression de $A_x - A_{x+n}$, par :

$$A_x - A_{x+n} = \ln(g) \cdot (c^n - 1) \cdot (c^x - c^{x+n})$$

Ainsi :

$$g = e^{\frac{A_x - A_{x+n}}{(c^n - 1) \cdot (c^x - c^{x+n})}}$$

Enfin, s est obtenu par l'expression de A_x obtenue précédemment :

$$s = e^{\frac{A_x - \ln(g) \cdot c^x (c^n - 1)}{n}}$$

Après avoir estimé c , g , et s , on peut déduire les valeurs de α et de β .

Ainsi :

$$\alpha = -\ln(s)$$

$$\beta = -\ln(g) \cdot \ln(c)$$

Puis β et c vont permettre de déduire les valeurs de b et γ .

En estimant p_x par la quantité $\hat{P}_x$ dans les égalités du haut de la page 36, ceci va donner une première estimation de d , e , f , que l'on notera : d_0 , e_0 , et f_0 .

2.1.3.2 Seconde estimation : le développement de Taylor

Ecrivons :

$$\begin{cases} \alpha = \alpha_0 - u \\ b = b_0 - v \\ \gamma = \gamma_0 - w \end{cases} \iff \begin{cases} \alpha - \alpha_0 = -u \\ b - b_0 = -v \\ \gamma - \gamma_0 = -w \end{cases}$$

Si u , v , w sont des écarts minimes, on aura approximativement :

$$\begin{cases} d = d_0 - ug - vh - wi \\ e = e_0 - uh - vj - wk \\ f = f_0 - ui - vk - wl \end{cases}$$

Ceci découle directement de l'application du développement de Taylor au second ordre aux fonctions e , d , et f aux points α_0 , b_0 et γ_0 , en posant :

$$\begin{cases} g = \frac{\partial^2 \Phi}{\partial \alpha^2} = \sum -\frac{p_x}{(1-p_x)^2} \cdot D_x \\ h = \frac{\partial^2 \Phi}{\partial \alpha \partial b} = \sum -\frac{e^{\gamma x} \cdot p_x}{(1-p_x)^2} \cdot D_x \\ i = \frac{\partial^2 \Phi}{\partial \alpha \partial \gamma} = b \sum -\frac{x \cdot e^{\gamma x}}{(1-p_x)^2} \cdot D_x \\ j = \frac{\partial^2 \Phi}{\partial \beta^2} = \sum -\frac{e^{2\gamma x} \cdot p_x}{(1-p_x)^2} \cdot D_x \\ k = \frac{\partial^2 \Phi}{\partial b \partial \gamma} = \sum -x \cdot e^{\gamma x} \cdot N_x + \sum \left(\frac{x \cdot e^{\gamma x}}{1-p_x} - b \frac{x \cdot e^{2\gamma x} \cdot p_x}{(1-p_x)^2} \right) \cdot D_x \\ l = \frac{\partial^2 \Phi}{\partial \gamma^2} = -b \sum x^2 \cdot e^{\gamma x} \cdot N_x + b \sum \left(\frac{x^2 \cdot e^{\gamma x}}{1-p_x} - b \frac{x^2 \cdot e^{2\gamma x} \cdot p_x}{(1-p_x)^2} \right) \cdot D_x \end{cases}$$

Sous forme matricielle, posons :

$$J = \begin{pmatrix} g & h & i \\ h & j & k \\ i & k & l \end{pmatrix}$$

Nous obtenons alors :

$$\begin{bmatrix} d \\ e \\ f \end{bmatrix} = \begin{bmatrix} d_0 \\ e_0 \\ f_0 \end{bmatrix} - J \times \begin{bmatrix} u \\ v \\ w \end{bmatrix}$$

La résolution du système d'équations aux dérivées partielles nulles impose la condition :

$$\begin{bmatrix} d \\ e \\ f \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}$$

Cette condition est satisfaite si et seulement si :

$$\begin{bmatrix} d \\ e \\ f \end{bmatrix} = J^{-1} \times \begin{bmatrix} d_0 \\ e_0 \\ f_0 \end{bmatrix}$$

2.1.3.3 Convergence du procédé et solution

Le calcul fait au paragraphe précédent, en estimant les coefficients de J avec les valeurs estimées $\hat{P}_x$, conduit à une seconde approximation :

$$\begin{cases} \alpha_1 &= \alpha_0 - u \\ b_1 &= b_0 - v \\ \gamma_1 &= \gamma_0 - w \end{cases}$$

Il suffit d'itérer ce procédé qui converge vers la solution à condition de partir d'une valeur de départ pas trop éloignée de l'optimum, chose qui a été réalisée par la méthode de King & Hardy.

Ainsi, les approximations qui convergent vers la solution permettront de calculer le taux instantané de mortalité, et donc les $\tilde{q}_x$ lissés.

2.1.4 Intervalle de confiance des $\tilde{q}_x$ lissés

Dans cette section, nous nous attacherons à calculer l'intervalle de confiance des q_x lissés en vue de mesurer avec un niveau de confiance fixé la plage dans laquelle ces taux de décès ajustés doivent se situer au regard des estimations effectuées et du modèle retenu.

Pour un nombre d'observations assez grand, la méthode du maximum de vraisemblance conduit à ce que les écarts entre les valeurs de (α, b, γ) avec l'estimation la plus vraisemblable, c'est à dire $(\hat{\alpha}, \hat{b}, \hat{\gamma})$, soient des variables normales centrées dont la matrice des variances-covariances est, au signe près :

$$I^{-1} = E[J]^{-1} = K$$

où I est la matrice d'information de Fisher, ce qui correspond au théorème de convergence classique, énoncé page 27 pour les estimés $\hat{q}_x$:

$$\sqrt{n}(\theta_n - \theta) \sim \mathcal{N}(0, I^{-1})$$

Comme $E[D_x] = N_x \cdot (1 - p_x)^2$, la matrice $E[J]$ est de la forme :

$$\begin{pmatrix} \hat{g} & \hat{h} & \hat{i} \\ \hat{h} & \hat{j} & \hat{k} \\ \hat{i} & \hat{k} & \hat{l} \end{pmatrix}$$

avec :

$$\begin{cases} \hat{g} &= \sum -\frac{p_x}{(1-p_x)} \cdot N_x \\ \hat{h} &= \sum -\frac{e^{\gamma x} \cdot p_x}{(1-p_x)} \cdot N_x \\ \hat{i} &= b \cdot \sum -\frac{x \cdot e^{\gamma x}}{(1-p_x)} \cdot N_x \\ \hat{j} &= \sum -\frac{e^{2\gamma x} \cdot p_x}{(1-p_x)} \cdot N_x \\ \hat{k} &= b \cdot \sum -\frac{x \cdot e^{2\gamma x} \cdot p_x}{(1-p_x)} \cdot N_x \\ \hat{l} &= b \cdot \sum -\frac{x^2 \cdot e^{2\gamma x} \cdot p_x}{(1-p_x)} \cdot N_x \end{cases}$$

On a par hypothèses :

$$\ln(p_x) = -\alpha - b \cdot e^{\gamma x}$$

Aux alentours de la valeur la plus vraisemblable, en estimant les coefficients de $E[J]$ à l'aide des $\hat{P}_x$, on aura approximativement pour la valeur lissée $\tilde{q}_x$:

$$\ln(\tilde{p}_x) = -\hat{\alpha} - \hat{b} e^{\hat{\gamma} x} - d\alpha - db \cdot e^{\hat{\gamma} x} - \hat{b} x \cdot e^{\hat{\gamma} x} \cdot d\gamma$$

où $d\alpha$, db , et $d\gamma$ sont des variables normales.

Dès lors :

$$Var[\ln(\tilde{p}_x)] = \begin{bmatrix} \hat{1} & e^{\hat{\gamma} x} & \hat{b} x e^{\hat{\gamma} x} \end{bmatrix} \times K \times \begin{bmatrix} 1 \\ e^{\hat{\gamma} x} \\ \hat{b} x e^{\hat{\gamma} x} \end{bmatrix}$$

Un intervalle de confiance à 95% pour $\ln(\tilde{p}_x)$ est par conséquent :

$$\left[-\hat{\alpha} - \hat{b} e^{\hat{\gamma} x} - 1,96 \times \sqrt{Var[\ln(\tilde{p}_x)]} ; -\hat{\alpha} - \hat{b} e^{\hat{\gamma} x} + 1,96 \times \sqrt{Var[\ln(\tilde{p}_x)]} \right]$$

²cf : D_x suit une loi binomiale

A partir de cet intervalle, on peut facilement déterminer un intervalle de confiance pour $\tilde{q}_x$, à l'intérieur duquel devrait se trouver la valeur estimée $\hat{Q}_x$.

Vu que l'on a $\tilde{q}_x = 1 - e^{\ln(\tilde{p}_x)}$, cet intervalle s'écrit :

$$\left[1 - e^{\left(\hat{\alpha} - \hat{b} e^{\hat{\gamma}x} - 1,96 \times \sqrt{\text{Var}[\ln(\tilde{p}_x)]}\right)} ; 1 - e^{\left(\hat{\alpha} - \hat{b} e^{\hat{\gamma}x} + 1,96 \times \sqrt{\text{Var}[\ln(\tilde{p}_x)]}\right)} \right]$$

2.2 Ajustement par lissage statistique

Nous avons vu dans la section précédente une méthode dont le principal intérêt est de fournir une loi analytique de mortalité, permettant alors d'homogénéiser les calculs d'assurance effectués sur ces bases. Très fréquemment, la miniaturisation de l'informatique permet désormais de substituer à la lecture de tarifs volumineux la consultation immédiate d'un fichier pour laquelle il n'y a pas de difficulté à entrer plusieurs âges, ni à effectuer un calcul exact, en se passant de l'hypothèse de Makeham.

De nombreuses méthodes sont donc apparues, visant à ne plus appliquer un modèle théorique (avec une hypothèse de mortalité) à la série des valeurs estimées $\hat{Q}_x$, mais plutôt à "remonter des données" en essayant de les lisser au mieux statistiquement. Parmi celles-ci, nous allons maintenant exposer la méthode utilisant la moyenne de Whittaker-Henderson. Cette méthode se base sur un lissage par série chronologique, dont nous allons donc effectuer quelques légers rappels.

2.2.1 Rappels

2.2.1.1 Différences avant

Les différences avant d'ordre 1 à z sont définies, pour un intervalle unitaire, comme suit :

$$\begin{aligned} \Delta f(x) &= f(x+1) - f(x) \\ \Delta^2 f(x) &= \Delta(\Delta f(x)) \\ &= \Delta(f(x+1) - f(x)) \\ &= f(x+2) - 2f(x+1) + f(x) \\ \\ \Delta^z f(x) &= \Delta(\Delta^{z-1} f(x)) \\ &= \Delta(\Delta(\Delta^{z-2} f(x))) \\ &= \dots \end{aligned}$$

2.2.1.2 Différences arrières

Les différences arrières d'ordre 1 à z sont définies, pour un intervalle unitaire, comme suit :

$$\begin{aligned}
 \nabla f(x) &= f(x) - f(x-1) \\
 \nabla^2 f(x) &= \nabla(\nabla f(x)) \\
 &= \nabla(f(x) - f(x-1)) \\
 &= f(x-2) - 2f(x-1) + f(x-2) \\
 \\
 \nabla^z f(x) &= \nabla(\nabla^{z-1} f(x)) \\
 &= \nabla(\nabla(\nabla^{z-2} f(x))) \\
 &= \dots
 \end{aligned}$$

2.2.1.3 Différences centrales

Les différences centrales d'ordre 1 à z sont définies, pour un intervalle unitaire, comme suit :

$$\begin{aligned}
 \delta f(x) &= f(x + \frac{1}{2}) - f(x - \frac{1}{2}) \\
 \delta^2 f(x) &= \delta(\delta f(x)) \\
 &= \delta(f(x + \frac{1}{2}) - f(x - \frac{1}{2})) \\
 &= f(x+1) - 2f(x) + f(x-1) \\
 \\
 \delta^z f(x) &= \delta(\delta^{z-1} f(x)) \\
 &= \delta(\delta(\delta^{z-2} f(x))) \\
 &= \dots
 \end{aligned}$$

2.2.1.4 Notations utilisées

La notation $\hat{Q}_x$ désignera les estimations pour l'âge x .

Les w_x seront les poids associés aux estimations.

Les $\tilde{q}_x$ seront quant à eux les valeurs ajustées.

2.2.2 La moyenne de Whittaker-Henderson

La moyenne de Whittaker-Henderson se base sur deux critères : un critère de fidélité et un critère de régularité.

2.2.2.1 Critère de fidélité

Le critère de fidélité mesure la distance euclidienne classique entre la mortalité lissée $\tilde{q}_x$ et la mortalité estimée $\hat{Q}_x$, chaque distance étant pondéré par w_x :

$$F = \sum_{x=1}^n w_x (\tilde{q}_x - \hat{Q}_x)^2$$

n étant l'âge de fin de l'étude.

La distance a pour fonction évidente de mesurer l'écart aux valeurs estimées. Les poids, eux, sont ajustables selon différentes modalités. On peut par exemple définir :

- *a priori*, si l'on remarque des incohérences pour certaines estimations ou des absences de données (en général au début et à la fin du lissage),

$$\begin{cases} w_x = 0 & \text{si } \hat{Q}_x = 0 \\ w_x = 1 & \text{sinon} \end{cases}$$

- *a posteriori*, si l'on remarque que le lissage tient trop compte de certaines valeurs estimées au regard de l'effectif présent à cet âge (typiquement en fin de lissage, âges élevés),

$$w_x = \frac{N_x}{\bar{N}}$$

où N_x représente l'effectif observé d'âge x et $\bar{N} = \frac{\sum_{x=1}^n N_x}{n}$ l'effectif moyen sur la plage d'âges de l'étude.

On peut évidemment composer ces deux définitions. Par contre on ne pourra pas définir de distance du type Khi-deux³, vu que l'on n'a pas de forme explicite pour la série des valeurs lissées $\tilde{q}_x$, et donc pas de variance exprimable des résidus du lissage.

2.2.2.2 Critère de régularité

Le critère de régularité mesure une distance (dépendante de z) entre valeurs lissées par différences avant :

$$R = \sum_{x=1}^{n-z} (\Delta^z \tilde{q}_x)^2$$

Où z fixe le degré du polynôme utilisé pour le critère de régularité. Par exemple, si l'on prend $z = 1$, on retrouvera la distance euclidienne entre deux valeurs lissées consécutives : $(\tilde{q}_{x+1} - \tilde{q}_x)^2$.

Pour tenir compte de suffisamment d'informations consécutives, z est généralement compris entre 2 et 5. De plus, vu la méthode par différences avant, on doit arrêter le critère à une plage d'âges de longueur $(n - z)$.

2.2.2.3 Moyenne de Whittaker-Henderson

La moyenne de Whittaker-Henderson s'obtient par une combinaison linéaire de la fidélité et de la régularité en mettant plus ou moins l'accent sur la régularité au

³cf paragraphe 3.2.1

moyen du paramètre h :

$$M = F + h R = \sum_{x=1}^n w_x (\tilde{q}_x - \hat{Q}_x)^2 + h \sum_{x=1}^{n-z} (\Delta^z \tilde{q}_x)^2$$

2.2.3 Lissage par Whittaker-Henderson

Les valeurs ajustées $\tilde{q}_x$ pour $x = 1, 2, \dots, n$ seront celles qui minimisent la mesure composite M qui est fonction des n valeurs inconnues de $\tilde{q}_x$.

Pour trouver les $\tilde{q}_x$, une condition nécessaire est que les n équations provenant des dérivées partielles de M par rapport à chacun des $\tilde{q}_x$ soient nulles, et elle est suffisante étant donnée la convexité des composantes de M donc de M .

Il faut donc résoudre :

$$\forall x, \frac{\partial M}{\partial \tilde{q}_x} = 0$$

Pour trouver la solution, nous allons utiliser la solution matricielle. Nous posons les vecteurs et matrices suivants :

$$Q = \begin{bmatrix} \hat{Q}_1 \\ \hat{Q}_2 \\ \dots \\ \dots \\ \dots \\ \hat{Q}_n \end{bmatrix} ; \quad V = \begin{bmatrix} \tilde{q}_1 \\ \tilde{q}_2 \\ \dots \\ \dots \\ \dots \\ \tilde{q}_n \end{bmatrix} ; \quad W = \begin{pmatrix} w_1 & 0 & \dots & \dots & 0 \\ 0 & w_2 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots \\ 0 & \dots & \dots & 0 & w_n \end{pmatrix}$$

Où W est une matrice diagonale, matrice des poids $(w_1, w_2, \dots, w_n)$.

Nous pouvons réécrire F en notation matricielle

$$F = (V - Q)^T \times W \times (V - Q)$$

Nous devons également définir le vecteur suivant, qui représente les différences avant d'ordre z de V , vecteur colonne de $n - z$ lignes :

$$\Delta^z V = \begin{bmatrix} \Delta \tilde{q}_1 \\ \Delta \tilde{q}_2 \\ \dots \\ \dots \\ \dots \\ \Delta \tilde{q}_{n-z} \end{bmatrix}$$

ce qui nous permet d'écrire R en notation matricielle :

$$R = (\Delta^z V)^T \cdot (\Delta^z V)$$

Pour trouver $\Delta^z V$, réécrivons ce vecteur en fonction de V , en utilisant une matrice spéciale K_z qui contient les coefficients binomiaux d'ordre z avec alternance de signe. La dimension d'une telle matrice est $(n - z) \times n$. Par exemple, pour $z = 3$, on aura :

$$K_3 = \left(\underbrace{\begin{pmatrix} -1 & 2 & -1 & 0 & \dots & \dots & 0 \\ 0 & -1 & 2 & -1 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & \dots & \dots & 0 & -1 & 2 & -1 \end{pmatrix}}_{n \text{ colonnes}} \right) \left. \vphantom{\begin{pmatrix} -1 & 2 & -1 & 0 & \dots & \dots & 0 \\ 0 & -1 & 2 & -1 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & \dots & \dots & 0 & -1 & 2 & -1 \end{pmatrix}} \right\} (n - 3) \text{ lignes}$$

Nous avons alors : $\Delta^z V = K_z \cdot V$

La mesure M devient alors :

$$\begin{aligned} M &= F + h R \\ &= (V - Q)^T \cdot W \cdot (V - Q) + h (\Delta^z V)^T \cdot (\Delta^z V) \\ &= (V - Q)^T \cdot W \cdot (V - Q) + h V^T \cdot K_z^T \cdot K_z \cdot V \end{aligned}$$

Après développement puis regroupement des termes, nous obtenons :

$$M = V^T \cdot W \cdot V - 2 V^T \cdot W \cdot Q + Q^T \cdot W \cdot Q + h V^T \cdot K_z^T \cdot K_z \cdot V$$

Dérivons M vectoriellement :

$$\frac{\partial M}{\partial V} = 2 W \cdot V - 2 W \cdot Q + 2h K_z^T \cdot K_z \cdot V = 0$$

Après résolution, nous trouvons :

$$W \cdot V + h K_z^T \cdot K_z \cdot V = W \cdot Q$$

En posant $C = W + h K_z^T \cdot K_z$, nous pouvons écrire :

$$C \cdot V = W \cdot Q$$

Si C est inversible, nous pouvons obtenir la solution suivante :

$$V = C^{-1} \cdot W \cdot Q$$

Au terme de ce calcul, nous aurons donc abouti à une série de valeurs lissées $\tilde{q}_t$ de manière statistique. Par rapport au lissage par Makeham exposé précédemment, il n'existe pas d'intervalles de confiance, étant donné que la méthode ne se base pas

sur des paramètres estimés.

Le choix des coefficients $\{\{w_x\}_{x=1,n}, h, z\}$ du modèle s'effectue souvent selon l'appréciation humaine, jusqu'à obtenir la courbe qui paraît la plus adéquate. Comme nous le verrons en application, nous avons retenu deux critères favoris de choix : d'une part une allure de courbe conservant une forme exponentielle, et d'autre part selon des critères d'adéquation statistique résultant des taux de mortalité lissés.

Reste maintenant à pouvoir mesurer l'efficacité des deux méthodes présentées dans cette partie, et à pouvoir les comparer à partir d'indicateurs pertinents. Ceci fait l'objet du prochain chapitre.

Chapitre 3

Adéquation des ajustements obtenus

Les tests statistiques en matière de tables de mortalité sont utilisés de façon classique pour comparer les observations faites sur une population avec une table de mortalité de référence. On peut également les employer dans le cadre d'un lissage, à condition de respecter une certaine cohérence vis à vis du but recherché, et donc de ne pas les interpréter de manière trop rigide.

On retiendra notamment l'application du principe de prudence sur toute la table, ou sur un intervalle particulier. Par exemple, pour une population affiliée à une assurance décès, la mortalité ne devra pas être sous-estimée, et a contrario on ne devra pas surestimer la mortalité d'une assurance de rente.

3.1 Tests d'adéquation de l'ajustement

Dans la suite, nous noterons F_n (respectivement F) la fonction de répartition empirique de la mortalité brute ou estimée (respectivement de la mortalité lissée), définie comme au paragraphe 1.4.3.1 par :

$$F_n(x) = P[X \leq x] = 1 - S_n(x) = 1 - \frac{\hat{L}_x}{l_0}$$

avec $\hat{L}_x$ nombre estimé de survivants à l'âge x (respectivement : $F(x) = 1 - \frac{\tilde{L}_x}{l_0}$) et l_0 le nombre exact de survivants en 0, c'est à dire la population totale observée. On a donc les relations :

$$\begin{aligned} F_n(x+1) - F_n(x) &= \frac{\hat{L}_x - \hat{L}_{x+1}}{l_0} \\ &= \frac{\hat{D}_x}{l_0} \\ &= \frac{\hat{L}_x \cdot \hat{Q}_x}{l_0} \\ &= \hat{Q}_x \cdot \prod_{y=0}^{x-1} (1 - \hat{Q}_y) \end{aligned}$$

la dernière égalité venant du fait que :

$$\hat{L}_x = l_0 \cdot \prod_{y=0}^{x-1} (1 - \hat{Q}_y)$$

(mêmes relations avec la fonction F de mortalité lissée)

3.1.1 Principe statistique

L'examen de la loi de probabilité empirique de mortalité associée à une population dont la loi parente est inconnue, permet de choisir parmi les lois usuelles celle qui lui "ressemble" le plus, ou ici de confirmer l'adéquation d'une loi particulière. Si notre choix s'oriente vers une certaine loi P de fonction de répartition F , on pourra retenir l'hypothèse que l'échantillon provient de cette loi si la distance entre la fonction de répartition théorique F et la fonction de répartition empirique F_n est faible.

Ayant fait le choix d'une certaine distance d entre fonctions de répartition, on se fixera une règle de décision qui s'énonce ainsi :

"Si l'événement $d(F_n ; F) < C$ est réalisé, alors je retiens l'hypothèse qu'il s'agit d'un échantillon de la loi de fonction de répartition F ."

On peut cependant se tromper en rejetant cette hypothèse alors que F est bien la fonction de répartition des variables de l'échantillon ; cette erreur se produit avec une probabilité qui s'exprime par :

$$\alpha = P[d(F_n ; F) > C]$$

Si on veut que ce risque d'erreur soit faible, on fixera une valeur α faible à cette probabilité (par exemple 5% ou 1%) et cette valeur permettra alors de préciser la valeur de la constante C qui apparaît dans la règle de décision, si on connaît la loi de probabilité de la variable aléatoire $d(F_n ; F)$.

Nous aurons ainsi réalisé un test d'adéquation, ou d'ajustement, entre une loi théorique donnée et une loi empirique associée à un échantillon d'observations. La fixation du risque α déterminera alors la valeur du seuil d'acceptation, ou seuil critique C .

Nous allons maintenant présenter et faire la critique de deux tests couramment utilisés en la matière. Ils sont associés à deux distances entre fonctions de répartition, permettant de déterminer la loi approchée de la variable aléatoire $d(F_n ; F)$ pour toute fonction de répartition F , le premier étant plutôt destiné aux lois discrètes (Test du Khi deux) et l'autre aux lois continues (Test de Kolmogorov-Smirnov).

Dans un souci de clarté, nous avons reporté en annexe 2 les énoncés détaillés de

ces trois tests, et nous nous concentrerons sur leur application et les limites de leur utilisation pour valider les ajustements obtenus.

3.1.2 Test du Khi-deux

3.1.2.1 Application du test du Khi-deux

La distance utilisée ici est la distance du Khi-deux, comme définie en annexe 2, paragraphe 1. En reprenant les notations utilisées précédemment, il nous faut calculer la quantité :

$$\chi_{cal}^2 = \sum_{x=1}^n \frac{(N_x \cdot \hat{Q}_x - N_x \cdot \tilde{q}_x)^2}{N_x \cdot \tilde{q}_x}$$

Ayant défini un degré de confiance $\alpha \in]0, 1[$ on a une approximation du seuil critique C avec la loi du χ_{n-1-r}^2 , pour $r = 3$ dans le cas du lissage par Makeham (3 paramètres), et $r = 0$ dans le cas du lissage par Whittaker-Henderson.

La règle de test s'énoncera par "on rejette l'hypothèse que χ_{cal}^2 suit une loi du χ_{n-1-r}^2 si χ_{cal}^2 est supérieur à C ".

En d'autres termes, on acceptera l'ajustement pour $\chi_{cal}^2 < C$.

3.1.2.2 Limites du test du Khi-deux

La première limitation est intrinsèque au calcul effectué. En effet, si on se réfère à la définition de la distance du χ^2 , nous n'avons pas calculé la distance $d_{\chi^2}(F_n, F)$, mais plutôt la distance entre deux 'simili' fonctions de répartition G_n et G telles que :

$$\begin{aligned} G_{N_x}(x+1) - G_{N_x}(x) &= \hat{Q}_x \\ G(x+1) - G(x) &= \tilde{q}_x \end{aligned}$$

Quelle est l'erreur commise ? D'abord le test n'est pas faux en lui-même vu qu'il mesure la distance euclidienne entre l'effectif lissé et l'effectif estimé des décès, et que la variable aléatoire $\frac{N_x \cdot \hat{Q}_x - N_x \cdot \tilde{q}_x}{\sqrt{N_x \cdot \tilde{q}_x}}$ est une variable aléatoire centrée convergent, si l'ajustement est correct, vers la loi $\mathcal{N}(0, \sqrt{1 - q_x})$.

Mais pour ce faire, on a commis une approximation des effectifs soumis au risque de décès pour chaque plage d'âges $(x, x+1]$. En effet, si on écrit la définition *stricto sensu*, on a :

$$d_{\chi^2}(F_n, F) = \sum_{x=1}^n \frac{(l_0 \cdot \hat{Q}_x \cdot \prod_{y=0}^{x-1} (1 - \hat{Q}_y) - l_0 \cdot \tilde{q}_x \cdot \prod_{y=0}^{x-1} (1 - \tilde{q}_y))^2}{l_0 \cdot \tilde{q}_x \cdot \prod_{y=0}^{x-1} (1 - \tilde{q}_y)}$$

Or l_0 , la population totale observée, est strictement inaccessible à la mesure dans le cadre d'observations transversales (on ne suit pas une population du début jusqu'à son 'extinction'). Nous sommes donc forcés d'en rester à l'estimation dans le

premier terme de la différence euclidienne :

$$l_0 \cdot \prod_{y=0}^{x-1} (1 - \hat{Q}_y) = N_x$$

ce qui n'est pas trop gênant vu que cela revient simplement à se positionner avec un effectif initial différent. Mais il faut alors se placer sur les mêmes bases pour approximer le terme de droite dans la différence...

Si on se réfère, à chaque classe d'âge, à un effectif initial dépendant de N_x (effectif observé d'âge x) et s'écrivant :

$$l_0 = \frac{N_x}{\prod_{y=0}^{x-1} (1 - \hat{Q}_y)}$$

on doit donc estimer dans le terme de droite :

$$l_0 \cdot \prod_{y=0}^{x-1} (1 - \tilde{q}_y) = N_x \cdot \frac{\prod_{y=0}^{x-1} (1 - \tilde{q}_y)}{\prod_{y=0}^{x-1} (1 - \hat{Q}_y)}$$

Et finalement, la distance du χ^2 s'exprime de manière correcte par :

$$d_{\chi^2}(F_n, F) = \sum_{x=1}^n \frac{(N_x \cdot \hat{Q}_x - N_x \cdot \tilde{q}_x \cdot \prod_{y=0}^{x-1} \frac{1-\tilde{q}_y}{1-\hat{Q}_y})^2}{N_x \cdot \tilde{q}_x \cdot \prod_{y=0}^{x-1} \frac{1-\tilde{q}_y}{1-\hat{Q}_y}}$$

En conclusion, l'approximation faite couramment dans le calcul du χ_{cal}^2 et qui consiste à considérer l'effectif des survivants selon la loi de lissage $N_x \cdot \prod_{y=0}^{x-1} \frac{1-\tilde{q}_y}{1-\hat{Q}_y}$ par l'effectif observé N_x conduit à une perte des informations du passé du lissage et à un critère de test plus faible qu'avec la distance décrite ci-dessus.

Pourtant, nous verrons que dans notre étude cette différence se révèle infime et soulignerons la pertinence des critiques faites à son égard, notamment quand ce test échoue à détecter :

- l'existence d'un certain nombre de grands écarts contrebalancés par un très grand nombre de petits écarts ;
- un grand écart cumulé sur une partie ou sur la totalité de l'intervalle ;
- un excès d'écarts positifs (ou négatifs) sur une partie ou sur la totalité de l'intervalle ;
- un grand excès d'écarts du même signe.

3.1.2.3 Intérêt du test du Khi-deux

Ce test aura pour nous deux intérêts principaux :

- sa facilité de mise en oeuvre, et pour Makeham, et pour Whittaker-Henderson ;
- son utilisation pour comparer les tables lissées obtenues entre elles et avec la table certifiée PF 60/64

Nous allons maintenant voir de façon succincte deux tests plus "pointus".

3.1.3 Tests d'adéquation de Kolmogorov-Smirnov

3.1.3.1 Application du test de Kolmogorov-Smirnov

Considérant les deux fonctions de répartition F_n et F décrites précédemment :

$$F_n(x) = P[X \leq x] = 1 - \frac{\hat{L}_x}{l_0} = 1 - \prod_{y=0}^{x-1} (1 - \hat{Q}_y)$$

$$F(x) = P[X \leq x] = 1 - \frac{\tilde{l}_x}{l_0} = 1 - \prod_{y=0}^{x-1} (1 - \tilde{q}_y)$$

nous devons calculer la distance 'sup' comme définie en annexe 2, paragraphe 2, par :

$$d_{sup}(F_n, F) = \max \left(\sup_{x \in \mathbb{R}} [F_n(x) - F(x)], \sup_{x \in \mathbb{R}} [F(x) - F_n(x)] \right)$$

Et on définit pour ce faire les statistiques (n = nombres d'âges concernés par l'ajustement) :

$$d^+(F_n ; F) = \sup_{x \in \mathbb{R}} [F_n(x) - F(x)] = \max_{1 \leq k \leq n} (F_n(k) - F(k))$$

$$d^+(F ; F_n) = \sup_{x \in \mathbb{R}} [F(x) - F_n(x)] = \max_{1 \leq k \leq n} (F(k) - F_n(k))$$

Ayant défini un degré de confiance $\alpha \in]0, 1[$ on a une approximation du seuil critique C par la table de la fonction de répartition asymptotique (ou en utilisant l'approximation de Feller, pour n assez grand, on peut avoir une approximation de C par : $C = \sqrt{-\frac{1}{2} \times \ln \alpha}$).

La règle de test s'énoncera par "on rejette l'hypothèse que $\sqrt{n} \cdot d_{sup}(F_n, F)$ suit une loi limite de type $K = 1 - 2 \times \sum_{k=-\infty}^{\infty} (-1)^{k+1} e^{-2k^2 x^2}$ si $\sqrt{n} \cdot d_{sup}(F_n, F)$ est supérieur à C ".

En d'autres termes, on acceptera l'ajustement pour $\sqrt{n} \cdot d_{sup}(F_n, F) < C$.

3.1.3.2 Intérêt et limites du test de Kolmogorov-Smirnov

Le principal intérêt de ce test, de mise en oeuvre plus délicate que le test du χ^2 est de venir confirmer de manière plus fine l'adéquation des lissages obtenus. En revanche, nous ne l'utiliseront pas pour comparer les lissages entre eux et avec la table PF 60/64, et nous préférons un indicateur plus pragmatique, l'espérance de vie, pour juger de la bonne fin des tables obtenues.

3.2 Méthodes de comparaison des ajustements obtenus

Nous rentrons ici dans une partie plus descriptive, où nous cherchons à définir des critères pertinents de comparaison entre les ajustements obtenus. Outre la liste non exhaustive des tests statistiques cités dans la section précédente, utilisables pour juger de l'adéquation des différents lissages entre eux, il faut se pencher sur l'utilisation future de la table de mortalité qui, rappelons le, sera choisie pour l'estimation des différents engagements pris par l'institution vis à vis de sa population affiliée.

Nous avons retenu, dans le cadre de notre étude, trois critères :

- l'adéquation statistique au moyen du test du χ^2
- la pertinence des espérances de vie obtenues par les lissages
- la comparaison des deux lissages à la table certifiée PF 60/64

Nous allons maintenant détailler ces méthodes et expliquer leur choix.

3.2.1 critère statistique

La méthode de comparaison statistique passera par deux étapes :

- le classement des deux lissages selon leur adéquation statistique aux estimations
- l'adéquation des deux lissages au moyen de la distance $d_{\chi^2}(F_{Mak}, F_{Whit})$ ¹

Il est *a priori* intuitif qu'une méthode de lissage statistique, telle que la moyenne de Whittaker-Henderson, obtiendra des meilleurs résultats en termes statistiques qu'une méthode appliquant une loi analytique de forme prédéfinie, telle que le lissage par la loi de Makeham, avec laquelle on ne dispose que de trois paramètres pour "coller" au mieux aux données. Nous verrons en application que cette première étape qui apparaît déjà comme une formalité ne l'est en fait absolument pas .

¹cf paragraphe 3.2.1, test classique

Afin de vérifier une certaine cohérence des lissages par Makeham et par Whittaker-Henderson, objet de la deuxième étape, nous avons gardé comme critère statistique le test du χ^2 . Ce choix est simplement dû à la simplicité de mise en oeuvre de ce test, étant donné que le résultat d'un tel type de test n'aura d'autre interprétation possible qu'une adéquation entre les deux lissages, sans référencement à ce que devrait être la loi inconnue de mortalité exacte.

Nous chercherons donc un critère de choix plus adéquat en mesurant quel peut être l'impact des tables obtenues sur l'évaluation des engagements pris par l'institution.

3.2.2 critère de l'espérance de vie

Si l'on s'intéresse à la mortalité d'une population affiliée à un contrat d'assurance décès, il ne faut pas sous-estimer la mortalité. En effet le risque pour l'assureur est de constater un trop grand nombre de sinistres, et ici de décès, par rapport à ce qui était prévu. Inversement, s'il s'agit d'un contrat d'assurance de rente, le risque pour l'assureur est de constater un trop grand nombre de sinistres, cette fois-ci de 'survie', par rapport à ce qui était prévu.

Considérant cette logique de sinistralité, nous nous sommes appuyés sur l'indicateur de l'espérance de vie pour définir ce que devrait être une table prudente, et pour comparer les deux lissages sur ces bases. Concrètement, en définissant l'espérance abrégée de vie à l'âge x comme :

$$e_x = \sum_{j=1}^{\omega} \frac{l_{x+j}}{l_x}$$

où ω est le dernier âge observé, on utilise la règle prudentielle suivante :

- si le sinistre encouru est le décès, on cherchera à ordonner $\{\tilde{e}_x^{mak}, \tilde{e}_x^{whit}\}$ de manière croissante, le lissage le plus prudent étant celui qui donne une espérance de vie plus grande ;
- *a contrario*, si le sinistre encouru est la survie, on cherchera à ordonner $\{\tilde{e}_x^{mak}, \tilde{e}_x^{whit}\}$ de manière décroissante, le lissage le plus prudent étant celui qui donne une espérance de vie plus faible ;

Nous signalons que cette règle n'est valable qu'à condition d'avoir vérifié préalablement l'adéquation des lissages aux estimations. En effet, il ne s'agit pas de majorer à l'excès les risques de sinistres, ce qui engendrerait des surtarifications notoires.

De plus, dans le cas où des prestations couvrant les deux natures de risque sont fournies par le régime ou le contrat d'assurance, la meilleure solution reste encore d'élaborer deux tables de mortalité distinctes pour chaque risque.

3.2.3 Comparaison à la table PF60/64

La comparaison à une table de référence aura pour nous un double intérêt. D'abord dans le cadre de l'étude effectuée, une question était de savoir si la table PF60/64 (makehamisée), table certifiée, pouvait être utilisée comme table de mortalité pour le régime, et servir de référence pour élaborer à terme une table d'expérience prospective². D'autre part pour la critique des deux lissages obtenus, une telle table pourra nous servir de valeur de référence.

Nous avons donc rajouté dans la règle de comparaison des espérances de vie³, les valeurs 'références' $\tilde{e}_x^{PF60/64}$.

Ce qu'il faut retenir ici, c'est l'importance du jugement humain dans la sélection de la table de référence, qui doit tout de même épouser l'idée que l'on se fait *a priori* sur l'allure de la courbe.

²escomptant l'augmentation de l'espérance de vie

³cf paragraphe précédent

Application à la population d'un régime de retraite marocain

Chapitre 4

Présentation de l'étude

Nous abordons ici la mise en pratique des méthodes de construction vues précédemment, sur la base d'une mission que j'ai contribué à réaliser courant août 2003, au sein du cabinet JWA-International, pour le compte d'un régime de retraite marocain. Nous allons présenter succinctement les paramètres du régime et les enjeux de l'étude, puis s'en suivra la partie relative au traitement des données.

4.1 Présentation du régime

Le régime étudié est un régime à prestations définies. Nous ne rentrerons pas dans les caractéristiques d'affiliation vu le caractère non confidentiel de ce mémoire, et nous ne donnerons que le mécanisme global d'évaluation des engagements du régime.

4.1.1 Ressources du régime

Ce régime est financé de façon classique par des cotisations salariales et des contributions patronales, composées les unes et les autres d'une part fixe et d'une part variable.

4.1.2 Prestations du régime

Les prestations servies aux affiliés sont de deux natures :

- pension de retraite
- pension d'ayant cause ou pension de réversion

Pension de retraite

Sous condition d'un minimum d'ancienneté requis, un affilié a droit à un montant de retraite déterminé comme dans tout régime à prestations définies en fonction du :

- SAMCR¹
- nombre d’années d’activité

La pension se calcule comme un pourcentage fixé du produit de ces deux termes :

$$Pension = \% \times SAMCR \times Ancienneté$$

Pension de réversion

La pension de réversion à laquelle a droit l’éventuel conjoint survivant se calcule comme un pourcentage de la pension de retraite du conjoint décédé :

$$Pension\ de\ reversion = tx_{réversion} \times Pension$$

Revalorisation des pensions

Le régime prévoit une revalorisation indexée sur l’évolution du salaire moyen annuel global :

$$Pension_n = Pension_{n-1} \times i_n$$

où $i_n = \frac{SM_n}{SM_{n-1}}$ est le taux de revalorisation de l’année n avec SM_n est le salaire moyen du régime de l’année n .

Transfert

Lorsqu’un affilié adhère à un autre régime de retraite, la valeur capitalisée de ses cotisations salariales et contributions patronales fixes est transférée au nouveau régime. Notre régime est alors libéré de tout engagement à l’égard de l’affilié. Ceci signifie qu’il faudra prendre en compte un turnover dans l’évaluation des engagements du régime.

4.1.3 Evaluation des engagements

L’évaluation se fait séparément selon que l’on considère l’effectif des actifs, des retraités, ou des ayants-droits (réversion). Nous présentons ici le calcul actuariel d’évaluation selon ces différents cas.

L’important ici est de voir les différents paramètres rentrant en jeu dans l’évaluation, en portant une attention toute particulière à la mortalité.

Cas d’un actif

Si l’on considère le cas d’un actif d’âge x affilié au régime, on représente les engagements à son égard comme une rente viagère indexée et démarrant à la date de retraite de l’individu (cf figure).

¹Salaire Annuel Moyen de Carrière Revalorisé

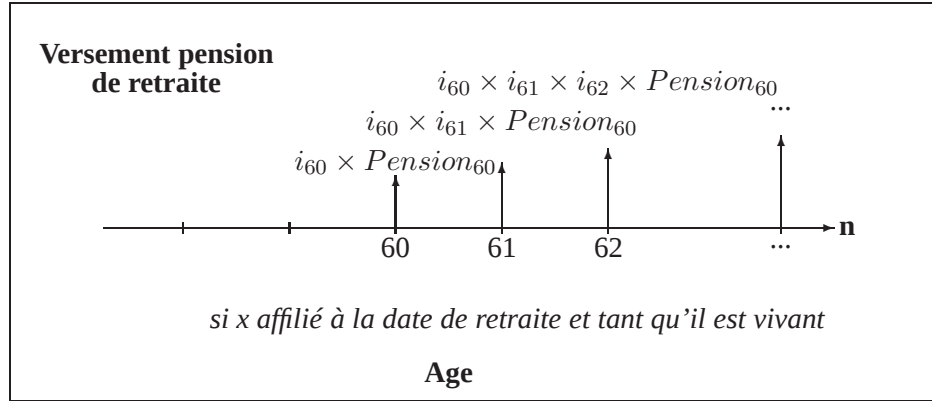


FIG. 4.1 – Engagement, cas d'un actif

On calcule dans un premier temps la valeur de l'engagement à la date de retraite, ici à 60 ans :

$$Rente_{60} = \sum_{k=0}^{\omega-60} v^{k+1} \times \frac{l_{60+k+1}}{l_{60+k}} \times pension_{60+k}$$

avec :

- $pension_{60+k} = pension_{60} \times \prod_{j=0}^k i_{60+j}$
- $pension_{60} = \% \times SAMCR \times anciennete$ (comme vu précédemment)
- $v = (1 + t)^{-1}$ où t est le taux technique ou taux d'actualisation
- ω limite de l'âge de vie biologiquement envisageable

Puis on actualise cette somme à l'âge x correspondant à la date d'évaluation, en tenant compte de la mortalité et du turnover :

$$VAP^2 = \underbrace{v^{60-x}}_{\text{actualisation}} \times \underbrace{\frac{l_{60}}{l_x} \times \prod_{j=x}^{60} (1 - r_j)^{-1}}_{\substack{\text{mortalité} \\ \text{turnover}}} \times Rente_{60}$$

probabilité de présence de l'individu à la date de retraite

où r_x est le taux de départ à l'âge x , caractéristique de la population active affiliée, et tabulé de façon statistique au même titre que les taux de survie l_x .

Cas d'un retraité

Pour le calcul de l'engagement vis à vis d'un retraité d'âge x , on reprend simplement le début du calcul précédent, en écrivant que ($x \geq 60$) :

$$VAP = Rente_x = \sum_{k=0}^{\omega-x} v^k \times \frac{l_{x+k+1}}{l_{x+k}} \times pension_{x+k}$$

²Valeur Actuelle Probable de l'engagement de l'assureur à la date d'évaluation

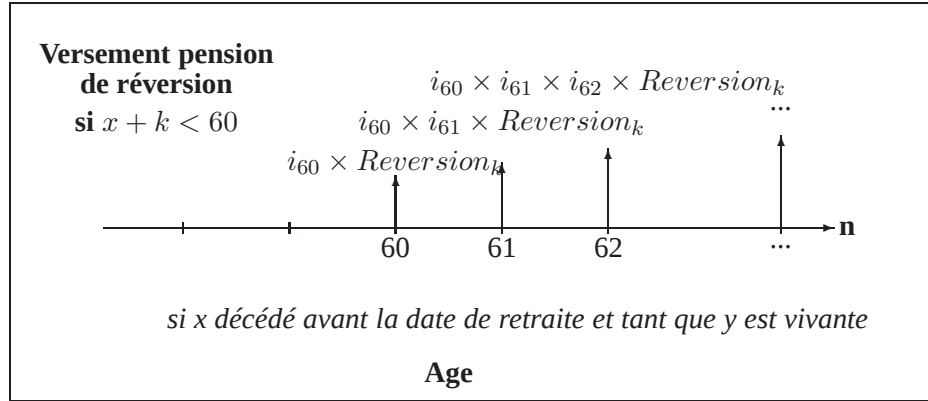


FIG. 4.2 – Engagement, cas d'une pension de réversion

Cas d'une pension de réversion

Si l'on considère un affilié A et son conjoint B d'âge respectivement x et y , le montant des engagements envers le conjoint si A décède dépend de l'âge $x + k$ du décès de A. Deux cas peuvent se présenter.

Si $x + k < 60$

Ce cas se présente comme une rente de survie viagère indexée et différée (à l'âge virtuel $x = 60$). On calcule donc l'engagement à l'âge de retraite virtuel de A, dépendant de l'indice k :

$${}_k Rente - rev_{60} = \sum_{s=0}^{\omega-60} v^{s+1} \times \frac{l_{(y+x-60)+s+1}^B}{l_{(y+x-60)+s}^B} \times \prod_{j=0}^s i_{60+j} \times Reversion_k$$

où l'on retrouve les mêmes termes que précédemment sauf que la mortalité est appliquée au conjoint B et que la pension de réversion dépend de k : $Reversion_k = tx_{reversion} \times {}_k Pension_{60}$ avec ${}_k Pension_{60}$ dépendant de l'ancienneté et du SAMCR valable au décès virtuel de A à l'âge $x + k$.

On doit alors sommer l'ensemble des probabilités de décès de A sur la plage d'âges $[x, 60]$, en tenant compte de la mortalité de B et du turnover de A :

$$VAP = \sum_{k=0}^{60-x} v^{k+1} \times \underbrace{\frac{l_{x+k}^A - l_{x+k+1}^A}{l_{x+k}^A} \times \prod_{j=1}^{k+1} (1 - t_j)^{-1}}_{\text{probabilité de décès de A affilié à l'âge } x+k} \times \frac{l_{y+k+1}^B}{l_{y+k}^B} \times {}_k Rente - rev_{60}$$

Si $x + k \geq 60$

Dans cette configuration virtuelle, on considère que le décès de A arrivera après la date de retraite (toutes les autres possibilités de décès étant prises en compte dans le cas précédent), et donc on a affaire à une rente de survie viagère indexée immédiate à l'âge de décès virtuel $x + k$.

$$Rente - rev_{60} = \sum_{s=0}^{\omega-60} v^{s+1} \times \frac{l_{60+s}^A - l_{60+s+1}^A}{l_{60+s}^A} \times \frac{l_{(y+x-60)+s+1}^B}{l_{(y+x-60)+s}^B} \times \ddot{a}_{(y+x-60)+s+1}$$

avec :

$$\ddot{a}_{(y+x-60)+s+1} = \sum_{j=0}^{\omega-(y+x-60)} v^{j+1} \times \frac{l_{(y+x-60)+s+j+1}^B}{l_{(y+x-60)+s+j}^B} \times \prod_{f=0}^j i_{60+s+f} \times Reversion_{60}$$

engagement, à la date de décès, pour la rente versée au conjoint B, et $Reversion_{60} = \%Pension_{60}$ telle que définie précédemment.

On actualise alors en tenant compte de la survie de A, du turnover, et de la survie de B :

$$VAP = \underbrace{v^{60-x}}_{\text{actualisation}} \times \underbrace{\frac{l_{60}^A}{l_x^A}}_{\substack{\text{survie de A} \\ \text{probabilité de présence de A à la date de retraite}}} \times \underbrace{\prod_{j=x}^{60} (1 - r_j)^{-1}}_{\text{turnover}} \times \underbrace{\frac{l_{y+x-60}^B}{l_y^B}}_{\text{survie de B}} \times Rente - rev_{60}$$

L'engagement à la date d'évaluation à l'égard du conjoint est la somme des deux VAP calculées ci-dessus.

Ces formules d'engagements, que nous avons développées pour expliciter les phénomènes temporels intervenant dans la méthode prospective d'estimation, sont simplifiées et automatisées de manière classique pour obtenir l'engagement global de la caisse.

4.1.4 Premières conclusions

On peut d'abord énumérer l'ensemble des paramètres intervenant dans le calcul des engagements du régime vis à vis de ses affiliés :

- date d'évaluation (âge de l'affilié)
- taux technique
- profil de carrière (évaluation du SAMCR)

- turnover (évaluation de l'ancienneté)
- dérive du salaire annuel moyen du régime (évaluation de la revalorisation)
- table de mortalité
- différence d'âge entre l'affilié et son conjoint

En ce qui concerne plus particulièrement la mortalité, les prestations du régime sont analogues à des rentes de survie, pour lesquelles on utilise normalement des tables prospectives³. Compte tenu des premiers résultats⁴, nous nous sommes restreints à une table d'expérience simple, la plus caractéristique possible du portefeuille étudié, cette table devant faire l'objet d'un suivi d'adéquation par la suite.

Dans le même cadre, le lecteur aura remarqué la contradiction venant de l'utilisation d'une seule table de mortalité pour l'affilié et le conjoint bénéficiaire. En effet, en termes de prudence, la mortalité de l'affilié devrait être plutôt sousestimée que surestimée dans le cas d'une rente directe, et l'inverse dans le cas d'une rente de réversion. Mais alors c'est faire preuve d'un excès de prudence, étant donné que les deux engagements évoluent conjointement, selon un même risque de décès de l'affilié.

Finalement il faut souligner que le risque d'autosélection⁵ n'existe pas dans un régime de retraite, vu qu'on ne peut pas définir d'utilité pour l'affilié si ce n'est celle de pouvoir bénéficier d'une retraite durement méritée. Ceci nous conforte dans l'idée qu'il n'y a pas lieu de considérer des tables séparées de vie et de décès, et que l'essentiel est d'être le plus fidèle possible à la mortalité caractéristique de la population affiliée en veillant à ne pas la surestimer.

4.2 Présentation de l'étude

4.2.1 Enjeux

L'élaboration d'une table d'expérience spécifique au régime se situe dans un cadre plus vaste d'études actuarielles, avec pour objectif d'avoir accès à tous les paramètres pour piloter le régime, paramètres qui seront autant de clés pour éviter à terme que ce régime ne périclite.

Un décalage important entre la mortalité constatée et la mortalité escomptée avait été relevé lors de surveillances menées antérieurement. Notre travail consiste en

³tenant compte de l'évolution de l'espérance de vie, cf §1.1

⁴cf traitement des données, §4.3

⁵phénomène d'autosélection en assurance-vie :

- les garanties en cas de décès attirent *a priori* les personnes en mauvaise santé
- les garanties en cas de vie attirent *a priori* les individus ayant l'espoir de vivre suffisamment de temps

la reconstitution d'une table de mortalité caractéristique de la population affiliée, et pour ce faire nous allons appliquer pas à pas les méthodes détaillées dans les chapitres précédents.

4.2.2 Présentation des données

Le base de données initiale est une base historique (unique) des pensionnés du régime sur une période d'observation de dix ans, des années 1990 à 2000. Elle se présente sous format texte, et nous en avons effectué la conversion et le traitement au moyen du logiciel ACCESS.

Cette base présente les huit informations suivantes :

- un code d'enregistrement (= 0 si vivant, = 1 si décédé au 31/12/2000)
- un n° d'identification
- la date de naissance (format JJMMAAAA)
- la date éventuelle de décès (format JJMMAAA)
- la nature de la pension ('D' pour Directe, 'I' pour Indirecte)
- le sexe ('M' ou 'F')
- le trimestre d'effet de la pension
- l'année d'effet de la pension (format AAAA)

Notons que la qualité de cette base (nombre de lignes > 70.000) est assez remarquable, ce que confirme le peu d'anomalies décelées lors du traitement de données que nous allons décrire maintenant.

4.2.3 Traitement des données

La première étape du traitement de données consiste en l'extraction et la conversion de celles-ci, puis de la détection des anomalies. Au moyen de requêtes SQRL simples, nous avons vérifié la complétude des différents champs, et détecté deux anomalies :

- > 3.000 individus indiqués décédés mais ayant une date de décès manquante. Il s'est avéré que ces personnes étaient décédées avant 1986 et ne rentraient donc pas dans le périmètre de l'étude.
- < 20 individus avec une date de naissance manquante. Au vu du faible nombre de personnes concernées, ces personnes ont été écartées de l'étude.

La deuxième étape a consisté en un découpage de ladite base (de 1990 à 1999), effectué en vue de disposer de dix périodes d'observation et de pouvoir dégager des statistiques annuelles ainsi qu'une tendance éventuelle d'augmentation de l'espérance de vie.

Nous avons effectué ce découpage pour chaque année 199(X) en utilisant la règle

suivante :

$$\left\{ \begin{array}{ll} \text{date de décès} & > 31/12/199(X - 1) \text{ ou 'vide'} \\ \text{et} & \\ \text{année d'effet de la pension} & < 199(X + 1) \end{array} \right.$$

ce qui revient à isoler les effectifs des personnes pensionnées de l'année 199(X).

Sans rentrer dans le détail des statistiques annuelles par souci de confidentialité, il faut savoir que globalement :

- les effectifs masculins se concentrent entre 50 et 80 ans
- les effectifs féminins se concentrent entre 40 et 80 ans
- les décès masculins se concentrent entre 60 et 85 ans
- les décès féminins se concentrent entre 60 et 80 ans

La prochaine étape est l'estimation de la population soumise au risque, sans distinction de sexe, réalisée pour chaque période d'observation de un an. La contribution de chaque individu de la population à la période d'observation sera évaluée en nombre de mois. Le total des contributions à un âge x ramené en années déterminera le nombre d'individus d'âge x soumis au risque décès, soit N_x . On distingue alors quatre cas :

- si l'adhérent est présent et a atteint l'âge x avant la date de début d'observation, sa contribution à N_x est le nombre de mois séparant la date de début d'observation et son prochain anniversaire.
- si l'adhérent a atteint l'âge x durant la période d'observation et est présent à la date de fin d'observation, sa contribution à N_x est le nombre de mois entre son anniversaire et la date de fin de période d'observation.
- si l'adhérent a intégré le régime à l'âge x pendant la période d'observation, sa contribution à N_x est le nombre de mois séparant sa date d'affiliation au régime (trimestre) et son prochain anniversaire.
- si l'adhérent décède à l'âge x pendant la période d'observation, sa contribution à N_x est le nombre de mois entre la date de début d'observation et sa date de décès.

On peut désormais procéder à l'estimation des taux bruts de mortalité sur chaque plage d'observation.

Remarque importante pour la suite :

A partir des résultats obtenus par lissage de Makeham sur les taux bruts de chaque période d'observation, nous avons essayé de déceler une tendance d'évolution de l'espérance de vie sur la période 1990-1999⁶.

⁶cf § 8.4

Compte tenu des résultats difficilement exploitables, nous avons procédé à la sommation des effectifs N_x des dix périodes (idem pour les effectifs de décédés D_x), ce qui nous a permis d'obtenir des taux bruts de mortalité et d'élaborer, dans un premier temps, une table unique sur la période 1990-1999.

Nous ne présenterons dans les deux prochains chapitres que les résultats obtenus à partir de ces effectifs globaux.

Chapitre 5

Estimation des taux bruts de mortalité

Vu le traitement de données effectué, sur des périodes d'observation de un an, nous avons fait l'hypothèse la plus simple, celle de distribution uniforme des décès, ce qui nous amène à l'estimateur pour la table globale 1990-1999 :

$$\hat{Q}_x = \frac{D_x}{N_x} = \frac{\sum_{i=1990}^{1999} D_x^i}{\sum_{i=1990}^{1999} N_x^i}$$

Nous présentons dans ce chapitre les estimations obtenues sur la plage d'âges allant de 30 ans à 105 ans, puis les intervalles de confiance des taux estimés selon la variance estimée et la variance exacte.

5.1 Taux bruts de mortalité

Le graphique figure 5.1 présente les taux bruts de mortalité par âge $\hat{Q}_x$ obtenus sur la population étudiée. Alors que sur la plage d'âges [30, 85] on observe une tendance exponentielle très nette, se prêtant *a priori* à un ajustement par loi analytique, on remarque des fluctuations importantes pour les âges plus élevés, dues essentiellement à la relative faiblesse des effectifs observés à ces âges.

Il ne faut pas donner à ces fluctuations un sens qu'elles n'ont pas. Notamment, nous ne pourrions pas justifier une chute de la mortalité à un âge avancé, contraire à l'évidence au phénomène de mortalité.

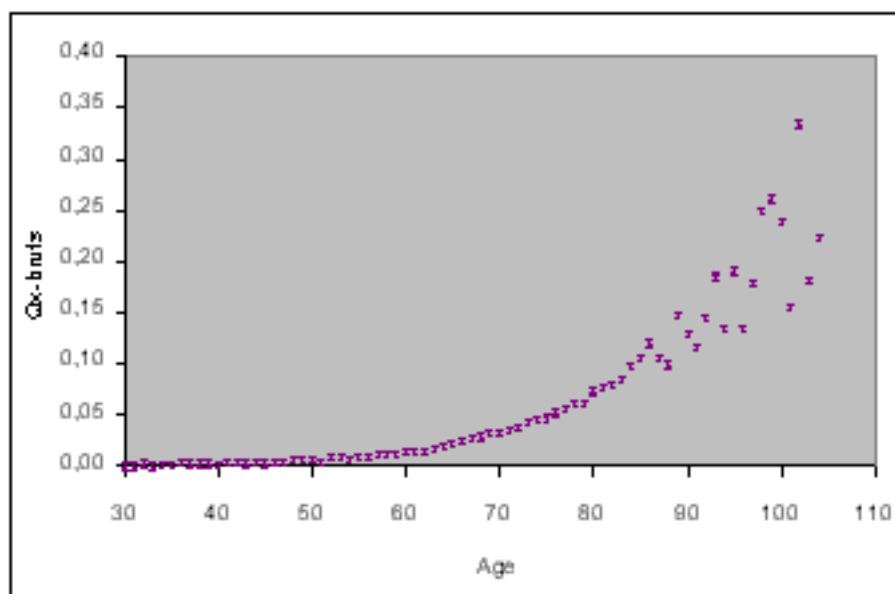


FIG. 5.1 – Taux bruts de mortalité par âge

5.2 Intervalles de confiance - avec la variance estimée

Les intervalles de confiance à 95 % obtenus avec la variance estimée (cf chapitre 1) sont présentés figure 5.2. Vu que le taux de décès ne peut être négatif, nous avons pris le maximum entre la valeur calculée et 0.

Ces intervalles confirment très nettement l'impact de la faiblesse des effectifs pour les âges les plus élevés : passé 90 ans, les observations ne sont plus vraiment représentatives de la réalité tant la possibilité de s'en écarter est importante.

5.3 Intervalles de confiance - avec la variance exacte

Nous présentons de même les intervalles de confiance à 95 % obtenus avec la variance exacte, figure 5.3.

Nous pouvons observer deux différences minimales :

- des bornes inférieures toujours distinctes de 0, même pour les âges les plus élevés
- des intervalles de confiance moins larges

La conclusion n'en reste pas moins la même : il faudra tenir compte du caractère approximatif des observations passé un certain âge.

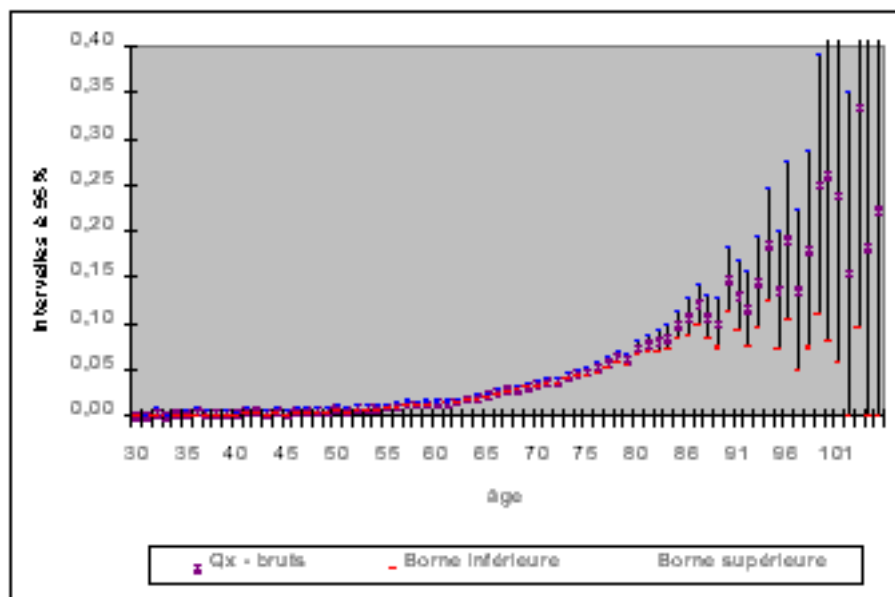


FIG. 5.2 – Intervalles de confiance, variance estimée

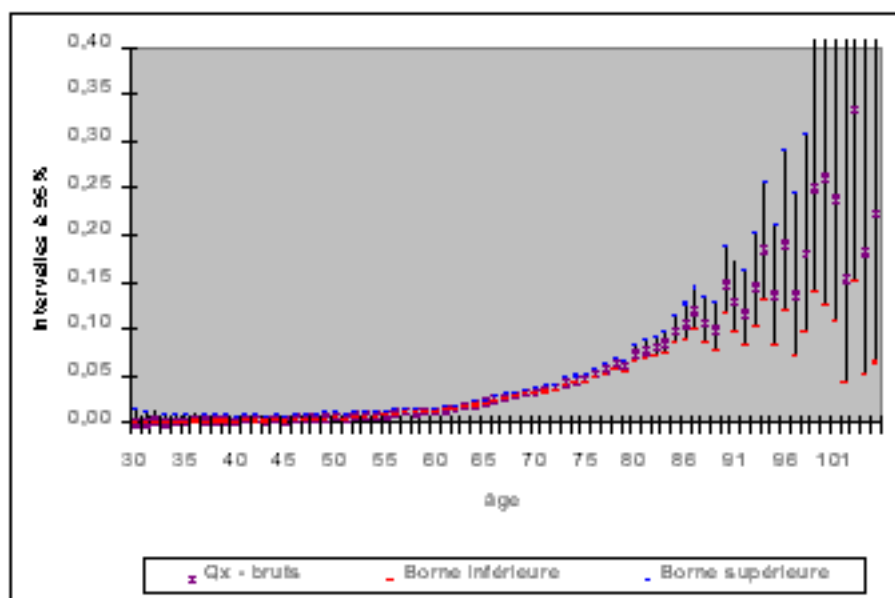


FIG. 5.3 – Intervalles de confiance, variance exacte

Chapitre 6

Ajustement des taux bruts de mortalité par la loi de Makeham

Après avoir défini un intervalle plausible d'ajustement makehamien, nous exposons dans ce chapitre les résultats de ce lissage et de son adéquation aux taux bruts, suivant la méthode exposée section 2.1.

6.1 Ajustement par Makeham

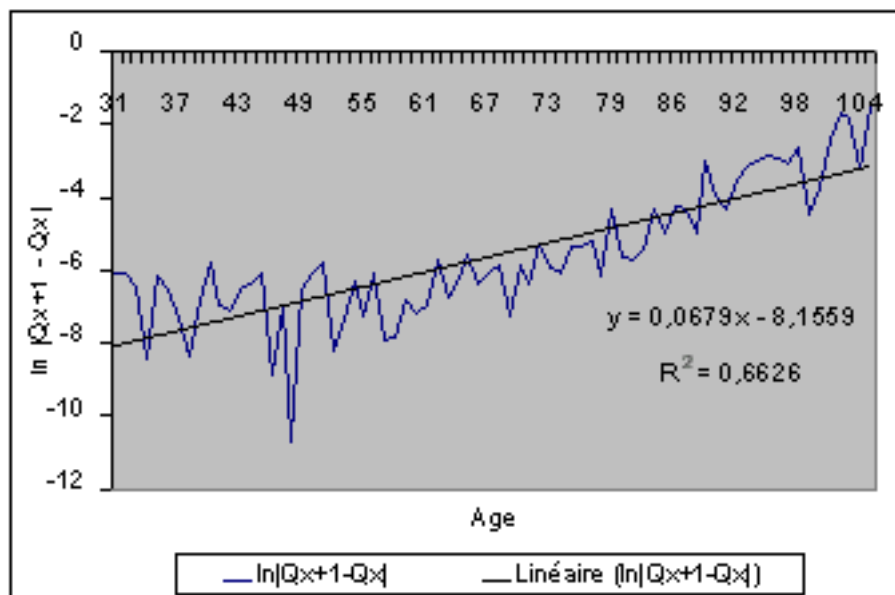


FIG. 6.1 – Intervalle d'ajustement

La fonction $\ln|\hat{Q}_{x+1} - \hat{Q}_x|$ ¹ décrite figure 6.1 présente une tendance linéaire assez forte ($\approx 0,07$), avec une adéquation relativement correcte ($R^2 \approx 0,7$). Nous avons effectué une régression linéaire sur l'intervalle d'âges [30, 90] qui donne sensiblement les mêmes résultats.

Compte tenu du fait que l'algorithme d'apprentissage des coefficients du modèle tient compte de la taille des effectifs à chaque âge, nous avons donc réalisé l'ajustement sur toute la plage d'âges observés.

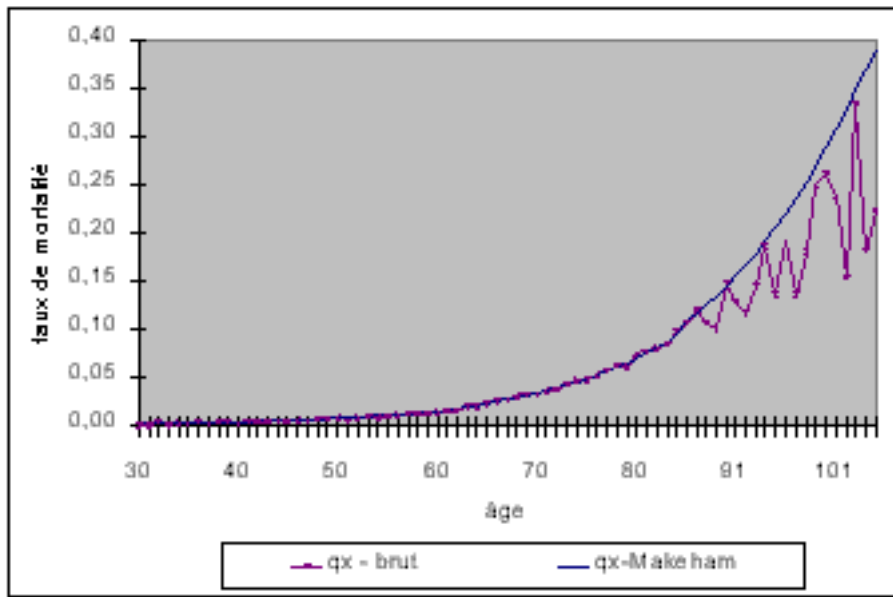


FIG. 6.2 – Lissage des taux bruts par Makeham

Nous avons alors utilisé sous EXCEL un programme d'analyse basé sur la méthode du maximum de vraisemblance, et abouti aux paramètres de la loi de Makeham :

$$\hat{\mu}_x = \hat{\alpha} + \hat{\beta} \cdot \hat{c}^x$$

avec :

$$\begin{cases} \hat{\alpha} \approx 0,000100 \\ \hat{\beta} \approx 0,000140 \\ \hat{c} \approx 1,080575 \end{cases}$$

Ces coefficients nous ont permis de déduire la loi supposée de mortalité :

$$\tilde{q}_x = 1 - \tilde{p}_x = 1 - e^{(\hat{\alpha} - \hat{\beta} e^{\hat{c}^x})}$$

¹cf § 2.1.1

avec :

$$\begin{cases} \hat{\alpha} \approx 0,000100 \\ \hat{b} \approx 0,000146 \\ \hat{\gamma} \approx 0,077493 \end{cases}$$

Les résultats du lissage sont présentés figure 6.2. Graphiquement, l'ajustement paraît très bon jusqu'à l'âge de 90 ans puis se détériore pour les âges plus élevés.

6.2 Intervalles de confiance des taux lissés

Nous avons ensuite calculé les intervalles de confiance de l'ajustement².

Rappelons ici que les valeurs estimées devraient, si l'ajustement est bon, se trouver entre les bornes de l'intervalle de confiance du lissage, et ce à chaque âge. Nous présentons les résultats de cette adéquation figure 6.3.

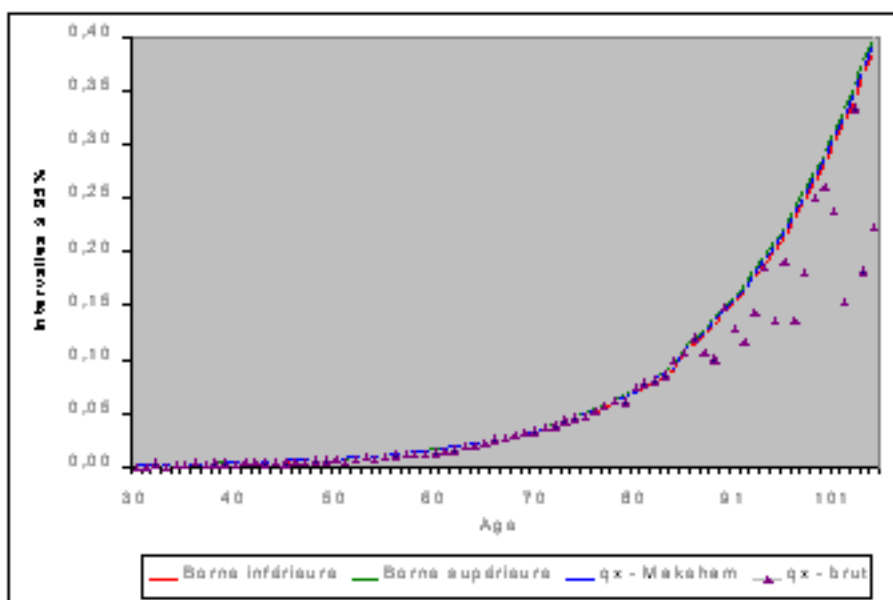


FIG. 6.3 – Intervalle de confiance du lissage par Makeham

Visiblement l'adéquation est excellente au départ, puis dérape à partir de 87 ans, vu que les taux bruts sortent de l'intervalle aux âges plus élevés. Nous allons pourtant voir que ce problème reste relativement mineur étant donné les approximations commises pour ces âges.

²comme étudié au § 2.1.4

6.3 Tests d'adéquation statistiques

Nous avons tout d'abord reporté les résidus du lissage figure 6.4, résidus qui sont évidemment loin d'être assimilable à une loi normale. La sur-estimation apparente de la mortalité aux âges supérieurs à 87 ans est assez dérangeante en soi, mais on doit faire preuve de prudence dans les conclusions à donner à un tel graphique.

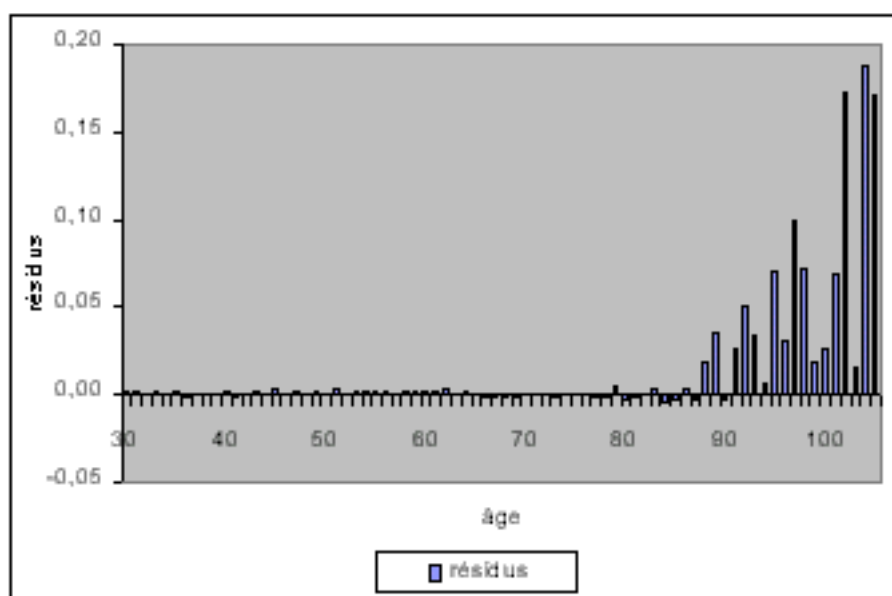


FIG. 6.4 – Résidus du lissage par Makeham - sur les $\hat{q}_x$

Pour vérifier la validité de l'ajustement, nous avons effectué les deux tests statistiques décrits au chapitre 3, dont voici les résultats :

Test du Khi2

Intervalle	Distance approximée	Distance exacte	Degré de liberté	Seuil critique
jusque 60 ans	26,9258	22,1080	27	40,1132
jusque 90 ans	64,8032	67,8527	57	75,6237
jusque 105 ans	81,2103	82,0469	71	91,6702

TAB. 6.1 – Lissage par Makeham - test du Khi2

Les distances calculées sont inférieures au seuil critique du test. Le test du Khi2 ne rejette donc pas l'ajustement.

Test de Kolmogorov-Smirnov

Intervalle	Distance	Seuil critique
jusque 60 ans	0,0595	0,2520
jusque 90 ans	0,4059	1,2239
jusque 105 ans	0,4530	1,2239

TAB. 6.2 – Lissage par Makeham - test de Kolmogorov-Smirnov

Les distances calculées sont toujours inférieures au seuil critique du test. Le test de Kolmogorov-Smirnov ne rejette donc pas l'ajustement. Graphiquement, l'ajustement de la fonction de répartition paraît très bon (cf figure 6.4), mais les résidus sont de même signe sur des plages d'âges étendues (cf figure 6.5) :

$$\begin{cases} x \in [30, 70) & \text{résidus} = F(x) - F_n(x) > 0 \\ x \in [70, 90) & \text{résidus} = F(x) - F_n(x) < 0 \\ x \in [90, 105) & \text{résidus} = F(x) - F_n(x) > 0 \end{cases}$$

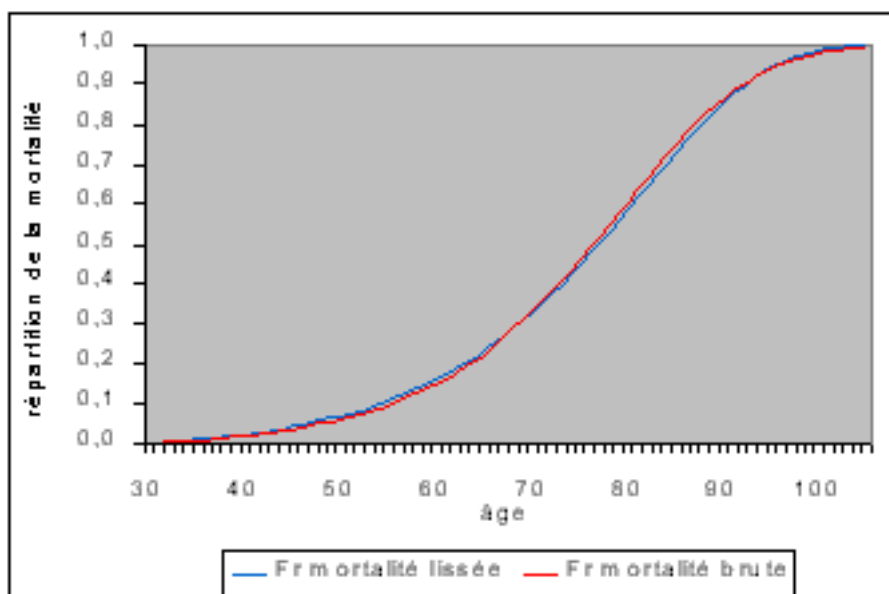


FIG. 6.5 – Fonctions de répartition brutes et lissées

On peut donc considérer si l'on est optimiste que l'ajustement est valable, et si l'on est pessimiste que les tests effectués échouent à détecter l'inadéquation de l'ajustement.

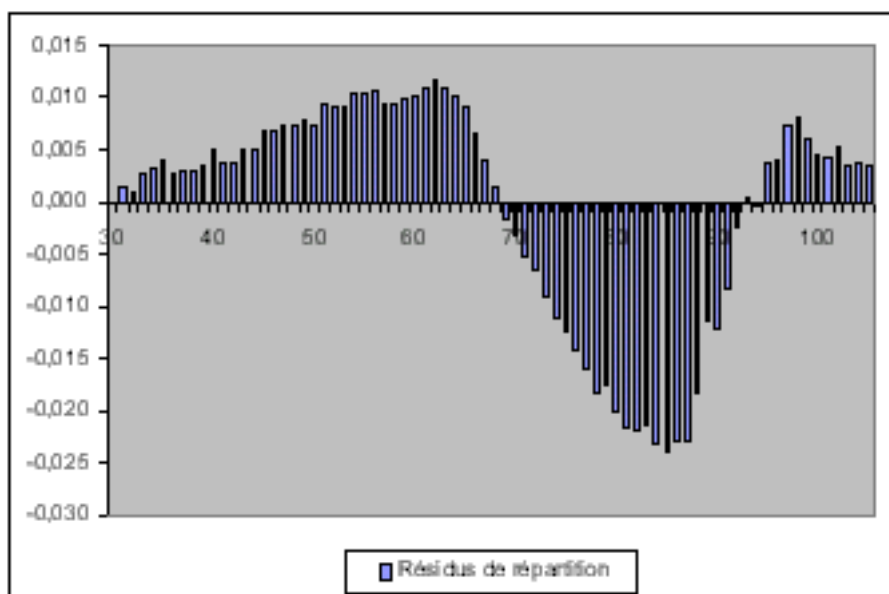


FIG. 6.6 – Résidus du lissage par Makeham - sur les fonctions de répartition

Nous avons effectué un test complémentaire sous MINITAB, où l'on étudie l'hypothèse : $F = F_n + \epsilon$ c'est à dire où l'on teste si le résidu ϵ est de loi normale, sur la base de la statistique d'Anderson-Darling :

$$A_n = n \int_{-\infty}^{+\infty} \frac{(F_n(x) - F(x))^2}{F(x)(1 - F(x))} dF(x)$$

Ce test plus puissant, ou plus contraignant, rejette très nettement l'hypothèse de normalité (le test d'Anderson-Darling attribue une probabilité quasi-nulle de normalité des résidus), ce à quoi on était en droit d'attendre.

Maintenant, il faut faire preuve d'une relative souplesse dans la prise en compte de ces tests, étant donné les remarques que nous avons pu faire précédemment pour ce qui est de l'estimation des taux bruts de mortalité. Rappelons-nous que l'on recherche une loi de mortalité dont les taux bruts historiques soient susceptibles d'être des observations.

Nous pouvons donc définir un test de cohérence, beaucoup plus intéressant à nos yeux, en vérifiant que la courbe lissée se trouve dans l'intervalle de confiance de l'estimation des taux bruts de mortalité, et ce pour tout âge.

Le résultat de ce test, présenté graphiquement figure 6.6, prouve que la courbe lissée se situe dans les intervalles de confiance pour tout âge, et donc que les observations historiques $\hat{Q}_x$ sont fortement susceptibles d'être des observations de loi

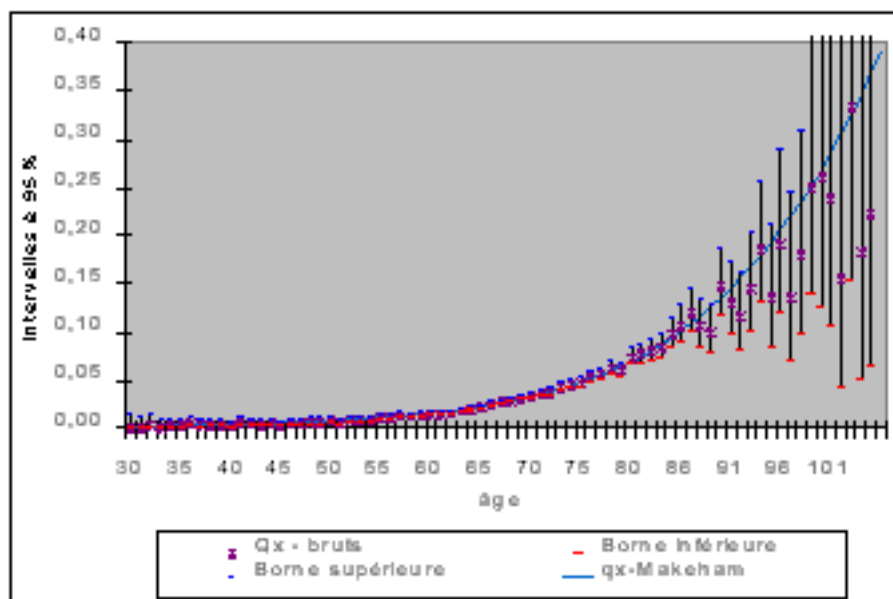


FIG. 6.7 – Test de cohérence du lissage par Makeham

de mortalité réelle $\tilde{q}_x$, y compris pour les âges extrêmes (plages d'âges $[30, 40)$ et $[87, 105)$). Un tel test permet de tenir compte de la faiblesse des effectifs à certains âges et des fluctuations injustifiables observées en conséquence. Il nous permet en outre de qualifier le lissage par Makeham effectué de cohérent, à défaut d'être absolument adéquat.

Finalement, pour améliorer cette méthode de lissage peu maniable, mais qui a depuis longtemps fait ses preuves, il reste à corriger les estimations en amont du lissage. Ces corrections peuvent être l'objet d'un lissage préalable par moyenne de Whittaker-Henderson, au même titre qu'un lissage global. Cette méthode est appliquée au prochain chapitre.

Chapitre 7

Ajustement des taux bruts de mortalité par la moyenne de Whittaker-Henderson

Nous abordons ici la méthode décrite section 2.2. Nous avons effectué ce lissage par un programme écrit sous MATLAB dont nous donnons la substance en annexe 3, étant donné que EXCEL ne permet pas de calcul matriciel sur de telles dimensions¹.

7.1 Ajustement par Whittaker-Henderson

Nous présentons dans cette section les différents résultats obtenus, sur le support de graphiques plus parlants que les tables correspondantes. Compte tenu de la dispersion des données étudiées, nous avons essayer de jouer au maximum sur le critère de régularité, sans utiliser de pondération particulière. Nous exposons successivement les lissages effectués avec le paramètre de régularité $h = 100 \ 1000 \ 10000$. Nous avons conservé pour chaque graphique la courbe du lissage par Makeham afin de comparer les résultats des deux méthodes.

¹ 76×76 , cf section 2.2

paramètre de régularité $h = 100$

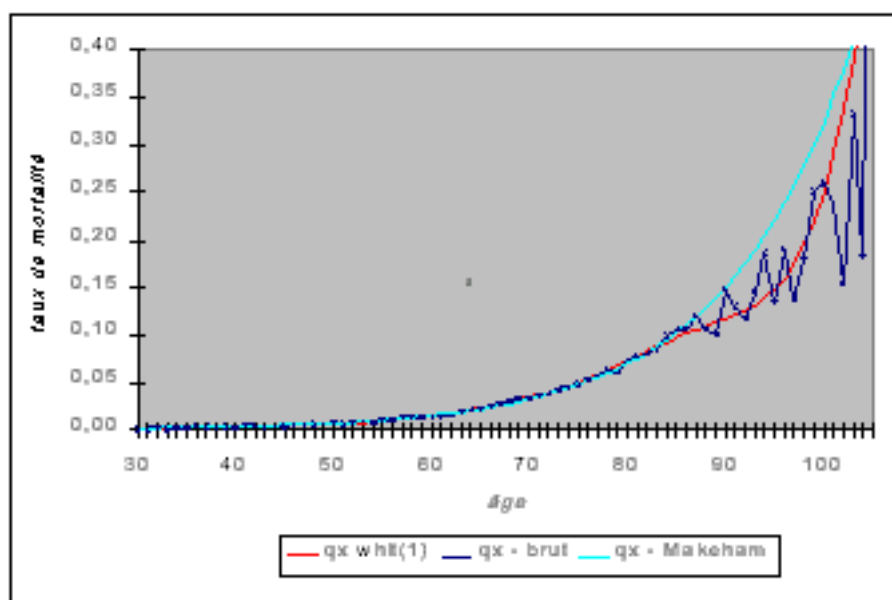


FIG. 7.1 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 100$; $z = 3$)

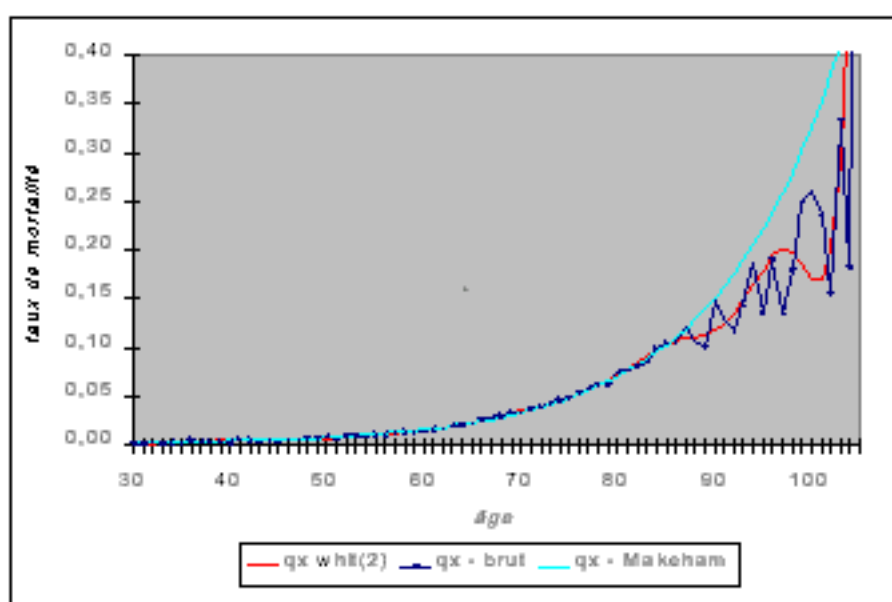


FIG. 7.2 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 100$; $z = 5$)

paramètre de régularité $h = 1000$

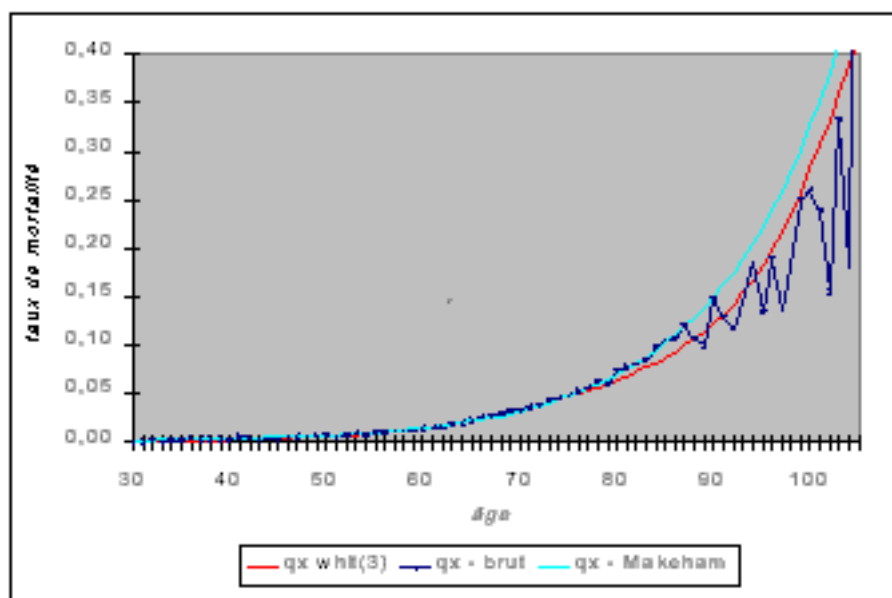


FIG. 7.3 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 1000$; $z = 3$)

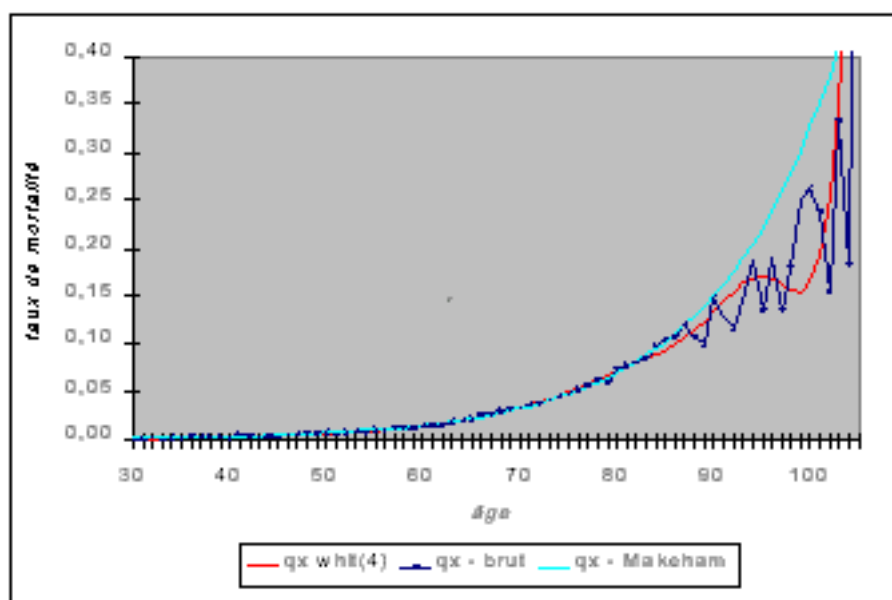


FIG. 7.4 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 1000$; $z = 5$)

paramètre de régularité $h = 10000$

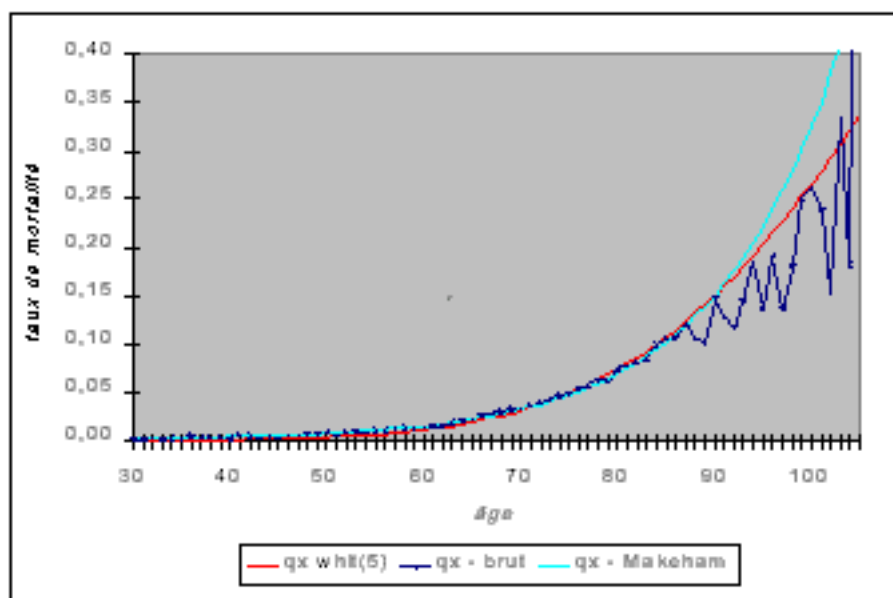


FIG. 7.5 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 10000$; $z = 3$)

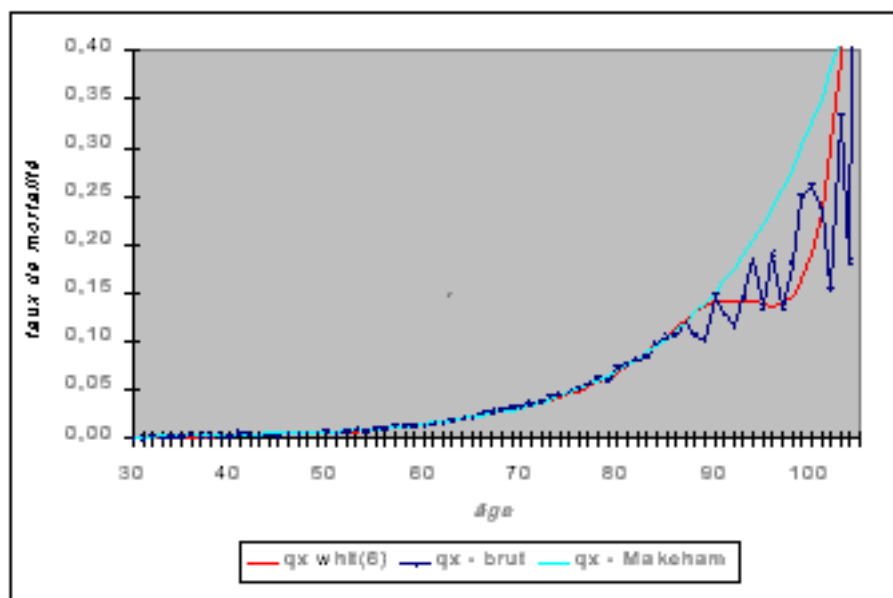


FIG. 7.6 – Lissage par Whittaker-Henderson (poids 0/1 ; $h = 10000$; $z = 5$)

7.2 Choix de la courbe la plus adaptée

Le choix du lissage peut se baser sur différents critères, jusqu'à obtenir le lissage le plus adéquat.

Comme nous l'avons déjà souligné, les fluctuations intempestives d'une courbe de mortalité ne sont nullement explicables, *a fortiori* les variations à la baisse sont difficilement imaginables d'un point de vue conceptuel pour la plage d'âges observée en dehors d'une conjoncture de guerre ou d'épidémie. Ceci nous permet d'écrémer l'échantillon de courbes exposé dans la section précédente en écartant d'ores et déjà les lissages basés sur des différences avant avec plus de trois valeurs consécutives ($z > 3$). En effet ces types de lissage collent trop aux estimations, pour ainsi dire, en conservant leurs aberrations.

Pour les mêmes raisons, nous avons ensuite écarté la courbe restante pour $h = 100$, qui présente une variation de trajectoire non justifiable vers l'âge 90.

Pour les courbes restantes, "whit(3)" et "whit(5)", nous avons effectué les mêmes tests que pour le lissage par Makeham, présentés dans la prochaine section.

7.3 Tests d'adéquation statistiques

Pour vérifier la validité de l'ajustement, nous avons effectué les deux tests statistiques décrits au chapitre 3, en voilà les résultats :

Test du Khi2

Intervalle	Distance approximée	Distance exacte	Degré de liberté	Seuil critique
jusque 60 ans	27,2392	27,8411	30	43,7729
jusque 90 ans	85,2034	68,4537	60	79,0819
jusque 105 ans	94,2829	82,0878	90	95,0814

TAB. 7.1 – Lissage par Whittaker-Henderson, whit(3) - test du Khi2

Intervalle	Distance approximée	Distance exacte	Degré de liberté	Seuil critique
jusque 60 ans	562,2495	504,7479	30	43,7729
jusque 90 ans	707,8955	633,1495	60	79,0819
jusque 105 ans	723,1342	647,9178	90	95,0814

TAB. 7.2 – Lissage par Whittaker-Henderson, whit(5) - test du Khi2

Suite à ces résultats, nous n'avons conservé que le lissage "whit(3)". Nous pouvons remarquer que l'essentiel de l'écart par la distance du Khi2 est acquis avant

90 ans, ce qui paraît plutôt problématique. On a confirmation de ces résultats sur le graphique présentant les résidus $F - F_n$, figure 7.7, où l'on peut observer un fossé grandissant entre les deux répartitions sur la plage d'âges [30,70) ans. Ce fossé signifie en termes concrets que le lissage surestimera le nombre de survivants à l'âge de retraite, et donc les engagements.

Test de Kolmogorov-Smirnov

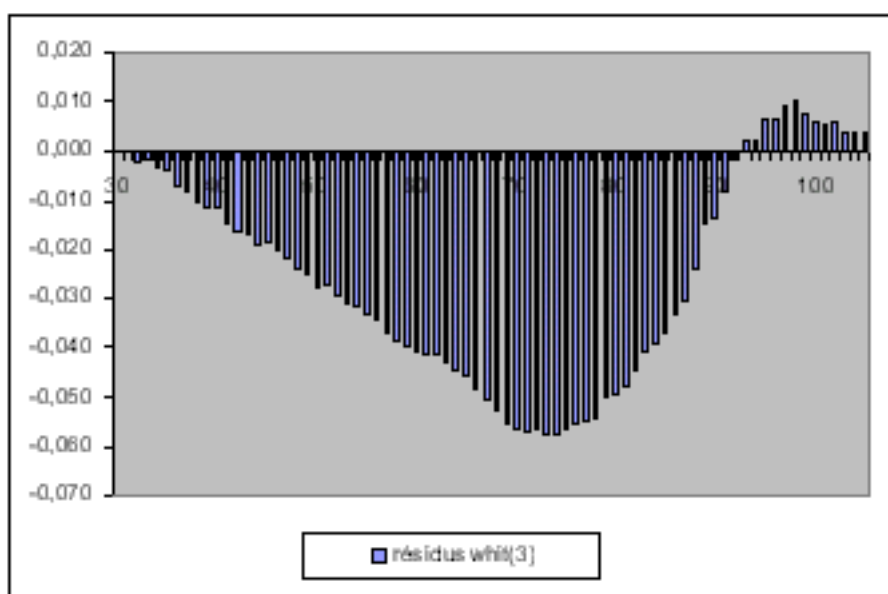


FIG. 7.7 – Résidus de répartition - whit(3)

Intervalle	Distance	Seuil critique
jusque 60 ans	0,1910	0,2520
jusque 90 ans	0,6232	1,2239
jusque 105 ans	0,6956	1,2239

TAB. 7.3 – Lissage par Whittaker-Henderson, whit(3) - test de Kolmogorov-Smirnov

Le test d'Anderson-Darling défini au chapitre précédent rejette d'ailleurs l'ajustement. La question reste de savoir si cette surestimation est inadéquate ou juste prudente, ce dont nous jugerons au chapitre suivant.

D'un point de vue de cohérence avec l'estimation faite des données, on observe

une relative cohérence (cf figure 7.8) sauf entre 80 et 90 ans, où le lissage flirte avec la borne inférieure de l'estimation.

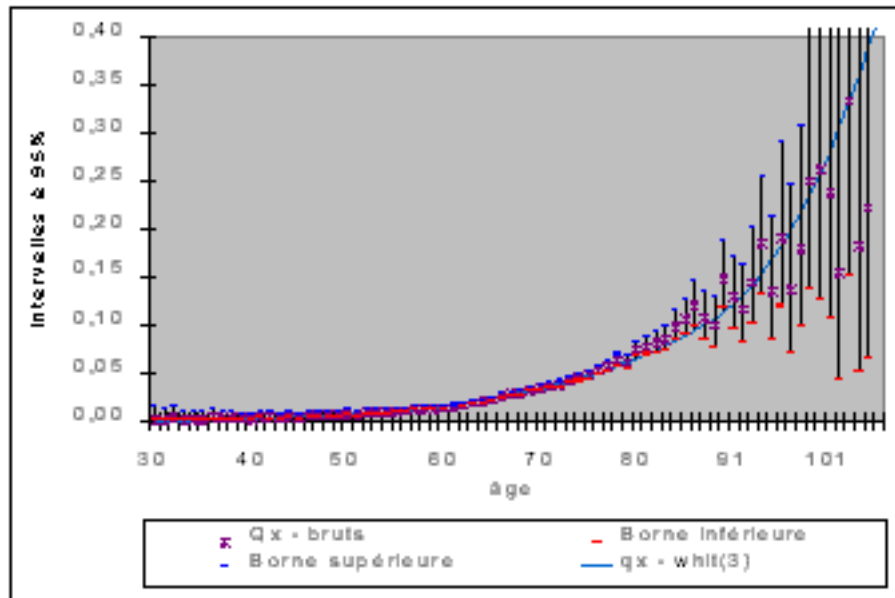


FIG. 7.8 – Test de cohérence - whit(3)

Chapitre 8

Comparaison des deux méthodes retenues

L'intérêt d'une telle comparaison¹ réside plus sur la place à donner à chacune des méthodes au cours de l'élaboration d'une table d'expérience qu'à une confrontation des résultats.

Nous allons dans ce chapitre dresser un bilan des tests d'adéquation, puis effectuer la comparaison sur la variable la plus parlante : l'espérance de vie, en se référant à la table PF60-64, et finalement étudier l'impact du choix de la table sur les effectifs de retraités du régime.

8.1 Limite des tests statistiques

Nous avons reporté sur le tableau 8.1 les résultats des tests du χ^2 effectués sur les deux lissages retenus, en ajoutant le test d'adéquation de l'ajustement par Whittaker-Henderson à l'ajustement par Makeham.

Ces valeurs nous indiquent que le lissage par Whittaker-Henderson est moins performant en termes d'adéquation pour les âges peu élevés, là où les effectifs sont les plus nombreux. Cela est peu surprenant vu le processus de lissage qui, pour Makeham, tient compte des effectifs, effectifs que nous n'avons pas pris en compte pour l'autre ajustement.

Ce que nous retenons de cette première comparaison, c'est qu'il ne faut pas se

¹un récapitulatif des taux de mortalité obtenus est présenté en annexe 3

Intervalle	Distance approximée	Distance exacte	Degré de liberté	Seuil critique
<i>Makeham</i>				
jusque 60 ans	26,9258	22,1080	27	40,1132
jusque 90 ans	64,8032	67,8527	57	75,6237
jusque 105 ans	81,2103	82,0469	71	91,6702
<i>Whittaker-Henderson</i>				
jusque 60 ans	27,2392	27,8411	30	43,7729
jusque 90 ans	85,2034	68,4537	60	79,0819
jusque 105 ans	94,2829	82,0878	90	95,0814
<i>Distance inter-lissages</i>				
jusque 60 ans	48,6700	42,3122	30	43,7729
jusque 90 ans	81,8467	56,0691	60	79,0819
jusque 105 ans	86,3976	72,3607	90	95,0814

TAB. 8.1 – Comparaison des lissages obtenus - test du Khi2

priver d'une loi analytique générale quand on peut faire jouer la loi des grands nombres, et de se réserver l'utilisation d'une méthode de lissage par série chronologique soit pour de faibles effectifs, soit éventuellement pour corriger et lisser certaines plages d'âge, typiquement ici les âges élevés.

8.2 Comparaison sur les espérances de vie

Les espérances de vie obtenues sont présentées pour les plus significatives d'entre elles dans le tableau 8.2, et graphiquement figure 8.1. Nous y adjoignons les valeurs de la table PF 60-64, table certifiée (makehamisée) dont la mortalité est la plus proche de celle observée sur la population des pensionnés étudiée.

On observe deux plages d'âges distinctes :

Espérance de vie :	à 60 ans	à 65 ans	à 70 ans	à 90 ans
PF 60-64	19,40	15,56	12,08	2,86
Makeham	19,31	15,88	12,78	4,03
Whittaker-Henderson	19,73	16,35	13,39	4,80

TAB. 8.2 – Comparaison des lissages obtenus - espérances de vie

$$\begin{cases} x \in [30, 60) & e_{\text{Makeham}} < e_{\text{Whittaker}} < e_{\text{PF60-64}} \\ x \in [60, 105) & e_{\text{PF60-64}} < e_{\text{Makeham}} < e_{\text{Whittaker}} \end{cases}$$

Le critère de prudence, comme nous l'avons montré au chapitre 4, est de ne pas surestimer la mortalité (les engagements étant sensiblement des rentes viagères),

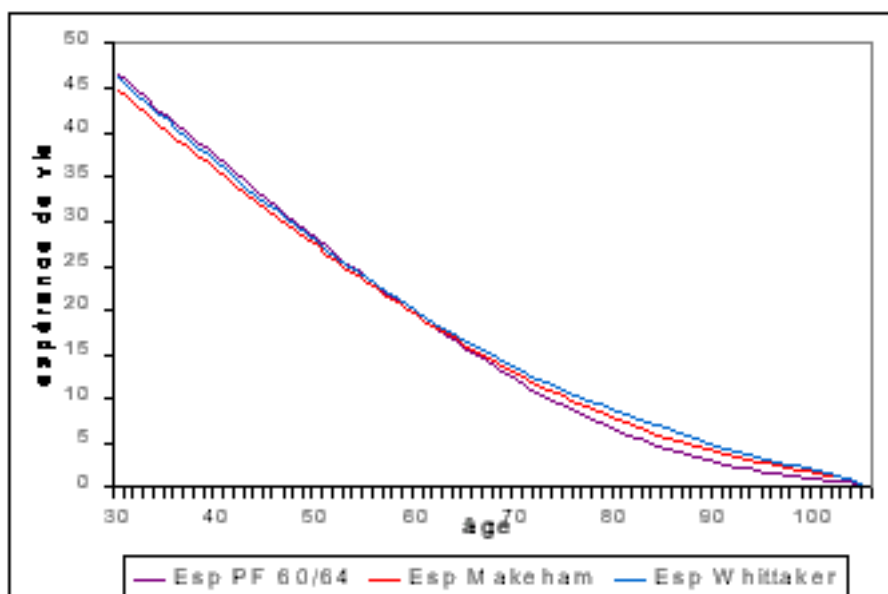


FIG. 8.1 – Comparaisons des espérances de vie

dans le cas d'un actif comme dans celui d'un retraité. Du moins sur la base de cette règle la courbe de Whittaker-Henderson serait donc à privilégier, si l'on ne tenait pas compte de la qualité de l'ajustement.

8.3 Impact sur le régime de retraite

Dans le cadre de l'étude de la viabilité du régime de retraite, nous avons effectué une projection comparative des effectifs de retraités escomptés d'ici à 2040, projection qui met en place tous les paramètres démographiques influant sur les engagements de la caisse de retraite.

Outre la date et la durée de projection, les hypothèses actuarielles utilisées sont :

- pour la population active :
 - Effectif initial (connu pour l'année 2000)
 - Turnover
 - Taux d'évolution des effectifs (entrants/sortants)
 - Âge de départ à la retraite (connu)
 - Table de mortalité (comparaison : Makeham / Whittaker-Henderson / PF 60-64)
- pour la population retraitée :
 - Effectif initial (connu pour l'année 2000)

- Les effectifs résultant de cette projection sont présentés tableau 8.3. Nous avons normé les effectifs initiaux à 100000 pour la population active, et 50000 pour la population retraitée. Il ne faut donc pas tenir compte du rapport entre actifs et retraités qui est faux, mais s'intéresser aux variations observées sur chaque projection d'effectifs, et comparer les résultats obtenus selon la table de mortalité choisie.

Projections des populations actives et retraitées						
	Makeham		Whittaker-Henderson		PF 60-64	
Année	Actifs	Retraités	Actifs	Retraités	Actifs	Retraités
2001	100 000	50 000	100 000	50 000	100 000	50 000
2002	101 089	54 639	101 139	54 776	101 180	54 624
2003	103 581	60 243	103 664	60 522	103 726	60 222
2004	105 747	66 646	105 863	67 056	105 948	66 629
2005	107 595	72 934	107 746	73 476	107 854	72 929
2006	108 943	80 493	109 128	81 142	109 259	80 512
2007	110 268	86 499	110 490	87 288	110 644	86 558
2008	111 535	93 548	111 792	94 448	111 972	93 652
2009	112 733	100 542	113 027	101 540	113 231	100 694
2010	113 846	108 410	114 178	109 455	114 407	108 606
2011	114 944	115 513	115 314	116 609	115 568	115 753
2012	115 966	123 413	116 375	124 495	116 655	123 681
2013	116 978	130 909	117 426	131 956	117 732	131 198
2014	117 909	138 428	118 396	139 394	118 728	138 719
2015	118 814	145 618	119 340	146 475	119 700	145 896
2016	119 673	153 060	120 240	153 742	120 626	153 296
2017	120 516	160 063	121 124	160 538	121 537	160 237
2018	121 323	167 129	121 972	167 336	122 413	167 209
2019	122 109	173 770	122 799	173 880	123 268	173 946
2020	122 862	180 046	123 595	180 544	124 091	180 833
2021	123 604	185 823	124 379	186 745	124 903	187 262
2022	124 313	191 483	125 130	192 885	125 683	193 651
2023	125 001	196 754	125 861	198 672	126 443	199 692
2024	125 658	201 881	126 561	204 356	127 171	205 638
2025	126 300	206 677	127 246	209 728	127 886	211 268
2026	126 909	211 262	127 899	214 903	128 568	216 697
2027	127 494	215 440	128 527	219 652	129 225	221 677
2028	128 046	219 429	129 123	224 182	129 851	226 413
2029	128 584	223 205	129 706	228 458	130 462	230 867
2030	129 102	226 837	130 268	232 520	131 055	235 068
2031	129 582	230 297	130 793	236 317	131 609	238 957

	Makeham		Whittaker-Henderson		PF 60-64	
Année	Actifs	Retraités	Actifs	Retraités	Actifs	Retraités
2032	130 059	233 725	131 315	239 982	132 160	242 667
2033	130 521	237 085	131 822	243 454	132 697	246 130
2034	130 977	240 381	132 323	246 724	133 229	249 334
2035	131 432	243 605	132 823	249 773	133 760	252 260
2036	131 923	246 802	133 361	252 647	134 328	254 956
2037	132 428	249 918	133 912	255 289	134 911	257 371
2038	132 949	252 924	134 480	257 666	135 510	259 470
2039	133 494	255 818	135 073	259 814	136 135	261 305
2040	134 060	258 602	135 688	261 742	136 782	262 892
Fin des projections						

Ces projections sont conformes aux espérances de vie trouvées précédemment, si l'on tient compte du taux d'évolution de la population active qui escompte une assez forte croissance (entrée d'affiliés pour l'ouverture du régime à d'autres cotisants) sur les dix premières années.

En effet, compte tenu de cette croissance escomptée et de l'importance de la population active proche de l'âge de retraite, le phénomène d'arrivée des actifs survivants à l'âge de retraite prend rapidement le pas sur la mortalité des retraités (après 60 ans). Du coup les effectifs de retraités estimés avec la PF60-64 sont plus nombreux que ceux estimés par le lissage makehamien dès l'année 2006, et dès 2019 par rapport au lissage par Whittaker-Henderson.

Notons que les différences observées entre les effectifs selon la table choisie sont minimales : à terme (année 2040), l'effectif des retraités évalué en utilisant la table makehamisée est seulement de 1,5% inférieur à celui obtenu avec la PF60-64 et de 1% à celui obtenu avec la table de Whittaker-Henderson.

En bilan, et compte tenu des remarques que nous avons faites, c'est la table d'expérience par Makeham que nous retiendrons. En plus de la qualité indéniable de l'ajustement makehamien, force est de constater que sa mise en place pour l'évaluation des engagements semble toujours plus commode, y compris au niveau des coûts de développement. Compte tenu toutefois de la relative inadaptation d'une telle table à prévoir la mortalité dans le futur, il a été prévu la mise en place d'une table d'expérience prospective utilisant une version adaptée de la PF60-64 déclinable par génération.

Remarquons d'ailleurs qu'il n'existe pas, à notre connaissance, de table de génération certifiée marocaine, et que la réglementation marocaine à ce sujet est inexistante. Il est donc *a priori* légal à défaut d'être prudent d'utiliser une table de moment, contrairement à ce qui se passe en France pour les engagements assimilés à des rentes viagères, obligatoirement évalués avec des tables prospectives.

Nous présentons succinctement dans la section suivante les prémices d'élaboration d'une table prospective sur ces bases.

8.4 Vers une table d'expérience prospective

A partir des dix périodes d'observation initiales, nous avons effectué dix lissages makehamiens en vue d'étudier une éventuelle dérive de la mortalité. Sur la base de ces lissages, nous avons calculé l'espérance de vie à différents âges, puisqu'il s'agit de la variable la plus parlante. Ces résultats sont présentés sur le tableau 8.4.

Espérance de vie :	à 60 ans	à 65 ans	à 70 ans	à 90 ans
Table 1990	18,6	15,4	12,5	4,0
Table 1991	18,9	15,7	12,8	4,4
Table 1992	17,8	14,4	11,4	2,8
Table 1993	18,6	15,2	12,1	3,6
Table 1994	18,6	15,2	12,0	3,3
Table 1995	18,8	15,5	12,4	3,7
Table 1996	19,1	15,7	12,7	3,9
Table 1997	19,2	15,7	12,6	3,9
Table 1998	19,6	16,0	12,7	3,7
Table 1999	21,1	17,4	14,1	4,4
Table 90/99	19,3	15,9	12,8	4,0

TAB. 8.4 – Espérances de vie selon l'année d'observation

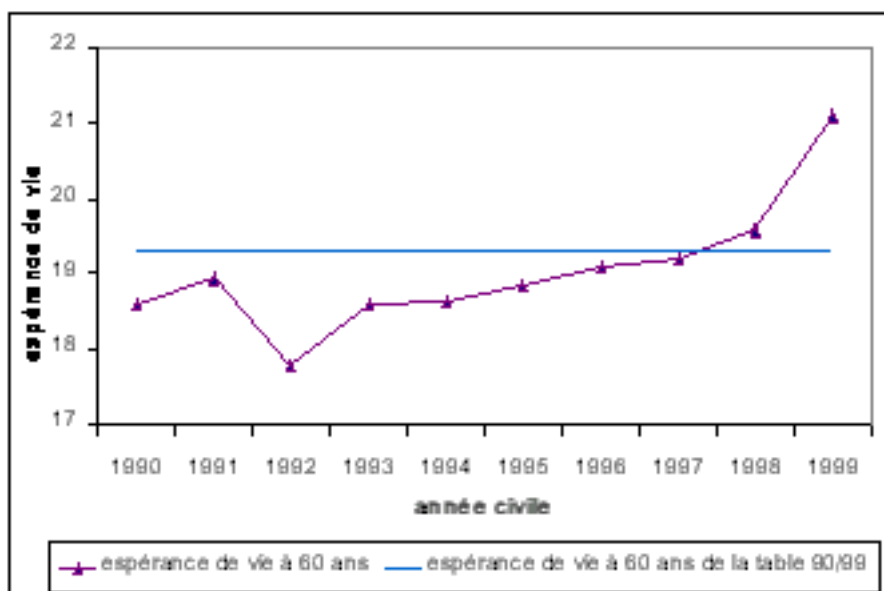


FIG. 8.2 – Espérances de vie à 60 ans selon l'année d'observation

Le graphique de l'espérance de vie à 60 ans au cours des périodes d'observations 1990 à 1999 est présenté figure 8.2, afin de rendre plus lisibles les informations contenues dans le tableau précédent.

Les années 1992 et 1999 apparaissent comme étant des années atypiques suite à des événements inexpliqués. Ces années mises à part, on observe une relative stabilité de l'espérance de vie de la population concernée. Le tableau 8.5 résume l'évolution de l'espérance de vie durant la période observée.

espérance de vie	Evolution entre 1990 et 1999	Evolution moyenne par année
à 60 ans	2 ans et demi	1 trimestre et 2 mois
à 65 ans	2 ans	1 trimestre
à 70 ans	1 an et demi	1 trimestre
à 90 ans	< 1/2 année	< 1 mois

TAB. 8.5 – Dérive de l'espérance de vie

Si la table moyenne réalisée sur les dix années d'observations s'avère être la table de référence à prendre en considération, une table prospective devra se baser sur cette dérive assez faible de la mortalité pour la population concernée.

Conclusion

Au terme de cette étude, pour laquelle nous avons abordé les principaux points ayant trait aux différentes techniques d'ajustement de courbes de mortalité de la population affiliée à un régime de retraite, nous avons abouti à diverses conclusions.

D'abord, s'agissant de la technique pure d'élaboration d'une table d'expérience, on ne mettra jamais assez l'accent sur la qualité des données historiques à la base de l'étude, et du traitement de données, quelque soit la méthode de lissage utilisée par la suite. Notamment la faiblesse des effectifs de la population observée conduit à des variations de taux de mortalité importantes, qui se traduisent par des difficultés pour obtenir un lissage adéquat.

En ce qui concerne les méthodes non exhaustives de lissage que nous avons pu mettre en oeuvre, il nous semble que l'application d'une loi analytique est plus adaptée à ce type de données, et qu'une méthode telle que celle de Whittaker-Henderson est meilleure pour lisser en amont, ou corriger certaines plages d'âges sur lesquelles les taux bruts présentent de trop grandes fluctuations.

Sur le plan de la prévision actuarielle, nous soulignerons d'abord l'importance de cadrer l'impact de la mortalité sur les engagements, ce qui permet en outre de dépasser les simples considérations statistiques. Nous considérons également primordial d'effectuer de nombreuses projections, afin de se faire une meilleure idée de cet impact et de juger de la réelle prudence de la table envisagée.

Finalement, cette étude m'a beaucoup appris, au niveau statistique bien sûr, mais surtout en termes d'expérience actuarielle, prise dans le sens d'une science de prévision économique, au niveau des concepts de modélisation et de la perception des divers phénomènes pouvant entrer en jeu dans un modèle économique, perception qui n'est pas toujours intuitive.

Bibliographie

- [1] Yoann GOUYEN. *Le risque de longévité : application aux rentes viagères en France*. Mémoire EURIA, juin 2002.
- [2] Robert LANGMEIER. *Etudes de différentes méthodes d'ajustement de tables de mortalité : application aux données d'une compagnie d'assurance*. Ecole HEC, Université de Lausanne, octobre 2000.
- [3] Abderrahim JANANI. *Elaboration d'une table d'expérience*. Mémoire ISUP, 1999.
- [4] Philippe SCHMITT. *Mise en place d'une table prospective de la mortalité au Luxembourg*. Mémoire de l'Université Louis Pasteur de Strasbourg, 1998.
- [5] Didier AUJOUX, Gilles CARBONEL. *Mortalité d'expérience et risque financier pour un portefeuille de rentes*. Mémoire du CEA Dauphine, 94-95.
- [6] Pierre PÉTAUTON. *Théorie et pratique de l'assurance-vie*. DUNOD, 2e édition, 2000.
- [7] Jean-Pierre LECOUTRE. *Statistique et probabilités*. DUNOD, 1998.
- [8] ORGANISATION MONDIALE DE LA SANTÉ, INSTITUT NATIONAL D'ÉTUDES DÉMOGRAPHIQUES, sous la direction de Roland PRESSAT. *Manuel d'analyse de la mortalité*. Imprimerie Louis Jean, octobre 1985.

Annexes

Annexe 1 : Taux instantané de mortalité

On peut mieux appréhender la notion de taux instantané par le raisonnement suivant : si on considère un individu de la population soumise au risque de mortalité (un assuré) pris en observation à l'âge x , et supposé vivant jusqu'à l'âge $x + t$, la probabilité pour qu'il décède au temps T_x dans l'intervalle de temps $(t, t + \Delta t)$ s'écrit :

$$P[t < T_x < t + \Delta t \mid T_x > t] = \frac{P[t < T_x < t + \Delta t]}{P[T_x > t]} = \frac{{}_t|\Delta tq_x}{{}_tp_x}$$

Or, en supposant la fonction de vie en temps continu ${}_tp_x$ dérivable par rapport à t , on peut écrire :

$$\lim_{\Delta t \rightarrow 0} \left[\frac{1}{\Delta t} \cdot {}_t|\Delta tq_x \right] = \lim_{\Delta t \rightarrow 0} \left[\frac{1}{\Delta t} \cdot ({}_tp_x - {}_{t+\Delta t}p_x) \right] = - {}_tp'_x$$

De plus, on a la relation :

$${}_tp_x = \frac{l_{x+t}}{l_x}, \text{ donc } {}_tp'_x = \frac{l'_{x+t}}{l_x}$$

Par conséquent, la définition du taux instantané de mortalité à l'âge $(x + t)$, μ_{x+t} , peut s'écrire comme suit :

$$\mu_{x+t} = \lim_{\Delta t \rightarrow 0} \left[\frac{1}{\Delta t} \cdot P[t < T_x < t + \Delta t \mid T_x > t] \right] = \frac{l'_{x+t}}{l_{x+t}}$$

Annexe 2 : Tests statistiques

Enoncé du test du Khi-deux

Ce test est à retenir si les données sont discrètes, avec des valeurs possibles notées x_i , de probabilité p_i pour $1 \leq i \leq k$, ou si les données individuelles ne sont pas fournies, mais ont été réparties en classes (a_i, a_{i+1}) dont les fréquences théoriques sont calculées à partir de la loi théorique postulée :

$$p_i = P[X \in (a_i, a_{i+1})] = F(a_{i+1}) - F(a_i)$$

Si N_i est le nombre (aléatoire) d'observations x_i , ou appartenant à la classe (a_i, a_{i+1}) , nous allons le comparer à l'effectif théorique qui est np_i . La distance euclidienne classique entre F_n représentée par les k effectifs observés N_i , et la fonction de répartition F , représentée par les k effectifs observés np_i , serait

$$\sum_{i=1}^k (N_i - np_i)^2$$

Cependant, comme cette distance ne permet pas de déterminer la loi asymptotique de cette variable aléatoire, on préfère retenir une autre distance. Cette dernière sera déterminée à partir de la remarque que les variables aléatoires N_i suivent des lois binomiales de paramètres n et p_i et que les variables centrées $\frac{N_i - np_i}{\sqrt{np_i}}$ convergent vers la loi $\mathcal{N}(0, \sqrt{1 - p_i})$. On retient donc la distance :

$$d(F_n, F) = \sum_{i=1}^k \frac{(N_i - np_i)^2}{np_i}$$

et cette somme de carrés de variables aléatoires centrées qui sont asymptotiquement normales et liées par la relation $\sum_{i=1}^k (N_i - np_i) = 0$ converge vers une loi χ^2_{k-1} . La valeur de C est alors déterminée approximativement, en utilisant cette loi asymptotique, comme le fractile d'ordre $1 - \alpha$ de la loi du χ^2 à $k - 1$ degrés de liberté, défini par :

$$P[\chi^2_{k-1} < C] = 1 - \alpha$$

Cette approximation est justifiée si n est assez grand et p_i pas trop petit, avec comme règle empirique $np_i \geq 5$. Si ce n'est pas le cas à cause d'une valeur de p_i

trop petite, on doit regrouper des classes (ou des valeurs) contiguës.

Cependant, si le type de loi est précisé, la loi dépend de paramètres dont la valeur n'est pas spécifiée et qu'il va falloir estimer pour pouvoir calculer les fréquences théoriques p_i . Si on doit estimer r paramètres, cela diminue d'autant le nombre de degrés de liberté qui devient alors $k - 1 - r$.

Enoncé du test de Kolmogorov-Smirnov

Dans le cas d'une variable continue pour laquelle on dispose des données individuelles, il est préférable d'utiliser toute l'information disponible et de ne pas regrouper les observations en classes. On retient alors la distance de Kolmogorov, ou distance de la convergence uniforme, définie par :

$$K_n = d(F_n, F) = \sup_{x \in \mathbb{R}} |F_n(x) - F(x)|$$

Là encore, on retiendra l'hypothèse que la loi parente admet F comme fonction de répartition si cette distance est faible, c'est à dire plus précisément si l'événement $d(F_n, F) < C$ est réalisé. La valeur de C sera déterminée par la fixation du risque d'erreur $\alpha = P[d(F_n, F) < C]$ et en utilisant la loi limite de la variable aléatoire $\sqrt{n}K_n$ qui admet pour fonction de répartition la fonction K définie pour $x > 0$ par :

$$K(x) = \sum_{k=-\infty}^{\infty} (-1)^k e^{-2k^2 x^2} = 1 - 2 \sum_{k=-\infty}^{\infty} (-1)^{k+1} e^{-2k^2 x^2}$$

Les valeurs de K sont tabulées, permettant de déterminer les fractiles de la loi. Les valeurs de C sont données en fonction de α dans la table suivante :

n	$\alpha = 0,10$	$\alpha = 0,05$	$\alpha = 0,01$
5	0,509	0,563	0,669
10	0,369	0,409	0,486
15	0,304	0,338	0,404
20	0,265	0,294	0,352
25	0,238	0,264	0,317
30	0,218	0,252	0,290
40	0,189	0,210	0,252
$n > 40$	$\frac{1,22}{\sqrt{n}}$	$\frac{1,36}{\sqrt{n}}$	$\frac{1,63}{\sqrt{n}}$

Pour le calcul pratique de cette distance, on utilise la définition de F_n faisant intervenir l'échantillon ordonné $X_{(1)} < X_{(2)} < \dots < X_{(n)}$. L'expression de $F_n = P_n[\cdot] - \infty; x]$ s'écrit alors :

$$F_n(x) = \begin{cases} 0 & \text{si } x \leq X_{(1)} \\ \frac{i-1}{n} & \text{si } x \in [X_{(i-1)}; X_{(i)}] \\ 1 & \text{si } x \geq X_{(n)} \end{cases}$$

On calcule au préalable les statistiques :

$$d^+(F_n ; F) = \sup_{x \in \mathbb{R}} [F_n(x) - F(x)] = \max_{1 \leq i \leq n} \left(\frac{i}{n} - F(X_{(i)}) \right)$$

$$d^+(F ; F_n) = \sup_{x \in \mathbb{R}} [F(x) - F_n(x)] = \max_{1 \leq i \leq n} \left(F(X_{(i)}) - \frac{i}{n} \right)$$

car F_n est constante sur chacun des intervalles délimités par les points de l'échantillon ordonné :

$$\sup_{x \in]X_{(i)} ; X_{(i+1)}]} [F_n(x) - F(x)] = \frac{i}{n} - \inf_{X_{(i)} < x \leq X_{(i+1)}} F(x) = \frac{i}{n} - F(X_{(i)} + 0) = \frac{i}{n} - F(X_{(i)})$$

On calcule ensuite :

$$d(F_n ; F) = \max \left(d^+(F_n ; F) ; d^+(F ; F_n) \right)$$

Annexe 3 : Récapitulatif des taux de mortalité

Table des taux de mortalité obtenus				
Age	q _x brut	q _x Makeham	q _x whit(3)	q _x whit(5)
30	0,0000	0,0016	0,0001	0,0001
31	0,0000	0,0017	0,0001	0,0001
32	0,0024	0,0018	0,0003	0,0001
33	0,0000	0,0020	0,0004	0,0001
34	0,0016	0,0021	0,0006	0,0001
35	0,0013	0,0023	0,0008	0,0001
36	0,0035	0,0025	0,0011	0,0001
37	0,0021	0,0027	0,0013	0,0002
38	0,0027	0,0029	0,0016	0,0002
39	0,0025	0,0031	0,0018	0,0003
40	0,0014	0,0033	0,0021	0,0004
41	0,0045	0,0036	0,0024	0,0005
42	0,0036	0,0039	0,0027	0,0007
43	0,0028	0,0042	0,0030	0,0008
44	0,0042	0,0045	0,0033	0,0010
45	0,0025	0,0049	0,0037	0,0012
46	0,0048	0,0052	0,0040	0,0014
47	0,0047	0,0057	0,0044	0,0017
48	0,0056	0,0061	0,0049	0,0020
49	0,0056	0,0066	0,0054	0,0024
50	0,0071	0,0071	0,0059	0,0028
51	0,0048	0,0077	0,0064	0,0032
52	0,0079	0,0083	0,0071	0,0038
53	0,0082	0,0089	0,0077	0,0043
54	0,0075	0,0096	0,0085	0,0050
55	0,0094	0,0104	0,0093	0,0057
56	0,0101	0,0112	0,0102	0,0065
57	0,0124	0,0121	0,0111	0,0074
Suite de la table à la page suivante ...				

Age	q _x brut	q _x Makeham	q _x whit(3)	q _x whit(5)
58	0,0121	0,0131	0,0122	0,0084
59	0,0125	0,0141	0,0133	0,0095
60	0,0136	0,0153	0,0146	0,0107
61	0,0143	0,0165	0,0159	0,0121
62	0,0153	0,0178	0,0174	0,0135
63	0,0186	0,0192	0,0189	0,0151
64	0,0198	0,0207	0,0206	0,0169
65	0,0217	0,0223	0,0224	0,0188
66	0,0254	0,0241	0,0243	0,0208
67	0,0271	0,0260	0,0263	0,0231
68	0,0294	0,0281	0,0284	0,0255
69	0,0323	0,0303	0,0306	0,0281
70	0,0330	0,0327	0,0330	0,0310
71	0,0358	0,0353	0,0354	0,0340
72	0,0374	0,0381	0,0380	0,0373
73	0,0425	0,0411	0,0407	0,0409
74	0,0451	0,0443	0,0436	0,0447
75	0,0474	0,0478	0,0466	0,0487
76	0,0522	0,0515	0,0497	0,0531
77	0,0567	0,0555	0,0529	0,0577
78	0,0622	0,0598	0,0564	0,0626
79	0,0600	0,0645	0,0600	0,0679
80	0,0734	0,0695	0,0637	0,0735
81	0,0769	0,0749	0,0677	0,0794
82	0,0801	0,0807	0,0720	0,0857
83	0,0844	0,0869	0,0765	0,0923
84	0,0980	0,0935	0,0813	0,0994
85	0,1047	0,1006	0,0866	0,1067
86	0,1059	0,1083	0,0923	0,1145
87	0,1200	0,1165	0,0985	0,1227
88	0,1064	0,1252	0,1053	0,1313
89	0,0994	0,1346	0,1130	0,1403
90	0,1476	0,1446	0,1215	0,1498
91	0,1290	0,1553	0,1311	0,1596
92	0,1160	0,1667	0,1417	0,1699
93	0,1444	0,1789	0,1537	0,1806
94	0,1854	0,1918	0,1671	0,1917
95	0,1351	0,2056	0,1819	0,2033
96	0,1899	0,2202	0,1984	0,2152
97	0,1356	0,2356	0,2165	0,2275
98	0,1800	0,2520	0,2364	0,2401
Suite de la table à la page suivante ...				

Age	q _x brut	q _x Makeham	q _x whit(3)	q _x whit(5)
99	0,2500	0,2693	0,2580	0,2531
100	0,2609	0,2875	0,2813	0,2664
101	0,2381	0,3067	0,3062	0,2799
102	0,1538	0,3269	0,3325	0,2936
103	0,3333	0,3480	0,3601	0,3075
104	0,1818	0,3701	0,3887	0,3215
105	0,2232	0,3931	0,4179	0,3356
Fin de la table				

Annexe 4 : Code Matlab utilisé pour la méthode de Whittaker-Henderson

Calcul de la matrice des poids :

```
function W = mat.poids(w)
(A,B,C,D,E,F,G,H)=textread('données.txt','%f %f %f %f %f %f %f %f','headerlines',1);
for i = 1 :length(F)
    if F(i)>0.002
        v(i,1)=w;
    else v(i,1)=0;
    end
end
W = diag(v);
W = full(W)
```

Calcul de la matrice des coefficients binomiaux d'ordre z :

```
function K = mat.binomial(z)
(A,B,C,D,E,F,G,H)=textread('données.txt','%f %f %f %f %f %f %f %f','headerlines',1);
PA1 = transpose(pascal(z,2));
v=[PA1(1, :)];
long = [0 :z-1];
for i = 1 :length(F)
    PA2(i, :) = v;
end
K=full(spdiags(PA2,long,length(F),length(F)))
```

Calcul du vecteur de Whittaker-Henderson :

```
function V(w,h,z)
(A,B,C,D,E,F,G,H)=textread('données.txt','%f %f %f %f %f %f %f %f','headerlines',1);
```

Annexe 4 : Code Matlab utilisé pour la méthode de Whittaker-Henderson100

```
C1=mat.poids(w)+h*mat.binomial(z)*transpose(mat.binomial(z)) ;  
C2 = inv(C1) ;  
V=C2*mat.poids(w)*F  
plot(V)
```