



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

le 25 Septembre 2015

Par : GUENNEUGUES Alexandre

Titre : Assurance-crédit : Impact de l'intégration des sinistres des polices mondiales sur les paramètres du modèle interne.

Confidentialité : Oui - (Durée: 2 ans)

Les signataires s'engagent à respecter la confidentialité indiquée ci-dessus

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

Mr Damien Tremel

Signature :

Entreprise :

Euler Hermes

Signature :

Membres présents du jury de l'EURIA : Directeur de mémoire en entreprise :

Mme Karine Arzur

Mr Pierre Aillot

Francisco PINA PENA

Signature :

Invité :

Signature :

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

(après expiration de l'éventuel délai de confidentialité)

Signature du responsable entreprise :

Signature du candidat :

Secrétariat :

Bibliothèque :

Résumé

Mots clefs: Assurance-crédit ; Modèle interne ; Probabilités de défaut ; Paramètres ; *Backtesting* ; Test de Kolmogorov-Smirnov ; Test de Wilcoxon

Solvabilité 2 impose de nouvelles normes prudentielles aux sociétés d'assurance. A partir du 1^{er} janvier 2016, celles-ci devront calculer le *Solvency Capital Requirement* (SCR). Cet indicateur représente le niveau de fonds propres à conserver afin d'assurer la solvabilité de la société dans l'année qui vient dans 99.5% des cas. Pour cela, il est nécessaire de construire par simulation une distribution des pertes dans 1 an. L'assurance-crédit est un secteur de l'assurance bien particulier, et bien souvent les compagnies d'assurance se détachent de la formule standard pour calculer un SCR "plus approprié" en développant un modèle interne.

EH étudie actuellement son implémentation de son modèle et l'objectif de ce mémoire est de revoir la segmentation des paramètres d'une entité particulière d'Euler Hermès (EH) disposant de polices mondiales afin que le modèle interne produise un SCR de qualité. Ces contrats sont gérés par les entités locales et, jusqu'à présent, ce sont les paramètres de ces dernières qui sont utilisés pour calculer le SCR lié à cette entité.

Le modèle prend en entrée les données portefeuilles ainsi que les paramètres fournis par les entités d'Euler Hermès. A l'aide de tests non-paramétriques, nous étudierons deux différentes manières d'intégrer ces polices dans le processus de calculs des paramètres. La première consiste en l'intégration de la sinistralité des polices mondiales dans le calcul des paramètres de l'entité. La seconde consiste à établir une entité à part pour ces contrats. La combinaison de différents tests statistiques à différents niveaux de segmentation a permis de présenter la deuxième solution comme candidate pour les années à venir. En effet, pour le moment les résultats des tests ne permettaient pas d'envisager des modifications sur le modèle interne d'EH.

Abstract

Keywords: Credit insurance ; Internal model ; Probability of default ; Parameters ; Backtesting ; Kolmogorov-Smirnov's test ; Wilcoxon's test

Solvency 2 sets new prudential standards for insurance companies. From 1st January 2016, these one will have to calculate the Solvency Capital Requirement. This indicator represents the level of equity to maintain in order to ensure the solvency of the company in the coming year in 99.5 % of the cases. For this, it is necessary to construct by simulation a loss distribution within 1 year. Credit insurance is a particular insurance industry and often the insurance companies stand out from the standard formula to calculate an SCR "more appropriate" by developing an internal model.

EH is currently studying its model implementation and the objective of this paper is to review the segmentation of the parameters in a particular entity which has world policies, so that the model Internal produce a quality SCR. These contracts are managed by local entities, and so far, the parameters of the latters are used to calculate the SCR related to this entity.

The model takes into portfolios input datas and parameters provided by the entities of Euler Hermes. Using nonparametric tests, we will explore two different ways to integrate these policies in the process of parameters' calculation. The first is the integration of the world policies' losses in the calculation of the entity settings. The second is to establish a separate entity for these contracts. The combination of various statistical tests at different levels of segmentation was used to present the second solution as a candidate for the coming years. Indeed, for the moment the test results did not allow to consider modifications on EH internal model.

Note de synthèse

Mots clefs: Assurance-crédit ; Modèle interne ; Probabilités de défaut ; Paramètres ; *Backtesting* ; Test de Kolmogorov-Smirnov ; Test de Wilcoxon

Introduction

L'assurance-crédit a pour objet de garantir les assurés contre les risques de défaillance de leurs clients (que nous appellerons acheteurs). Ce mémoire a été effectué au sein de la compagnie d'assurance-crédit Euler Hermes (EH) qui a décidé de développer un modèle interne qui englobe le risque d'assurance-crédit et de caution pour calculer son exigence de fonds propres (SCR *Solvency Capital Requirement*). Allianz est l'actionnaire majoritaire du groupe EH et partage ses modules de risques pour permettre à EH d'établir un modèle interne global. Pour faire fonctionner le modèle interne, il y a de nombreux paramètres à calibrer. C'est au niveau groupe que sont calculés les paramètres pour obtenir une certaine homogénéité dans le processus d'obtention de ces paramètres. De nombreux échanges avec l'ensemble des entités légales (EL) permettent de confirmer la validité des paramètres calculés par le groupe.

EH dispose de toutes les ressources nécessaires pour calibrer les paramètres de l'ensemble des EL modélisées par le modèle interne (31 EL). Cependant, dans le cadre de ce mémoire, nous traitons le cas d'une entité particulière. Il existe une filiale qui joue le rôle de courtier, ce n'est pas une société d'assurance car elle n'a pas d'agrément. Créée en 2007, la World Agency (WA) est un intermédiaire pour les clients importants présents dans chaque pays où réside une EL d'EH. Cette entité négocie les contrats puis les attribue aux EL concernées. Ces EL sont des sociétés d'assurance et elles reçoivent les primes associées à ces polices mondiales et indemnisent les assurés.

Actuellement, au niveau de chaque EL, il y a deux types de contrats dans les portefeuilles, les contrats de l'entité WA lié à cet EL et les contrats classiques de cette EL. Dans chaque EL, les paramètres sont calculés uniquement sur les données des contrats classiques. Les paramètres du risque crédit de WA ne sont pas estimés, ce sont les paramètres des EL correspondantes (selon le pays de l'assuré) qui sont utilisés. Ainsi l'exposition de WA est séparée en différentes parties (une par EL) et à chaque partie, nous appliquons le jeu de paramètre correspondant. WA ne dispose pas de paramètres propres à sa sinistralité.

La problématique de ce mémoire est de savoir si les acheteurs souscrits par les assurés

de WA ont un comportement similaire à ceux des filiales concernées et s'il est pertinent d'utiliser les mêmes paramètres pour ces deux types d'acheteurs en entrée du modèle interne. Dans le cas contraire, nous étudierons deux propositions pour intégrer les polices mondiales de WA au processus de calcul des paramètres.

Nous considérons alors deux solutions possibles :

- La première solution consiste à intégrer les données de sinistres et d'exposition de WA dans le calcul des paramètres des EL. Cela ne changerait pas le modèle, mais juste les paramètres d'entrée. Au lieu d'appliquer les paramètres des EL aux différents segments de WA, nous calibrons sur chacun de ces segments un jeu de paramètre basé sur la combinaison des données de WA et de l'EL liée à ce segment. Cette solution est envisageable dans un délai court, elle engendre peu de changement et ne touche pas le modèle.
- La seconde solution est de calculer les paramètres pour WA comme s'il s'agissait d'une unique EL. De cette manière, les paramètres seront adaptés au portefeuille WA. Mettre en place le calcul d'un jeu de paramètre pour WA est semblable à l'ajout d'une nouvelle EL. L'ajout de nouvelles EL est un sujet qui a déjà été traité, nous avons donc de la visibilité sur la méthodologie à suivre. Ensuite, la complexité sera d'enlever le traitement de WA tel qu'il est fait actuellement par le modèle interne. Il y aura donc des modifications à effectuer dans le code du modèle interne et à justifier auprès du régulateur. De plus, Il n'est pas évident que les paramètres globaux soient adaptés à tous les segments.

Travaux

Nous avons commencé par présenter les paramètres du modèle interne d'EH. Ensuite, nous avons expliqué les fondements du modèle afin de mettre en avant la façon dont ils interviennent dans le processus de calcul du SCR. Par la suite, nous avons établi une comparaison entre les polices mondiales de WA et les polices classiques des EL. Puis, nous avons procédé des tests de validation rétroactifs sur les paramètres actuellement utilisés pour calculer le SCR. Ensuite, nous avons étudié la sensibilité du SCR face à des chocs sur les différents paramètres d'entrée du modèle interne, afin de décerner l'impact d'une faible variation des paramètres sur le SCR. Puis nous avons comparé les comportements des paramètres impactant le plus le SCR, sous différentes segmentations afin de définir des ressemblances à l'aide de tests statistiques. Nous avons d'abord tenté d'accorder les nouvelles données avec celles existantes pour obtenir un jeu de paramètres plus adapté. Puis, nous avons étudié la construction d'une base propre aux nouvelles données. Pour finir, nous mettrons en avant la stratégie choisie ainsi que les arguments qui nous auront poussés à choisir cette stratégie.

Une particularité importante du secteur du risque de crédit est la prédiction mathématique qu'un défaut ait lieu. C'est pourquoi, pour illustrer notre présentation des paramètres, nous définissons les Probabilités de Défaut (PD).

D'abord, commençons par définir un événement de défaut. Dans la calibration des paramètres, un acheteur en défaut est un acheteur pour lequel a été déposée au moins une déclaration de sinistre de la part d'un assuré.

La PD à la date t représente la probabilité qu'au moins un assuré déclare un sinistre sur un acheteur connu, pendant la période d'observation de 12 mois débutant à la date t .

$$PD(t) = \frac{\text{Nombre d'acheteurs dénommés avec une exposition positive en } t \\ \text{et au moins un sinistre déclaré dans les 12 mois suivant } t}{\text{Nombre d'acheteurs dénommés avec une exposition positive en date } t} \quad (1)$$

Concrètement, il faut prendre une photo du portefeuille au début de la période d'observation, puis compter le nombre d'acheteurs dans le portefeuille (le dénominateur). Sur les 12 mois suivants, nous observons parmi ces acheteurs, ceux qui font défaut (le numérateur). Pour qu'il y ait vraiment défaut, il faut un montant assuré strictement positif.

La date du défaut dépend du type de contrat, il y a deux cas possibles :

- Soit c'est la date de facturation. Dans ce cas, nous parlons d'attachement au risque.
- Soit c'est la date de déclaration de sinistre. Dans ce cas, nous nous rattachons à la date d'apparition de la perte.

Un point important est la segmentation, en effet suite à un long travail, EH a choisi sa segmentation pour calculer ses PD. Cette segmentation se retrouve dans le calcul des autres paramètres. Nous appellerons PARS le deuxième paramètre le plus important du modèle dans le reste de cette synthèse.

Nous avons ensuite décrit les particularités de l'entité étudiée. WA offre des produits simples et flexibles à des situations parfois complexes. Elle dispose du réseau mondial des experts EH tant en prévention qu'en recouvrement. Elle met en place une plate-forme de décision centralisée. Elle offre un service dans la langue et la devise du client. Par la suite, nous avons calculé les différents paramètres associés à cet entité, puis nous avons mené une étude sur les corrélations des PD ainsi que des tests de validation rétroactifs (*backtesting*) sur les paramètres du modèle interne ainsi que sur la perte attendue à l'horizon d'un an.

- La première étude nous montre que les corrélations entre WA et ses principaux segments sont élevées. Nous supposons que la probabilité globale de WA et celles qui sont séparées par pays pourraient suivre la même tendance. En regardant les corrélations entre les segments fortement corrélés avec WA, nous observons qu'ils sont également fortement corrélés entre eux.

- Le *backtesting* des PD nous permet de conclure que la méthodologie actuellement utilisée fonctionne pour certains segments de WA et qu'un effet de compensation entre les segments de WA sous-estimés et surestimés pourrait donner des résultats cohérents. Cependant, les prédictions ne reflètent pas fidèlement le comportement des segments de WA.
- Le *backtesting* des PARS nous fait remarquer une tendance générale à la baisse convergeant vers un niveau commun. Cependant, même si les écarts sont plus faibles, il reste toujours un écart conséquent entre les EL les plus éloignées. Même si le niveau des PARS n'est pas homogène entre les EL, nous pouvons observer que les PARS de WA sont plus élevés que ceux des EL.
- Le *backtesting* de la perte attendue nous montre un manque évident de régularité : Sur certains segments, nous avons d'importantes surestimations, sur d'autres, nous avons des sous-estimations non négligeables.

L'étude de la situation actuelle montre que la méthodologie utilisée n'est pas vraiment adaptée pour calculer les paramètres de WA.

Ensuite, pour mettre en avant les paramètres dont les variations ont l'impact le plus fort sur le SCR, nous avons procédé à des tests de sensibilité sur le modèle interne. Cela permettra de déterminer quels paramètres doivent être les mieux calibrés et quels paramètres sont suffisamment proches pour être agrégés. Nous avons dû calibrer ces chocs à l'aide des données correspondant à un moment précis dans le temps, ce ne pouvait pas être un écart-type historique. L'étude sur la sensibilité des paramètres du modèle nous a permis d'observer que les paramètres qui ont le plus d'influence sur le SCR lorsqu'ils sont choqués, sont les PD ainsi que les PARS. Ils ont une sensibilité majeure que nous pouvons qualifier de linéaire. Nous fonderons donc notre étude des différentes configurations possibles sur les PD et les PARS

Enfin, nous testerons les deux solutions envisagées. Tout d'abord, la solution consistant à intégrer les données WA au sein du processus de calcul existant pour chaque EL. Puis, nous constituerons une EL à part entière avec les données WA. Les études descriptives ont montré que la part de WA dans l'exposition totale augmente de manière régulière. La croissance de WA nous permet d'obtenir davantage de données pour effectuer des calculs. Cela permet d'envisager une intégration en tant qu'EL virtuelle au sein du modèle interne, et si l'exposition continue d'augmenter de manière soutenue, un calcul de paramètres pour WA segmenté par les pays où elle exerce son activité d'assurance deviendra envisageable.

Afin de procéder aux différentes comparaisons des paramètres, nous avons utilisé différents tests. Les tests rencontrés dans notre formation supposent généralement que la variable aléatoire X étudiée suit une loi normale dans les populations considérées. Cette hypothèse n'est pas toujours vérifiée, notamment dans le cas des PARS qui n'ont que des valeurs positives et semble plutôt suivre une loi gamma. Ainsi, nous travaillons avec les tests non-paramétriques de Kolmogorov-Smirnov et Wilcoxon. Après avoir présenté

les tests retenus, nous avons testé si les échantillons des différentes bases ont le même comportement et s'ils peuvent être issus d'une même population. Ce qui est équivalent à savoir si les différences entre les échantillons sont significatives.

Nous synthétisons les résultats des tests non-paramétriques dans ce tableau :

Segment de WA	méthode 1		méthode 2	
	PD	PARS	PD	PARS
ACI	Refusé	Accepté	Refusé	Accepté
HERM	Refusé	Accepté	Refusé	Accepté
SFAC	Refusé	Refusé	Accepté	Accepté
SIAC	Refusé	Refusé	Refusé	Accepté
TI	Refusé	Refusé	Refusé	Refusé

TABLE 1 – Récapitulatif des résultats des tests

En ce qui concerne la première méthode, il s'agit de comparer les EL avec les segments WA correspondants.

Au niveau des PD, nous avons vu que le test sur données appariées de Wilcoxon rejette toutes nos hypothèses de comportements similaires au niveau global.

Au niveau des PARS, l'EL de l'Italie révèle le plus de différences entre WA et EL. Cependant, nous observons que globalement, il n'y a pas de différences significatives entre l'échantillon ALL et l'échantillon EL, et ceci pour l'ensemble des EL.

Les PD et les PARS sont les éléments les plus importants dans le processus de calcul de la perte attendue, il faut éviter qu'ils varient de manière trop importante. Ainsi, même si sur certains segments, l'incorporation des données de WA dans le processus de calcul des paramètres d'entrée du modèle interne serait possible et refléterait mieux la réalité, dans l'ensemble il y a trop de variabilité entre les données WA et EL. Il va donc falloir considérer la deuxième méthode : Calculer des paramètres pour WA comme si c'était une EL à part entière. Cette méthode semble être une meilleure solution. En effet, nous observons au niveau des PARS qu'il y a de nombreuses correspondances entre les différents segments de WA et WA.

Cependant au niveau des PD, les résultats ne sont pas satisfaisants, Il n'y a que la France qui a des similarités suffisantes pour envisager l'application de cette méthode.

La conclusion de ces tests est que pour le moment nous ne pouvons pas changer notre méthode de prise en compte des données de WA dans notre modèle. Nous avons montré que la façon actuelle de modéliser le SCR sur ce segment n'est pas très fidèle à la sinistralité. Cependant, notre méthode actuelle se révèle très prudente sur les dernières années (*Expected Loss* bien supérieure au sinistres observées). Cependant, il va falloir traiter WA d'une manière plus adapté dans le modèle interne. En effet l'ensemble des tests menés ont montré que le comportement de WA est significativement différent de celui des EL. Cependant, la solution consistant à prendre en compte les données de WA dans le calcul des paramètres n'est pas satisfaisante, cela a un impact trop fort sur les paramètres robustes que l'on applique aux EL. Celle consistant à considérer WA comme

une EL donne de meilleurs résultats au niveau des UGD, mais pas les résultats des tests sur les PD ne sont toujours pas suffisants.

Conclusion

Après avoir remis en cause la non prise en compte actuelle des polices mondiales dans le calcul des paramètres du modèle interne, nous avons testé deux nouvelles méthodes permettant d'intégrer ces polices. La première consistant à mélanger les polices mondiales et les polices locales afin d'obtenir un jeu de paramètres par EL. La seconde consistant à former une EL avec les polices mondiales. Pour sélectionner la méthode la plus cohérente avec les données, nous avons effectués des tests non-paramétriques (Wilcoxon, Kolmogorov-Smirnov). Nous avons dû poser certaines hypothèses qui nous permis d'obtenir les résultats du tableau 1. Ces tests ont été appliqués aux PD et aux PARS. L'hypothèse de la première méthode est que le comportement de ces deux paramètres est similaire entre les EL et les segments de WA correspondants. Pour la deuxième méthode, nous avons pris pour hypothèse que les paramètres (PD et PARS) des segments de WA ont des comportements similaires à WA.

WA n'est pas encore assez développée pour pouvoir calculer des paramètres propres à elle sur chacun des segments reliés à une EL, cette solution sur-mesure n'est pas envisageable. En effet, pour le moment, il n'y a pas assez de données pour les grands segments, et il est difficile de trouver une méthode de regroupement satisfaisante pour les petits segments. C'est surtout au niveau de la sinistralité qui se révèle faible que nous manquons de données.

Nous avons calculé des paramètres pour l'entité virtuelle que formerai WA, pour le moment, nous observons une perte attendue significativement plus élevée que celle que nous avions auparavant. Il reste quelques ajustements possibles (adaptation des classes d'exposition,...). Mettre cette mesure en place va requérir beaucoup de temps, notamment au niveau de la programmation du modèle interne. Cela requiert des études de segmentation afin de définir de nouvelles classes d'exposition spécifique à cette EL. La mise en place d'un modèle de projection des PD pour WA est envisageable. Cependant, il serait fondé sur des indicateurs externes mondiaux difficiles à choisir, car les données internes sont pauvres. Une PD définie par jugement d'expert peut être utilisée en attendant. Cela va impacter de nombreux codes. Il va falloir présenter cela au régulateur afin qu'il valide notre étude. Il reste quelques axes qui n'ont pas été exploré, notamment l'étude sur les derniers paramètres qui selon les tests de sensibilité ont un effet moindre sur le SCR.

Pour le moment, il n'est pas envisageable d'impacter notre modèle, il faut attendre au moins une année pour confirmer la convergence des PD d'une part, et d'autre part, il faut commencer à envisager toutes les modifications que cela va entraîner. Ce mémoire a mis en avant qu'il n'était clairement pas envisageables de regrouper les segments de WA aux EL correspondantes. Il a également permis de mettre en avant une méthode qui ne peut pas être mise en place pour le moment, mais qui se sera à réexaminer dans les années à venir.

Executive summary

Mots clefs: Credit insurance ; Internal model ; Probability of default ; Parameters ; Backtesting ; Kolmogorov-Smirnov's test ; Wilcoxon's test

Introduction

Credit insurance insures the policyholder against the risk of default of their customers (called buyers). This paper has been written in Euler Hermes (EH), a credit insurance company, who has decided to develop an internal model which includes the credit insurance risk and surety, to calculate its solvency capital requirement (SCR). Allianz is the majority shareholder of the EH group, and it shares some risks' modules to allow EH to develop a comprehensive internal model. To operate the internal model, there are many parameters to calibrate. It's at group level that are calculated these parameters to maintain a certain homogeneity in the parameters calculation process. Many meetings with legal entities (EL) confirm the validity of the parameters calculated by the group.

EH has all the resources needed to calibrate the parameters of the EL modeled by internal model (31 EL). However, as part of this thesis, we will treat the case of a particular entity. There is a subsidiary that acts as a broker, this is not an insurance company because it has no license. Created in 2007, the World Agency (WA) is an intermediate for major customers in each country where there is an EL of EH. This entity negotiates contracts and assigns them to EL concerned. These EL are insurance companies and they receive the premiums associated with these world policies and compensate policyholders.

Currently, in each EL, there are two types of contracts in the portfolios, contracts of WA entity related to that EL and classics EL contracts. In each EL, the parameters are calculated only on data of conventional contracts. WA credit risk parameters are not estimated, EH use the parameters of the corresponding EL (according to the country of the policyholder). Thus the exposure of WA is separated into different parts (one per EL) and for each part, the corresponding set of parameters is applied. WA does not have its own loss parameters.

The issue of this paper is to know whether policyholders' buyers subscribed by WA have the same behaviors as those of the subsidiaries concerned and whether it is appropriate to use the same parameters for both types of buyers in the internal model. Otherwise, we will discuss of two proposals for WA policies integration in the calculation process.

We then consider two possible solutions :

- The first solution is to integrate the loss and exposure data of WA in the calculation of EL parameters. This does not change the model, but just the input parameters. Instead of applying the EL parameters to different segments of WA, we calibrated for each one of these segments one set of parameters based on the combination of data from EL and from WA segment linked to this EL. This solution could be set up in short term, it doesn't affect the model and generates few changes.
- The second solution is to calculate the parameters for WA as if it was a single EL. In this way, the parameters would be adapted to WA portfolio. Implement the calculation of a set of parameters for WA is similar to adding a new EL, so there is visibility on what there is to do. Then, the complexity would be to remove the current method used for WA. There would be changes to be made in the internal model code and to be defended in front of the insurance regulators. Moreover, it is not evident that the global settings would be adapted to all segments.

Methodology

We first presented the parameters of the EH internal model. Then we explained the fundamentals of the model to highlight how they are involved in the SCR calculation process. Subsequently, we established a comparison between world policies and those of EL. Then we made retroactive validation tests on the parameters currently used to calculate the SCR. Then we studied the SCR sensitivity to shocks on different input parameters of the internal model in order to discern the impact of a small variation of parameters on the SCR. Then we focused on the parameters that have the biggest impact on the SCR. We compared the behavior of those parameters under different segmentations to define similarities by using statistical tests. First we tried to mix the new data with the existing ones to obtain a set of more suitable parameters. Then we studied the construction of a proper data base for the WA. Finally, we defined the final strategy and the arguments that have led us to choose this strategy.

An important feature of the credit risk sector is the mathematical prediction that a default occurs. Therefore, to illustrate our parameters presentation, we defined the probabilities of Default (PD).

First, we start by defining an event of default. In the parameters calibration, a defaulting buyer is a buyer for which at least one claim has been filed by a policyholder.

The PD at time t is the probability that at least one policyholder declares a claim on a named buyer, during the observation period of 12 months beginning on the date t .

$$PD(t) = \frac{\text{number of named buyers with a positive exposure in } t \text{ with at least one claim reported in the next 12 months after } t}{\text{number of named buyers with a positive exposure at time } t} \quad (2)$$

Specifically, we must take a picture of the portfolio at the beginning of the observation period, then we have to count the number of buyers in the portfolio (the denominator). On the following 12 months, we observe among these buyers, the defaulting ones (the numerator). In order to avoid default that are not covered we define that the insured amount must be strictly positive.

The date of default depends on the type of contract, there are two possible cases :

- Either this is the invoice date. In this case, it's risk attaching.
- Either this is the declaration date of claim. In this case, it's loss occurring

An important point is the segmentation, indeed after a long study, EH has chosen its segmentation to calculate its PD. This segmentation is reflected in the calculation of other parameters. We will call PARS the second most important parameter of the model in the rest of this synthesis. It's the parameters which has the biggest impact on the SCR after the PD.

We then described the specificities of the entity studied. WA offers simple and flexible products to complex situations. It has a global network of EH experts in prevention and recovery. It establishes a centralized decision platform. It offers a service in the language and currency of the customer.

Subsequently, we calculated the parameters associated with that entity, then we conducted a correlation study on the PD and some retroactive validation tests (backtesting) on the parameters of the internal model, as well as the expected loss within one year.

- The first study shows that the correlations between WA and its main segments are high. We assume that the overall probability of WA and those that are separated by countries could follow the same trend. Looking at the correlations between the segments strongly correlated with WA, we observed that they are also highly correlated.
- The backtesting of PD enables us to conclude that the current methodology works for some segments of WA and a compensation effect between segments of WA underestimated and overestimated could give consistent results. However, the predictions do not accurately reflect the behavior of segments of WA.
- The backtesting of PARS pointed out to us a general downward trend converging to a common level. However, even if the differences are smaller, there is still a substantial gap between the farthest EL. Although the level of PARS is not uniform between the EL, we can observe that the PARS WA are higher than the EL.
- The backtesting of the expected loss shows a clear lack of regularity : On some segments, we have significant overestimates and on others, we have significant underestimates.

The study of the current situation shows that the methodology used is not really adapted to calculate the WA parameters.

In order to highlight the parameters whose variations have the greatest impact on the SCR, we performed sensitivity tests on the internal model. This determined which parameters should be best calibrated and which parameters are close enough to be aggregated. We had to calibrate these shocks with the data corresponding to a specific point in time, it could not be a historical standard deviation. The study on the sensitivity of the model parameters allowed us to observe that the parameters that have the most influence on the shocked SCR are the PD and the PARS. They have a major sensitivity that we can qualify of linear. We based our study of different possible configurations on PD and PARS

Finally, we will test the two proposed solutions. First, the solution consisting in integrating WA's data within the existing calculation process for each EL. Then, we will build a full EL with data WA. Descriptive studies have shown that the share in the total exposure WA increases regularly. WA growth allows us to get more data to perform calculations. This allows to consider integration of WA as one virtual EL within the internal model. If the exposure still increasing steadily, a parameters set for WA segmented by countries where it conducts its insurance business could be developped.

To make the different comparisons between the bases of the parameters that have been calculated, we used different tests. The tests encountered in our educational background generally assume that the studied random variable X follows a normal distribution in considered populations. This assumption is not always true, particularly in the case of PARS that only have positive values and moreover follow a gamma distribution. So, we worked with the non-parametric tests of Kolmogorov-Smirnov and Wilcoxon. After presenting the selected tests, we tested whether the samples different bases behave the same way and if they can be derived from the same population. Which is equivalent to know whether the differences between the samples are significant.

We synthesize the results of the nonparametric tests in this table :

Segment de WA	méthode 1		méthode 2	
	PD	PARS	PD	PARS
ACI	Rejected	Accepted	Rejected	Accepted
HERM	Rejected	Accepted	Rejected	Accepted
SFAC	Rejected	Rejected	Accepted	Accepted
SIAC	Rejected	Rejected	Rejected	Accepted
TI	Rejected	Rejected	Rejected	Rejected

TABLE 2 – Summary of test results

As regards the first method, we compared the EL with the corresponding WA segments.

For the PD, we saw that the Wilcoxon test on paired data rejects all our tests.

In terms of PARS, El Italy reveals more differences between WA and EL than the others entities. However, it is observed that overall there was no significant difference between the EL sample and the ALL sample, and this for all EL.

The PD and PARS are the most important elements in the process of calculating the expected loss, we have too avoid too big variations. So while in some segments, the incorporation of data WA in the process of calculating the internal model input parameters would be possible and reflect better the reality, overall there is too much variability between WA datas and EL datas. So we have to consider the second method : Calculate the parameters for WA as if it were an EL. This method seems to be the best solution. Indeed, we observed that the PARS have several connections between WA and its different segments.

However at the PD, the results are not satisfactory, Only France has sufficient similarity to consider the application of this method.

The conclusion of these tests is that currently we can not change our method of data integration in our model. While we have shown that the current way of modeling the SCR in this segment is not very faithful to the related claims. However, our current method proves to be very prudent in recent years (Expected Loss well above the observed losses).

It is to be dealt WA a more suitable way in the internal model. Indeed all the conducted tests showed that the behavior of WA is significantly different from that of EL. However, the option of taking into account the WA data in the calculation of the parameters is not satisfactory, this has a too strong impact on robust parameters that apply to EL. That of considering as a WA EL gives better results at the PUD, but the results of tests on PD are still not sufficient.

Conclusion

After challenging the not consideration of the current global policies in the calculation of the internal model parameters, we tested two new methods to integrate these policies. The first consisting in merging global policies and local policies to obtain a parameter set for each EL. The second consisting to forming an EL with global policies. To select the most consistent method with the data, we performed nonparametric tests (Wilcoxon, Kolmogorov-Smirnov). We had to make certain assumptions that allow us to get the results in Table 2 on the previous page. These tests were applied to PD and PARS. The hypothesis of the first method is that the behavior of these two parameters is similar between the EL and segments corresponding WA. For the second method, we assumed that the parameters (PD and PARS) of WA's segments have similar behaviors to WA.

WA is not yet developed enough to be able to calculate specific parameters to each of the segments connected to one EL, this customized solution is not possible. Indeed, for the moment, there is not enough data for several segments, and it is difficult to find a satisfactory method of consolidation for small segments. It is especially due to the low number of observations.

We calculated the parameters for the virtual entity that will form WA, for the moment, we have a significantly higher expected loss than we had before. There are some

possible adjustments (adaptation of exposure classes, ...). To put this measure in place will require considerable time, particularly in terms of programming the internal model. This requires segmentation studies to define new classes of specific exposure to the EL. The establishment of a projection model PD for WA is possible. However, it would be based on global external indicators because the internal data are very poor. A PD determined by expert judgment may be used in the meantime. This will impact many codes. We'll have to submit it to the regulator validation. There are still some areas that have not been explored, including the study on the last parameters which have a lower effect on the SCR according to the sensitivity tests.

For now, it is not possible to impact our model, on the one hand we have to wait at least a year to confirm the convergence of PD, and on the other hand, we must begin to consider any changes that will train. This thesis put forward that it was clearly not possible to group the segments of the corresponding WA EL. It has also put forward a method that can not be implemented for now, but which will review in the coming years.

Remerciements

Avant de développer cette expérience dans le milieu professionnel, il semble approprié de commencer ce mémoire par remercier les personnes qui m'ont consacré de leur temps et m'ont fait apprendre de nombreuses notions pendant ce stage. Je tiens également à remercier les personnes qui ont eu la gentillesse de faire de ce stage un moment très agréable.

Je tiens à remercier vivement mon maître de stage, Francisco PINA PENA, modélisateur de risque au sein de l'entreprise EULER HERMES, pour son accueil, le temps passé ensemble et le partage de son expertise au quotidien. Il m'a formé et accompagné tout au long de cette expérience professionnelle avec beaucoup de patience et de pédagogie. Grâce aussi à sa confiance, j'ai pu m'accomplir totalement dans mes missions avec son aide précieuse dans les moments les plus délicats.

Je remercie également Anna PORYADYNSKA, responsable de l'équipe modélisation, qui a contribué à la mise en place du sujet de ce mémoire, et m'a encadrée pour de nombreuses analyses.

Je remercie M DAVIT, responsable du risque du groupe (*Chief Risk Officer*), et M CORNU responsable pilier 1, pour m'avoir accueilli au sein de leurs équipes. Merci à toute l'équipe solvabilité 2 qui a contribué au succès de mon stage par leur accueil et leur esprit d'équipe et notamment M DELMOTTE, consultant externe, qui m'a beaucoup aidé à comprendre les problématiques de l'ORSA.

Mes remerciements s'adressent aussi au directeur de l'EURO Institut d'Actuariat de Brest, monsieur Marc Quincampoix, ainsi qu'à l'ensemble du corps enseignant pour m'avoir permis de suivre une formation de qualité.

Enfin, je tiens à remercier toutes les personnes qui m'ont conseillé et relu lors de la rédaction de ce rapport de stage : ma tutrice attribuée par mon école Mme LEDONGE, mes collègues de promotion, ma famille.

Table des matières

Résumé	i
Abstract	ii
Note de synthèse	iii
Executive summary	ix
Remerciements	xv
Introduction	1
1 L'assurance-crédit chez Euler Hermes	3
1.1 L'assurance-crédit	3
1.2 Le choix du modèle interne	4
1.3 Les paramètres d'entrée	8
1.3.1 Probabilités de défaut	8
1.3.2 Perte réelle de l'assuré en cas de défaut	11
1.3.3 Corrélations	14
1.3.4 Acheteurs non dénommés	15
1.3.5 Acheteurs non liés	16
1.4 Description du fonctionnement du modèle interne	17
1.4.1 Modélisation du facteur de crédit par une approche de type Merton	18
1.4.2 RSQ et matrice de variance covariance	22
1.4.3 Simulation stochastique des UGD	24
1.4.4 Application des autres paramètres de réduction de la perte	25
2 Cas d'étude : la World Agency	27
2.1 Présentation de la World Agency	27
2.1.1 Principe	27
2.1.2 Aperçu des produits	28
2.1.3 Similarité avec les EL classiques	29
2.1.4 Contexte	29
2.2 Analyses préliminaires	31

2.2.1	Répartitions de l'exposition entre les pays	31
2.2.2	Corrélations	32
2.2.3	Comparaisons des PD de WA par pays avec celles des EL correspondantes.	34
2.2.4	Tests de validation rétroactifs des PD de WA par pays par rapport à celles qui sont prédites pour les EL correspondantes	35
2.2.5	Test de validation rétroactifs des UGD de WA par pays par rapport aux UGD des EL correspondantes	37
2.2.6	Test de validation rétroactif de l' <i>Expected Loss</i> de WA par rapport à celle calculée avec l'exposition de WA et les paramètres des EL correspondantes	38
2.3	Méthodes envisageables	39
3	Tests de sensibilité sur les paramètres du modèle interne	41
3.1	Présentation des tests de sensibilité	41
3.1.1	Choc sur les PD	42
3.1.2	Choc sur les UGD	43
3.1.3	Choc sur les RR et SR	44
3.1.4	Choc sur les expositions et acheteurs non dénommés et non liés	45
3.2	Résultats	46
3.3	Impact dans notre étude	47
4	Etude de l'intégration des données de WA dans le calcul des paramètres	49
4.1	Qualité des données	49
4.1.1	Complétude	49
4.1.2	Conformité	50
4.1.3	Cohérence	50
4.1.4	Précision	50
4.1.5	Duplication	51
4.1.6	Intégrité	51
4.1.7	Exactitude Temporelle	51
4.1.8	Répartition des différents tests	51
4.1.9	Qualité des données World Agency	53
4.2	Préparation des bases de données	54
4.3	Description des bases utilisées	55
4.3.1	Etude des expositions	55
4.3.2	Etude de la sinistralité	57
4.4	Description des tests utilisés	60
4.4.1	Test de Box-Pierce	61
4.4.2	Tests des <i>runs</i>	61
4.4.3	Test d'indépendance du Khi-carré	62
4.4.4	Test d'indépendance de Fisher	65
4.4.5	Test de Kolmogorov-Smirnov	65
4.4.6	Test de Wilcoxon	66

4.5	Etude de dépendance sur les PD	69
4.5.1	Création de classes de PD	69
4.5.2	Dépendance au sein des échantillons	70
4.5.3	Dépendance entre les échantillons	70
4.6	Etude de l'intégration de WA dans les bases des EL	71
4.6.1	Exemple concret d'application des test de comparaison	72
4.6.2	Etude sur les PD	74
4.6.3	Etude sur les UGD	75
4.7	Etude des paramètres de WA en tant qu'une entité propre	78
4.7.1	Etude sur les PD	78
4.7.2	Etude sur les UGD	79
	Conclusion	83
	Bibliographie	87
	Glossaire	90
	Annexes	90
	A Backtesting des PD WA avec celles des EL	91
	B Etude des expositions de L'EL et de WA pour différents pays	92
	C Test de bruits blancs et d'indépendance	93
	C.1 Tests de bruits blancs	93
	C.2 Tests d'indépendances	94
	D Tests de KS et de Wilcoxon	94
	D.1 Codes R des Applications numériques	94
	D.2 Résultats des tests sur PD	98
	D.3 Arbre illustrant les résultats sur les UGD avec une segmentation plus fine	98
	E Qualité des données	98

Introduction

L'objectif de ce mémoire est de trouver la meilleure solution de calibration des paramètres pour une entité particulière du modèle interne développé par Euler Hermès (EH). EH développe un modèle interne pour le risque d'assurance-crédit et de caution que nous appelons le modèle TCI&S (Trade Credit Insurance and Surety). Les autres risques sont modélisés via le modèle interne d'Allianz. C'est au niveau groupe que sont calculés les paramètres pour obtenir une certaine homogénéité dans le processus d'obtention de ces paramètres. Il y a donc de nombreux échanges avec l'ensemble des entités pour confirmer la validité des paramètres calculés par le groupe.

EH dispose de toutes les ressources nécessaires pour calibrer les paramètres de l'ensemble des entités légales (EL) modélisées par le modèle interne (31 EL). Cependant, il existe une filiale qui ne dispose pas d'agréments d'assurance, mais joue un rôle assimilable à celui d'un courtier, il s'agit de la World Agency (WA). Créée en 2007, WA est un intermédiaire pour les clients importants présent dans chaque pays où réside une EL d'EH. Elle offre des produits simples et flexibles à des situations parfois complexes. Elle dispose du réseau mondial des experts EH tant en prévention qu'en recouvrement. Elle met en place une plate-forme de décision centralisée. Elle offre un service dans la langue et la devise du client. Cette entité négocie les contrats puis les attribue aux EL concernées. Ces EL sont des sociétés d'assurance et elles reçoivent les primes et indemnisent les assurés. La problématique est de savoir si les acheteurs souscrits par les assurés de WA ont les mêmes comportements que ceux des filiales concernées et s'il est pertinent d'utiliser les mêmes paramètres pour modéliser le comportement de ces deux types de d'acheteur avec le modèle interne. Jusqu'à présent, EH ne distinguait pas la provenance des contrats et appliquait les paramètres de l'EL locale à tous les contrats qu'ils proviennent de WA ou non.

Nous présenterons dans un premier temps les différents paramètres de ce modèle, ensuite nous expliquerons le fonctionnement du modèle interne afin de voir à quel niveau ces paramètres rentrent en jeu. Dans un second temps, nous établirons une comparaison entre WA et les EL, pour cela, (EL) nous procéderons à de nombreux tests. Par la suite, nous étudierons la sensibilité du SCR face à des chocs sur les différents paramètres d'entrée du modèle interne. Puis nous comparerons les comportements des paramètres impactant le plus le SCR, sous différentes segmentations afin de définir des ressemblances à l'aide de tests statistiques. D'abord, nous tenterons d'accorder les nouvelles données avec celles existantes pour obtenir un jeu de paramètres plus adapté. Puis nous étudierons

la construction d'une base propre aux nouvelles données. Pour finir, nous mettrons en avant la stratégie choisie ainsi que les arguments qui nous auront poussé à choisir cette stratégie.

Chapitre 1

L'assurance-crédit chez Euler Hermes

1.1 L'assurance-crédit

L'assurance-crédit protège les entreprises contre les risques d'impayés en provenance de clients domestiques ou internationaux (on parle alors d'assurance-crédit export). Elle assure le fournisseur contre le risque d'impayés des biens livrés ou des services fournis à son client. En outre, l'assureur-crédit informe le fournisseur sur la solvabilité de ses clients.



FIGURE 1.1 – Relation triangulaire en assurance-crédit

Un contrat d'assurance-crédit peut être considéré comme une relation en triangle. C'est représentable comme sur la figure 1.1 ci-dessus.

La relation contractuelle implique deux parties, l'assuré et l'assureur (EH). Cependant, il y a également une troisième partie qui intervient dans cette relation et qui joue un grand rôle, le client de l'assuré que nous appellerons l'acheteur. Si les contrats d'assurance-crédit existent, c'est parce que les compagnies (assurés) souscrivent de tels contrats pour se protéger contre le défaut de l'un de leurs clients (l'acheteur). Il est important de savoir qu'un acheteur peut faire parti de plusieurs contrats, à partir du moment où il peut être client de plusieurs assurés qui ont souscrit un contrat d'assurance-crédit concernant cet acheteur. Si un acheteur fait défaut et que l'assuré déclare le sinistre correspondant, l'assureur doit déclencher la procédure d'indemnisation et après avoir pris en compte les différents paramètres du contrat, il pourrait rembourser l'assuré. Ces paramètres permettent à l'assureur de diminuer les effets d'un défaut.

Une mesure effective du risque de crédit dans un portefeuille implique quatre indicateurs :

- La probabilité individuelle de défaut de chaque acheteur,
- La perte due au défaut de l'acheteur,
- La corrélation entre les probabilités de défaut des différents acheteurs,
- Les différents critères du contrat.

Nous expliquerons pourquoi EH a décidé de mettre en place un modèle interne, puis nous vous présenterons comment mesurer ces indicateurs, ensuite nous regarderons comment traiter les sinistres liés aux acheteurs non dénommés et ceux qui sont liés aux acheteurs non liés.

1.2 Le choix du modèle interne

Si EH a décidé de développer un modèle interne, c'est parce que la formule standard a un certain nombre de failles pour la représentation du risque de prime sur l'assurance-crédit et de cautionnement. Pour une entreprise focalisée sur les produits de crédit et de caution, ces faiblesses conduisent à avoir des doutes sur le SCR découlant de la formule standard. Ces lacunes peuvent être traitées par la modélisation du risque de prime sur le crédit et le cautionnement en utilisant un modèle de risque de crédit, ainsi EH a décidé d'opter pour un modèle interne.

Les éléments suivants sont considérés comme des défauts attachés au calcul du SCR sur le risque de prime rattaché à l'assurance-crédit et de cautionnement sous la formule standard :

- (i) La séparation du SCR prime et du SCR CAT (catastrophe)
- (ii) L'application de traités de réassurance dans la formule standard, menant à une sous-estimation ou une surestimation du SCR
- (iii) La surestimation du SCR prime en intégrant les primes acquises nettes versées aux réassureurs pour les traités non proportionnels

- (iv) La sous-estimation des caractéristiques présentes dans les contrats atténuant les risques non-linéaires
- (v) Les valeurs rétrospectives représentant l'effet des mesures d'atténuation des risques mises en œuvre dans le contrat ne permettant pas de prendre en compte les changements de caractéristique d'atténuation des risques
- (vi) Le calcul sur la base des primes n'est pas la mesure la plus adéquate du risque pris par l'assureur

Ces défauts ont été communiqués à divers organismes de réglementation, à l'EIOPA (Autorité Européenne des Assurances et des Pensions Professionnelles) et à la commission européenne, soit directement ou par l'intermédiaire de l'ICISA (Association Internationale d'Assurance-Crédit et Cautionnement). Des discussions ont eu lieu et certaines améliorations seront apportées, soit sur la formule standard, soit sur son calibrage. Cependant, tous les défauts mentionnés ne seront pas pris en compte, tout particulièrement parce que l'objectif est d'avoir une formule unique pour toutes les entreprises d'assurances. Cette volonté d'avoir une formule unique conduit à un processus de calcul non approprié pour une société dont les activités sont principalement l'assurance-crédit et le cautionnement.

La séparation du SCR prime et du SCR CAT

Afin d'avoir une approche uniforme dans toute l'industrie, la formule standard fait une distinction entre le niveau de pertes "normal extrême" (SCR primes) et le niveau de pertes "extrême extrême" (SCR CAT basé sur des scénarios d'événements sélectionnés par l'homme). Cette distinction ne représente pas assez les activités d'assurance-crédit et caution en créant une séparation artificielle ces deux catégories de pertes alors qu'il y a une continuité des scénarios de pertes.

En outre, avoir ces deux modules séparés présente une discussion supplémentaire sur la corrélation entre ces deux modules qui peut être évité en n'ayant qu'un seul module.

L'application de traités de réassurance dans la formule standard mène à une sous-estimation ou une surestimation du SCR

La structure de réassurance du groupe EH intègre des traités proportionnels (*Quota-Share* (QS)) et des traités non proportionnels (l'excès de perte (*Exces of Loss* = XoL), *Stop Loss* (SL) et Grand Perte combinée avec *Stop Loss* (LLSL)).

S'il peut être considéré que le QS est correctement appliquée dans le calcul (l'hypothèse sous-jacente est que la variation de la commission de la réassurance ou de la commission de profit n'a pas d'impact matériel sur les résultats), le principal problème réside sur l'application des traités non proportionnels.

En ce qui concerne l'application du SL, compte tenu de la formule standard actuel, le SL peut seulement être intégrée dans le risque SCR CAT pour les multiples événements de perte. L'autre approche serait de l'intégrer dans les facteurs d'ajustement pour

réassurance non-proportionnelle (NP), ceci réduisant l'écart-type. Cependant, le modèle actuellement proposé couvre les traités XoL et non les traités NP.

En forçant l'application du SL pour le calcul du SCR prime, le SCR obtenu est de l'ordre de 1% à 2% plus fort que lors de l'intégration du SL dans le SCR CAT[9].

En ce qui concerne l'application du XoL, il peut être pris en compte par la composante d'événement de perte unique du SCR CAT ou par le facteur NP attachée au risque de primes. Comme il ne peut pas être pris en compte à plusieurs reprises, EH a fait le choix de l'appliquer à l'événement de perte unique et non pas sur les facteurs NP.

Cependant, cette approche n'est pas parfaite, car elle conduit à sous-estimer l'allègement possible venant des traités XoL. Comme les scénarios de risque SCR CAT sont fondés sur trois grandes pertes, cela signifie que toutes les couches XoL avec plus de deux rétablissements ne sont pas utilisées entraînant ainsi potentiellement à une surestimation du SCR prime.

La surestimation du SCR prime en intégrant les primes acquises nettes versées aux réassureurs pour les traités non proportionnels

Les primes nettes acquises sur lesquelles se fonde le calcul de la formule standard intègrent les primes versées aux réassureurs pour l'XoL, le SL et les traités de LLSL créant ainsi un SCR prime comprenant une part sur un risque qui a été redistribué aux réassureurs. Ce fait conduit à une surestimation de la base pour le SCR prime dans un niveau de 4-6%[9].

La sous-estimation des caractéristiques présentes dans les contrats atténuant les risques non linéaires

Les caractéristiques d'atténuation des risques attachés aux contrats sont prises en compte par le calcul de la volatilité des primes. Ce dernier est fondé sur l'analyse de l'évolution du ratio de perte sur une période historique d'au moins 5 ans. Certaines des caractéristiques d'atténuation des risques intégrés dans les contrats ont été conçues pour être déclenchées en cas de niveau extrême de pertes (par exemple : le total annuel). En conséquence, plus l'exposition au défaut est élevée et plus l'impact de ces éléments non linéaires devrait être élevé. Toutefois, la formule standard actuelle ne prend pas en compte cette relation, car la volatilité des primes est calculée sur scénario de perte historique qui peut être considérée comme non extrême. Cela signifie que le SCR calculé grâce à la formule standard sous-estimera l'impact positif de la fonction d'atténuation des risques non linéaires attachés aux contrats d'assurance.

Les valeurs rétrospectives représentant l'effet des mesures d'atténuation des risques mises en œuvre dans le contrat ne permettent pas de prendre en compte les changements de caractéristique d'atténuation des risques

Comme indiqué ci-dessus, la volatilité des primes est calculée sur des données historiques. En conséquence, c'est une vue rétrospective de l'impact de l'effet des clauses

d'atténuation du risque qui sont intégrées dans les contrats.

Lorsqu'elle fait face à une crise, une société d'assurance-crédit et de cautionnement utilise les caractéristiques d'atténuation des risques comme un levier pour réduire ses pertes. Même si le changement dans les caractéristiques d'atténuation des risques n'est pas immédiat (12 mois sont nécessaires pour examiner et renouveler les contrats), la volatilité des primes calculées sur des données historiques n'intègre pas ces changements.

En conséquence, la volatilité des primes utilisée dans le calcul conduira à une surestimation du SCR vu que l'impact des mesures d'atténuation du risque sera sous-estimé.

Le calcul sur la base des primes n'est pas la mesure la plus adéquate du risque pris par l'assureur

La prime à payer pour un assuré est égale au taux de prime stipulé dans le contrat appliqué à la somme de toutes les factures qui ont été assurées.

Les primes ne sont pas une mesure adéquate du profil de risque lié au portefeuille pour plusieurs raisons :

- * Si nous prenons le cas théorique dans lequel, deux souscripteurs ont la même liste de factures assurées, mais deux taux de primes différents, l'assuré ayant le taux de prime le plus élevé générera un SCR supérieur à l'autre malgré le fait qu'ils aient le même profil de risque.
- * Si nous prenons deux portefeuilles d'assurance avec le même montant de primes, ils génèrent la même quantité de SCR prime, mais comme leurs profils de risque peuvent différer, le risque supporté par ces portefeuilles est en fait différent, ce qui ne se reflète pas dans la formule standard.

Solution

L'activité d'assurance-crédit d'EH consiste à prendre une part du risque de crédit supporté par l'assuré sur son acheteur. La modélisation par approches classiques du risque de crédit, en particulier celles utilisées dans le monde de la banque, sont pleinement applicables à la situation d'EH afin de simuler une distribution de perte représentant la perte attendue sur un horizon d'un an.

Une modélisation du risque de crédit se résume en une approche en deux étapes :

- (i) la simulation de l'exposition qui va faire défaut conduisant à définir la valeur exposée au risque (exposition au moment du défaut= EAD)
- (ii) l'application des facteurs d'atténuation présents dans les contrats et traités de réassurance ou d'autres clauses d'atténuation menant à définir la perte finale (UL= *Ultimate Loss*) portée par la compagnie d'assurances.

Ce cadre de modélisation permet d'aborder toutes les questions mentionnées auparavant :

- Production d'une distribution de perte couvrant tous les types de scénarios de perte et d'événements de perte.
- Les caractéristiques d'atténuation des risques (présents dans le contrat ou d'un traité de réassurance) peuvent être modélisées de la façon dont ils fonctionnent et non estimés.
- Les caractéristiques d'atténuation des risques attachés aux contrats représentent l'état actuel du portefeuille.
- Le modèle est basé sur l'exposition qui est la mesure clé de la prise de risque.
- Les paramètres sont définis afin de représenter le risque supporté par EH sur un horizon d'un an.

En tant que membre du groupe Allianz, EH a décidé d'utiliser les éléments développés par Allianz pour tous les autres risques que ceux pris en compte dans le portefeuille TCI&S. Pour chaque module développé par Allianz, EH a fait des analyses afin de vérifier que le modèle lui était applicable. Si des limites ont été trouvées, elles ont été mises en avant et traité avec Allianz afin de décider des solutions à adapter. Allianz a créé un service pour gérer ses relations avec EH.

1.3 Les paramètres d'entrée

1.3.1 Probabilités de défaut

Une particularité importante du secteur du risque de crédit est la prédiction mathématique qu'un défaut ait lieu.

D'abord, nous définissons l'événement de défaut. Dans le modèle TCI&S et donc dans la calibration de ses paramètres, un acheteur en défaut est un acheteur pour lequel nous avons reçu au moins une déclaration de sinistre de la part d'un de nos assurés disposant d'un montant assuré positif.

La probabilité de défaut (PD) à la date t représente la probabilité qu'au moins un assuré déclare un sinistre sur un acheteur connu, pendant la période d'observation de 12 mois débutant à la date t .

$$PD(t) = \frac{\begin{array}{c} \text{Nombre d'acheteurs dénommés avec une exposition positive en } t \\ \text{et au moins un sinistre déclaré dans les 12 mois suivant } t \end{array}}{\text{Nombre d'acheteurs dénommés avec une exposition positive en date } t} \quad (1.1)$$

Concrètement (Figure 1.2), nous prenons une photo du portefeuille au début de la période d'observation (exemple : 1^{er} février 2012), nous comptons le nombre d'acheteurs dans le portefeuille (le dénominateur). Sur les 12 mois suivants, nous observons parmi ces acheteurs, ceux qui font défaut (le numérateur). Pour qu'il y ait vraiment défaut, il faut un montant assuré strictement positif.

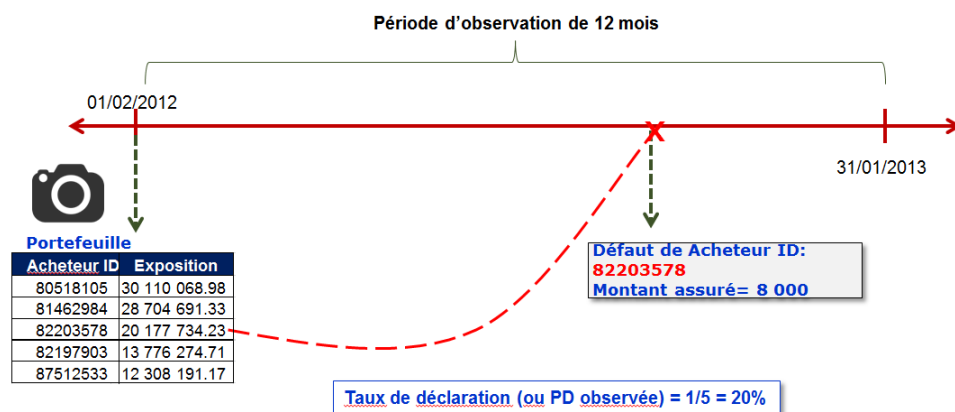


FIGURE 1.2 – Exemple simplifié de calcul de PD

La date du défaut dépend du type de contrat, il y a deux cas possibles :

- Soit c'est la date de facturation. Dans ce cas, nous parlons d'attachement au risque et c'est le cas le plus fréquent.
- Soit c'est la date de déclaration de sinistre. Dans ce cas, nous nous rattachons à la date d'apparition de la perte. Ce cas assez exceptionnel est tout de même utilisé pour certaines EL comme l'Allemagne.

Un point important est la segmentation, en effet suite à un long travail, EH a choisi de calculer ses PD en utilisant la segmentation suivante :

- Entité. Le modèle interne comprend 31 EL correspondant à des pays (France, Allemagne, ...) ou à des regroupements de pays (APAC pour l'Asie pacifique). Il y a également des EL virtuelles pour traiter le *bonding* (produit de caution), il y a par exemple UK et UK *bonding* pour le Royaume-Uni. Dans ce mémoire, nous travaillons sur le traitement d'une entité qui n'est pas explicitement présente dans cette liste.
- Type de pays de l'acheteur (domestique ou export). Chaque contrat d'assurance est rattaché à l'EL de l'assuré. Si l'acheteur qui fait défaut réside également sur le territoire couvert par cette même EL, alors la transaction ayant conduit à ce défaut est qualifiée de domestique et nous qualifions également ce défaut de domestique. Dans le cas contraire, nous qualifierons ce défaut d'export.
- Type de défaut. Il y a deux différents types d'événements de défaut :

- Le défaut que nous qualifions d'"*insolvency*" (insolvabilité). Cet événement de défaut est l'insolvabilité légale de l'acheteur. Cela se traduit par la mise en place de procédures collectives à son encontre, c'est-à-dire des procédures de redressement et de liquidation judiciaire de l'entreprise, prononcées par un tribunal, lorsque l'acheteur se trouve en état de cessation des paiements.
 - Le défaut que nous qualifions de "*protracted*". L'événement de défaut est le plus souvent un retard de paiement ou un litige. C'est lorsque l'acheteur faisant défaut dispose d'un état de solvabilité suffisant pour faire face à ses engagements, notamment de payer.
- Note (*grade*). EH est la première compagnie d'assurance à avoir obtenu un agrément européen pour sa société de notation. EH met en place un projet d'harmonisation du système automatique de notation. Jusqu'à présent de nombreux processus automatiques pour définir les notes ont été définis par les EL à partir d'un processus de notation manuel. Le but de ce nouveau projet est de remplacer ce processus automatique par une notation centrale dirigée par le groupe et paramétrée au niveau local. La note de l'acheteur est la note interne à EH évaluée sur une échelle de 1 à 10[10], liée à cet acheteur, cela reflète sa situation financière actuelle et sa capacité à régler ses dettes.
- Taille d'exposition. Des classes d'expositions sont établies afin de catégoriser les assurés selon leurs expositions. Ainsi, les classes d'expositions regroupent les acheteurs ayant les mêmes niveaux d'expositions. Elles sont différentes selon les EL et calibrées en fonction de la composition du portefeuille en terme d'exposition.

Le modèle interne d'EH prend en input une estimation des PD sur les 12 mois qui viennent, elles sont établies chaque trimestre. EH utilise différentes méthodes pour les obtenir :

- Les PD sont estimées en utilisant des modèles de PD ou des courbes de développement spécifique pour chaque EL : Les modèles de PD sont actuellement utilisés en Allemagne, Royaume-Uni, Belgique, France, Pays Nordiques, Italie, Amérique du Nord et Pays-Bas. Ces modèles sont des régressions basées sur des données internes et économiques.
- Les EL ne disposant pas d'assez de données historiques pour construire un modèle statistique, utilisent la méthode de Chain-Ladder. Les courbes de développement de sinistres sont aussi utilisées pour faire un niveau de référence aux sorties du modèle de PD.
- Les experts locaux peuvent ajuster les PD selon leurs attentes, par exemple : afin de refléter l'impact des décisions de gestion ou les changements récents de la sinistralité qui ne sont pas capturés par les statistiques. Dans tous les cas, ce sont eux qui valident les PD choisies.

Une fois que l'on connaît la PD, il faut savoir estimer le montant perdu en cas de défaut, pour cela nous nous intéressons désormais à la perte réelle de l'assureur en cas de défaut.

1.3.2 Perte réelle de l'assuré en cas de défaut

Si un acheteur fait défaut alors la totalité des encours de l'assuré ne sera pas à la charge de l'assureur (sauf cas exceptionnels). Chaque contrat d'assurance comporte des clauses permettant de maîtriser le risque porté par l'assureur. Nous allons présenter quelques mesures préventives que l'assureur peut prendre lorsque cela est stipulé dans le contrat et d'autres mesures envisageables après que le contrat soit signé. Pour diminuer le montant à indemniser au client, l'assureur dispose de plusieurs manières de réduire ses pertes :

- Tout d'abord, le contrat peut inclure des clauses de part non assurée, franchises, responsabilité maximale ou total annuel. Alors que les deux premiers termes peuvent sembler familiers, les notions de responsabilité maximale et total annuel doivent être définies :

La responsabilité maximale correspond au montant total maximal d'indemnisation que l'assureur paiera dans un horizon de temps donné.

Le total annuel est une franchise qui s'applique à la somme des sinistres ou indemnisations d'un assuré sur une certaine durée. Si la somme est inférieure au montant de franchise alors l'assureur n'indemnise pas, sinon, c'est le montant au-dessus de cette franchise qui est indemnisé.

Ces restrictions s'assimilent à des produits de réassurance : *Quota-Share* (QS), franchises, *Annual Agregate Deductible* (AAD) et *Annual Agregate Limit* (AAL) appliqués à l'assuré.

- L'assureur peut choisir de partager le risque avec d'autres assureurs, c'est-à-dire se réassurer, ce qui est une technique bien plus classique.
- Les recouvrements provenant des récupérations avant et après indemnisation sont également un moyen de réduire les pertes dues aux différents défauts.

UGD

L'un des paramètres les plus importants du contrat est sa limite. La limite est définie comme le montant maximal couvert par EH à un assuré pour un acheteur, c'est la couverture maximale. Les limites de crédit peuvent être réduites ou annulées à tout moment selon la solvabilité de l'acheteur, ceci en respectant les clauses contractuelles comprenant le plus souvent un délai allant de 30 à 90 jours pour que la modification soit effective. Cela permet à l'assuré soit de mettre fin aux transactions avec cet acheteur, soit de trouver une autre assurance sur le marché. Dans tous les cas, c'est la limite au moment du défaut (facturation) qui vaut et non la limite à la date de déclaration du

défaut. C'est pourquoi EH peut modifier les limites en prévenant son client, car cela ne modifiera en rien l'assurance des transactions en cours.

La limite des contrats intervient dans la définition de l'*Exposure At Default* (EAD). L'EAD est définie comme le montant d'en cours au moment du défaut, qui représente la quantité utilisée de la limite du défaut. En général pendant la période de maturité d'une année, le montant des encours est inférieur à la limite accordée elle-même. Cependant, nous pouvons observer qu'en général, la limite accordée diminue plusieurs fois par an en fonction de la détérioration de la solvabilité des acheteurs dont les performances financières se dégradent. L'EAD peut être facilement estimée en faisant des hypothèses sur la proportion de la limite utilisée en cas de défaut en fonction de la limite accordée. Cette proportion est appelée *Usage Given Default* (UGD). De manière formelle nous avons :

$$\text{EAD} = \text{limite accordée} \times \text{UGD} \quad \text{avec} \quad \text{UGD} = \frac{\text{limite utilisée au moment du défaut}}{\text{limite accordée au 1}^{\text{er}} \text{ janvier}} \quad (1.2)$$

Nous devons préciser que la proportion de réduction de limite montre l'efficacité de la gestion des contrats. Dans le meilleur des cas, la limite accordée est annulée, c'est-à-dire que le contrat est annulé, avant le défaut de l'acheteur. Dans ce cas, l'assureur est libéré de toute obligation légale et le défaut de l'acheteur n'aura aucune conséquence. Le plus souvent, quand l'assureur abaisse la limite accordée, l'assuré va rationnellement diminuer ses activités avec cet acheteur. Il y a deux raisons qui mènent à ça. La première raison est que si l'assureur abaisse la limite, l'assuré va prendre cette mesure comme un message de prévention disant que l'acheteur est moins capable de rembourser ses dettes, ainsi les chances qu'il ne rembourse pas les montants facturés augmentent. La seconde raison est qu'à partir du moment où la limite accordée sur un acheteur diminue, si l'assuré ne réduit pas son exposition sur cet acheteur, le montant qu'il perdra en cas de défaut de cet acheteur sera plus fort, car la prise en charge par EH sera moindre.

Une fois l'échantillon des UGD établi, nous prenons sa moyenne et son écart-type en tant que paramètre d'entrée du modèle interne. Ils serviront à paramétrer un UGD stochastique basé sur une loi gamma dont les paramètres sont calculés à partir des moments d'ordre 1 et 2 (moyenne et variance(=écart-type²)). Plus de détails seront données au point 1.4.3.

LGD

Parmi les trois dernières mesures de prévention, nous avons deux types de recouvrement qui sont agrégées en un seul paramètre appelé *Loss Given Default* (LGD). Le recouvrement avant indemnisation est appelé *recollection rate* (taux de recouvrement RR). La récupération après indemnisation est appelée *salvages rate* (le taux de "sauvetage" SR)

$$RR = \frac{\sum \text{montant récupéré avant indemnisation}}{\sum \text{Montant assuré}} \quad (1.3)$$

$$SR = \frac{\sum \text{Montant récupéré après indemnisation}}{\sum \text{Montant indemnisé}} \quad (1.4)$$

$$LGD = (1 - RR)(1 - SR) \quad (1.5)$$

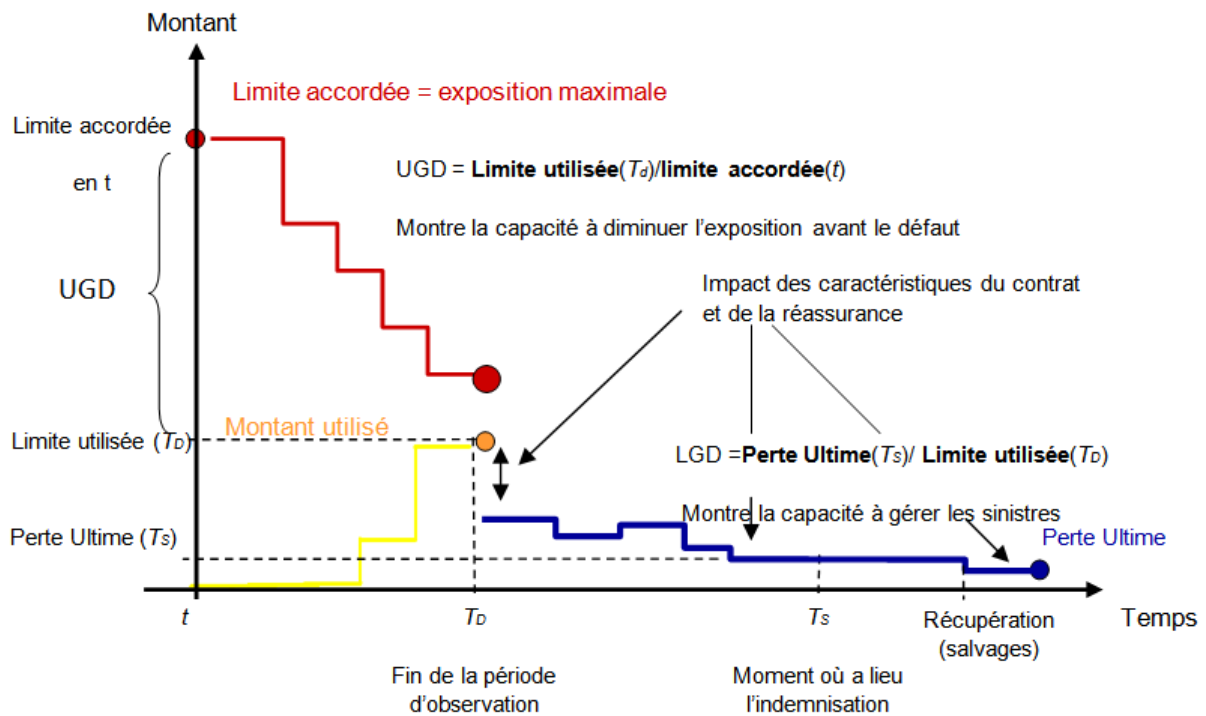


FIGURE 1.3 – Gestion des sinistres

L'image ci-avant résume comment un contrat est typiquement géré et plusieurs des notions définies ci-dessus. On observe en rouge l'évolution de la limite accordée (le sinistre maximal), celle-ci diminue avant le défaut lorsque l'assureur dégrade l'acheteur. En jaune, nous avons le montant utilisé qui augmente car l'assuré développe son activité avec l'acheteur. Au moment du défaut, nous avons un UGD qui correspond au montant utilisé au défaut sur la limite accordée en début de période. Puis, les montants récupérés avant le paiement nous permettent d'obtenir le RR. Et après l'indemnisation des sinistres nous obtenons le second taux de recouvrement (SR). Au final, nous observons que la perte ultime est bien inférieure à la limite accordée à l'assuré.

Avec les paramètres que nous venons de présenter, nous obtenons la perte attendue (*Expected Loss*). Ceci est synthétisé dans le schéma 1.4.

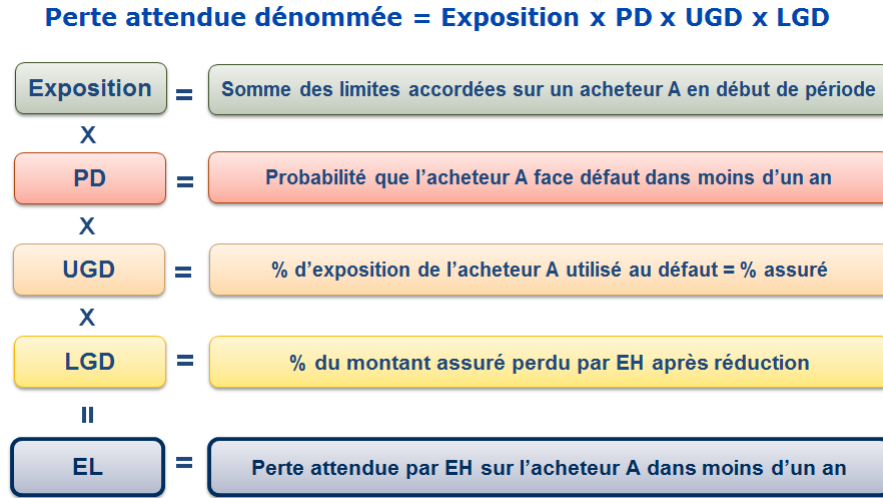


FIGURE 1.4 – Relation entre les paramètres présentés

Dans ce qui précède, nous avons présenté comment les contrats d'assurance sont managés, et nous avons donné une idée du montant des pertes auquel l'assureur doit faire face en cas de défaut d'un acheteur spécifique. Nous avons également vu au préalable le processus d'obtentions des PD. Il reste d'autres facteurs très importants à prendre en compte dont la corrélation entre les PD.

1.3.3 Corrélations

Il est important de connaître les PD individuelles, car elles sont utilisées dans le calcul du capital requis pour un acheteur individuel. Cependant, dans le cas de portefeuille de crédit, connaître le risque lié à chacun des acheteurs n'est pas suffisant pour évaluer le risque total. En effet, en faisant cela, nous ne prendrions pas en compte le fait que l'événement de défaut d'une compagnie est probablement lié au défaut d'une autre compagnie dans le même pays et dans le même secteur (décomposition utilisée détaillé en 1.4.2). Etudier et intégrer les corrélations entre les événements de défaut dans le modèle de risque se révèle être crucial afin d'évaluer correctement le risque supporté par l'assureur crédit. Un portefeuille composé d'entreprises avec un faible risque de crédit, mais dont les valeurs d'actifs sont très corrélées avec une autre entreprise, est très risqué.

Exemple simplifié : Soit 1 et 2 deux compagnies dans le même secteur. Soit D_i la variable de Bernoulli qui indique si l'acheteur $i=1,2$ fait défaut. Soit ρ le coefficient de corrélation annuel entre les deux compagnies. Soit p_i tel que $p_i = P(D_i = 1)$ for $i = 1,2$ et P_{12} tel que $p_{12} = P(D_1 = 1, D_2 = 1)$. Le coefficient de corrélation entre les événements de défaut des deux compagnies est donné par la formule suivante :

$$\rho = \frac{Cov(D_1, D_2)}{\sqrt{V(D_1)}\sqrt{V(D_2)}} = \frac{\mathbb{E}[D_1 D_2] - \mathbb{E}[D_1]\mathbb{E}[D_2]}{\sqrt{V(D_1)}\sqrt{V(D_2)}} = \frac{p_{12} - p_1 p_2}{\sqrt{p_1(1-p_1)}\sqrt{p_2(1-p_2)}}$$

Supposons par exemple que les PD p_1 et p_2 pour l'année N sont toutes les deux égales à 2% et que le coefficient de corrélation entre ces défauts soit de 0.1. Alors la probabilité que les deux compagnies fassent défaut est $p_{12}=0.24\%$. Soit deux scénarios pour l'année N+1 :

- (a) les probabilités de défauts individuels p_1 et p_2 augmentent et deviennent égales à 0.4 et que la corrélation mesuré par ρ reste la même.
- (b) les probabilités de défauts individuels p_1 et p_2 augmentent et deviennent égales à 0.4 et que le coefficient de corrélation augmente et devienne égal à $\rho = 0.2$. Ce scénario a plus de chances d'arriver que le premier, car en général si la qualité de défaut se détériore, la corrélation entre les défauts a tendance à augmenter.

Dans le premier scénario (a) la probabilité jointe de défaut arrive à $p_{12} = 0.544\%$ et dans le scénario (b) $p_{12}=0.928\%$. Ainsi, la non prise en compte du coefficient de corrélation entre les défauts sous-estime la probabilité jointe de défaut de près de la moitié, ce qui est considérable.

Les coefficients de corrélations sont très complexes à calculer. Un moyen d'observer ces corrélations est de connaître les probabilités individuelles de défauts et les corrélations des actifs. Cela reste très difficile car une compagnie est considéré en défaut à partir du moment où la valeur de ses actifs tombe en dessous d'un certain seuil, donc deux compagnies font défaut simultanément si la valeur de leurs actifs respectifs tombe simultanément en dessous d'un certain seuil. La simultanéité de ces deux baisses et son niveau définissent la corrélation des actifs des deux compagnies. Un moyen de traiter ce problème est de poser l'hypothèse que la distribution jointe de défaut des deux compagnies est une copule gaussienne dont les coefficients sont les corrélations entre les valeurs des actifs de ces compagnies. Ce développement fait déjà l'objet d'un mémoire[5].

Nous nous attardons désormais sur deux types de sinistres traités différemment et qui représentent part importante des sinistres.

1.3.4 Acheteurs non dénommés

Dans un contrat d'EH, il peut être défini une limite discrétionnaire. C'est une couverture pour l'assuré sur les factures impayées d'acheteurs pour lesquels aucune limite spécifique n'a été définie. Ces acheteurs sont qualifiés de non dénommés. Les pertes sur ces acheteurs ne sont pas simulées. La différence se trouve au niveau de la préparation des données car elles sont par définition inconnues d'EH.

Par définition, il n'est pas possible de connaître la situation des acheteurs attachés à ces contrats. Vu que les pertes qui proviennent d'un acheteur non dénommé peuvent représenter une part significative de la perte ultime, EH a conçu une solution afin d'intégrer les acheteurs non dénommés dans le calcul de la mesure de risque.

Les éléments d'entrée suivants sont calibrés pour la préparation des données :

- Le nombre d'acheteurs non dénommés dont le défaut est attendu.
- La note attachée à l'exposition non dénommée.
- La classe d'exposition attachée à l'exposition non dénommée.
- La perte ultime attendue avant réassurance, attaché aux acheteurs dénommés.

En se fondant sur ces informations ainsi que sur les paramètres d'entrée pour les acheteurs dénommés, une limite artificielle agrégeant les expositions des acheteurs non dénommés est créée pour chaque EL. A haut niveau, la simulation du modèle joint les estimations de pertes pour acheteurs non dénommés cf. figure 1.6.

1.3.5 Acheteurs non liés

Nous avons défini comment obtenir les PD (cf. 1.2). En effet, celles-ci sont utilisées pour simuler les événements de défauts dans une période de 12 mois pour les limites connues des acheteurs dénommés à un moment fixé dans le temps. Ainsi, les acheteurs sans exposition au début de la période d'observation de douze mois, mais obtenant une limite pendant cette période d'observation, ne sont pas pris en compte dans cette évaluation des PD.

Il faut noter qu'à partir du moment où il y a une exposition positive, et malgré le fait que l'exposition est nulle au début de la période d'observation, un risque est porté par EH sur ces acheteurs.

Ainsi, en plus de la modélisation pour les limites connues à un instant et en plus de la prise en compte des acheteurs non dénommés dans l'estimation des pertes, quelques modélisations supplémentaires doivent être introduites pour refléter les pertes possibles sur des acheteurs dénommés absents du portefeuille au début de la période d'observation, mais dont l'exposition devient positive pendant la période d'observation. Cela est appelé par EH comme la modélisation des acheteurs non liés.

Dans l'exemple suivant (figure 1.5), supposons que nous sommes le 01/01/2015, la limite accordée à l'acheteur X dans les douze mois suivants va générer une exposition et un risque qui ne sera pas pris en compte par la modélisation normale des acheteurs dénommés, à partir du moment où l'exposition au 01/01/2015 est nulle.

Dans le calcul des paramètres, nous prenons souvent des données sur un an. Pour cela, nous prenons les expositions en début de période et nous suivons leur comportement sur l'année suivante. Cependant, les expositions changent, en effet des assurés peuvent demander à rajouter de nouveaux acheteurs. De plus ces nouveaux acheteurs peuvent faire défaut avant la fin de la période d'observation.

Des recherches ont été réalisées sur les données de sinistralité historique d'EH. Pour plusieurs EL, la perte découlant des acheteurs non liés représente une part significative de la perte ultime. Cela prouve l'importance de modéliser les acheteurs non liés pour le calcul de la mesure de risque.

Etant donné que nous ne connaissons pas ces acheteurs au début de la période d'observation, il a été choisi d'appliquer une modélisation similaire à celle utilisée pour les

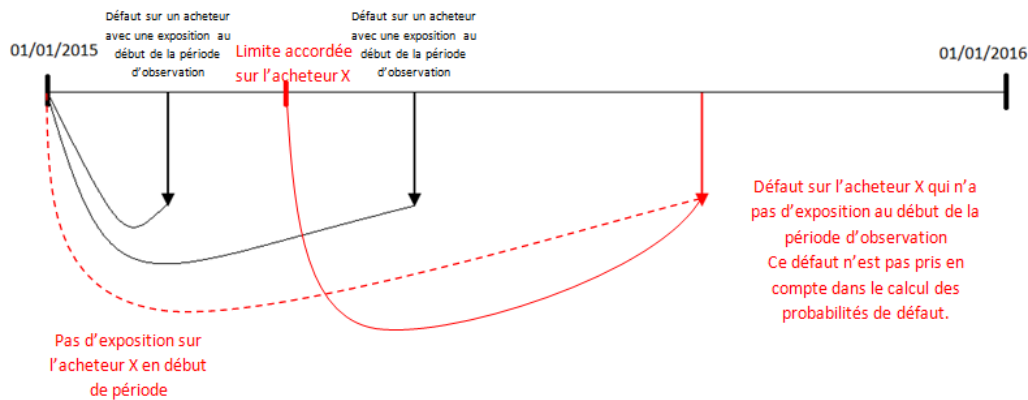


FIGURE 1.5 – Schématisation d'un sinistre sur un acheteur non lié

acheteurs non dénommés. Il faut donc pour chaque EL, calibrer les éléments suivants, pertinents pour les 12 mois suivants :

- Le nombre d'acheteurs non liés faisant défaut attendu.
- La note attachée à l'exposition des acheteurs non liés.
- La classe d'exposition attribuée aux acheteurs non liés.
- La perte ultime attendue avant la réassurance, liée aux acheteurs dénommés.

En se fondant sur ces informations ainsi que sur les paramètres d'entrée pour les acheteurs dénommés, une limite artificielle agrégeant les expositions des acheteurs non liés est créée pour chaque EL. A haut niveau, la simulation du modèle joint les estimations de pertes pour acheteur non liés (cf. figure 1.6).

Une perte attendue est calculée afin d'avoir une valeur de référence de ce que le modèle TCI&S devrait donner.

$$\text{Perte attendue} = \text{Exposition courante} * \text{PD} * \text{UGD} * \text{LGD} + \text{Pertes non dénommées} + \text{pertes non liées} \quad (1.6)$$

Cette estimation doit refléter les changements et évolutions des paramètres.

1.4 Description du fonctionnement du modèle interne

Dans cette partie nous présentons les principaux axes du modèle d'EH. Cette partie prend pour source la documentation du modèle interne d'EH[9]. Dans un premier temps nous expliquerons l'établissement de la situation de défaut par une approche de type

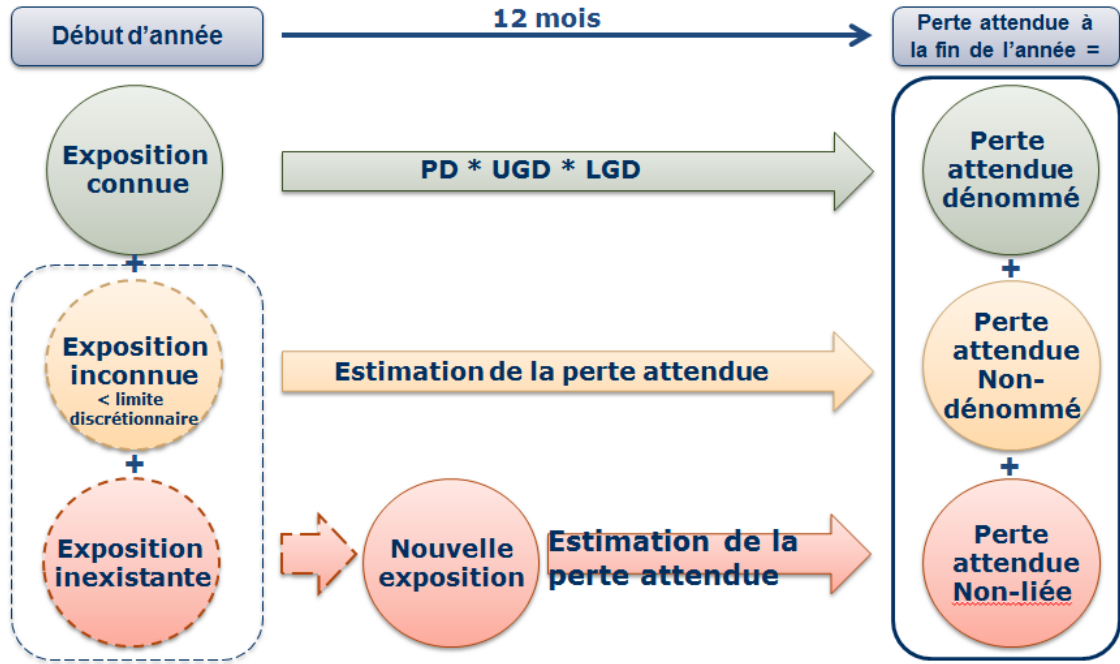


FIGURE 1.6 – Estimation de la perte attendue

Merton, puis nous donnerons quelques précisions sur la détermination des corrélations, ensuite nous présenterons la modélisation des UGD stochastiques et pour finir ce chapitre, nous expliquerons l'application des autres paramètres permettant de réduire les pertes.

1.4.1 Modélisation du facteur de crédit par une approche de type Merton

Pour le calcul de son risque d'assurance-crédit et de caution, EH applique une méthodologie de modèle bien connu de facteur de crédit[3] [15]. L'hypothèse principale est que la valeur de l'actif de chacune des sociétés de l'acheteur dans le portefeuille suit un mouvement brownien géométrique et qu'un défaut arrive dans le cas où la valeur de l'actif serait inférieure à un certain seuil (en théorie reflétant les passifs). Dans ce cadre, la simulation défaut / non défaut se résume à simuler une loi normale reflétant le rendement logarithmique de la valeur de l'actif de l'entreprise sur un an et la compare à l'hypothèse d'un mouvement brownien géométrique.

Les hypothèses de base du modèle de portefeuille choisi sont :

- Un acheteur fait défaut dans l'intervalle de temps $[0, T]$ si la valeur de ses actifs est inférieure à un certain seuil.
- La valeur de ses actifs A suit un mouvement brownien géométrique :

$$\frac{dA_t}{A_t} = \mu dt + \sigma dW_t$$

où W_t est un processus de mouvement brownien. La dérive μ et la volatilité σ sont des constantes spécifiques à l'entreprise.

Appliquant le lemme d'Itô, cela donne ce qui suit pour le $\log(A_t)$:

$$d[\log(A_t)] = \sigma A_t \frac{1}{A_t} dW_t + \left[\mu A_t \frac{1}{A_t} + \frac{\sigma^2 A_t^2}{2} \left(-\frac{1}{A_t^2} \right) \right] dt = \left[\mu - \frac{\sigma^2}{2} \right] dt + \sigma dW_t$$

En fait, cela signifie que la valeur de l'actif à l'instant T , c'est à dire A_T , satisfait

$$\log(A_T) - \log(A_0) = \left(\mu - \frac{\sigma^2}{2} \right) T + \sigma \sqrt{T} ATP$$

où ATP ("capacité à payer") est une variable distribuée suivant une loi normale centrée réduite. Autrement dit, les rendements logarithmiques de la valeur de l'actif sont normalement distribués.

En pratique, soit PD_j la PD attribuée à un acheteur particulier j sur une période T , où $T = 1$ an dans le cas de la modélisation EH. Nous générons une simulation ATP_j issue d'une loi normale centrée réduite et nous la comparons au seuil

$$c_j = N^{-1}(PD_j),$$

où N^{-1} est l'inverse de la loi normale centrée réduite. Un défaut de la contrepartie au sein de $[0, T]$ est déclenché si la simulation est inférieure à ce seuil. Faire cela une fois pour chaque acheteur conduit à un scénario de perte plausible pour le portefeuille. Plusieurs scénarios de portefeuille de ce genre donnent à la fin une simulation de la distribution de la perte du portefeuille. Ceci est connu comme une méthode de Monte Carlo. Pour approximer N^{-1} EH utilise l'algorithme de Moro[17].

En général, nous observons une corrélation entre les entreprises d'un portefeuille. Ainsi, simuler des valeurs indépendantes ATP pour les différentes entreprises, conduirait à une sous-estimation de la queue de la distribution des pertes. La corrélation des défauts peut se référer à :

- L'influence de l'économie : Pendant les périodes de stress ou les crises économiques, les défauts simultanés se produisent plus souvent. Ce genre de mouvement de cycle peut être particulier pour un secteur industriel et par pays.
- La relation entre les entreprises : Un défaut de certaines entreprises peut augmenter la PD des autres entreprises connexes.

Une approche commune, remontant à Merton et Vasicek [20], modélise l' ATP comme la somme de deux composantes :

- La partie systémique, qui est commune pour un certain nombre d'entreprises.
- La partie idiosyncratique ou spécifique, qui est indépendante d'une entreprise à l'autre.

Suivant cette approche, EH implémente la formule suivante pour modéliser l'ATP :

$$ATP_j = \sqrt{RSQ_j} \frac{\vec{\omega}_j^T \vec{\phi}^T}{\sqrt{\vec{\omega}_j^T \Sigma \vec{\omega}_j}} + \sqrt{1 - RSQ_j} \varepsilon_j, \quad j = 1, \dots, N \quad (1.7)$$

La première partie de l'équation est le risque systématique (pays et l'industrie) et la deuxième partie est le risque spécifique par ATP. Plus en détail,

- N est le nombre d'acheteurs (dans le portefeuille évalué).
- ε_j est la partie idiosyncratique pour l'acheteur j . Ce sont des réalisations d'une variable aléatoire normale centrée réduite indépendante et identiquement distribuée, indépendante des facteurs systémiques.
- $\vec{\phi}^T = (\phi_1, \dots, \phi_s)$ est un vecteur de facteurs systémiques. Ceux-ci suivent une distribution normale multivariée $N(\vec{0}, \Sigma)$ où Σ est la matrice de covariance.
- $\vec{\omega}_j^T = (\omega_{j,1}, \dots, \omega_{j,s})$ est un vecteur de constantes, soit paramètres déterministes. Ils sont appelés «poids» parce qu'ils montrent l'intensité de la dépendance entre le facteur systémique de l'ATP de l'acheteur j et les autres facteurs.
- RSQ_j est un paramètre déterministe. Il représente la corrélation entre l'ATP de l'acheteur j et son facteur systémique particulier (étant une combinaison linéaire de l'ensemble général des facteurs systémiques).

Remarquons que l'on a alors $\frac{\vec{\omega}_j^T \vec{\phi}^T}{\sqrt{\vec{\omega}_j^T \Sigma \vec{\omega}_j}} \sim N(0, 1)$ et du coup $ATP_j \sim N(0, 1)$ ceci étant l'hypothèse stochastique pour l'ATP mentionné précédemment.

EH a intégré dans sa modélisation les deux principaux événements déclencheurs de la couverture d'un contrat : défaut *protracted* et défaut *insolvency*.

Il faut savoir que pour l'assurance-crédit, le calibrage des UGD et des paramètres d'atténuation des pertes dépend de façon significative du fait que l'entreprise défaillante est dans un état ou non d'insolvabilité. Un cas *insolvency* est en général pire qu'un cas *protracted*. Par conséquent, dans les scénarios simulés, nous devons savoir si un sinistre simulé se réfère à un acheteur insolvable ou non.

Ainsi, Euler Hermes a étendu l'approche du modèle de Merton afin d'intégrer un second seuil de défaut comme décrit dans le graphique 1.7. Comme le cas *insolvency* est pire que le cas *protracted*, le seuil du statut de défaut *insolvency* est attaché au déclenchement de *insolvency* afin de pour effectuer des simulations.

On a pour cela :

1. $PD_j = PD_{j,insolvency} + PD_{j,protracted}$
2. $c_{j,insolvency} = N^{-1}(PD_{j,insolvency})$ et $c_j = N^{-1}(PD_j)$
3. Si ATP_j simulé $< c_{j,insolvency}$ nous avons un événement de défaut avec statut *insolvency* sur l'acheteur j .

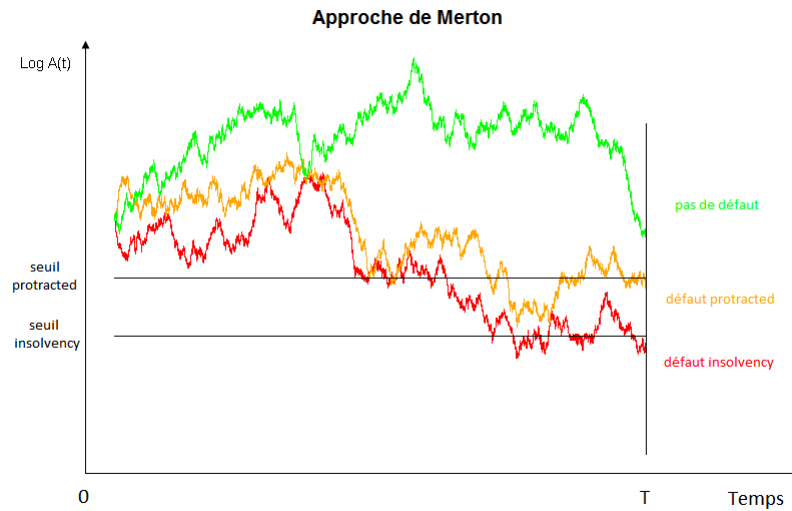


FIGURE 1.7 – Illustration de l'approche de Merton

4. Si $c_{j,insolvency} < ATP_j \text{ simulé} < c_j$ nous avons alors un événement de défaut avec statut *protracted* sur l'acheteur j .

Comme mentionné précédemment, nous calculons les paramètres pour chaque EL. Cela signifie que, s'il existe une limite pour un acheteur particulier dans différentes EL, alors plusieurs PD sont affectées à cet acheteur. De cette manière, EH applique une simulation de défaut pour chaque combinaison "Acheteur EL". Cela est également considéré au niveau de l'ATP et ainsi un ATP sera simulé pour chaque acheteur par EL.

Intuitivement, cela pourrait être incompatible, car il serait étrange qu'un acheteur soit simulé comme «insolvable» dans une seule EL et non dans les autres. C'est pourquoi, quelques règles supplémentaires ont été implémentées. Dans le cas d'un portefeuille agrégé (plusieurs EL), certaines règles d'agrégation sont nécessaires pour avoir une approche de simulation cohérente. En particulier, parce que la calibration des PD est faite par EL et en conséquence, la simulation de l'ATP est ensuite effectuée à ce niveau, nous pourrions nous retrouver avec des scénarios où un acheteur particulier est simulé comme insolvable dans une seule EL et non dans d'autres.

EH a mis en œuvre deux règles pour éviter ces situations :

- (i) L'homogénéisation du paramètre ATP. Supposons qu'un acheteur ait des expositions dans différentes EL. Alors, une de ces EL est choisie au hasard et l'ATP simulé de cet acheteur dans cette EL est attribué à toutes ses expositions.
- (ii) Homogénéisation de la PD de l'insolvabilité. Supposons qu'un acheteur a des expositions dans différentes EL. Alors, une de ces EL est choisie au hasard et la PD *insolvency* attribuée à l'acheteur est attribuée à toutes ses expositions. Si la première règle ci-dessus est activée, l'EL choisie au hasard est la même pour les deux

règles.

Le tableau suivant illustre ces deux règles :

Acheteur	EL	ATP	PD Insolvency	PD Protracted
A	EL1	ATP1	PD1	PD1
A	EL2	ATP2	PD2	PD2
A	EL3	ATP3	PD3	PD3
A	EL4	ATP4	PD4	PD4

EL choisié aléatoirement règle i règle ii

FIGURE 1.8 – Exemple règles d'agrégation

1.4.2 RSQ et matrice de variance covariance

Cette section traite la corrélation présentée dans la formule du modèle de facteur de crédit. Nous allons expliquer les paramètres de corrélations.

Dans le modèle de facteur de crédit, un ensemble de facteurs systémiques $\vec{\phi}^T = (\phi_1, \dots, \phi_s)$ suivant une loi normale est défini. La matrice Σ est la matrice de covariance de ces facteurs. En particulier l'élément (i,j) est égal à :

$$\Sigma_{i,j} = COV(\phi_i, \phi_j) \quad (1.8)$$

Il faut noter que la partie idiosyncratique de l'ATP d'un certain acheteur (dans une EL) est indépendant de tous les autres. De cette manière, nous pouvons faire différentes interprétations :

- $\sqrt{RSQ}$ représente la corrélation entre l'ATP d'un acheteur dans l'EL et sa part systémique spécifique (normalisé). C'est-à-dire sa dépendance avec l'effet systémique.

$$\sqrt{RSQ_j} = CORR(ATP_j, S_j) \quad (1.9)$$

Où $S_j = \frac{\vec{\omega}_j^T \vec{\phi}}{\sqrt{\vec{\omega}_j^T \Sigma \vec{\omega}_j}}$ est la part systémique normalisée et $\vec{\omega}_j^T = (\omega_{j,1}, \dots, \omega_{j,s})$ est le vecteur avec les poids déterministes indiquant dans quelles proportions les facteurs systémiques contribuent à cette partie systémique de l'acheteur j dans l'EL donnée.

- Supposons que l'acheteur i et j aient la même partie systémique et le même RSQ, c'est-à-dire $S_i = S_j$ et $RSQ_i = RSQ_j$. Ainsi, RSQ représente la corrélation entre leurs ATP.

$$RSQ_i = RSQ_j = CORR(ATP_i, ATP_j) \quad (1.10)$$

La corrélation des actifs mesure le degré auquel les valeurs de deux entreprises ont tendance à bouger en même temps. Plus les facteurs économiques en commun de deux en-

treprises sont liés et plus il est probable que leur fortune montent et descendent ensemble. La forme mathématique est la suivante :

$$CORR(ATP_i, ATP_j) = \sqrt{RSQ_i} \sqrt{RSQ_j} CORR(S_i, S_j) = \frac{\sqrt{RSQ_i} \sqrt{RSQ_j} \vec{\omega}_i^T \Sigma \vec{\omega}_j}{\sqrt{\vec{\omega}_i^T \Sigma \vec{\omega}_i} \sqrt{\vec{\omega}_j^T \Sigma \vec{\omega}_j}} \quad (1.11)$$

La corrélation entre les parts systémique S_i et S_j est une fonction de leurs poids sur les facteurs systémiques définis et la matrice Σ .

Afin d'estimer les corrélations entre facteurs systémiques (et RSQ), EH repose sur les séries chronologiques des PD observées. Celles-ci peuvent uniquement être observées sur un ensemble d'acheteurs, que nous appellerons échantillon. Par conséquent, il est choisi de diviser le portefeuille en différents échantillons et de joindre à chaque échantillon un facteur systémique correspondant. Dans ce contexte la formule 1.7 devient :

$$ATP_j = \sqrt{RSQ_j} \frac{\phi_{p(j)}}{\sqrt{\Sigma_{p(j),p(j)}}} + \sqrt{1 - RSQ_j} \varepsilon_j, \quad j = 1, \dots, N \quad (1.12)$$

où $p(j)$ réfère à l'échantillon auquel l'acheteur j appartient. En d'autres termes, les composants sont tous nuls sauf pour la place représentant le segment de l'acheteur qui sera égal à 1.

Le niveau sur lequel les facteurs systémiques et leurs corrélations sont définies est généralement établi en suivant à une grille de région et d'industrie. D'un point théorique plus de segmentation conduit à une modélisation plus précise dans le sens de la diversification. Cependant, dans la pratique, une segmentation trop fine conduit à un plus petit nombre d'acheteurs dans les échantillons et par conséquent une incertitude statistique et un risque plus élevé sur le bruit idiosyncratique des PD observés. Il est alors difficile de savoir si les tendances systémiques observées, et en particulier la corrélation entre les échantillons, forment un bon point de référence pour l'avenir. Pour résumer, il faut être prudent et ne pas segmenter trop, sinon cela conduit à un sur-apprentissage sur les données historiques (c'est-à-dire à une non robustesse du modèle).

Sur la base des critères de segmentation disponibles et en tenant compte du nombre d'acheteurs, EH a choisi d'appliquer les critères suivants pour définir les échantillons :

- EL.
- Type de transaction.
- 4 classes sectorielles :
 - o Produits finis.
 - o Produits bruts.
 - o Finance / Immobilier.
 - o Services / Commerce.

1.4.3 Simulation stochastique des UGD

Dans cette partie nous allons montrer comment les paramètres calculés pour les UGD sont utilisés dans le modèle. Supposons un segment particulier d'acheteurs ayant des caractéristiques similaires. Même au sein de ce segment, différents UGD seront observés dans un scénario réaliste. Ne pas tenir compte de cette variation induit une imprécision de la distribution des pertes et surtout dans la queue de distribution.

Pour cela, dans chaque EL, pour chaque acheteur faisant défaut, l'UGD qui doit être appliqué dans un scénario particulier est échantillonné aléatoirement par une distribution de probabilité. Certaines analyses ont été effectuées sur les données observées historiquement et il a été conclu que la loi Gamma à 2 paramètres, conduit à un bon ajustement. Ci-dessous, un exemple est donné sur les données historiques d'une EL d'EH.

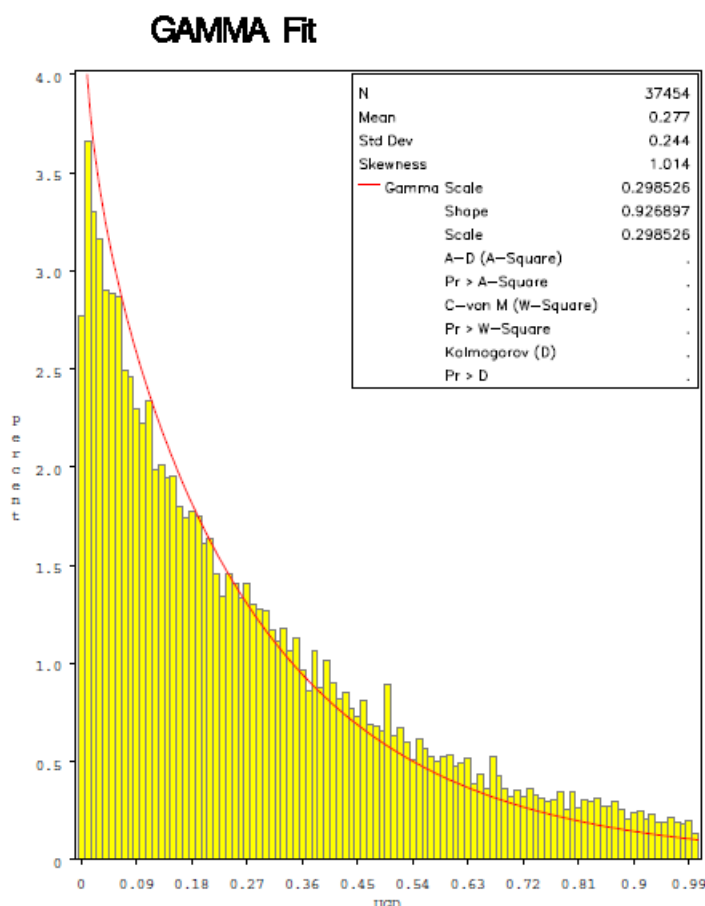


FIGURE 1.9 – Loi gamma ajusté sur les UGD d'une EL

Pour une simulation aléatoire Gamma, un étalonnage de deux paramètres est nécessaire. Ainsi, pour chaque EL, il faut calibrer les moyennes des UGD (sur l'ensemble des acheteurs ayant une exposition au moment de référence et défaillants dans la pé-

riode d'observation de 12 mois suivants) et les écarts-types liés sur chaque axe de la segmentation. La segmentation des UGD est la même que celle des PD.

Si nous appelons $UGD_1, \dots, UGD_n$ les n observations d'UGD des acheteurs défaillants dans le segment particulier, nous avons avec les définitions classiques de moyenne (m) et écart-type (s) :

$$m(UGD) = \overline{UGD} = \frac{1}{n} \sum_{i=1}^n UGD_i \quad \text{et} \quad s(UGD) = \sqrt{\frac{1}{n} \sum_{i=1}^n (UGD_i - \overline{UGD})^2} \quad (1.13) \quad (1.14)$$

A partir de ces deux paramètres, et connaissant les moments d'ordre 1 et 2 de la loi gamma, nous obtenons les paramètres d'intensité (β) et de forme (α) de la loi gamma :

$$\begin{cases} \mathbb{E}[UGD] = \frac{\alpha}{\beta} \\ Var[UGD] = \frac{\alpha}{\beta^2} \end{cases} \Leftrightarrow \begin{cases} m(UGD) = \frac{\alpha}{\beta} \\ m(UGD) = \frac{\alpha}{\beta^2} \end{cases} \Leftrightarrow \begin{cases} \alpha = \frac{m(UGD)^2}{s(UGD)^2} \\ \beta = \frac{m(UGD)}{s(UGD)^2} \end{cases} \quad (1.15)$$

Remarque : La loi gamma permet de simuler un UGD-dessus de 100%. Cela se produit régulièrement, en particulier pour les acheteurs ayant une faible exposition au moment de référence. Leurs limites sont susceptibles d'être augmentées au sein de l'année d'observation.

1.4.4 Application des autres paramètres de réduction de la perte

Une fois que l'on dispose des défauts et des UGD associés à ces défauts, il faut appliquer les autres paramètres propres aux contrats et aux groupes de contrats entre EH et l'assuré. Cela fait le lien entre l'EAD et la perte ultime brute. Sont également appliqués les LGD rentrés dans le modèle interne.

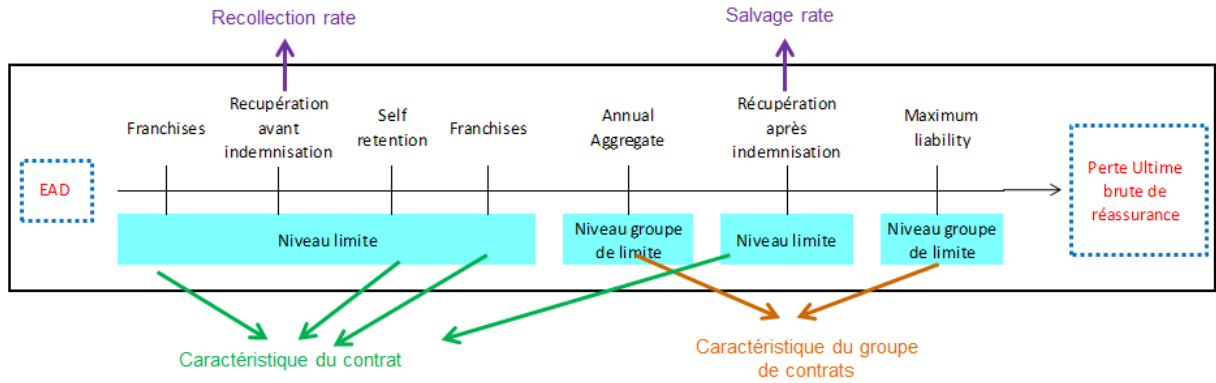


FIGURE 1.10 – Schématisation de l'application des caractéristiques des contrats

Ce schéma 1.10 permet de bien visualiser les différentes caractéristiques utilisées par le modèle ainsi que l'ordre d'application. Il y a deux types de caractéristiques : celles

qui sont propres au contrat d'assurance liant un assuré, un acheteur et EH et celles du groupe de contrats liant l'assuré et EH.

- Au niveau du contrat, nous avons :
 - Différentes franchises.
 - Les taux de récupération.
 - La *self retention*, qui s'assimile au produit de réassurance *Quota-Share*, c'est le partage des pertes entre l'assuré et l'assureur, ainsi il y a un pourcentage fixé des pertes que l'assuré devra porté par lui-même sur chaque sinistre, c'est souvent de l'ordre de 1 à 10 %.
- Au niveau du groupe de contrats, les pertes simulées doivent être additionnées pour appliquer les clauses à ce niveau, et ensuite redistribuées sur les limites. Il y a deux principales clauses :
 - L'*Annual Aggregate* : Les sinistres individuels sont additionnés et seul le montant au-dessus de cet *Annual Aggregate* est à payer. En dessous de ce montant, c'est à la charge de l'assuré.
 - *Maximum liability* : C'est le montant maximal assuré par le groupe de contrat. Au-delà de ce montant, les pertes sont à la charge de l'assuré.

Chapitre 2

Cas d'étude : la World Agency

Dans ce chapitre, nous allons présenter l'entité dont nous voulons modifier l'intégration au modèle interne. L'objectif de ce mémoire est de montrer qu'il existe une manière plus adaptée de prendre en compte cette entité. Nous procéderons à de nombreux *backtesting* (tests de validation rétroactifs) sur les paramètres actuellement utilisés pour prendre en compte la contribution de cette entité dans le SCR. Pour finir ce chapitre, nous formulerons les différentes modifications envisageables.

2.1 Présentation de la World Agency

Nous présentons l'entité sur laquelle le mémoire est basé : la World Agency. Tout d'abord, nous mettrons en avant sa raison d'être et ses particularités, puis nous présenterons ses produits, enfin nous évoquerons comment elle est intégrée au modèle interne.

2.1.1 Principe

EH WA est une compagnie fondée, en 2007, pour fournir au groupe la possibilité de délivrer des polices multinationales à ses plus grands clients. Les missions d'EH WA ont été formées par ses deux principaux objectifs :

- Eviter la concurrence intra-groupe pour les clients multinationaux d'EH.
- Délivrer un haut niveau de services intégrant une couverture multinationale.

EH WA est un courtier interne qui crée des contrats mondiaux et internationaux qui sont portés par les EL. WA définit les termes généraux du contrat et ensuite délègue la gestion aux EL locales, c'est un intermédiaire. Ce sont donc les EL qui gèrent les contrats (exposition, sinistres) et qui reçoivent les primes. WA offre des produits sur-mesure et permet aux assurés d'obtenir la couverture dont ils ont besoin. Ainsi, l'assuré dispose de tous les avantages des contrats classiques mais également de quelques services sur-mesures :

- Les contrats proposés par WA sont simples et flexibles, permettant à l'assuré de comprendre les moindres détails du contrat et d'ajuster le contrat à ses besoins.
- Titulaire d'un contrat avec EH WA, l'assuré dispose du réseau d'experts d'EH dans la prévention et dans le recouvrement.
- Ces assurés sont présents dans au moins deux différents pays et ont la chance d'être face à une plate-forme de décision centrale unique. De plus, ils disposeront d'un service dans leurs langues et dans leurs monnaies.

2.1.2 Aperçu des produits

WA propose trois types de produits. Le premier est le programme mondial (*World Program* = WP), c'est un contrat d'assurance-crédit de court terme pour les groupes internationaux (500M€ de chiffre d'affaires). Le second est le XoL, c'est une assurance moyen terme avec des franchises plus fortes et des limites très segmentées. Le dernier type de produits est la couverture de transaction (*Transactional cover Unit* =TCU), c'est un contrat d'assurance de court et moyen terme, concernant une transaction avec un seul acheteur et requérant un très haut niveau de spécialisation dans les termes du contrat. Nous détaillons ces catégories de produits.

Le WP est un produit d'assurance-crédit à court terme pour lequel les termes généraux du contrat sont définis par WA et sont exécutés par l'EL locale. Les termes généraux du contrat sont négociés entre les équipes régionales de WA et la société-mère cliente (siège social). Ensuite, les branches de la société-mère peuvent souscrire un contrat avec la branche d'EH qui leur est la plus proche, avec les conditions générales du contrat établi avec la société-mère et des conditions particulières dues au business local de la branche.

Le XoL est un produit d'assurance-crédit à moyen terme, qui cible les clients ayant déjà un contrat d'assurance-crédit et/ou une bonne gestion de leur crédit. Ce produit apporte de la protection face à des pertes de crédit exceptionnelles. Il y a toutefois quelques spécificités. Cela concerne l'assurance-crédit à moyen terme, les limites sont fortement segmentées, la couverture d'assurance n'est pas annulable pendant les 12 premiers mois et le niveau des franchises est élevé, notamment le total annuel.

Le TCU est un contrat personnalisé permettant de réduire et gérer les risques liés aux transactions telles que les interruptions de contrats, le non-paiement ou les confiscations par violence politique. La durée de ces contrats varie entre 1 mois et 8 ans. Le montant maximal couvert est de 100 millions d'euros. Ces contrats personnalisés ne couvrent qu'une transaction par contrat. Ainsi, seule la transaction spécifiée dans le contrat est assurée.

Dans le cadre de notre étude, nous nous concentrons sur les contrats WP qui représentent 91 % de l'exposition du portefeuille de WA.

2.1.3 Similarité avec les EL classiques

Commençons par mettre en avant les points communs entre WA et les EL classiques. Ce sont des contrats de courts termes (inférieurs à 12 mois) et annulables. Dans les deux situations, le processus de collecte est assuré par les équipes locales.

Cependant, de nombreuses différences sont à noter. C'est ce qui d'ailleurs nous invite à modifier le système actuel. Au niveau de la taille des assurés, le contrat WP comporte uniquement des grandes compagnies ou des filiales de grandes compagnies, tandis que les contrats classiques, sont conçus pour une gamme d'assuré beaucoup plus vaste.

Au niveau de la localisation des acheteurs des assurés, les contrats WP comportent une très large diversification d'exposition à travers de nombreux pays, tandis que les EL classiques disposent d'une exposition principalement domestique.

Concernant la qualité des acheteurs, dans les contrats WP, nous observons une note moyenne de 3.7. Sur la partie France nous avons une note moyenne de 4 pour les contrats WP contrairement à 4.2 en contrat classique, pour l'Allemagne 3.7 au lieu de 3.8 et au Royaume-Uni 3.3 au lieu de 4.1. On constate donc que les assurés souscrivant aux contrats d'assurance WP font des transactions avec des acheteurs qui ont de meilleures solvabilités, situations financières et capacités à payer leurs dettes.

Là où une différence moins évidente pourrait être mise en avant, c'est sur le prix de l'assurance, en effet lorsque l'on regarde combien d'euros d'expositions sont couverts pour un euro de primes, nous obtenons de grandes divergences entre le WP et les contrats classiques des EL. Pour la partie France, nous obtenons une exposition plus forte pour chaque euro de prime chez WA par rapport à l'EL d'EH France. Cet écart est conservé en Allemagne et amplifié au Royaume-Uni. L'idée est que les acheteurs de ces assurés sont plus solides, leurs défauts sont moins fréquents, et même si les montants sont plus importants, dans l'ensemble, il y a moins de montants en sinistres dans WA que dans les EL classiques, en considérant un ratio S/P (montant des sinistres/montant des primes) égal.

Dans le domaine des conditions des déclarations de sinistres, pour le contrat WP, les assurés doivent déclarer immédiatement les événements de non-paiements si la date n'est pas respectée, car cela est spécifié dans les clauses du contrat. Dans les contrats classiques des EL, il n'y a pas de clause obligeant les assurés à déclarer l'événement de non-paiement directement après la date due.

Chez WA, la procédure en cas de sinistre est la même pour tous les pays, alors qu'au niveau des EL, la procédure pour les sinistres peut varier en fonction de l'EL.

2.1.4 Contexte

Actuellement, au niveau de chaque EL, il y a deux types de contrats dans le portefeuille, les contrats de l'entité WA lié à cet EL et les contrats classiques de cette EL. Pour chaque EL, les paramètres sont calculés de la manière que nous avons présentée dans le chapitre 1. Cependant, ils sont calibrés uniquement sur les données des contrats classiques. Les paramètres du risque crédit de WA (PD/UGD/LGD) ne sont pas estimés, ce sont les paramètres de l'EL correspondante (selon le pays de l'assuré) qui sont

utilisés. Ainsi l'exposition de WA est séparée en différentes parties (une par EL) et à chaque partie, nous appliquons le jeu de paramètre correspondant. WA ne dispose pas de paramètres propres à sa sinistralité.

Pour comprendre cette approche, il faut savoir qu'au moment de la création du modèle interne, WA était une entité très récente et ne disposait pas d'un historique de sinistralité assez grand pour pouvoir calibrer ses propres paramètres. C'est pourquoi EH a formulé l'hypothèse que l'exposition et la sinistralité de WA ont des comportements très similaires à ceux des EL locales. Cette hypothèse prend effet au sein de chaque EL, c'est-à-dire par exemple que la sinistralité de la France et du segment France de WA sont supposés assez similaires pour que nous leurs appliquons les mêmes paramètres.

La proposition principale de notre analyse est de contrôler si l'approche courante est cohérente avec la sinistralité observée. Nous utiliserons les fichiers de sinistres récemment reçus de WA afin de mener nos analyses et d'envisager des axes d'amélioration. Dans un premier temps, nous mènerons des études sur la répartition des sinistres pour tester l'homogénéité et challenger la segmentation. Ensuite, les PD WA vont être calculées et étudiées en fonction des branches où la police a été signée. Puis nous effectuerons une comparaison entre les taux de défaut des segments de WA et les taux de défaut des EL respectives. Par la suite, nous mènerons de nombreux tests de validité rétroactifs (*backtesting*) sur les PD, UGD et pertes attendues afin d'évaluer la consistance des hypothèses prises par le passé. Le *backtesting* ou test rétro-actif de validité consiste à tester la pertinence d'une modélisation en s'appuyant sur un large ensemble de données historiques réelles. Pour clôturer ce chapitre, nous proposerons des alternatives à la méthode actuellement en place.

2.2 Analyses préliminaires

2.2.1 Répartitions de l'exposition entre les pays

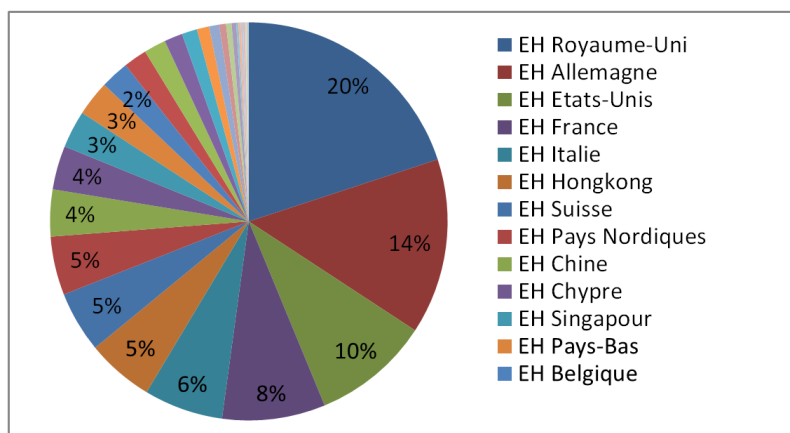


FIGURE 2.1 – Répartition de l'exposition de WA auprès des différentes EL

Nous remarquons que les 4 principales EL (Royaume-Uni, Allemagne, Amérique du Nord et France) représentent plus de 50% de l'exposition totale. L'entité EH Pays Nordiques regroupe 4 EL (EH Suède, EH Finlande, EH Danemark, EH Norvège). Cet agrégat représente 5% de l'exposition totale. En considérant l'EL Pays Nordiques, nous observons que les 8 plus grandes EL représentent 75%. C'est-à-dire que les 20 autres segments de WA se partagent 25% d'exposition.

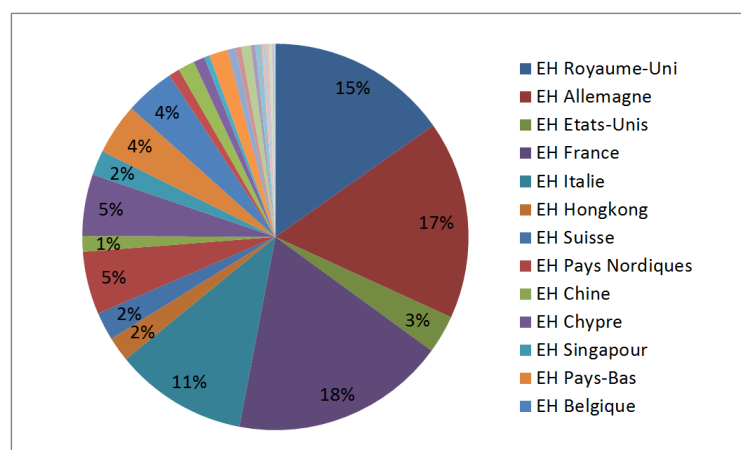


FIGURE 2.2 – Répartition des acheteurs de WA auprès des différentes EL

S'agissant du nombre d'acheteurs, nous observons des différences. Ainsi, il y a plus d'acheteurs de WA à Chypre qu'aux Etats-Unis. Les segments ayant les plus grands nombres d'acheteurs sont la France, l'Allemagne, le Royaume-Uni et l'Italie.

Par la suite, nous avons procédé aux calculs des PD. Afin d'étudier le comportement des contreparties de WA, nous avons observé les corrélations entre les PD des différents pays.

2.2.2 Corrélations

L'échelle de couleur explique la force de la corrélation, le bleu foncé représente les corrélations fortes positives et le rouge foncé montre les corrélations importantes négatives.

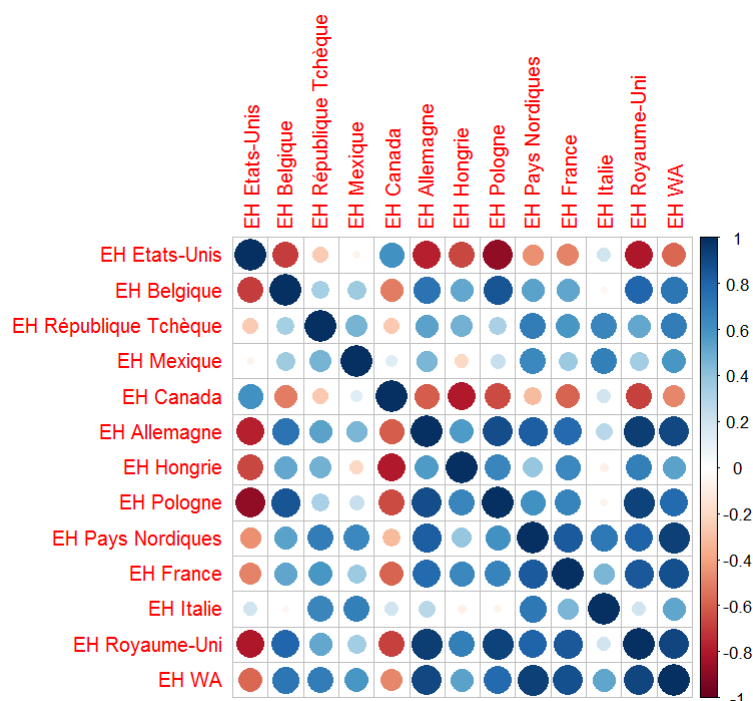


FIGURE 2.3 – Corrélations entre les PD au sein de WA

Nous observons tout d'abord les fortes corrélations entre WA et ses segments. On observe que les segments France (88 %), Allemagne (90%), Royaume-Uni (91%) et Pays Nordiques (94%) ont de fortes corrélations positives avec WA. Ce sont les principaux pilotes du portefeuille car, unis ils représentent une part importante de l'exposition. Nous supposons que la PD globale de WA et celles qui sont séparées par pays pourraient suivre la même tendance.

On observe également que les segments Etats-Unis (-60%) et Canada (-50%) ont des corrélations négatives également fortes. Ces deux EL sont souvent regroupées dans une entité virtuelle : Amérique du Nord (EH *North America*). C'est donc logique que la corrélation entre ces deux segments soit positive et assez forte (60%). Par contre, le fait que ces deux segments évoluent dans le sens contraire de WA n'est pas aussi logique. On peut penser que cela est dû au fait que c'est une région économique puissante et assez

autonome et ne suit pas la tendance générale de WA.

En regardant les corrélations entre les segments fortement corrélés avec WA, nous observons qu'ils sont également corrélés fortement entre eux : Royaume-Uni et France (85%), Royaume-Uni et Allemagne (95%), Allemagne et France (78%), Pays Nordiques et Allemagne (82%).

On conserve également le comportement opposé du segment Etats-Unis, en effet il a des corrélations négatives très fortes avec les segments Allemagne (-77%), Royaume-Uni (-80%) et Pologne (-90%).

Dans la suite de ce mémoire, les notations suivantes pourront être utilisées pour désigner les EL :

- ACI (EH Etats-Unis)
- TI (EH Royaume-Uni)
- SIAC (EH Italie)
- SFAC (EH France)
- NORDIC (EH Pays Nordiques)
- HERM (EH Allemagne)

Nous allons désormais comparer le comportement des PD des segments de WA avec celles des EL correspondantes.

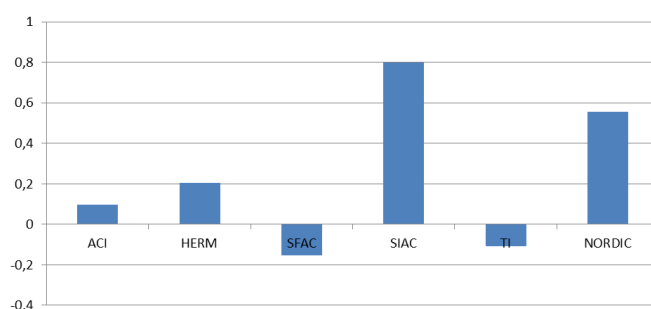


FIGURE 2.4 – Corrélations entre les PD au sein de WA et les PD au sein des EL correspondantes

Nous remarquons sur la figure 2.4 que seule l'Italie et de manière moins convaincante, les Pays nordiques ont une corrélation positive significative indiquant que ces pays ont un comportement similaire entre les polices WA et non WA.

Sur le graphique 2.5, nous observons les PD de la WA par pays et la globale de la WA. On se concentre sur les six plus grands contributeurs : Allemagne, France, Royaume-Uni, Etats-Unis, Italie et les pays Nordiques. La tendance bien particulière des Etats-Unis est opposée à la tendance commune. Les Etats-Unis représentant une importante part de l'exposition totale ne peuvent pas être mis à part, dans notre étude il faudra faire attention aux résultats sur ce segment de WA. Nous observons également une forte hausse des PD de l'Italie en 2011. De plus, pour les autres EL, qui ont une tendance similaire, nous pouvons voir que la convergence n'est pas stable et le niveau le plus récent des PD laisse un écart de 0.65% ce qui est important.

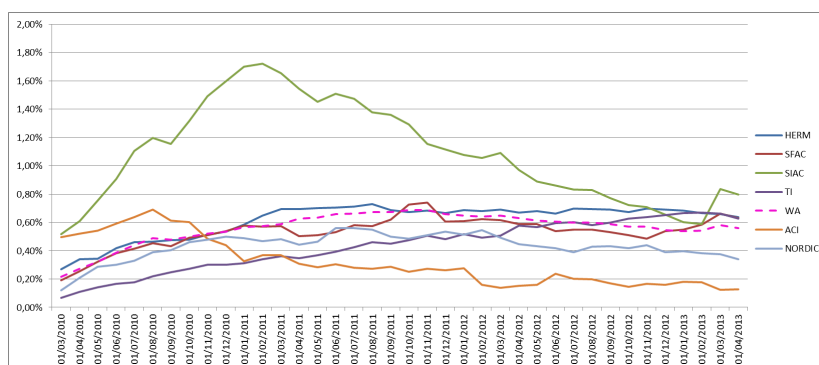


FIGURE 2.5 – PD de WA et des grands segments de WA

Par la suite, nous avons mené de nombreux *backtesting*, nous avons utilisé les PD observés, es UGD observés et les pertes observées.

2.2.3 Comparaisons des PD de WA par pays avec celles des EL correspondantes.

Cette étude est complémentaire à celle des tests de corrélation réalisés au préalable. Nous étudions désormais le comportement des courbes.

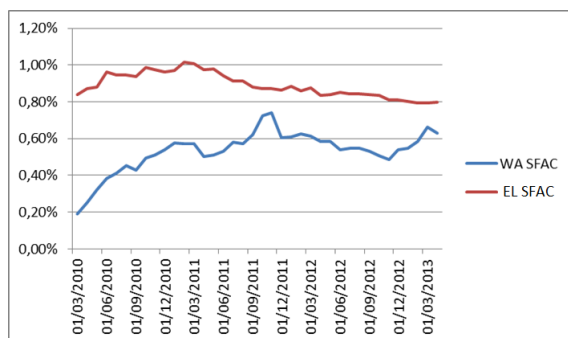


FIGURE 2.6 – Exemple de la France

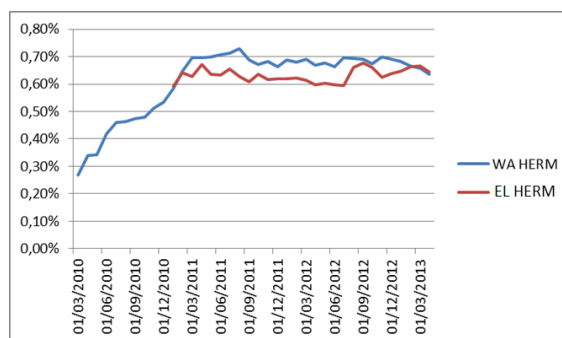


FIGURE 2.7 – Exemple de l'Allemagne

On constate, qu'il y a de nombreuses disparités dans les observations. Comme nous pouvons le constater, même si les écarts sont importants et irréguliers, les Etats-Unis (annexe AA.4), la France et l'Italie(annexe AA.1) ont un taux de défaut observé sur le portefeuille WA supérieur à celui des EL classiques correspondantes.

Au contraire, pour l'Allemagne, au Royaume-Uni (annexe AA.3) et les Pays Nordiques (annexe AA.2), nous voyons que les PD sont plus élevées au niveau de WA qu'au niveau des EL. Pourtant, la prise en compte de la différence de la note moyenne et de la taille d'exposition, nous ne pouvons pas conclure que les PD de WA ne sont pas comparables avec les PD des autres EL d'EH. Pour la suite de nos analyses, nous ne considérons

plus l'EL NORDIC, en effet, sachant que c'est une agrégation de 4 EL (EH Suède, EH Finlande, EH Danemark, EH Norvège) qui ne représente pas une part suffisante de l'exposition (5%) pour mener les études qui vont suivre.

2.2.4 Tests de validation rétroactifs des PD de WA par pays par rapport à celles qui sont prédites pour les EL correspondantes

Nous comparons désormais les PD établies un an auparavant pour les EL avec les taux de défaut constaté sur WA par pays (cf. figure 2.8).

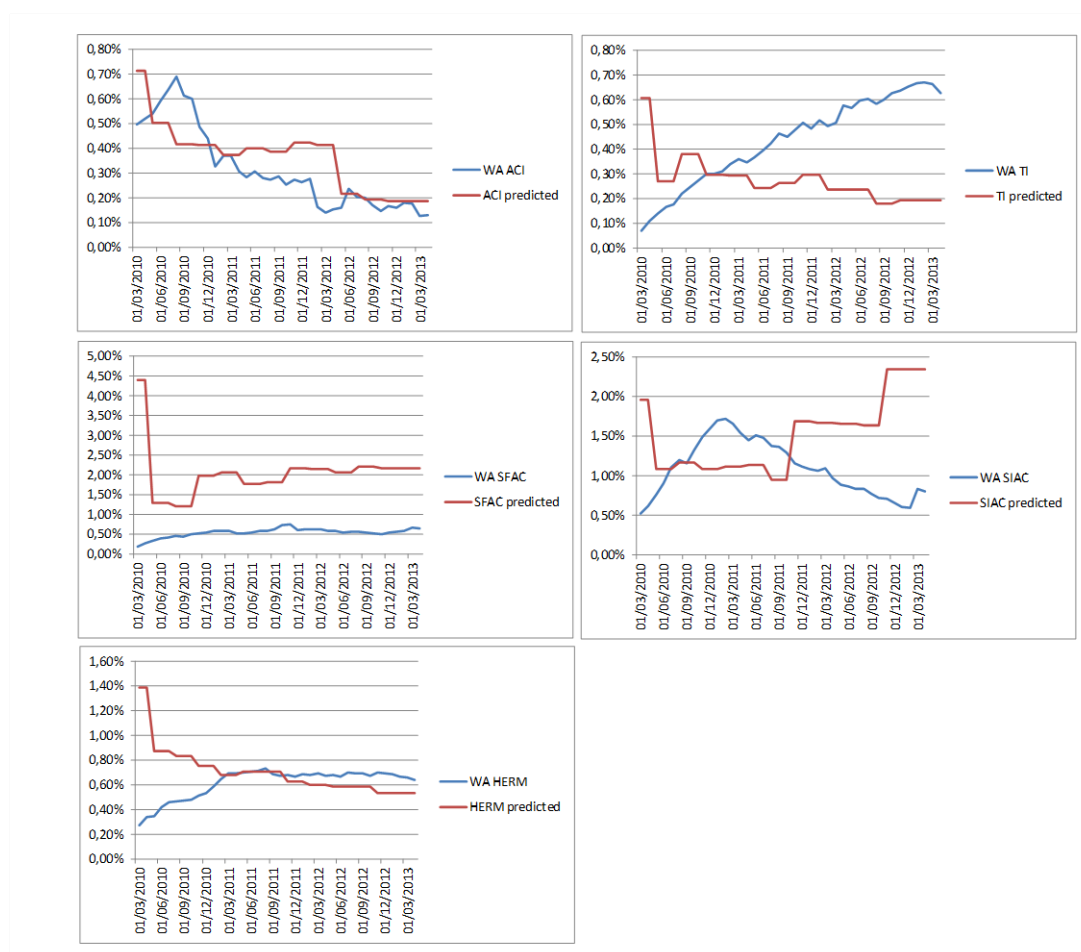


FIGURE 2.8 – Backtesting des PD prédites avec les PD observées

Pour les Etats-Unis, nous constatons que la prédiction s'améliore avec le temps et que la sinistralité au préalablement sous-estimée puis surestimée se retrouve juste en dessous de la prédiction sur l'année 2013, les prédictions sont recevables pour ce segment de WA. Concernant la part rattachée à la France, la sinistralité de WA est clairement en dessous de celle de l'EL France. Les PD sont donc surestimés pour ce segment. De même pour

la part rattachée à l'Italie, nous voyons dans un historique récent que la sinistralité de WA est inférieure à celle de l'EL. On retrouve le pic de 2011 où les PD étaient plus fortes pour la WA, mais cela ne reflète pas la tendance globale. En Allemagne, nous observons qu'au début, les PD de WA étaient fortement inférieures à celles prédites pour l'EL allemande, cependant elles ont convergé l'une vers l'autre et dans les années les plus récentes la sinistralité est légèrement sous-estimée en prenant les prédictions pour les PD de l'Allemagne. Pour ce qui est du Royaume-Uni, nous observons une sinistralité en forte croissance pour WA alors que les prédictions pour la partie EL ont diminué et se sont stabilisées. La sinistralité est nettement sous-estimée.

On peut conclure que dans l'ensemble, cette technique fonctionne pour certaines EL et qu'un effet de compensation entre les segments sous-estimés et surestimés pourrait donner des résultants cohérents. Cependant, les prédictions ne reflètent pas fidèlement le comportement des segments de WA.

2.2.5 Test de validation rétroactifs des UGD de WA par pays par rapport aux UGD des EL correspondantes

La comparaison entre les UGD prédits pour les EL et ceux de WA par pays nous montrent que la dernière est plus volatile. Il faut préciser que les premiers sont disponibles par trimestres, tandis que ceux qui sont calculés pour WA sont disponibles mensuellement. C'est une raison pour laquelle nous observons plus de variation avec les UGD de la WA. Cette forte volatilité observée sur les UGD de WA particulièrement en début de la période d'observation est principalement due à un manque de données sur cette période. En effet, pour la partie UK de WA, nous ne disposons que de 30 sinistres sur l'année 2010.

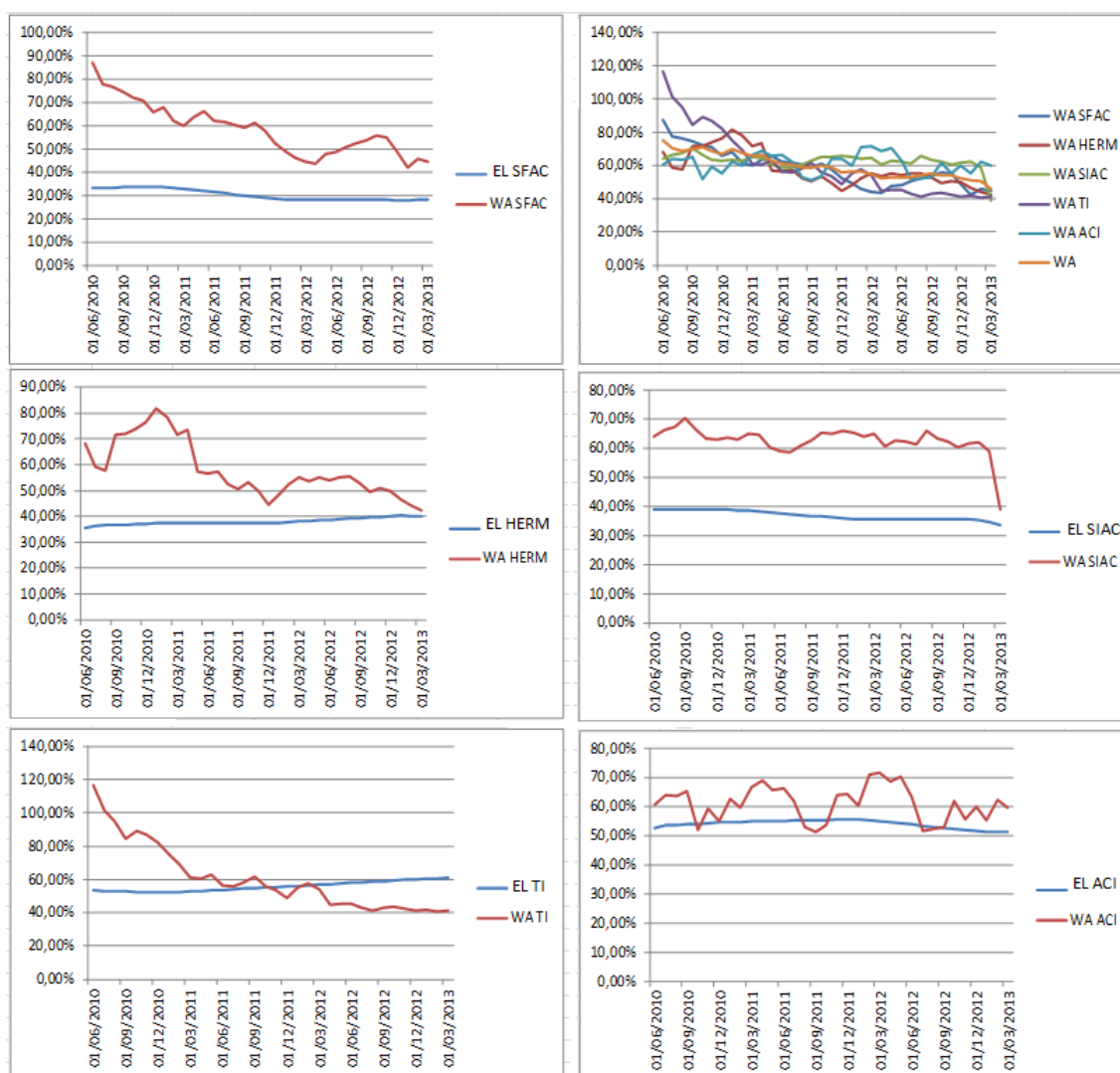


FIGURE 2.9 – Backtesting des UGD prédits avec les UGD observés

Nous pouvons remarquer une tendance générale à la baisse de l'UGD convergeant à 45%. Cependant, même si les écarts sont plus faibles, il reste toujours 20% d'écart entre l'Italie (40%) et le Royaume-Uni (60%) ce qui est conséquent.

Même si le niveau des UGD n'est pas homogène entre les EL, nous pouvons observer qu'à part le Royaume-Uni récemment, les UGD de WA sont plus élevés que celui des EL.

2.2.6 Test de validation rétroactif de l'*Expected Loss* de WA par rapport à celle calculée avec l'exposition de WA et les paramètres des EL correspondantes

La perte attendue (*Expected Loss*) dépend beaucoup de l'exposition. Nous supposons désormais que les paramètres des EL sont utilisables pour WA et nous attribuons donc l'exposition de WA par EL avec les différents paramètres de chaque EL. Dans un premier temps, nous ne prenions pas en compte les acheteurs non dénommés et les acheteurs non liés. Puis constatant que les EL étaient trop faibles, il a été établi qu'on devait les prendre en compte. On a donc supposé que ces catégories d'acheteurs seraient dans les mêmes proportions pour WA que pour chaque EL en appliquant le prorata d'exposition.

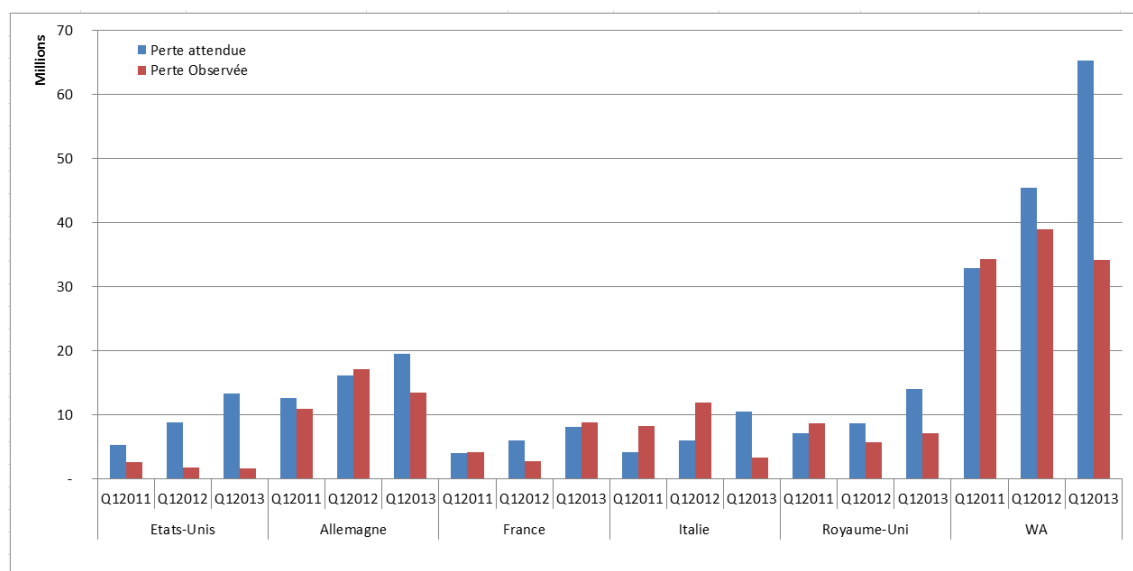


FIGURE 2.10 – Backtesting de l'Expected Loss

Nous pouvons observer un manque évident de régularité : d'une part pour WA, la perte attendue pour les Etats-Unis est considérablement surestimée, d'autre part, la perte attendue de WA en Italie a été sous-estimée, en 2011, et en 2012. Globalement, nous pouvons observer que, dans certains cas, nous sommes très conservateurs en utilisant les paramètres des EL. Cependant, il y a 3 cas où la perte attendue est dépassée de plus de 10%. Les surestimations sont importantes : dans 38% des cas, la perte attendue est presque le double de la perte observée. Spécialement aux Etats-Unis (ACI) où l'EL est

toujours au-dessus de la perte réelle. Le plus important est au niveau de la WA, nous observons qu'en 2011, la perte attendue est dépassée mais seulement de 5% et par la suite elle est bien au-dessus des pertes observées 86% et 52%. On voit clairement que ce n'est pas fidèle. L'approximation de cette manière de la perte attendue ne peut être maintenue, car dans certains cas, elle n'est pas assez forte (2011) et dans d'autres cas, elle est beaucoup trop restrictive (2013). Pour 2011, c'est principalement dû à une sous-estimation de l'Italie et du Royaume-Uni.

Certes, ce processus se révèle prudent pour 2012 et 2013. Mais nous observons qu'il est trop éloigné de la réalité, il y a une différence de 30 millions en 2013, l'estimation est deux fois plus grande que ce qui est observé par la suite.

2.3 Méthodes envisageables

L'étude de la situation actuelle montre que la méthodologie utilisée n'est pas vraiment adaptée pour calculer les paramètres de WA. Devrions-nous garder, pouvons-nous l'améliorer ou proposer une alternative ? Nous considérons quatre solutions possibles :

- La première est de conserver la méthodologie actuelle qui consiste à utiliser les paramètres des EL avec l'exposition de WA. Globalement, la perte attendue est très conservatrice, donc, nous pouvons accepter ce conservatisme pour éviter tout changements dans l'évolution de la méthodologie se révéleraient très coûteux en termes de temps. Cependant les paramètres utilisés ne sont pas bien adaptés aux limites de WA, et il faudrait sûrement calibrer un ajustement forfaitaire. Cette méthode est celle qui sera retenu, si aucune des autres n'arrive à se démarquer.
- La seconde solution sera étudiée plus en détail, elle consiste à intégrer les données de sinistres et d'exposition de WA dans le calcul des paramètres. Cela ne changerait pas le modèle, mais juste les paramètres d'entrée. Au lieu d'appliquer les paramètres des EL aux différents segments de WA, nous calibrons sur chacun de ces segments un jeu de paramètre basé sur la combinaison des données de WA et de l'EL liée à ce segment. Cette solution est envisageable dans un délai court, elle engendre peu de changement et ne touche pas le modèle.
- La troisième solution est de calculer les paramètres pour WA comme s'il s'agissait d'une unique EL. De cette manière, les paramètres seront adaptés au portefeuille WA. Mettre en place le calcul d'un jeu de paramètre pour WA est semblable à l'ajout d'une nouvelle EL. L'ajout de nouvelles EL est un sujet qui a déjà été traité, nous avons donc de la visibilité sur la méthodologie à suivre. Ensuite, la complexité sera d'enlever le traitement de WA tel qu'il est fait actuellement par le modèle interne. Il y aura donc des modifications à effectuer dans le code du modèle interne et à justifier auprès du régulateur. De plus, Il n'est pas évident que les paramètres globaux soient adaptés à tous les segments.
- La dernière solution est de calculer les paramètres spécifiques à WA mais ventilées par pays de l'assureur (gestion confiée à l'assureur). C'est une solution sur mesure,

qui semble la plus appropriée. Cependant, elle requiert un groupement par pays afin d'obtenir une base de données cohérente. Car si en considérant WA comme une EL, elle devient la plus grande en matière d'exposition, la répartition de ses expositions ne permet pas de faire des calculs sur chacun des segments, nous aurions donc une segmentation moins granulaire et des paramètres instables. Comme cela est montré dès le début de ce chapitre, 21 EL se partagent 25% de l'exposition. Sachant que par la suite, il y a encore des segmentations, il faudrait regrouper certains pays et choisir une segmentation moins fine que pour les autres EL. Le processus est le même que dans le troisième cas, mais pour chaque groupe (éventuellement Europe occidentale et centrale, Italie et méditerranée, et autres) Cette solution serait la plus coûteuse en temps et en argent et ne garantit pas un résultat satisfaisant.

Avant de mener des études sur les différents paramètres d'entrée du modèle interne, nous procéderons à des tests de sensibilités afin de mettre en avant les paramètres dont les variations ont l'impact le plus fort sur le SCR. Cela permettra d'avoir une vision sur la conséquence d'une variation de chaque paramètre, et de déterminer quels paramètres doivent être les mieux calibrés et suffisamment proches pour être agrégés.

Chapitre 3

Tests de sensibilité sur les paramètres du modèle interne

Dans ce chapitre, nous présentons les tests de sensibilité que nous avons menés sur le modèle interne. Ils ont été menés afin d'observer comment réagit le SCR face aux chocs sur chacun des paramètres.

3.1 Présentation des tests de sensibilité

Nous avons procédé à des tests sur six différents paramètres d'entrée du modèle interne suivants :

- Les PD.
- Les UGD.
- Les RR.
- Les SR.
- L'exposition.
- La part des acheteurs non liés et non dénommés.

Le but de ces tests de sensibilité est d'établir les réactions sur le SCR du modèle interne face à des chocs dépendant de la volatilité des paramètres. Soit σ l'écart-type de l'échantillon, entre $-\sigma$ et $+\sigma$ se trouvent 68% des valeurs pour une loi normale et entre -2σ et $+2\sigma$ se trouvent 95% des valeurs toujours dans le cadre d'une loi normale. Le choc retenu est l'écart-type relatif théorique correspondant à la formule :

$$\sigma = \frac{\frac{s}{m}}{\sqrt{N}} \quad (3.1)$$

- s : l'écart-type.
- m : la moyenne.
- N : le nombre d'observations.

Pour chacun de ces paramètres, les chocs δ , $\delta \in \{-2\sigma, -\sigma, \sigma, 2\sigma\}$ seront effectués.

Dans un premier temps, nous calibrons σ qui représente la force du choc à appliquer. Le choc sera calculé pour chaque EL. EH a fait le choix d'utiliser une calibration à un

moment fixe dans le temps pour calculer les paramètres du modèle TCI&S. Ainsi, nous ne travaillerons pas sur un historique de données, mais sur les paramètres à une date donnée. Le SCR que nous choquerons au cours de ces différents tests est celui de juin 2015.

3.1.1 Choc sur les PD

Pour calibrer le choc σ des PD, nous prenons la PD de décembre 2013, car il faut généralement attendre 18 mois pour observer tous les sinistres référant à un exercice de 12 mois. En effet sachant que les délais de paiement après facturation peuvent atteindre 6 mois et que le défaut est rattaché à la date de facturation, il est donc nécessaire de prendre 18 mois pour observer une PD complète. Nous ne disposons que du ratio du nombre de défaut sur le nombre de limites :

$$PD_{statut} = \frac{D_{statut}}{N} \quad (3.2)$$

Avec :

- Statut : *protracted* ou *insolvency*.
- D : le nombre de limite sur lesquelles un acheteur ayant une limite positive en décembre 2013 a fait défaut sur la période de décembre 2013 à novembre 2014.
- N : le nombre de limites positives en décembre 2013.

Il faut savoir que D est assimilable à une variable suivant une loi de Bernoulli qui prend la valeur 1 en cas de défaut avec probabilité p et 0 dans le cas contraire avec probabilité 1-p. On obtient la moyenne et la variance de PD à l'aide de celle de D. On en déduit l'écart-type s (racine de la variance).

$$\begin{cases} \mathbb{E}[D] = p \\ Var[D] = p(1-p) \\ s[D] = \sqrt{p(1-p)} \end{cases} \Leftrightarrow \begin{cases} \mathbb{E}[PD] = \mathbb{E}\left[\frac{D}{N}\right] = \frac{p}{N} \\ Var[PD] = Var\left[\frac{D}{N}\right] = \frac{p(1-p)}{N} \\ s[PD] = s\left[\frac{D}{N}\right] = \sqrt{\frac{p(1-p)}{N}} \end{cases} \quad (3.3)$$

le *sigma* est obtenu par la formule suivante :

$$\sigma_{statut} = \frac{\sqrt{\frac{p(1-p)}{N}}}{\frac{p}{N}} = \sqrt{\frac{1-p}{pn}} \quad (3.4)$$

Ceci est équivalent à :

$$\frac{\sqrt{\frac{PD_{statut}(1-PD_{statut})}{N}}}{PD_{statut}} \quad (3.5)$$

Nous avons pu calculer les chocs sur les PD de chaque EL à partir de cette formule. Nous avons fait des calculs par statut de défaut dans un premier temps, cependant dans un souci de lisibilité des résultats (rapporter l'impact au choc effectué) nous avons

besoin également du choc moyen que nous avons obtenu en appliquant cette formule à la PD globale de l'EL. Ceci a également été effectué pour les chocs sur UGD et taux de recouvrement.

3.1.2 Choc sur les UGD

Les paramètres d'UGD choisis en entrée du modèle interne sont la moyenne et l'écart-type des UGD sur l'année sur chaque axe de la segmentation. Nous avons choisi l'année 2014, nous avons calculé les UGD de chaque défaut, puis conservé uniquement ceux inférieurs à 10 car ceci est une condition qui a été fixée par l'équipe en accord avec le régulateur après de nombreuses études montrant l'impact significatif sur le SCR. Les UGD supérieurs à 10 sont des cas exceptionnels, cela veut dire que le sinistre indemnisé est 10 fois plus fort que la limite en début d'année, ces cas-là arrivent lorsque les limites en janvier sont très basse et augmentent significativement au cours de l'année. Ainsi les UGD qui sont censés rester entre 0 et 1 peuvent dépasser les 100 % car les limites bougent au cours de l'année. Il a été établi que ces UGD supérieurs à 1000 % avaient un impact négatif sur la calibration des paramètres. Par la suite, nous avons appliqué la formule 3.1 aux UGD afin d'obtenir les paramètres du modèle interne.

Nous nous sommes demandé si le choc que nous venons de calculer devait s'appliquer seulement à la moyenne ou si nous devions l'appliquer à l'écart-type également. Sachant que lorsque qu'une population subit un choc, il est directement retransmis sur la moyenne et sur l'écart-type. En effet si $y_i = px_i$ alors :

$$m = \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$$

$$m = \frac{1}{n} \sum_{i=1}^n px_i = p\bar{x}$$

(3.6)

$$s(y) = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2} =$$

$$\sqrt{\frac{1}{n} \sum_{i=1}^n (px_i - p\bar{x})^2} =$$

$$\sqrt{p^2 \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2} = ps(x)$$

(3.7)

Nous avons donc par la suite procédé à une étude pour savoir s'il fallait appliquer le choc sur l'écart-type des UGD. Pour cela nous avons étudié un cas avec des données moyennes présentes sur les figures 3.1 et 3.2.

Nous avons fait deux cas, le premier où l'on choque seulement la moyenne et le second où l'on choque également l'écart-type.

L'utilisation de la loi gamma dans le modèle interne est détaillée au point 1.4.3. Dans un premier temps nous observons que le cas 1 réagit de manière importante sur la partie où se trouve la majorité de la densité (3.3). En effet, nous observons visuellement de grands écarts entre les maximum des courbes de densité, surtout lors du choc en *insolvency*. Le cas numéro 2 montre également des variations au niveau du max mais de manière moins importante. Par contre lorsque l'on s'intéresse un peu plus à la queue de distribution nous observons que les courbes de densité du cas n°1 convergent entre elles beaucoup plus rapidement que celles du cas n°2. L'objectif de cette étude étant de

cas n°1 choc uniquement sur UGD				
	UGD	UGD_SD	alpha	beta
insolvency choc = 0,12				
sans choc	0,390	0,42	0,862	2,211
-2	0,296	0,42	0,498	1,680
-1	0,343	0,42	0,668	1,946
1	0,437	0,42	1,082	2,476
2	0,484	0,42	1,326	2,741
protracted choc = 0,06				
sans choc	0,350	0,38	0,848	2,424
-2	0,308	0,38	0,657	2,133
-1	0,329	0,38	0,750	2,278
1	0,371	0,38	0,953	2,569
2	0,392	0,38	1,064	2,715

FIGURE 3.1 – CAS 1 choc uniquement sur la moyenne (UGD)

cas n°2 choc sur UGD et UGD_SD				
	UGD	UGD_SD	alpha	beta
insolvency choc = 0,12				
sans choc	0,390	0,420	0,862	2,211
-2	0,296	0,319	0,862	2,909
-1	0,343	0,370	0,862	2,512
1	0,437	0,470	0,862	1,974
2	0,484	0,521	0,862	1,783
protracted choc = 0,06				
sans choc	0,35	0,380	0,848	2,424
-2	0,308	0,334	0,848	2,754
-1	0,329	0,357	0,848	2,579
1	0,371	0,403	0,848	2,287
2	0,392	0,426	0,848	2,164

FIGURE 3.2 – CAS 2 choc sur moyenne et écart-type(UGD_SD)

choquer le SCR, il nous faut donc choquer la queue de distribution. En effet :

$$SCR = CVAR99.5\%(Pertes Attendues). \quad (3.8)$$

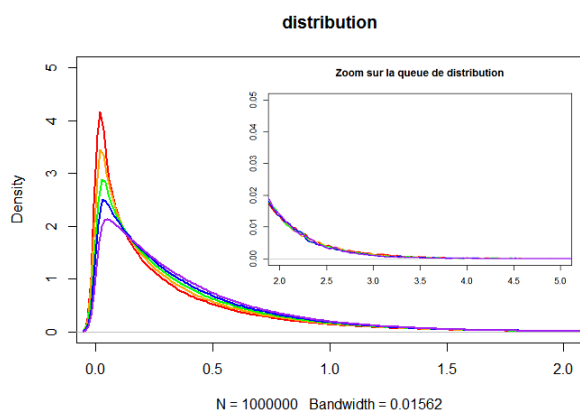
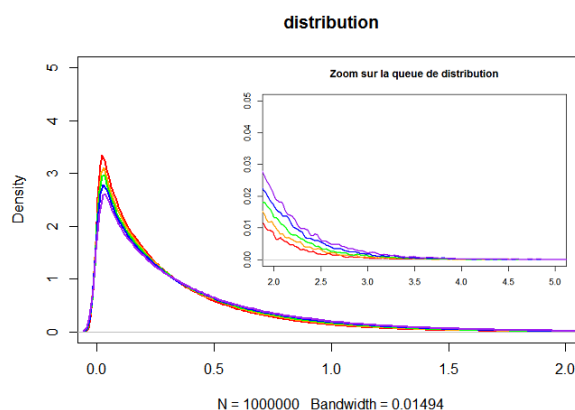
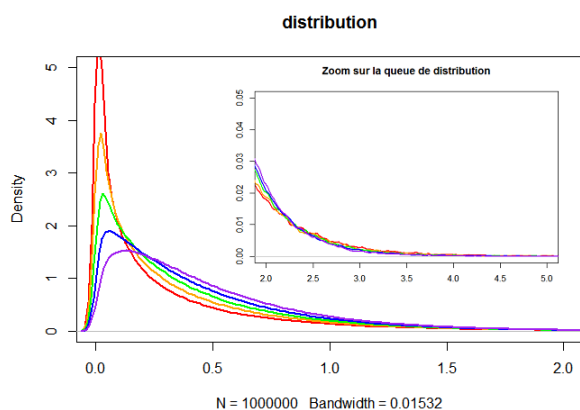
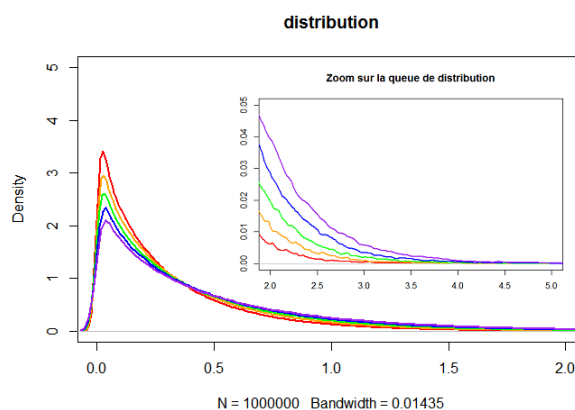
Chez EH, la notation CVAR ne réfère pas à la *Conditional Value at Risk* (CVaR) mais à l'écart entre la *Value at Risk* (VaR) et la perte attendue (*Expected Loss*) :

$$CVAR99.5\% = VaR_{99,5} - Expected Loss. \quad (3.9)$$

Le SCR est donc dépendant du quantile à 99,5% de notre distribution de perte finale et il est donc plus intéressant de choquer la queue de distribution. Cette étude nous a montré que lorsque nous choquions la moyenne et l'écart-type de l'échantillon, nous obtenions des variations plus importantes au niveau de la queue de distribution. La conclusion de cette étude est qu'il faut appliquer les chocs à la moyenne mais également à l'écart-type.

3.1.3 Choc sur les RR et SR

En ce qui concerne les taux de recouvrement avant et après indemnisation (RR et SR), la méthode de calcul de ces paramètres d'entrée dans le modèle interne est différente de celle des UGD. En effet, au lieu de prendre la moyenne comme paramètre, nous prenons le rapport entre deux sommes comme cela est précisé au point 1.5. La question à se poser est donc de savoir s'il est plus intéressant de choquer par rapport à la volatilité des taux de recouvrements ou celle des montants indemnisés. Le choc a pour cible le comportement du recouvrement, nous avons décidé d'appliquer la formule 3.1 aux montants récupérés avant et après indemnisation.

FIGURE 3.3 – Cas 1 *protracted*FIGURE 3.5 – Cas 2 *protracted*FIGURE 3.4 – Cas 1 *insolvency*FIGURE 3.6 – Cas 2 *insolvency*

3.1.4 Choc sur les expositions et acheteurs non dénommés et non liés

Pour les chocs sur l'exposition et sur les acheteurs non dénommés et non liés, nous ne disposons pas de méthode pour déterminer une volatilité, il s'agit de paramètres constants à moment fixé dans le temps. Nous nous sommes donc référés à la documentation de la formule standard de l'EIOPA [11] laquelle spécifie un écart-type combiné pour le risque de prime d'assurance non-vie et le risque de réserve. Pour l'assurance-crédit, il a été fixé à 12%. Nous avons donc pris cet écart-type forfaitaire pour effectuer nos chocs.

Par la suite, nous appliquons les différents chocs aux paramètres concernés, puis nous effectuons un calcul de SCR pour chacun des chocs sur chacun des paramètres. Dans le cadre de notre étude, il n'y aura pas de combinaison de chocs. Nous avons fait 640 000 simulations pour chaque test. Nous avons 24 calculs de SCR à effectuer. Nous illustrons les résultats dans la partie suivante.

3.2 Résultats

On obtient de nombreux résultats en sortie de cette étude. Dans un premier temps nous observons les résultats au niveau groupe, puis si nous mettrons en avant les EL importantes ayant un comportement différent.

Comme cela a été précisé au préalable, nous avons réalisé ses tests sur le SCR brut et net de réassurance. Pour chaque test de sensibilité, nous avons calculé le $SCR = CVAR_{99,5\%}$ (cf. équation 3.9) et nous le comparons au SCR non choqué. La conclusion du test de sensibilité est définie à l'aide de l'indicateur R qui est calculé pour chacun des tests à l'aide des sigma globaux que nous avons calculés en enlevant la distinction du statut au moment du défaut.

$$R = \left| \frac{CVAR_{99.5\%choqué} - CVAR_{99.5\%initial}}{CVAR_{99.5\%initial} \delta} \right| \quad (3.10)$$

$$avec \begin{cases} R < 5\%, & \text{sensibilité non matérielle} \\ 5\% \leq R < 25\%, & \text{sensibilité modérée} \\ 25\% \leq R < 75\%, & \text{sensibilité majeure "en dessous de linéaire"} \\ 75\% \leq R < 125\%, & \text{sensibilité majeure "linéaire"} \\ 125\% \leq R, & \text{sensibilité majeure "au-dessus de linéaire"} \end{cases}$$

Nous avons les résultats de ces tests au niveau du groupe sur la figure 3.7.

Nous observons un comportement similaire pour le SCR net et le SCR brut. Il subsiste quelques écarts significatifs au niveau des chocs -2σ notamment pour les UGD et expositions. On observe que l'impact des chocs sur les taux de récupération RR et SR est très faible comparé aux autres chocs. La sensibilité du SCR à ces deux taux est modérée voir non matérielle pour les taux de récupération après indemnisation. L'impact des chocs sur la proportion des sinistres sur acheteurs non liés ou non dénommés est également plus faible que les autres. Nous observons une sensibilité majeure mais bien en dessous de linéaire. L'exposition à une sensibilité majeure en dessous de la linéarité.

Par contre, la sensibilité sur les UGD et PD est majeure et linéaire, ce sont ces deux paramètres qui sont les plus sensibles aux variations et ils impactent linéairement le SCR. Nous allons nous consacrer sur ces deux paramètres à partir de maintenant.

Ayant constaté peu de différences entre les réactions du SCR net et du SCR Brut, nous avons décidé de nous consacrer uniquement aux résultats bruts par la suite. En effet, ce mémoire n'évoque pas le schéma de réassurance entrepris par le groupe, et au contraire évoque les étapes du calcul du SCR brut.

Au niveau des UGD, nous observons des spécificités sur certaines EL. En effet, l'Italie a une sensibilité bien au-dessus de la linéarité au choc des UGD. Nous observons également que le Royaume-Uni au contraire est moins sensible que le groupe à ce choc (cf. figure 3.8).

Au niveau des PD, nous observons que l'Allemagne a une sensibilité bien plus forte que la sensibilité linéaire. Sinon les autres EL suivent la tendance du groupe, le Royaume-

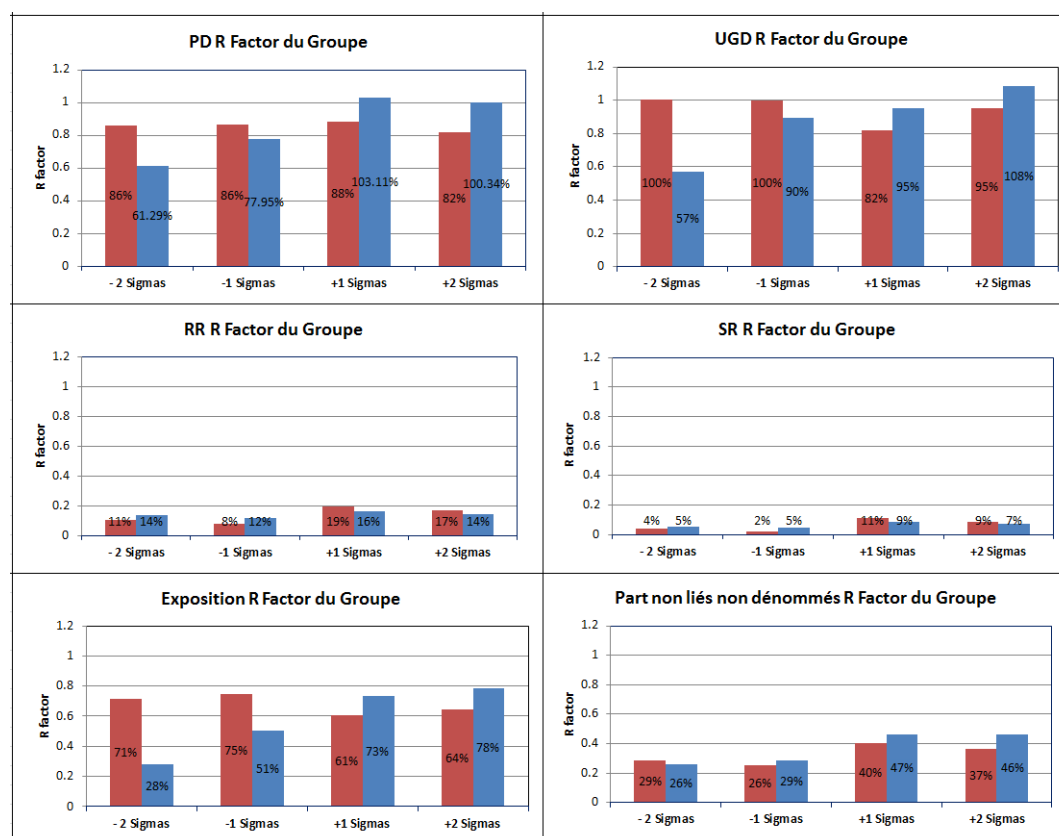


FIGURE 3.7 – Résultat des tests de sensibilité au niveau groupe (net en bleu et brut en rouge)

Uni se révèle encore l'EL la moins sensible, mais celle-ci reste tout de même majeure (cf. figure 3.9).

3.3 Impact dans notre étude

Nous avons observé que les paramètres qui ont le plus d'influence sur le SCR lorsqu'ils sont choqués, sont les PD et les UGD. Ils ont une sensibilité majeure que nous pouvons qualifier de linéaire. Les taux de récupération (SR et RR) ont quant à eux un impact modéré presque non matériel. Ils peuvent donc être négligés par rapport aux deux paramètres précédents. Concernant les tests sur l'exposition et les acheteurs non dénommés et non liés, pour lesquels nous ne disposons pas de manière simple de calibrer un choc, nous ne pouvons pas nous attendre à une logique sur l'évolution de ces paramètres en incorporant des sinistres en plus. Il s'agit de données quantitatives qui ne sont pas pondérées. Nous ne pouvons donc pas nous attendre à conserver ces paramètres intacts en incorporant de nouvelles données.

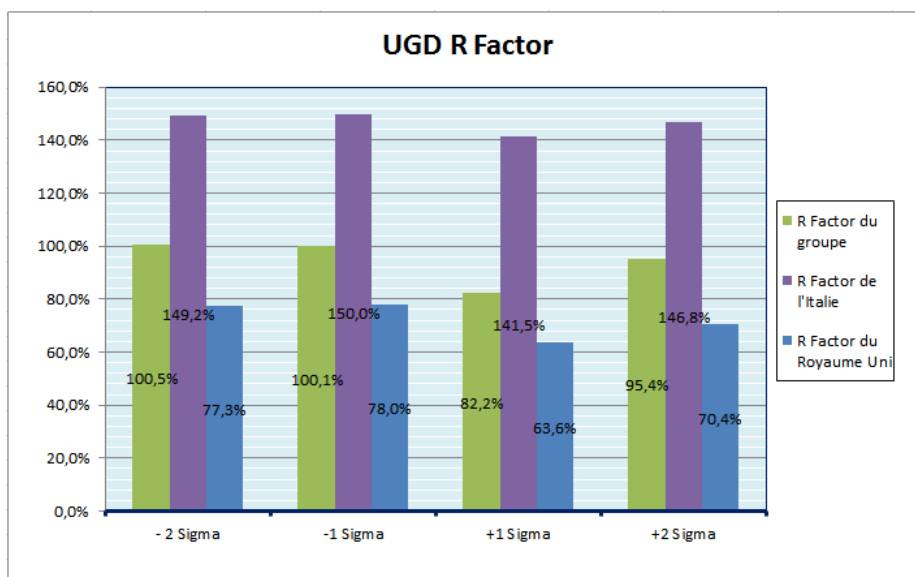


FIGURE 3.8 – R Factor sur UGD Groupe Italie Royaume-Uni

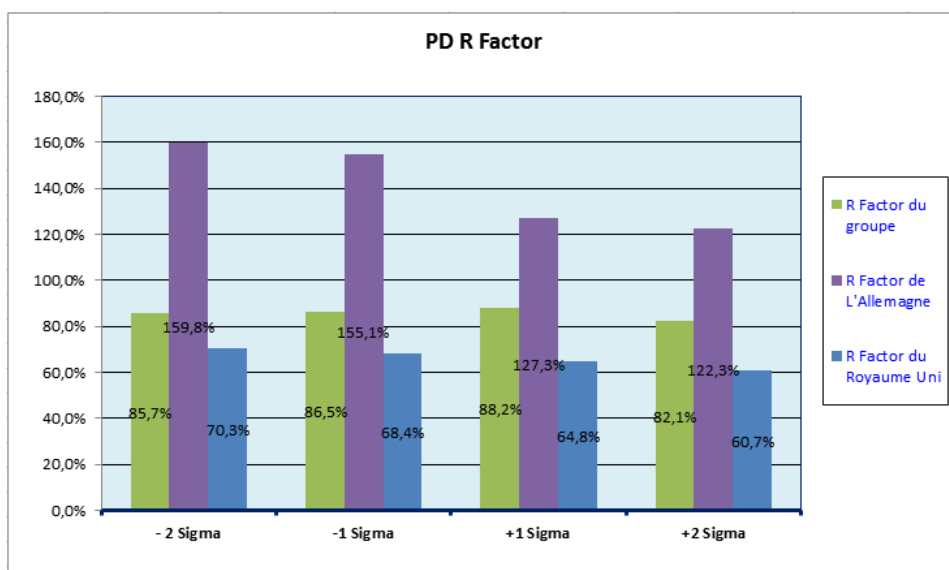


FIGURE 3.9 – R Factor sur PD Groupe Italie Royaume-Uni

Nous fonderons donc notre étude des différentes configurations possibles sur les PD et les UGD.

Chapitre 4

Etude de l'intégration des données de WA dans le calcul des paramètres

Dans ce chapitre, nous allons tester deux des solutions envisagées. Tout d'abord, la solution consistant à intégrer les données WA au sein du processus de calcul existant pour chaque EL. Puis, nous constituerons une EL à part entière avec les données WA. Avant de procéder à des tests, nous décrirons les fichiers sinistres et expositions dont nous disposons, puis nous calculerons les différents paramètres requis pour les tests statistiques à établir.

4.1 Qualité des données

Chez Euler Hermes de nombreux tests sont effectués pour assurer la qualité des données. En effet, Solvabilité rentre en vigueur à partir de janvier 2016 et il est très important de pouvoir mettre en avant la puissance des données utilisées pour réaliser l'intégralité de nos calculs.

EH a établis sept différents aspects des données à étudier :

4.1.1 Complétude

Complétude (*completeness*) : une attente de complétude indique que l'on doit affecter des valeurs à certains attributs du jeu de données. Complétude les règles peuvent être assignées à un ensemble en deux niveaux de contraintes données :

- Les caractéristiques obligatoires qui requièrent une valeur,
- les caractéristiques facultatives qui peuvent avoir une valeur basée sur un certain ensemble de conditions.

Exemple d'indicateur : «Pourcentage d'identifiants de contrats manquant dans le rapport sur les contrats envoyé sur une base mensuelle".

4.1.2 Conformité

Conformité (*conformity*) : Cette dimension se réfère à savoir si les données sont stockées, échangées, ou présentées dans un format qui est compatible avec le domaine des valeurs, ainsi que compatible avec d'autres valeurs d'attributs similaires. Chaque colonne a de nombreux attributs de métadonnées associées avec elle : son type de données, sa précision, les modèles de format, l'utilisation d'une énumération de valeurs prédéfinie, gammes de domaine, les formats de stockage sous-jacents, etc.

Exemple de l'indicateur : "Dans les fichiers de données, est ce que le taux de recouvrement est calculé à partir des montants recouverts des fichiers sinistres ?"

4.1.3 Cohérence

Cohérence (*Consistency*) : dans sa forme la plus basique, la cohérence se réfère à des valeurs de données dans un ensemble de données étant compatible avec les valeurs dans un autre ensemble de données. Une définition stricte de cohérence précise que deux valeurs de données tirées des ensembles de données distincts ne doivent pas entrer en conflit avec l'autre, bien que la cohérence ne signifie pas nécessairement l'exactitude. La cohérence peut être définie dans différents contextes :

- Entre un ensemble de valeurs d'attributs et un autre attribut défini dans le même enregistrement (cohérence niveau des enregistrements)
- Entre un ensemble de valeurs d'attributs et un autre ensemble d'attributs différents dossiers (cohérence entre enregistrements)
- Entre un ensemble de valeurs d'attributs et le même attribut mis dans le même enregistrement à différents points dans le temps (cohérence temporelle)
- La cohérence peut également prendre en compte la notion de «raisonnable», dans lequel une certaine gamme d'acceptabilité est imposée sur les valeurs d'un ensemble d'attributs.

Exemple de l'indicateur : "L'écart moyen en pourcentage entre les taux de recouvrements fournis par l'entité légale et les taux de recouvrements calculés par le groupe sur chaque segmentation (Grade * Catégorie d'exposition * domestique / Export)"

4.1.4 Précision

Précision (*Accuracy*) : la précision des données se rapporte au degré avec lequel les données représentent correctement les objets «vie réelle» qu'ils sont destinés à modéliser. Dans de nombreux cas, la précision est mesurée par la façon dont les valeurs concordent avec une source d'informations correcte (telles que des données de référence) identifiées. Il y a différentes sources d'information correcte : une base de données de l'enregistrement, un ensemble similaire de corroboration de valeurs de données d'une autre table, les valeurs calculées dynamiquement, ou peut-être le résultat d'un processus manuel.

Exemple de l'indicateur : "rapprochement des positions entre les actifs et les informations dépositaire de base de données interne (Soliam, SunGard ...)"

4.1.5 Duplication

Duplication (*Duplication*) : La dimension de la duplication est caractérisé en déclarant que les données existent plus d'une fois dans le jeu de données. Cette dimension peut être contrôlée de deux manières. Comme une évaluation statique, elle implique l'application d'une double analyse de l'ensemble afin de déterminer si les enregistrements en double des données existent, et comme un processus continu de surveillance, il suppose de fournir une adaptation d'identité et un service de résolution au moment de la création de l'enregistrement pour localiser les dossiers de correspondance exacte ou potentiels. Exemple de l'indicateur : «Pourcentage des sinistres dupliqués dans le fichier sinistre"

4.1.6 Intégrité

Intégrité (*Integrity*) : L'intégrité renvoie à des règles associées à l'intégrité référentielle qui sont souvent manifestés comme des contraintes contre la reproduction (pour garantir que chaque entité est représentée une fois, et une seule fois), et des règles d'intégrité de référence, qui affirment que toutes les valeurs utilisées se réfèrent à un fichier principal existant. Exemple de l'indicateur : «Pourcentage de mauvais codes d'industrie (non référencés dans la table ISO)"

4.1.7 Exactitude Temporelle

Exactitude temporelle (*Timeliness*) : L'exactitude temporelle se réfère à l'âge des données qui répond à l'exigence de l'utilisateur. Exemple de l'indicateur : "la date de réception du fichier sinistres"

4.1.8 Répartition des différents tests

L'ensemble des tests réalisés sur les données utilisées pour modéliser le risque d'assurance-crédit chez EH sont disponibles en annexe. Ces tests sont réalisés à différents niveaux par différents acteurs. Nous avons les services informatiques du groupe (*Group IT*) et des entités (*Local IT*) qui réalisent une grande partie de ces tests. Mais il y a également des équipes du risque locales (*Legal Entity*) qui vérifient de nombreux points. Pour finir au groupe (*GRC*), il y a une personne de l'équipe pilier II qui réalise les tests suivants :

#	Axis	Data to check	KPI Definition	Scale
KPI's from the Policy Report (Frequency: on a monthly basis, Reference date: end of last month)				
177	Timeliness	Policy Report	Reception Date of the file	On time
199	Timeliness	Group Policy Report	Reception Date of the file	On time
208	Timeliness	Claims Report	Reception Date of the file	On time
238	Timeliness	Unnamed Claims Report	Reception Date of the file	On time
248	Timeliness	CBIS Report	Reception Date D of the file	On time
249	Integrity	Country Code	Percentage of wrong country code (not referenced in the ISO table)	Not Applicable
250	Integrity	Industry Code	Percentage of wrong Industry code (not referenced in the ISO table)	Not Applicable
261	Timeliness	Data Requirements File	Reception Date D of the file	On time
262	Consistency	Exposure	Gap in percentage between the total exposure provided by the Business Unit and the total exposure computed from the CBIS exposure file	0.0%
263	Consistency	Exposure	Average Exposure relative Gap in percentage between the exposure provided by the Business Unit and the exposure computed from the CBIS exposure file on each segmentation (Grade * Exposure Class * Domestic / Export)	0.0%
264	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Domestic Protracted PD provided by the Business Unit and the Global Domestic Protracted PD computed by GRC	0.0%
265	Consistency	UGD	Average relative Gap between the UGD provided by the Business Unit and the UGD computed by GRC on each segmentation (Grade * Exposure Class * Domestic / Export)	0.1%
266	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Domestic Insolvency PD provided by the Business Unit and the Global Domestic Insolvency PD computed by GRC	0.0%
267	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Export Insolvency PD provided by the Business Unit and the Global Export Insolvency PD computed by GRC	0.0%
268	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Export Protracted PD provided by the Business Unit and the Global Export Protracted PD computed by GRC	0.0%
269	Consistency	Recollection Rate	Average relative Gap in percentage between the recollection rate provided by the Business Unit and the recollection computed by GRC on each segmentation (Grade * Domestic / Export)	0.0%
270	Consistency	Salvages Rate	Average relative Gap in percentage between the salvages rate provided by the Business Unit and the salvages rate computed by GRC on each segmentation (Grade * Domestic / Export)	0.2%
284	Timeliness	Sign-off of all Credit Insurance files	General sign-off of all Credit Insurance reports (Policy, Group Policy, Claims, Unnamed Claims, ReBus and Data Requirement) on the upload platform from the local responsible received on time?	Yes
285	Integrity	Credit Insurance files	Check whether new Credit Insurance reports (Policy, Group Policy, Claims, Unnamed Claims, ReBus and Data Requirement) have been uploaded on the upload platform	Yes
286	Timeliness	Sign-off of CI model Results	General sign-off of CI model Results from the local responsible received on time?	Not Applicable

C'est pourquoi, nous avons confiance quant à la qualité des données que nous utilisons dans nos études.

4.1.9 Qualité des données World Agency

Au niveau du calcul des paramètres cependant, il faut établir des paramètres pour WA. Pour cela, nous devons effectuer de nombreux échanges avec les équipes de WA afin de bien expliquer les données attendues. Nous avons reçu une première base de donnée en mars 2015 et une seconde en juillet 2015.

Dans la première, nous avons trois points sur lesquels nous ne sommes pas satisfait :

- la date de facturation
- la séparation entre sinistres dénommé / non dénommé
- la séparation entre les défauts *insolvency* / *protracted*

Pour remédier à ces problèmes nous avons dans un premier temps poser certaines hypothèses :

La date de facturation

Cette information ne nous a pas été communiqué au sein de la première base. Cependant, nous avons déjà eu ce problème sur l'EL Italie, après de nombreuses estimations, il a été accordé que retirer 195 jours à la date de déclaration de sinistres était une mesure prudente et recevable. Nous avons donc appliquer cette mesure aux sinistres de WA en attendant de disposer d'informations plus précises.

La séparation entre sinistre dénommé et non-dénommé

Concernant la distinction entre sinistres dénommés et non-dénommes, nous regardons pour chaque sinistre si le couple assuré/acheteur dispose d'une limite dans les 12 mois précédents la date de déclaration du sinistre. Si c'est le cas, nous considérons alors le sinistre comme dénommé. Nous obtenons les résultats suivants :

- En nombre de sinistres : 48.1% de dénommé.
- En montant assuré : 70.4% de dénommé.
- En montant indemnisé : 75.2% de dénommé.

Ces résultats sont satisfaisants, car ils sont cohérents avec ceux des EL.

la séparation entre les défauts *insolvency*/*protracted*

Concernant la séparation entre les défauts *insolvency* et *protracted*, nous nous sommes basés sur le délai avant indemnisation après déclaration du sinistre. L'idée est que les défauts *protracted* étaient indemnisés sous 30 jours, tandis que les défauts *insolvency* étaient indemnisés 6 mois après la date de déclaration de sinistres.

Délai de paiement après déclaration en nombre de sinistres dénommés :

- ≤ 30 Jours : 7%
- > 30 jours et ≤ 180 jours : 36%
- > 180 jours : 58%

Il persiste 36% d'incertitude, nous avons tenté d'ajuster le nombre de jours, mais cette méthode ne nous satisfaisait pas.

Après plusieurs échanges, WA a trouvé un indicateur qui traduisait le statut de défaut au sein de leur base. Ils ont pu nous fournir ces distinctions dans la dernière version que nous avons reçu.

Nous disposons ainsi de données de qualité qui ont été vérifiées par les équipes de la WA dans un premier temps et qui ont été ajustées pour nos besoins en se basant sur des hypothèses qui ont été retenues.

4.2 Préparation des bases de données

A ce stade de l'étude, nous travaillons à partir de la dernière version du fichier sinistre. Nous avons obtenu un nouveau fichier fin juillet 2015. Ce fichier est qualitativement meilleur que le premier, car WA a fait les ajustements que nous lui avons demandé, nous avons donc fait le choix de prendre ce nouveau fichier comme source pour l'ensemble des études qui vont suivre. La nouvelle base est sur une autre plage de données allant de janvier 2012 à juillet 2015, cela nous permet de créer nos bases de données pour calculer les paramètres sur la période allant de janvier 2012 à janvier 2014.

Afin de calculer les différents paramètres du modèle interne, il nous faut un fichier contenant l'exposition d'EH et un fichier contenant les sinistres.

Le fichier d'exposition est obtenu à partir de fichiers extraits d'un logiciel interne à EH. On obtient un fichier contenant l'ensemble des contrats avec le code l'EL, l'identifiant du contrat, l'identifiant de l'acheteur et l'exposition de ce contrat sur cet acheteur. On obtient également un fichier contenant les expositions globales par acheteur, dans ce fichier, nous disposons d'informations plus détaillées sur chacun des acheteurs, à partir de leur identifiant, nous avons le nom de la compagnie, son pays, son type d'industrie, son chiffre d'affaires, sa limite accordée, son nombre de limite, sa note, et son identifiant en tant que groupe. A partir du premier fichier, nous avons calculé l'exposition totale sur chacun des acheteurs en sommant les expositions des différents contrats correspondant à leur identifiant, puis ensuite nous récupérons les informations nécessaires dans le deuxième fichier (grade, industrie, pays,...). Nous obtenons ainsi une table des acheteurs

dans laquelle à partir de l'EL du client et le pays de l'acheteur, nous pouvons déduire s'il s'agit d'un contrat domestique ou export.

Les fichiers sinistres quant à eux sont fournis par les EL. Ils disposent d'un historique long (> 6 ans) voir très long pour certaines EL. Il y a deux types de sinistres à considérer, chacun formant un fichier, les sinistres dénommés et les sinistres non dénommés comme cela a été précisé auparavant. Les calculs seront réalisés uniquement à partir des sinistres dénommés. Ce fichier doit être retravaillé par notre équipe avant d'être utilisable. En effet, nous devons récupérer les limites accordées en début d'année. Nous devons également déduire le statut du défaut à l'aide du statut actuel et de celui de 6 mois auparavant. La règle est de prendre le pire des cas.

WA traite l'ensemble de ses données en euros, alors que les EL au Royaume-Uni et aux Etats-Unis traitent leurs opérations dans leurs devises locales. Ainsi, dans le cadre de cette étude, nous avons choisi de convertir l'ensemble des données en euros. Pour cela, nous avons utilisé les taux de change mensuels de l'OCDE.

Nous voulons étudier l'impact de la prise en compte de la WA dans le calcul des paramètres. Pour cela, nous disposons de plusieurs bases d'exposition et de sinistres : Pour chacun des 5 segments étudiés nous avons :

- Une base par EL.
- Une base par segment de WA de cette EL.
- Une base comprenant les deux précédentes.

Nous avons également, pour traiter le cas de WA comme une EL à part entière créé une base WA entière.

Une fois ces bases établies, nous avons pu calculer les différents paramètres (notamment : PD et UGD) à l'aide des codes établis par EH. Ces codes ont dû être retouchés, car nous sommes dans une architecture bien différente de l'architecture habituelle consistant à établir un jeu de paramètre unique par EL.

4.3 Description des bases utilisées

Dans cette partie, nous cherchons des points communs entre ces bases de données et établissons si elles sont concordantes. Il va également être possible d'observer, si WA dispose d'une part d'exposition suffisante pour être intégrée au processus de calcul des paramètres.

4.3.1 Etude des expositions

Tout d'abord, nous allons observer l'évolution des expositions de WA. En 4 ans, son exposition a été multiplié par 4. En effet cette entité croît de façon impressionnante au cours de ces dernières années et c'est ce qui nous pousse à étudier de nouvelles façons de calibrer ces paramètres.

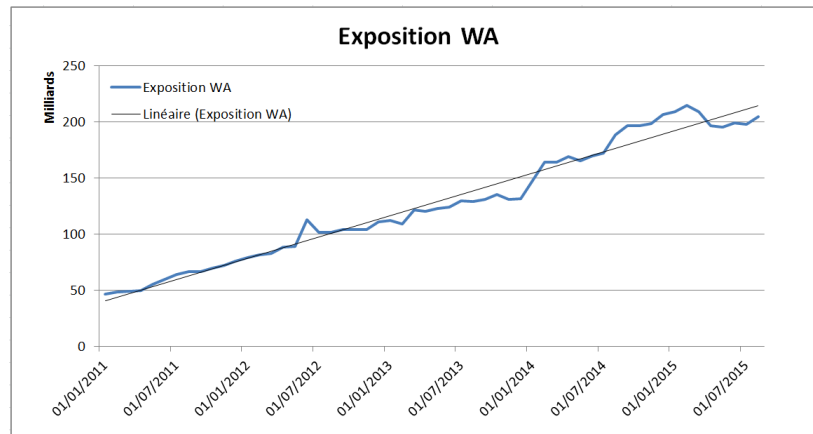


FIGURE 4.1 – Evolution de l'exposition WA

Lorsque l'on segmente WA par EL, il y a deux aspects intéressants à observer sur les expositions, d'un côté leur croissance, et d'un autre côté, la différence de croissance entre l'EL et WA, afin d'observer le comportement de la part de WA dans une entité virtuelle ALL regroupant l'EL et WA.

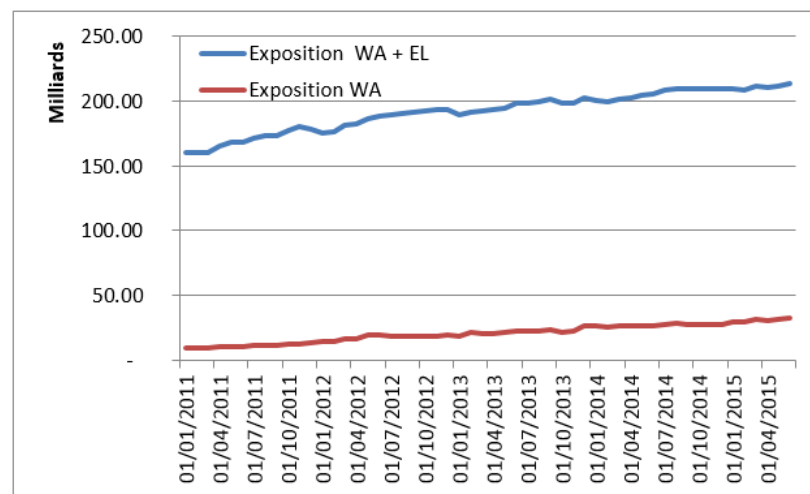


FIGURE 4.2 – Evolution de l'exposition EL et WA en Allemagne

Concernant l'évolution des expositions, nous constatons une croissance plus ou moins forte selon les EL, assez stable pour la France et très forte pour le Royaume-Uni (cf. annexe B.3). Concernant l'Allemagne, nous observons une croissance moyenne, il est difficile de se rendre compte si l'évolution de l'exposition de WA est plus forte que celle de l'EL. Sur le graphique de l'Allemagne, nous observons que l'exposition regroupant les deux échantillons croît de manière similaire à la croissance de WA.

Mais lorsque l'on regarde de plus près la part de WA dans L'exposition totale, nous

nous rendons compte que celle-ci est en forte croissance. Cela est normal pour les années les plus anciennes, car l'entité WA est jeune, mais cela se conserve sur ces dernières années comme nous pouvons le voir sur la figure 4.3. En traçant les tendances linéaires, nous observons même que cela augmente de manière régulière. Pour le Royaume Uni, suite à un pic d'exposition en 2014, le portefeuille WA est aussi grand que celui de l'EL. La croissance de WA nous permet d'obtenir plus de données et de faire des calculs qui auront plus de sens. Cela permet d'envisager une intégration de WA en tant qu'EL virtuelle au sein du modèle interne voire même si l'exposition augmente toujours de manière soutenue, un calcul de paramètres pour WA segmenté par les pays où elle exerce son activité d'assurance.

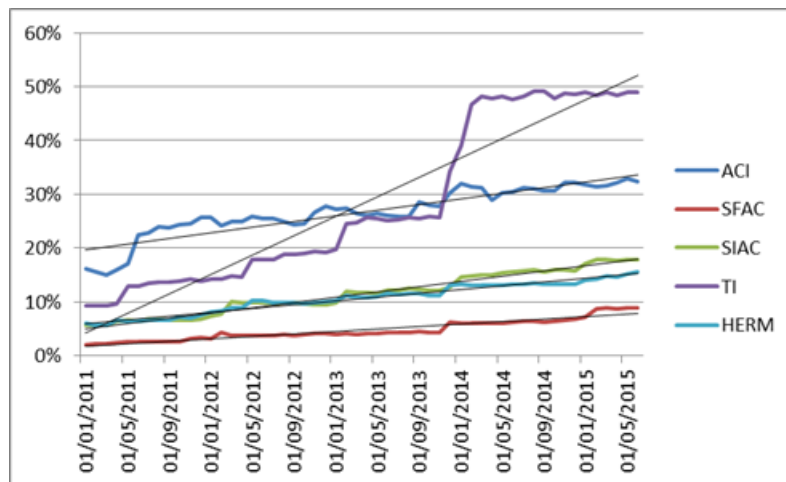


FIGURE 4.3 – Evolution de la part de WA dans les 5 principales EL

4.3.2 Etude de la sinistralité

Nous regardons la répartition des sinistres dans les différents portefeuilles afin d'établir des points communs ou des différences permettant d'utiliser des méthodes similaires dans le processus de calcul des paramètres.

Nous commençons par observer les proportions de sinistres sur acheteurs non dénommés dans les différents portefeuilles. En effet, les calculs étant exclusivement effectués sur les sinistres sur acheteurs dénommés, il faut que ces derniers constituent une part conséquente de la base. Il y a deux moyens d'observer cette répartition, soit nous regardons le nombre de sinistres déclarés, soit nous observons les montants totaux.

Dans le premier cas, nous observons que sur les segments WA de la France et de l'Italie, il y a une faible proportion de sinistres sur acheteurs dénommés, en 2013, et en 2014. Sur le Royaume-Uni, nous observons l'effet inverse, en effet il y a peu de sinistres dénommé sur l'EL et le double sur WA. Sur les autres segments cela représente plus de la moitié de la base et est cohérent avec les pourcentages des EL correspondantes. Au niveau des EL, nous observons des répartitions intrigantes. En effet pour la France et le

EL	année	sinistres WA dénomés	sinistres EL dénomés	montant WA dénomés	montant EL dénomés
SFAC	2012	48%	29%	69%	65%
	2013	26%	28%	78%	62%
	2014	31%	28%	75%	62%
SIAC	2012	46%	92%	74%	97%
	2013	37%	90%	80%	96%
	2014	30%	89%	67%	94%
HERM	2012	75%	56%	95%	77%
	2013	73%	55%	71%	88%
	2014	66%	53%	88%	83%
TI	2012	71%	31%	78%	68%
	2013	65%	32%	47%	67%
	2014	60%	36%	64%	71%
ACI	2012	75%	58%	96%	90%
	2013	53%	60%	25%	91%
	2014	69%	65%	96%	94%

FIGURE 4.4 – Répartition des sinistres sur acheteurs dénomés et non dénomés

Royaume-Uni, le nombre de sinistre sur acheteur dénomé constitue jusqu'à moins d'un tiers de tous les sinistres.

Il faut prendre en compte le fait que le nombre de sinistre est une donnée sensible à la gestion des EL. Si l'on reçoit deux sinistres dénomés pour le même contrat entre un acheteur et un assuré sur une courte période, nous aurons tendance à les rassembler en un seul sinistre. Tandis que dans le cas des sinistres sur acheteurs non dénomés, ce rassemblement ne peut pas être effectué. C'est pourquoi nous privilégierons dans la suite de nos études les répartitions en montant de sinistres.

En matière de montant, nous observons des proportions beaucoup plus raisonnables. En effet, les EL ont au minimum 62 % de leurs sinistres (en montant à partir de maintenant) sur des acheteurs dénomés et cela va jusqu'à 96% en moyenne sur les trois années en Italie. En ce qui concerne WA, il y a un segment qui ne dispose pas d'une part de sinistres sur acheteurs dénomés suffisante. Sur l'année 2013, les Etats-Unis ont une part de sinistres sur acheteurs dénomés de seulement 25%. Nous devons regarder plus en détails ce segment, car sur les autres années les parts de sinistres sur acheteurs dénomés en montant sont toutes supérieures à 90%. On remarque que ce segment est très petit, nous avons peu de sinistres que ce soit en nombre ou en montant. De plus, nous observons un sinistre extrême sur acheteurs non dénomés. Lorsque nous retirons uniquement ce sinistre, nous arrivons à une proportion d'acheteurs dénomés de 47% ce qui est déjà plus acceptable. Nous faisons donc l'hypothèse que nous pouvons calculer des paramètres sur ce segment.

Dans un second temps, nous nous sommes concentrés sur les sinistres sur acheteurs dénomés, car c'est à partir de ces sinistres que seront établis les différents paramètres du modèle interne. Nous avons étudié les répartitions selon les deux premières segmentations d'EH à savoir la localisation de l'acheteur et le statut du défaut.

Sur la répartition entre domestique et export, nous observons qu'en dehors du segment

EL	Echantillon	Année	domestique	Domestique Protracted	Export Protracted
ACI	WA	2012-2014	11%	43%	100%
	EL	2012	84%	70%	85%
		2013	82%	71%	100%
		2014	80%	72%	100%
HERM	WA	2012-2014	46%	27%	54%
	EL	2012	70%	13%	47%
		2013	78%	10%	52%
		2014	72%	15%	48%
SFAC	WA	2012-2014	73%	17%	65%
	EL	2012	64%	31%	80%
		2013	63%	26%	67%
		2014	75%	16%	47%
SIAC	WA	2012-2014	57%	77%	65%
	EL	2012	64%	88%	83%
		2013	67%	86%	84%
		2014	46%	90%	91%
TI	WA	2012-2014	70%	60%	81%
	EL	2012	83%	8%	100%
		2013	70%	13%	100%
		2014	78%	23%	100%

FIGURE 4.5 – Répartition des sinistres selon la localisation et le type de défaut

France, WA a beaucoup plus de sinistres à l'export. C'est assez logique, car WA concerne des clients internationaux qui exercent des activités dans de nombreux pays alors qu'ils ne sont rattachés qu'à un seul pays. Ainsi, nous observons qu'ils ont bien moins de sinistres domestiques. WA étant déjà segmenté par pays, nous ne disposons pas d'assez de données de sinistres pour avoir des chiffres robustes par année. Cela montre déjà une faiblesse au niveau de la quantité de données disponible sur WA. Sur les Etats-Unis nous observons même un chiffre fort de 89% des sinistres à l'export, celui-ci confirmant le fait que WA assure des clients internationaux faisant des transactions partout à travers le monde notamment en dehors du périmètre de l'EL correspondante.

Concernant la répartition sur le statut du défaut *insolvency* ou *protracted*, il faut tout d'abord préciser que le Royaume-Uni ne couvre pas les sinistres *protracted* des acheteurs à l'export. On observe que cette répartition est assez différente selon les EL et la localisation. En effet pour la France la majorité des sinistres domestiques sont *protracted* et ceux qui sont à l'export sont majoritairement *insolvency*. Alors que pour l'Italie que ce soit en domestique ou à l'export, la proportion des sinistres *insolvency* est très majoritaire (+ de 80%). En excluant l'Allemagne, nous observons sur les sinistres à l'export, une très forte concentration sur le statut *insolvency*. Si l'on se concentre sur les différences entre WA et les EL. Une différence notable est observable sur le Royaume Uni, en 2012, en domestique, nous observons une répartition des défauts inversée. Sinon dans la globalité, il n'y a pas de cas alarmant. Sur le segment export des Etats-Unis de

WA, nous sommes presque à 100 % sur *insolvency*, mais cela peut être dû aux clauses de leurs contrats ou sur la gestion de leurs sinistres. Au niveau de l'EL, nous retrouvons également cette tendance à avoir les sinistres concentré sur le segment *insolvency* lorsque l'on est à l'export.

Globalement, nous n'avons pas observé de différences significatives pouvant avoir un impact sur le calcul des paramètres. Nous avons une part de sinistre sur acheteurs dénommés suffisante pour calculer des paramètres à appliquer au portefeuille entier. Maintenant nous décrivons les tests qui vont par la suite être utilisés afin de comparer les paramètres calculés dans les différentes bases.

Nous avons également procédé à cette étude sur la base sinistre WA totale, nous avons eu une proportion dénommée / non dénommée suffisante cf. figure 4.6. Concernant la répartition sur une segmentation plus fine, il n'est pas possible de définir le statut de domiciliation de l'acheteur. En effet WA couvrant de nombreuses entités, il y a des acheteurs qui pourraient avoir le statut domestique sur certaines transactions et export sur d'autres, hors ce statut est censé être propre à l'acheteur et non seulement à la transaction. C'est pourquoi nous considérons l'ensemble des acheteurs de WA comme acheteurs à l'export. Au niveau de la répartition *insolvency* / *protracted*, nous observons qu'en 2013 avec seulement 38% des sinistres en *protracted*, il y a eu une répartition inverse à celle des deux autres années.

année	sinistres WA dénommés	montant WA dénommés	Montant protracted
2012	57%	82%	62%
2013	40%	77%	38%
2014	48%	69%	73%

FIGURE 4.6 – Répartition des sinistres WA

4.4 Description des tests utilisés

Afin de procéder aux différentes comparaisons des paramètres, nous avons utilisé différents tests. Les tests rencontrés dans notre formation supposent généralement que la variable aléatoire X étudiée suit une loi normale dans les populations considérées. Cette hypothèse n'est pas toujours vérifiée, notamment dans le cas des PARS qui n'ont que des valeurs positives et semblent plutôt suivre une loi gamma. Ainsi, nous travaillons avec les tests non-paramétriques de Kolmogorov-Smirnov et Wilcoxon.

La *p-value* des tests donnera le niveau de rejet de l'hypothèse H_0 . Si la *p-value* est inférieure à 5%, nous rejetons l'hypothèse nulle en faveur de l'hypothèse alternative, et le résultat du test est déclaré « statistiquement significatif ». Dans le cas contraire, si la *p-value* est supérieure au seuil, nous ne rejetons pas l'hypothèse nulle, et nous ne pouvons rien conclure quant aux hypothèses formulées.

Dans un premier point, nous présenterons les tests retenus. Nous souhaitons tester si les échantillons des différentes bases ont le même comportement et s'ils peuvent être

issus d'une même population. L'objectif principal de ces tests est de voir si la différence entre les échantillons est significative.

4.4.1 Test de Box-Pierce

Ce test repose sur la Statistiques de Box-Pierce [13] :

Le test de Box Pierce permet d'identifier les processus de bruit blanc (suite de variables aléatoires de même distribution et indépendantes entre elles). Nous devons donc identifier si $\forall k, p_k = cov(y_k, y_{t-k}) = 0$

Un processus de bruit blanc implique que $p_1 = p_2 = \dots = p_h = 0$, soit les hypothèses :

- $H_0 : p_1 = p_2 = \dots = p_h = 0$
- $H_1 : \text{Il existe au moins un } p_i \text{ significativement différent de } 0.$

Pour effectuer ce test, on recourt à la statistique de BoX-Pierce Q qui est donnée par :

$$Q = n \sum_{k=1}^h \hat{P}_k^2 \quad (4.1) \quad \text{où} \quad \begin{aligned} &\bullet \text{ h = nombre de retards,} \\ &\bullet \text{ } P_k = \text{autocorrélation empirique d'ordre k,} \\ &\bullet \text{ n = nombre d'observations.} \end{aligned}$$

La statistique Q est distribuée de manière asymptotique comme une loi du Khi-carré à h degrés de liberté. Nous rejetons donc l'hypothèse de bruit blanc, au seuil α , si la statistique Q est supérieure à celle du Khi-carré lu dans la table au seuil $(1-\alpha)$ et h degrés de liberté. Ce test est appelé par les anglosaxons : « *portmanteau test* » soit littéralement test «fourre-tout».

4.4.2 Tests des *runs*

Le test des *runs* [4], peut être utilisé pour décider si un ensemble de données provient d'un processus aléatoire. Un *run* est défini comme une série de valeurs croissantes ou une série de valeurs décroissantes. Le nombre de valeurs augmentant, ou à la baisse, est la longueur du *run*. Dans un ensemble de données aléatoires, la probabilité que la $(i+1)^{ime}$ valeur soit supérieure ou inférieure à la i^{ime} valeur suit une loi binomiale, qui constitue la base du critère des tests des *runs*.

LA première étape dans le test des *runs* consiste à compter le nombre de *runs* dans la séquence de données. Il existe plusieurs façons de définir les *runs* dans la littérature, toutefois, dans tous les cas, la formulation doit produire une séquence de valeurs dichotomique. Par exemple, une série de 20 lancers de pièces pourrait produire la séquence de Pile (P) et de Face (F) suivante :

P P F F P F P P P F P P F F F F F P P

Le nombre de *runs* pour cette série est de neuf. Il ya 11 piles et 9 faces dans la séquence.

Nous allons coder des valeurs supérieures à la médiane comme positives et les valeurs inférieures à la médiane comme négatives. Un *run* est défini comme une série de valeurs positives (ou négatives) consécutives. Le test des *runs* est défini comme :

- H_0 : la série a été générée de manière aléatoire
- H_a : la série n'a pas été générée de manière aléatoire
- Statistique du test :

$$Z = \frac{R - \bar{R}}{s_R} \quad (4.2) \quad \text{où} \quad \begin{array}{l} - R \text{ est le nombre de } runs \text{ observé,} \\ - \bar{R} \text{ est le nombre de } runs \text{ attendu,} \\ - s_R \text{ est l'écart-type du nombre de } runs. \end{array}$$

$\bar{R}$ et s_R s'obtiennent par les formules suivantes :

$$\bar{R} = \frac{2n_1n_2}{n_1 + n_2} + 1 \quad \text{et} \quad s_R = \frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2(n_1 + n_2 - 1)} \quad (4.3)$$

Le test des runs rejette l'hypothèse nulle si :

$$|Z| > Z_{1-\frac{\alpha}{2}} \quad (4.4)$$

Pour un grand échantillon de tests des runs ($n_1 > 10$ et $n_2 > 10$), la statistique du test est comparée à une table normale standard. Autrement dit, au niveau de signification de 5%, une statistique de test avec une valeur absolue supérieure à 1,96 indique une génération non aléatoire. Pour un petit échantillon court essai, il ya des tables pour déterminer les valeurs critiques qui dépendent des valeurs de n_1 et n_2 [2].

4.4.3 Test d'indépendance du Khi-carré

Le test d'indépendance du Khi-carré (l'écriture anglaise est « chi-square ») a été développé par Pearson[16] L'expression test du Khi-carré recouvre plusieurs tests statistiques, dont trois principaux :

- le test d'ajustement ou d'adéquation, qui compare globalement la distribution observée dans un échantillon statistique à une distribution théorique, celle du Khi-carré.
- Le test d'indépendance du Khi-carré qui permet de contrôler l'indépendance de deux caractères dans une population donnée.
- le test d'homogénéité du Khi-carré qui teste si des échantillons sont issus d'une même population.

Le test qui nous intéresse ici est uniquement le test d'indépendance du Khi-carré. Ce test sert à apprécier l'existence ou non d'une relation entre deux caractères au sein d'une population, lorsque ces caractères sont qualitatifs où lorsqu'un caractère est quantitatif et l'autre qualitatif, ou bien encore lorsque les deux caractères sont quantitatifs mais que les valeurs ont été regroupées.

Ce test s'effectue sur des variables dont les valeurs ont été regroupées en plusieurs classes.

Soient deux variables aléatoires X et Y, dont on veut tester l'indépendance :

- X est à I modalités correspondants aux classes $C_1, \dots, C_I$
- Y est à J modalités correspondants aux classes $C_1, \dots, C_J$

On note :

- $P_X(i) = P(X \in C_i) \forall i = 1 \dots I$
- $P_Y(j) = P(Y \in C_j) \forall j = 1 \dots J$
- $P_{X,Y}(i, j) = P(X \in C_i, Y \in C_j) \forall i, j$

Les Hypothèses du test sont :

- H_0 : les variables X et Y sont indépendantes : $\Leftrightarrow P_{X,Y}(i, j) = P_X(i)P_Y(j), \forall i \in 1, \dots, I, \forall j \in 1, \dots, J$.
- H_1 X et Y sont liées : $\Leftrightarrow \exists i, j$ tels que $P_{X,Y}(i, j) \neq P_X(i)P_Y(j)$

Ce test prend en entrée un tableau de contingence des observations. C'est un tableau qui contient les effectifs observés pour chaque événement $\{X \in C_i, Y \in C_j\}$, $\forall i \in 1, \dots, I, \forall j \in 1, \dots, J$. N est l'effectif total de la population observée. Exemple : N_1 est

l'effectif de l'événement $\bigcup_{j=1}^T \{X = 1, Y = j\}$

Les tableaux de contingences dont nous disposons pour notre étude sont de la forme suivante :

TABLE 4.1 – tableau de contingence					
Classes	A	B	C	D	E
PD échantillon 1	N_{A1}	N_{B1}	N_{C1}	N_{D1}	N_{E1}
PD échantillon 2	N_{A2}	N_{B2}	N_{C2}	N_{D2}	N_{E2}

où N_{A1} correspond au nombre de PD de l'échantillon 1 se situant dans la classe A

Il faut ensuite construire des effectifs théoriques en considérant que l'hypothèse H_0 est vérifiée. L'effectif théorique pour l'événement $\{X = i, Y = j\}$ est alors :

$$N_{ij} = NP(X = i)P(Y = j) \quad (4.5)$$

	1	...	j	...	J	Effectifs marginaux
1	$ \begin{array}{ccccc} N_{11} & & & & N_{1J} \\ & & & & \\ & & N_{ij} & & \\ & & & & \\ N_{I1} & & & & N_{IJ} \end{array} $					$N_{i.}$
$\vdots$						
i						
$\vdots$						
I						
Effectifs marginaux	$N_{.j}$					N

FIGURE 4.7 – Tableau de contingence

On estime $P_X(i)$ par $\frac{N_{i.}}{N}$ et $P_Y(j)$ par $\frac{N_{.j}}{N}$

L'effectif théorique correspondant à l'événement $\{X = i, Y = j\}$ est donc :

$$E_{ij} = \frac{N_{i.}}{N} \frac{N_{.j}}{N} N = \frac{N_{i.} N_{.j}}{N} \quad (4.6)$$

Pour utiliser ce test, selon le critère de Cochran de 1954, il faut que toutes les classes $C_{i,j}$ doivent avoir un effectif théorique non nul ($E_{ij} \geq 1$), et que 80% des classes doivent avoir un effectif théorique supérieur ou égal à 5 : $E_{ij} \geq 5$

Lorsque le nombre de classes est petit, cela revient à dire que toutes les classes doivent contenir un effectif théorique supérieur ou égal à 5.

La statistique de ce test mesure l'écart entre les effectifs observés et les effectifs théoriques. Plus cet écart est grand, moins les variables sont indépendantes. En effet, les effectifs théoriques étant construits sous H_0 , un écart important signifie que les effectifs théoriques et observés sont éloignés et donc que les deux variables n'ont pas de lien entre elles.

La statistique de test s'écrit :

$$Q = \sum_{\substack{1 \leq i \leq I \\ 1 \leq j \leq J}} \frac{(N_{ij} - E_{ij})^2}{E_{ij}} = \sum_{\substack{1 \leq i \leq I \\ 1 \leq j \leq J}} \frac{(N_{ij} - \frac{N_{i.} N_{.j}}{N})^2}{\frac{N_{i.} N_{.j}}{N}} \quad (4.7)$$

On a Q qui suit une loi $\chi^2_{(I-1),(J-1)}$

On rejette H_0 au risque α si $Q > \chi^2_{1-\alpha,(I-1),(J-1)}$ avec $\chi^2_{1-\alpha,(I-1),(J-1)}$ le quantile d'ordre $1-\alpha$ de la loi $\chi^2_{(I-1),(J-1)}$:

$$P \left\{ X \leq \chi^2_{1-\alpha,(I-1),(J-1)} \right\} = 1 - \alpha, \text{ si } X \text{ suit la loi } \chi^2_{(I-1),(J-1)} \quad (4.8)$$

4.4.4 Test d'indépendance de Fisher

Le test exact de Fisher, tout comme celui du Khi-carré repose sur des tableaux de contingences. Il permet de tester l'indépendance dans les cas où les conditions de Cochran 4.4.3 ne sont pas vérifiées. En effet, dans notre cas, nous avons des échantillons ou les PD sont réparties seulement en deux classes. Cela nous fait des tableaux de la forme suivante :

TABLE 4.2 – Exemple de tableau de contingence 2×2

		X1	
		Classes	A B
X2	PD échantillon 1	a	b
	PD échantillon 2	c	d

L'idée du test exact de Fisher est de comparer la distribution observée avec l'ensemble des combinaisons possibles et issues (au sens statistique) d'une distribution aléatoire. Le test exact de Fisher est, à la base, développé pour les tableaux de taille 2×2

Nous pouvons alors présenter les effectifs croisés entre (X_1, X_2) :

Nous sommes donc dans le cadre d'une distribution hypergéométrique et la formule générale pour la statistique du test exact de Fisher est donnée par :

$$p = \frac{(a+b)!(c+d)!(a+c)!(b+d)!}{a!b!c!d!n!} \quad (4.9)$$

Le calcul de la p-valeur est assez laborieuse à obtenir. En effet, il faut comparer la statistique de test associée aux données observées à celles obtenues sur chacune des K configurations conservant les mêmes effectifs marginaux $(a+b)$, $(c+d)$, $(a+c)$, $(b+d)$ que ceux observés et de produit d'effectifs croisés (nommé également « éloignement »),

$$\left| \prod_{\text{config } k \in [1, \dots, K]} \right| = |a \cdot d - c \cdot b|_k \geq \left| \prod_p \right| = |a \cdot d - c \cdot b|_p \quad (4.10)$$

La p-valeur vaut alors, $p + \sum_{k=1}^K p_k$ Où p_k désigne la valeur de la statistique de test calculée sur la configuration numéro k. La p-valeur est comparée directement au seuil de confiance fixé α .

L'hypothèse H_0 est : » Il y a indépendance entre les deux variables « .

Pour les tableau de contingence $> 2 \times 2$ ne vérifiant pas la condition du Cochran, le logiciel se base sur l'extension développée par G. H. Freeman, J. H. Halton. [12]

4.4.5 Test de Kolmogorov-Smirnov

Le test de Kolmogorov-Smirnov (KS) est un test d'hypothèse utilisé pour déterminer si un échantillon suit bien une loi donnée connue par sa fonction de distribution continue,

ou si deux échantillons suivent la même loi [14] [6]. Ce test statistique non paramétrique repose sur les propriétés des fonctions de répartition empirique. Dans notre cas, nous testons les deux échantillons $(x_1, \dots, x_n), (y_1, \dots, y_m)$ ayant pour fonctions de répartitions empiriques :

$$F_n = \frac{1}{n} \sum_{i=1}^n I_{\{X_i \leq x\}} \text{ et } G_m = \frac{1}{m} \sum_{i=1}^m I_{\{Y_m \leq y\}}$$

Le test que nous menons considère $H_0 : F_n = G_m$ contre $H_1 : F_n \neq G_m$

Le test de KS est basé sur la distance :

$$d_{n,m}(F_n, G_m) = \sup_x \|F_n(x) - G_m(x)\|$$

L'hypothèse H_0 est rejetée lorsque : $d_{n,m} > c\sqrt{\frac{n+m}{nm}}$

On a la probabilité que cela arrive appelée la *p-value* à partir de laquelle nous concluons si l'hypothèse H_0 est validée.

$$P \left[d_{n,m} > c\sqrt{\frac{n+m}{nm}} \right] \rightarrow \alpha(c) = 2 \sum_{r=1}^{+\infty} (-1)^{r-1} \exp(-2r^2 c^2) \quad (4.11)$$

On obtient des difficultés pour calculer la *p-value* dans les cas où nous avons des ex-æquo dans nos observations. Lorsque ces cas se présenteront, nous prendrons la *p-value* du test *bootstrap de KS* [18]. Ce test utilise du *Bootstrapping* (rééchantillonnage). Le test de KS obtient la probabilité de la statistique de ce test par la distribution Kolmogorov (distribution expliquant comment les statistiques de tests sont distribuées lorsque deux échantillons sont vraiment issus de la même distribution). Le *bootstrap* obtient la probabilité de la statistique de ce test par la distribution empirique, provenant de l'hypothèse nulle. Autrement dit, les deux échantillons étudiés sont combinés ensemble et de cette union, nous tirons deux nouveaux échantillons au hasard avec remplacements. Ces nouveaux échantillons sont tirés de la même distribution, et leurs statistiques de tests sont notées. On répète plusieurs fois cette procédure, nous obtenons la distribution empirique des statistiques de tests sous l'hypothèse nulle.

On a $\alpha(c)=5\%$ pour $c=1.36$, ainsi pour que l'hypothèse H_0 soit rejetée il faut une *p-value* inférieure à 0.05

4.4.6 Test de Wilcoxon

Frank Wilcoxon a proposé dans un seul article [21] les deux principaux tests qui lui sont attribué, c'est-à-dire un test sur les données appariées et un test sur les données non appariées. Ces deux tests sont dits non paramétriques, car ils ne nécessitent pas d'estimation de la moyenne et de la variance. En fait, ils n'utilisent même pas les valeurs x_i recueillies dans les échantillons, mais seulement leur rang dans la liste ordonnée de toutes les valeurs.

Dans le cas des données non appariées (UGD), ce test est équivalent à celui de Mann-Whitney et compare la moyenne expérimentale des deux échantillons. C'est le "*Wilcoxon rank sum test*".

Dans le cas des données appariées (PD), ce test travaille sur la différence des deux échantillons. C'est le "*Wilcoxon signed rank test*".

Test sur les données non appariées

Ce test est également connu sous le nom de test de Mann-Whitney[8]. Nous détaillons son fonctionnement. On dispose des mesures des valeurs de X dans deux échantillons indépendants E_1 et E_2 , de tailles respectives n_1 et n_2 . On souhaite comparer les deux moyennes expérimentales, c'est-à-dire tester l'hypothèse nulle $H_0 : \mu_1 = \mu_2$.

Pour cela, nous commençons par réunir les deux échantillons.

Si les individus des deux groupes proviennent de la même population, les rangs des individus du premier groupe sont tirés au hasard dans l'ensemble des n ($n=n_1 + n_2$) premiers entiers. Si les moyennes des deux échantillons ne sont pas égales, les rangs du premier groupe auront tendance à être trop grands ou trop petits. La statistique utilisée est la somme des rangs.

Le raisonnement est symétrique sur les deux groupes. Nous choisissons de raisonner sur le premier échantillon. Les valeurs x_i des deux échantillons réunis sont classées dans l'ordre croissant $x_{(i)}$ notés r_i , ce sont les réalisations de variables aléatoires R_i de loi uniforme. La statistique du test est la somme des rangs SR du premier groupe et est appelée statistique de Wilcoxon.

$$SR = \sum_{i=1}^{n_1} R_i \quad (4.12)$$

La statistique de Wilcoxon SR suit approximativement une loi normale de moyenne et de variance ($n_1 > 10$ et $n_2 > 10$) :

$$\mathbb{E}[SR] = \frac{n_1(n+1)}{2} \quad (4.13) \quad \mathbb{V}[SR] = \frac{n_1 n_2 (n+1)}{12} \quad (4.14)$$

Il existe des tables donnant les seuils de signification pour les petits effectifs et les bons logiciels donnent la loi exacte connue s'il n'y a pas d'ex-æquo. Souvent, les logiciels utilisent la statistique U de Mann-Whitney :

$$U = SR - \frac{m(m+1)}{2} \quad (4.15)$$

D'après ce qui précède et sous les mêmes conditions, celle-ci suit approximativement une loi normale de moyenne et de variance :

$$\mathbb{E}[U] = \frac{n_1 n_2}{2} \quad (4.16) \quad \mathbb{V}[U] = \frac{n_1 n_2 (n+1)}{12} \quad (4.17)$$

Cela explique le fait que ces deux tests aient deux noms différents. Nous pouvons ainsi conclure, avec un niveau de confiance 95 % ($\alpha = 0.05$) :

$$\begin{cases} \text{si } |S| > Q_Z(1 - \frac{\alpha}{2}), & (H_0) \text{ doit être rejetée;} \\ \text{sinon,} & \text{nous conservons l'hypothèse nulle avec un risque de niveau } \alpha. \end{cases} \quad (4.18)$$

L'inconvénient de ce test est qu'il ne prend pas en compte l'importance de la différence entre les paires, information qui est pourtant souvent disponible.

Wilcoxon a proposé un autre test qui permet de prendre en compte le niveau de différence à l'intérieur des paires. Ce test est appelé test de rang signé de Wilcoxon (*Wilcoxon signed rank test*), car le signe des différences intervient aussi. Cependant, ce dernier test ne s'applique qu'aux échantillons appariés.

Test sur les données appariées

Pour ce test, nous disposons de deux échantillons appariés E1 et E2, c'est-à-dire que chaque valeur de E1 est associée à une valeur de E2. Ainsi $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ est un échantillon formé de n couples. (x_i, y_i) est un échantillon aléatoire simple d'une loi quelconque L_i . $z_i = x_i - y_i$ est la réalisation d'une variable aléatoire de médiane 0. L'écart z_i mesuré sur le couple i a deux propriétés, son signe et sa valeur absolue. Il se pourrait que le signe soit positif environ une fois sur deux, mais que les différences dans un sens soient plus faibles en valeur absolue que les différences dans l'autre sens.

Le test de Wilcoxon sur les données appariées donne plus de poids à une paire qui montre une large différence entre les deux échantillons, qu'à une paire ayant une faible différence. Cela implique que l'on puisse dire quel membre d'une paire est plus grand que l'autre (donner le signe de la différence)[19], mais aussi que l'on puisse ranger les différences en ordre croissant.

On peut améliorer le raisonnement. Si les différences en valeur absolue sont rangées par ordre croissant, chacune d'entre elles, quel que soit son rang, a une chance sur deux d'être négative (le même raisonnement peut être fait pour positif). Le rang 1 a une chance sur deux de porter le signe -. Le rang 2 a une chance sur deux de porter le signe -. Le rang n a une chance sur deux de porter le signe -.

La somme des rangs (S) qui portent le signe - vaut en moyenne :

$$\mathbb{E}[S] = \frac{1}{2} + 2\frac{1}{2} + \dots + n\frac{1}{2} = \frac{n(n+1)}{4} \quad (4.19)$$

La variance de la somme des rangs qui portent le signe - est :

$$\mathbb{V}[S] = \frac{n(n+1)(2n+1)}{24} \quad (4.20)$$

Dans le cas des « grands » échantillons, lorsque N est supérieur à 25, il peut être

démontré que la somme des rangs S est pratiquement normale ; nous utilisons alors l'approximation normale :

$$Z = \frac{S - \frac{n(n+1)}{4}}{\sqrt{\frac{n(n+1)(2n+1)}{24}}} \quad Z \sim N(0, 1) \quad (4.21)$$

Ainsi, nous pouvons établir les conclusions de la même manière qu'au point 4.18

4.5 Etude de dépendance sur les PD

4.5.1 Création de classes de PD

Pour la réalisation de certains tests, nous aurons besoin de données qualitatives. Pour cela nous avons créé des classes de PD. Dans un premier temps, nous les avons créées à partir de toutes les PD calculées pour notre étude.

Pour cela, nous avons rassemblé toutes les PD au sein d'un échantillon et nous avons pris les quantiles 20%, 40%, 60% et 80%.

Le soucis est que lorsque nous appliquons ensuite ces classes de PD, nous n'observons pas de distinction au sein des échantillons de chaque EL. Nous avons donc conclu que ces classes étaient inadaptées aux données. Ceci est dû au à des écarts de PD entre les différentes EL.

Ensuite au niveau le plus segmenté, nous avons créé des classes à partir des 3 échantillons (WA, EL et ALL). On obtient une plus grande variabilité dans les résultats, cependant les classes ne sont pas assez uniformes et cela biaise les résultats de nos tests.

Après application des tests présentés dans la partie précédente, nous avons pu conclure que la meilleure calibration des classes de PD était au niveau des EL. Nous avons donc calculé des classes de PD par EL. Celles-ci permettent d'obtenir des résultats intéressant en formant le moins de schéma de classes différents.

TABLE 4.3 – Classes de PD par EL

	ACI	HERM	SFAC	SIAC	TI
Borne 1	0.03%	0.20%	0.14%	0.07%	0.10%
Borne 2	0.14%	0.28%	0.23%	0.32%	0.28%
Borne 3	0.50%	0.39%	0.56%	0.64%	0.42%
Borne 4	0.65%	0.52%	0.72%	0.87%	0.65%

A partir du tableau ci-dessus, nous formons 5 classes :

- la classe A comprend le nombre de PD inférieures à la borne 1
- la classe B comprend le nombre de PD comprises entre la borne 1 et la borne 2
- la classe C comprend le nombre de PD comprises entre la borne 2 et la borne 3

- la classe D comprend le nombre de PD comprises entre la borne 3 et la borne 4
- la classe E comprend le nombre de PD supérieures à la borne 4

4.5.2 Dépendance au sein des échantillons

Dans cette partie, nous appliquons les test de runs et test de Box-Pierce définis précédemment afin d'établir d'une part si nos échantillons EL, WA et ALL sont indépendantes et identiquement distribuées (iid et d'autre part si les échantillons obtenus en procédant aux différences EL-WA et EL-ALL (ces échantillons interviennent dans les tests sur données appariées) sont iid ou non. Les *p-values* qui ressortent de ces tests rejettent l'hypothèse que les échantillons soient iid. Il y a trop peu d'échantillons pour lesquels les test donne une p-value supérieure à 0.05.

Les résultats obtenus pour l'Allemagne sont les suivants :

EL	Segmentation	Le test de Box-Pierce					Le test des runs				
		EL-WA	EL-ALL	EL	WA	ALL	EL-WA	EL-ALL	EL	WA	ALL
HERM	Global	6.34E-03	1.27E-02	4.41E-06	6.96E-05	5.58E-06	8.00E-03	6.67E-02	2.62E-06	1.04E-04	2.62E-06
	Domestique	5.00E-05	7.27E-04	5.15E-06	6.63E-05	6.44E-06	1.04E-04	2.49E-02	2.62E-06	2.49E-02	2.62E-06
	Export	9.06E-05	1.01E-03	4.60E-06	2.10E-05	5.96E-06	5.19E-04	6.67E-02	2.62E-06	1.79E-05	1.04E-04
	Domestique Insolvency	1.29E-04	1.82E-05	2.43E-06	1.89E-05	2.44E-06	5.19E-04	1.04E-04	2.62E-06	2.20E-03	2.62E-06
	Domestique Protracted	1.48E-04	1.85E-03	1.41E-04	1.57E-04	4.23E-04	2.20E-03	2.20E-03	2.20E-03	2.20E-03	2.20E-03
	Export Insolvency	2.16E-06	3.01E-05	1.46E-06	1.17E-04	1.68E-06	1.79E-05	1.79E-05	2.62E-06	5.19E-04	2.62E-06
	Export Protracted	7.89E-05	2.78E-03	1.96E-05	2.34E-05	3.28E-05	5.19E-04	8.00E-03	1.04E-04	1.79E-05	5.19E-04

FIGURE 4.8 – Tests de Box-Pierce et tests des runs sur les données de l'Allemagne

Les autres EL ont des réponses similaires à ces deux tests, les résultats sont disponibles en annexes C.1.

Pour le reste de notre étude, nous allons tout de même considérer que ces variables sont iid. Cette hypothèse est très forte, car les tests menés précédemment indiquent le contraire. Cette hypothèse nous permettra d'appliquer de nombreux tests statistiques. Le fait de savoir que cette hypothèse n'a pas été vérifiée sera prise en compte dans nos analyses. En effet, selon les tests, cela diminue les p-values sur les tests entre échantillons non iid.

4.5.3 Dépendance entre les échantillons

Dans cette partie, nous étudions l'indépendance entre les différents échantillons que nous voulons comparer. A chaque niveau de segmentation, nous souhaitons comparer les échantillons EL et ALL d'une part et EL et WA d'autre part, c'est pourquoi nous avons établi les tableaux de contingence de la manière décrite dans la partie décrivant les tests statistiques. Ainsi, nous avons appliqué le test d'indépendance du Khi-carré à de multiples reprises. Cependant, certains résultats paraissaient extrêmes, la conclusion fut que les échantillons n'étaient pas indépendants. C'est pourquoi, nous avons décidé d'appliquer le test de Fisher. En effet, de la façon dont a été implémenté ce test sous le logiciel R, pour les tableaux de contingences supérieurs à 2×2 , lorsque les conditions

du Cochran 4.4.3 sont vérifiées, c'est le test du Khi-carré qui est appliqué et dans le cas contraire c'est le test de Fisher.

Pour l'Allemagne, nous avons obtenus les résultats suivants :

EL	Segmentation	Test de Fisher		Test du khi-carré	
		EL vs WA	EL vs ALL	EL vs WA	EL vs ALL
HERM	Global	0.73	0.01	0.23	0.12
	Domestique	0.00	0.00	0.09	0.12
	Export	0.00	0.00	0.69	0.23
	Domestique Insolvency	0.00	0.00	0.23	0.04
	Domestique Protracted	0.00	0.00	0.69	0.08
	Export Insolvency	0.00	0.00	0.40	0.40
	Export Protracted	0.22	0.74	0.69	0.23

FIGURE 4.9 – Tests d'indépendance sur les données de l'Allemagne

Les autres EL ont des réponses similaires à ces deux tests, les résultats sont disponibles en annexes C.2. On observe bien que les résultats semblent opposés, ceci est dû à la taille des échantillons. Le test du Khi-carré semble accepté l'indépendance, tandis que le test de Fisher la refuse. Il faut prendre en compte que les données au sein de chaque échantillons ne sont pas indépendantes et identiquement distribuées (tests de bruits blancs), ainsi les p-values que nous obtenons pour nos test d'indépendance peuvent se voir sous-estimées.

Le manque de données nous a fait construire de nombreux tableaux de contingence 2×2 pour lesquels le test de Fisher est le mieux adaptés. Les résultats de ce tests seront donc retenus. Malgré le fait qu'il faut augmenté les p-values pour tenir compte du fait que les données ne soient pas iid dans les échantillons, les p-values sont vraiment très faibles et mettent en avant de la dépendance entre les différents échantillons de PD.

Dans le reste de notre étude, nous prendrons en compte cette dépendance en appliquant des test sur données appariés pour comparer les différents échantillons de PD.

4.6 Etude de l'intégration de WA dans les bases des EL

Dans cette partie, nous analyserons l'impact d'une intégration des données d'exposition et de sinistralité de la WA dans le processus de calcul des paramètres. Pour cela, nous nous focaliserons sur les 5 EL les plus importantes en matière d'exposition et de sinistralité

Un objectif de ce mémoire est de voir s'il serait pertinent d'intégrer les données de WA aux différentes EL dans le calcul des paramètres. Les paramètres actuellement calculés pour chaque EL sont attribués à la part de WA liée à cette EL. Nous avons auparavant montré que ces paramètres n'étaient pas toujours en ligne avec le comportement de WA. Pour notre étude, nous travaillons avec 3 échantillons pour chacune des 5 principales EL du portefeuille à savoir l'Allemagne, les Etats-Unis, le Royaume-Uni, l'Italie et la France. Le premier échantillon comprendra les données de l'EL, le second comprendra les données de WA associées à cette EL et le troisième échantillon sera composé des deux précédents.

Ainsi, à l'aide de tests statistiques, nous pourrions mettre en avant les différences entre les deux premiers échantillons dans un premier temps, puis par la suite nous pourrions établir si la différence entre l'échantillon 1 actuellement utilisé et l'échantillon 3 est significative.

Nous exercerons ces tests au niveau global dans un premier temps, puis selon les résultats des tests, nous explorerons plus en profondeur la segmentation actuellement utilisée chez EH.

Nous commencerons au niveau de l'EL, puis nous segmenterons selon le critère domestique/export puis dans un second temps, nous segmenterons selon le statut du défaut : *insolvency* / *protracted*.

Pour les tests que nous avons présentés, il n'y a pas de taille minimum bien définie, car en cas de petite taille, ils ont leurs propres tables et n'utilisent pas l'approximation de la loi normale. Cependant, il est conseillé d'avoir des échantillons d'au moins 4 observations pour appliquer le test de Wilcoxon[1]. Pour le test de KS[7], nous trouvons qu'à partir de 8 observations les résultats peuvent être exploités. Pour nos études nous nous fixons un minimum de 15 valeurs par échantillons. Dans le cas où nous ne disposerions pas de 15 observations dans un échantillon, nous n'appliquerons pas les tests. Pour les PD, nous disposons de données mensuelles sur 2 ans soit 25 PD par échantillon et il n'y a donc aucun problème pour appliquer les tests.

4.6.1 Exemple concret d'application des test de comparaison

Dans un premier temps nous illustrons les tests de comparaisons décrits plus haut par un cas pratique. Les codes R de ces exemples sont disponible en annexe D.1. Ces test sont illustrés sur les PD malgré le fait qu'appliquer un test sur données non-appariés ou un test de Kolmogorov-Smirnov sur des données appariées n'a pas de sens. La méthodologie est la même peu importe les données en entrée, cela permet donc de comprendre le fonctionnement des tests.

Application du test de Kolmogorov-Smirnov

Nous détaillons l'application du test de Kolmogorov-Smirnov entre les PD globales de SIAC (a) les PD globales du segment SIAC de WA (b).

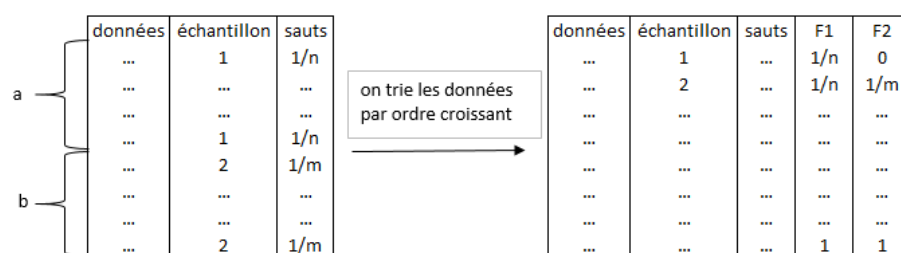


FIGURE 4.10 – Illustration de l'application du test de Kolmogorov-Smirnov aux PD de SIAC

Nous commençons par regrouper les deux échantillons. Nous créons une matrice avec dans la première colonne les observations et dans la deuxième colonne le groupe d'où elles proviennent, 1 pour a et 2 pour b. Ensuite, nous faisons une autre colonne dans laquelle nous indiquons selon l'échantillon le saut à effectuer (l'inverse de sa longueur). Dans le cas des tests sur les PD, les deux échantillons font la même longueur (25), nous mettons donc dans cette colonne des 0.04. Le rang d'une donnée est son ordre d'apparition lorsque les données sont triées dans l'ordre croissant. Ensuite nous trions la matrice en fonction du rang. Ainsi nous pourrions calculer les fonctions de répartition empirique que sont F1 et F2. Afin d'obtenir F1, nous évoluons dans l'ordre des ligne (l'ordre des rangs) et à chaque fois que l'on rencontrera une observation de l'échantillon 1 nous rajouterons 0.04. De la même façon, nous obtenons F2.

Pour obtenir la statistique D, il suffit de prendre le max de la valeur absolue de la différence entre F1 et F2. On obtient $D=0.48$. Pour obtenir la *p-value*, il faut poser :

$$c = D * \sqrt{n * m / (n + m)} \quad (4.22)$$

Nous avons $n=m=25$ et nous obtenons la *p-value* en faisant la formule du $\alpha(c)$ définie en 4.11. Il suffit de prendre la somme de $r = 1$ à 100 car l'exponentielle tend rapidement vers 0 lorsque x tend vers $-\infty$. On obtient une *p-value* de 0.05612483 et le test de R trouve une *p-value* de 0.05614.

Application du test de Wilcoxon apparié

Nous détaillons l'application du test de Wilcoxon sur données appariées entre les PD globales de SIAC (a) et les PD globales du segment SIAC de WA (b).

Lorsque nous procédons à la différence entre a et b nous obtenons 9 différences négatives sur les 25. Cela fait un taux de 36% qui ne nous permet pas de rejeter l'hypothèse qu'il y est une différence sur deux négative, car la probabilité d'obtenir ce résultat avec une loi binomiale est de 11.5% ($\text{pbinom}(9, 25, 0.5)$ sur R).

Nous sommes les rangs des différences positives et ceux des différences négatives nous obtenons 242 pour les écarts négatifs et 83 pour les écarts positifs. Ceci correspond à la statistique D en sortie du test, en sachant que la somme des rangs est de 325, plus nous sommes prêt de 162.5 et plus les échantillons ont un comportement similaire.

Ensuite nous calculons la moyenne et la variance de la loi normale des rangs négatifs cf. 4.19 et 4.20. Nous obtenons ainsi comme valeur normalisée $z=-2.139099$ pour $D=242$. Ce qui nous donne une *p-value* de $0.03242761 = 2*(1-P[Z < z])$ (avec Z loi normale). De même si l'on calcule z avec $D=83$ (les écarts positifs), nous obtenons $z=2.139099$ et de manière symétrique nous arrivons à la même *p-value*. Le test implémenté sous R (`wilcox.test()`) nous donne une p-value de 0.3243.

Application du test de Wilcoxon non apparié

De la même manière, nous détaillons l'application du test de Wilcoxon sur données non-appariées entre les PD globales de SIAC (a) et les PD globales du segment SIAC de

WA (b). La longueur des deux vecteurs est de 25. On les mets dans un même vecteur c de taille 50. On additionne les rang du premier vecteur : nous obtenons 723.

Ensuite nous calculons la moyenne et la variance de la loi normale des rang négatifs cf. 4.13 et 4.14. Nous obtenons ainsi comme valeur normalisée $z=1.658944$ pour $W=723$. Ce qui nous donne une *p-value* de $0.09712714 = 2*(1-P[Z < z])$ (avec Z loi normale). De même si l'on calcule z avec $W=552$ (les écarts positifs), nous obtenons $z=-1.658944$ et de manière symétrique nous arrivons à la même *p-value*.

En se rapportant à la statistique de Mann-Whitney, nous avons exactement les mêmes résultats. Les valeurs de U dans cet exemple sont 398 pour a et 227 pour b. La fonction `wilcox.test()` de R retourne la statistique de Mann-Whitney. Le test implémenté sous R (`wilcox.test()`) nous donne une p-value de 0.9713.

4.6.2 Etude sur les PD

Après avoir retravaillé les fichiers de sortie, nous ne conservons que les PD de janvier 2012 à janvier 2014. Nous disposons donc de 25 observations sur chaque segments, nous avons des PD pour l'ensemble de la base puis séparée en en domestique / export puis en *insolvency/protracted*. L'étape suivante consistant à une séparation en grade ayant peu d'intérêt au vu du nombre d'observation dont nous disposons pour WA.

Résultats

Nous regardons désormais les résultats de l'ensemble des tests.

Les échantillons ayant tous le même nombre d'observation, les statistiques sont visibles à travers les *p-values*. Ainsi, nous porterons nos analyses sur les PD sur les *p-values*.

Tout d'abord, il faut préciser que les valeurs sont arrondies au centième pour une meilleure lisibilité, ainsi les valeurs 0 et 1 ne sont jamais réellement atteintes. Nous nous concentrons sur les tests effectués sur les PD de l'Italie et sur celles l'Allemagne dans un premier temps, car ces segments permettent de voir la plupart des situations rencontrées. Puis nous présenterons les résultats globaux.

EL	Niveau	Wilcoxon apparié		impact
		EL vs WA	EL vs ALL	
SIAC	Global	0.03	0.00	104%
	Domestique	0.11	0.00	105%
	Export	0.01	0.13	101%
	Domestique Insolvency	0.00	0.00	111%
	Domestique Protracted	0.09	0.00	105%
	Export Insolvency	0.00	0.00	105%
	Export Protracted	0.01	0.75	101%

FIGURE 4.11 – Test de Wilcoxon sur données appariées appliqués aux PD de SIAC

Au niveau de l'Italie, nous retrouvons les résultats calculés au travers des exemples pour la comparaison EL vs WA. Au niveau global, nous remarquons que le test de Wilcoxon sur données appariées rejette l'hypothèse que les échantillons EL et ALL aient

des comportements similaires. Cela indique que ce test trouve que l'impact de l'intégration des données WA est assez fort pour modifier significativement le comportement de l'échantillon EL. Plus en détail, nous observons que les PD export ont un comportement similaire entre ALL et EL (p -value de 0.13). On observe sur cette EL que d'une part les PD de WA et de l'EL ont des comportements similaires (2 tests sur 7) et d'autre part que l'intégration de WA modifie le comportement de l'EL (5 tests sur 7)

En termes d'impact, nous observons une hausse globale des PD de 4% sur ce segment, d'après l'étude précédente, cela se refléterait par un impact similaire sur le SCR lié à ce segment, ce qui est acceptable.

Sur ce segment, nous ne pouvons pas envisager de mélanger les données.

EL	Niveau	Wilcoxon apparié		impact
		EL vs WA	EL vs ALL	
HERM	Global	0.00	0.00	110%
	Domestique	0.98	0.00	107%
	Export	0.00	0.00	121%
	Domestique Insolvency	0.00	0.00	101%
	Domestique Protracted	0.00	0.00	115%
	Export Insolvency	0.00	0.00	103%
	Export Protracted	0.00	0.00	130%

FIGURE 4.12 – Test de Wilcoxon sur données appariées appliqués aux PD de HERM

Au niveau de l'Allemagne, les résultats sont beaucoup moins satisfaisants. Au niveau des PD globales, nous observons que le test refuse l'hypothèse de similarité entre les échantillons EL et WA ainsi qu'entre EL et ALL. Le test accepte l'hypothèse que WA et EL est le même comportement sur les PD domestique, cependant ils refusent que ALL et EL aient le même comportement. Sur l'Allemagne, nous avons trop de tests refusant l'hypothèse de similarité des échantillons (13/14) pour considérer pouvoir mélanger les données. Si cela se faisait tout de même nous aurions une augmentation globale de 10% des PD.

Sur les autres EL, les résultats sont disponibles en annexe D.1. Tous les tests ont refusé l'hypothèse de comportements similaires sur le Royaume-Uni. Sur la France, tous les tests de similitude entre WA et EL ont été rejetés. Au niveau des Etats-Unis les résultats ne sont pas satisfaisants non plus, uniquement le segment Export *insolvency* obtient des résultats laissant l'hypothèse de similarité indécise avec une p -value de 3% pour le test de Wilcoxon sur données appariées. Les impacts dans le cas où l'on mélange les données sont neutre pour la France, une faible baisse des PD sur les Etats-Unis, et une forte hausse sur le Royaume-Uni qui voit ses PD doubler.

4.6.3 Etude sur les UGD

Les PD se calculent mensuellement alors que les UGD se calculent annuellement. Nous procéderons à des tests sur la distribution des UGD par année avant d'appliquer la segmentation d'EH. Les données n'étant pas appariées, nous n'appliquerons pas le test de Wilcoxon pour données appariées.

Nous avons un UGD par sinistre, ce qui fait que les échantillons ont une longueur variable. Ainsi le lien entre les statistiques et les p -values dépend de la longueur des échantillons. Nous conserverons donc les statistiques dans cette étude. Nous ne comparons uniquement les échantillons WA et EL car, la comparaison entre EL et ALL révèle un comportement similaire dans plus de 90% (191/210) des cas. En effet, le nombre de sinistres par segment étant significativement supérieur pour les EL, le rajout des sinistres WA n'est pas significatif.

Nous avons commencé par appliquer les tests de KS et de Wilcoxon au niveau global et nous avons obtenu les résultats de la figure 4.13.

EL	Année	Kolmogorov-Smirnov		Wilcoxon		impact
		p-value	statistique	p-value	statistique	
ACI	2012	0.982	0.103	0.789	8477	102%
	2013	0.368	0.247	0.931	6408	100%
	2014	0.655	0.117	0.680	22386	102%
HERM	2012	0.104	0.075	0.751	506936.5	105%
	2013	0.024	0.085	0.014	683005	101%
	2014	0.018	0.087	0.046	526889.5	102%
SFAC	2012	1.21E-08	0.257	3.43E-11	356635	103%
	2013	2.08E-10	0.329	6.70E-14	200409.5	103%
	2014	0.008	0.113	2.97E-05	489047.5	104%
SIAC	2012	0	0.331	7.72E-22	406714	104%
	2013	7.70E-11	0.285	4.22E-13	170457	105%
	2014	3.34E-09	0.285	3.48E-11	21926.5	138%
TI	2012	2.39E-05	0.220	0.002	54245	103%
	2013	8.71E-13	0.292	5.59E-16	85801	86%
	2014	1.07E-11	0.261	1.13E-09	94980	89%

FIGURE 4.13 – Test de Kolmogorov-Smirnov et Wilcoxon appliqués aux UGD au niveau global

Les résultats sont bien différents de ceux des PD. Les tests de KS et de Wilcoxon ont exactement les mêmes résultats. Les tests de KS ont dû faire une estimation de la p -value, car le fait qu'il y est des ex-æquo dans les données ne lui permettent pas de calculer la valeur exacte. Nous avons remédié à ce problème en utilisant la fonction `ks.boot()` sur R dont le fonctionnement est décrit au point 4.4.5

Nous observons que les tests acceptent la similarité entre WA et EL sur les Etats-Unis. Avec des p -values très élevées, et ceci malgré le fait qu'en 2013 les fonctions empiriques ont un écart maximal de 0.24 (la statistique). Nous avons tout le contraire sur le Royaume-Uni et l'Italie, où les échantillons ne sont pas du tout similaire d'après nos tests. On constate également que sur certaines années l'impact de prendre les sinistres WA dans la sinistralité peut avoir des conséquences importantes (à la hausse (Italie) comme à la baisse (Royaume-Uni)). Concernant l'Allemagne et la France, le test rejette globalement les similarités entre WA et EL. On observe sur certains segments que le test de KS donne

une statistique plus faible que celle des Etats-Unis alors que la p -value passe en dessous du seuil, mais n'est pas trop éloignée. Ceci est du à la longueur des échantillons qui rend plus strict le test. L'Allemagne peut accepter l'hypothèse de similarité, la France par contre uniquement sur l'année 2014 dans le meilleur des cas. Il est difficile d'interpréter la statistique de Mann-Whitney, elle dépend de la taille des différents échantillons et prendre en compte ces paramètres nous ramène à la p -value. Contrairement à la statistique de Wilcoxon elle ne correspond pas directement à la somme des rangs d'un échantillon et par conséquent nous ne statuerons pas sur cette statistique, mais plus sur la p -value.

Les résultats obtenus sont assez négatifs dans l'ensemble. Cependant, par mesure de prudence, nous conservons, la plus petite des p -value comme indicateur. Cela signifiera que les deux tests ne rejettent pas l'hypothèse que les échantillons sont issus d'une même population.

Afin d'avoir de la visibilité sur les résultats des tests, nous avons incorporé les résultats sous forme d'arbre pour chacun des segments. Lorsqu'il n'y a pas de p -values ou de couleur, c'est qu'il y avait un manque de données soit dans l'échantillon WA soit dans l'échantillon EL (moins de 15 sinistres sur ce segment).

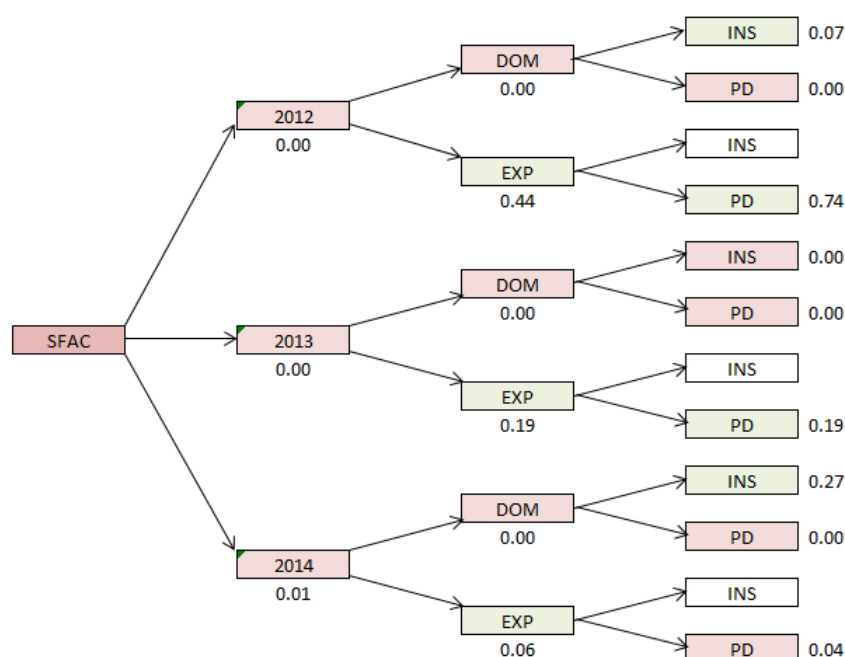


FIGURE 4.14 – Arbre des Test de Kolmogorov-Smirnov et Wilcoxon appliqués à la France

Pour la France, nous observons qu'au niveau global nous n'avons aucune concordance, cependant sur les segments Export, nous observons de nombreuses concordances entre les UGD WA et EL. Cependant la majorité des sinistres sur la France sont domestiques, il faut donc accorder plus de poids à ce segment. Si nous observons le niveau le plus segmenté nous avons la moitié des cas qui acceptent l'hypothèse de comportements simi-

lares. Les paramètres étant calculés de manière segmentée, il faut prendre en compte le fait que même si au niveau global des différences persistent, le niveau final reste le plus important. Il est difficile de dire si ces similarités sont suffisantes pour considérer que les UGD de l'EL France et du segment WA France soient issus de la même population. Une règle a été choisie : une EL accepte l'hypothèse de comportements similaires si 75% des tests sur cette EL acceptent l'hypothèse de similitude. Dans le cas de la France (38%), nous concluons que ces UGD WA et EL sont trop différents pour pouvoir être agrégés en une seule distribution.

Les autres arbres sont disponibles en annexe D.2 et suivantes. On observe que tous les tests sont positifs pour les Etats-Unis et cela à tout niveau. Il faut toutefois préciser que les Etats-Unis disposent d'une sinistralité très faible par rapport aux autres segments de WA. Ainsi, nous avons été assez rapidement limités pour effectuer nos tests. Pour l'Allemagne, il y a des différences significatives entre WA et les autres échantillons mais uniquement sur l'année 2013. Nous avons donc 81% de résultats positifs et pouvons considérer ce segment comme assimilable à son EL. Au Royaume-Uni, nous observons peu de correspondance entre les bases WA et EL (31%). L'Italie n'a aucun résultat positif sur l'ensemble des tests réalisés.

4.7 Etude des paramètres de WA en tant qu'une entité propre

Dans cette partie nous nous concentrerons sur les données de WA. Nous comparerons les données par segment aux données totales. D'abord nous nous concentrerons sur les PD puis ensuite sur les UGD. Nous avons effectué des tests similaires à ceux employés précédemment. Au lieu de comparer les distributions des segments de WA à celle des EL, nous les comparons à la distribution de WA entière. Comme nous l'avons précisé auparavant, en considérant WA comme une EL à part entière, nous perdons la notion de domestique export. Ainsi notre segmentation se limitera au statut du défaut.

4.7.1 Etude sur les PD

Pour WA, nous avons des PD assez stables, ainsi, nous observons que dans de nombreux cas, il y a toutes les données d'un échantillon qui sont en-dessous ou au-dessus de données de l'autre. Les deux dernières colonnes servent à identifier les cas où les échantillons sont disjoints. On observe ainsi que 7 de nos tests n'ont pas pu obtenir de résultats positifs et ont donc la *p-value* la plus petite. En effet soit les rangs d'un échantillon sont tous inférieurs à ceux de l'autre, ou soit la fonction de répartition empirique d'un échantillon est arrivé à 1 alors que l'autre est toujours à 0.

Nous observons quelques résultats surprenant lorsque l'on se pose sur les segment *insolvency* et *protracted*. En effet, alors que seule la France dispose de résultats positifs au global, en *protracted*, nous retrouvons le Royaume-Uni, tandis qu'à l'export, nous retrouvons l'Italie. Ainsi chaque segment de WA se retrouve dans WA entière dans au moins un secteur.

4.7. ETUDE DES PARAMÈTRES DE WA EN TANT QU'UNE ENTITÉ PROPRE79

Niveau	Segment de WA	Wilcoxon appariés	min(WA) > WA segment	max(WA) < WA segment
Global	ACI	0.00	17	6
	HERM	0.00	0	25
	SFAC	0.65	4	6
	SIAC	0.00	0	25
	TI	0.00	0	25
Protracted	ACI	0.00	25	0
	HERM	0.01	4	0
	SFAC	0.00	25	0
	SIAC	0.00	0	15
	TI	0.41	4	0
Insolvency	ACI	0.00	25	0
	HERM	0.00	0	21
	SFAC	0.02	0	10
	SIAC	0.44	9	10
	TI	0.00	25	0

FIGURE 4.15 – Résultats des tests sur PD entre WA entière et ses segments

A ce niveau nous ne pouvons pas nous faire un avis sur l'homogénéité de WA, en sachant que seule la France a des résultats satisfaisants au niveau global. Il y a trop de cas où les échantillons sont disjoint (7/15). Nous nous accordons donc sur le fait qu'uniquement le segment France a des PD similaire à celles de WA.

4.7.2 Etude sur les UGD

Dans cette partie, nous appliquons les tests aux UGD. Dans un premier temps au niveau global, puis ensuite nous regarderons au niveau *protracted* / *insolvency*.

Au niveau global, nous observons que les tests sur la France et les Etats-Unis s'accordent sur la concordance des distributions d'UGD. Pour le Royaume-Uni, c'est le contraire, nous observons que la distribution sur ce segment n'est pas similaire à celle de WA.

A un niveau plus poussé dans la granularité, tout en effectuant les tests sur les échantillons où le nombre de données est suffisant. Nous observons des résultats plus satisfaisants. Pour le Royaume-Uni, la majorité des sinistres étant *protracted*, les résultats positifs sur le segment *insolvency* ne sont pas conservés au niveau global. A ce niveau de segmentation, nous pouvons accepter que toutes les autres EL aient un comportement similaire avec WA. Au niveau des UGD, il est concevable d'établir des calculs pour WA en tant que EL à part entière.

Segment	Année	Kolmogorov-Smirnov		Wilcoxon	
		P-value	Statistique	P-value	Statistique
ACI	2012	0.76	0.15	0.53	12805.5
	2013				
	2014	0.06	0.21	0.05	29390.5
HERM	2012	0.01	0.11	0.01	208326
	2013	0.03	0.09	0.00	250019
	2014	0.34	0.06	0.24	315188.5
SFAC	2012	0.08	0.11	0.02	108122
	2013	0.10	0.12	0.16	67475
	2014	0.00	0.14	0.00	224251
SIAC	2012	0.00	0.13	0.00	125173
	2013	0.01	0.14	0.00	89020
	2014	0.48	0.06	0.46	167404.5
TI	2012	0.00	0.18	0.01	108060.5
	2013	0.00	0.20	0.00	211111
	2014	0.00	0.11	0.00	303608

FIGURE 4.16 – Résultats des tests sur les UGD entre WA entière et ses segments

4.7. ETUDE DES PARAMÈTRES DE WA EN TANT QU'UNE ENTITÉ PROPRE⁸¹

Segment	Statut	Année	Kolmogorov-Smirnov		Wilcoxon	
			P-value	Statistique	P-value	Statistique
ACI	Protracted	2012	0.57	0.20	0.41	8077
		2013				
		2014	0.06	0.22	0.03	23813
	Insolvency	2012				
		2013				
		2014				
HERM	Protracted	2012	0.00	0.14	0.00	144143
		2013	0.03	0.10	0.01	172073
		2014	0.47	0.05	0.37	238891
	Insolvency	2012	1.00	0.06	0.86	5893
		2013	0.15	0.16	0.04	7533
		2014	0.67	0.11	0.21	5192
SFAC	Protracted	2012	0.08	0.12	0.02	78946
		2013	0.19	0.14	0.16	35836
		2014	0.00	0.16	0.00	151214
	Insolvency	2012	0.56	0.17	0.48	2302
		2013	0.60	0.13	0.52	4025
		2014	0.36	0.14	0.10	6126
SIAC	Protracted	2012	0.00	0.14	0.00	91798
		2013	0.00	0.16	0.00	65229
		2014	0.40	0.07	0.37	145319
	Insolvency	2012	0.72	0.14	0.97	2541
		2013	0.79	0.15	0.75	1807
		2014				
TI	Protracted	2012	0.00	0.19	0.00	89040
		2013	0.00	0.20	0.00	172291
		2014	0.00	0.13	0.00	256788
	Insolvency	2012				
		2013	0.46	0.23	0.23	1156
		2014	0.07	0.27	0.05	1687

FIGURE 4.17 – Résultats des tests sur les UGD entre WA entière et ses segments au niveau de segmentation suivant

Conclusion

Pour conclure ce mémoire, nous ferons dans un premier temps, une synthèse des résultats obtenus dans le chapitre 4, puis ensuite nous récapitulerons le cheminement de notre raisonnement.

Synthèse des résultats

Après avoir remis en cause la non prise en compte actuelle des polices mondiales dans le calcul des paramètres du modèle interne dans le chapitre 2, nous avons testé deux nouvelles méthodes permettant d'intégrer ces polices dans le chapitre 4. La première consistant à mélanger les polices mondiales et les polices locales afin d'obtenir un jeu de paramètres par EL. La seconde consistant à former une EL avec les polices mondiales. Pour sélectionner la méthode la plus cohérente avec les données, nous avons effectué des tests non-paramétriques (Wilcoxon, Kolmogorov-Smirnov). Nous avons dû poser certaines hypothèses qui nous permis d'obtenir les résultats du tableau 4.4. Ces tests ont été appliqués aux PD et aux UGD. L'hypothèse de la première méthode est que le comportement de ces deux paramètres est similaire entre les EL et les segments de WA correspondants. Pour la deuxième méthode, nous avons pris pour hypothèse que les paramètres (PD et UGD) des segments de WA ont des comportements similaires aux paramètres globaux de WA.

Nous synthétisons les résultats de ces tests dans ce tableau :

Segment de WA	méthode 1		méthode 2	
	PD	UGD	PD	UGD
ACI	Refusé	Accepté	Refusé	Accepté
HERM	Refusé	Accepté	Refusé	Accepté
SFAC	Refusé	Refusé	Accepté	Accepté
SIAC	Refusé	Refusé	Refusé	Accepté
TI	Refusé	Refusé	Refusé	Refusé

TABLE 4.4 – Récapitulatif des résultats des tests

En ce qui concerne la première méthode, il s'agit de comparer les EL avec les segments WA correspondants.

Au niveau des PD, nous avons vu que le test sur données appariées de Wilcoxon rejette toutes nos hypothèses de comportements similaires au niveau global. Il est clair que les PD EL et WA n'ont pas un comportement similaire.

Nous avons observé que les échantillons WA et EL sont significativement différents. Cependant, l'échantillon ALL n'est pas conclu différent de celui EL, ceci pour quelques segments. Il y a donc concordance relative entre les PD ALL et EL sur l'Italie sur certains segments, mais ce n'est pas suffisant pour admettre au niveau de l'EL.

Au niveau des UGD, c'est le contraire, l'EL de l'Italie révèle le plus de différences entre WA et EL. Cependant, nous observons que globalement, il n'y a pas de différences significatives entre l'échantillon ALL et l'échantillon EL, et ceci pour l'ensemble des EL.

Les PD et les UGD sont les éléments les plus importants dans le processus de calcul de la perte attendue, il faut éviter qu'ils varient de manière trop importante. Ainsi, même si sur certains segments, l'incorporation des données de WA dans le processus de calcul des paramètres d'entrée du modèle interne serait possible et refléterait mieux la réalité, dans l'ensemble il y a trop de variabilité entre les données WA et EL. Il va donc falloir considérer la deuxième méthode : Calculer des paramètres pour WA comme si c'était une EL à part entière.

Cette méthode semble être la meilleure solution. En effet, nous observons au niveau des UGD qu'il y a de nombreuses correspondances entre les différents segments de WA et WA. Cependant au niveau des PD, les résultats ne sont pas satisfaisant, Il n'y a que la France qui a des similarité suffisantes pour envisager l'application de cette méthode.

La conclusion de ces tests est que pour le moment nous ne pouvons pas changer notre méthode de prise en compte des données de la WA dans notre modèle. Nous avons certes montré que la façon actuelle de modéliser le SCR sur ce segment n'est pas très fidèle à la sinistralité lié. Cependant, notre méthode actuelle se révèle très prudente sur les dernières années (*Expected Loss* bien supérieure au sinistres observées). Il va falloir traiter WA d'une manière plus adapté dans le modèle interne. En effet l'ensemble des tests menés ont montré que le comportement de WA est significativement différent de celui des EL. Cependant, la solution consistant à prendre en compte les données de WA dans le calcul des paramètres n'est pas satisfaisante, cela a un impact trop fort sur les paramètres robustes que l'on applique aux EL. Celle consistant à considérer WA comme une EL donne de meilleurs résultats au niveau des UGD, mais les résultats des tests sur les PD ne sont toujours pas suffisants.

Lorsque l'on étudiait WA (cf. 2.5), nous observions une convergence que nous avions qualifiée de non stable, et nous avons précisé qu'il persistait de gros écarts entre les PD (0.65 %). Avec la mise à jour des données, nous observons que la convergence continue et tend à se stabiliser dans un intervalle de 0.25 à 0.45 %.

On a représenté les régressions linéaires afin de visualiser si la convergence allait se confirmer dans les 10 mois suivants. La régression linéaire du Royaume-Uni tend à sortir de cet intervalle, mais si l'on observe bien les dernières PD, il y a un changement de tendance assez important qui nous permet d'accepter l'hypothèse d'une convergence des PD du Royaume-Uni vers celles de WA. Le seul segment s'écartant éventuellement de la tendance générale est la France dont les PD continuent à baisser alors qu'elles sont déjà

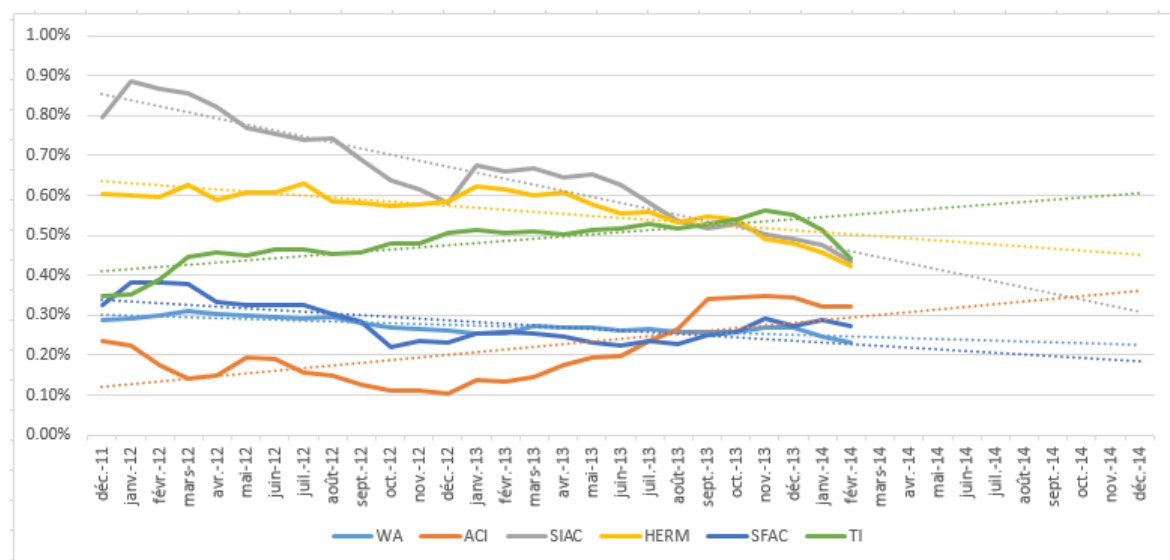


FIGURE 4.18 – PD WA et ses 5 principaux segments

les plus faibles. En nous fondant sur ces régressions linéaires, nous émettons l'hypothèse que l'Italie, les Etats-Unis et l'Allemagne auront d'ici un an des PD similaires. Cette convergence que nous observons sur les PD nous donne espoir de pouvoir considérer WA comme une EL à part entière dans un avenir proche. C'est pourquoi, il faudra refaire cette étude dans un an afin de constater ce rapprochement à l'aide des tests statistiques.

Bilan

A travers ce mémoire, nous avons décrit le fonctionnement du modèle interne d'Euler Hermès, nous avons détaillé la manière d'obtenir les paramètres. Nous avons ensuite présenté une entité différente des autres : la World Agency. Nous avons remis en question la manière dont elle est actuellement intégrée dans le modèle interne et nous avons proposé de nouvelles solutions. Puis, nous avons identifié les paramètres qui doivent être les mieux calibrés à travers des tests de sensibilité. Par la suite, nous avons présenté les tests non-paramétriques sur lesquels reposent nos résultats. Enfin nous avons effectué ces tests afin de comparer dans un premier temps, les PD et UGD de WA par segment aux EL liées à ces segments et dans un second temps, les PD et UGD de WA par segments aux PD et UGD de WA dans sa globalité.

WA n'est pas encore assez développée pour pouvoir calculer des paramètres propres à elle sur chacun des segments reliés à une EL, cette solution sur-mesure n'est pas envisageable. En effet, pour le moment, il n'y a pas assez de données pour les grands segments, et il est difficile de trouver une méthode de regroupement satisfaisante pour les petits segments. C'est surtout au niveau de la sinistralité qui se révèle faible que nous manquons

de données.

Nous avons calculé des paramètres pour l'entité virtuelle que formerai WA, pour le moment, nous observons une perte attendue significativement plus élevée que celle que nous avons auparavant. Il reste quelques ajustements possibles (adaptation des classes d'exposition,...). Mettre cette mesure en place va requérir beaucoup de temps, notamment au niveau de la programmation du modèle interne. Cela requiert des études de segmentation afin de définir de nouvelles classes d'exposition spécifique à cette EL. La mise en place d'un modèle de projection des PD pour WA est envisageable. Cependant, il serait fondé sur des indicateurs externes mondiaux difficiles à choisir, car les données internes sont pauvres. Une PD définie par jugement d'expert peut être utilisée en attendant. Cela va impacter de nombreux codes. Il va falloir présenter cela au régulateur afin qu'il valide notre étude. Il reste quelques axes qui n'ont pas été exploré, notamment l'étude sur les derniers paramètres qui selon les tests de sensibilité ont un effet moindre sur le SCR.

Pour le moment, il n'est pas envisageable d'impacter notre modèle, il faut attendre au moins une année pour confirmer la convergence des PD d'une part, et d'autre part, il faut commencer à envisager toutes les modifications que cela va entraîner. Ce mémoire a mis en avant qu'il n'était clairement pas envisageables de regrouper les segments de WA aux EL correspondantes. Il a également permis de mettre en avant une méthode qui ne peut pas être mise en place pour le moment, mais qui se sera à réexaminer dans les années à venir.

Bibliographie

- [1] Jean-Yves Baudot. Test de wilcoxon-mann-whitney. <http://www.jybaudot.fr/Inferentielle/mannwhitney.html>, 2014.
- [2] B.M. Beaver and R.B. Miller. *Instructor's Resource Manual for Statistics for Management and Economics, Fourth Edition, by William Mendenhall & James E. Reinmuth*. Duxbury Press, 1982.
- [3] Fischer Black and Myron Scholes. The pricing of options and corporate liabilities. *Journal of Political Economy*, 81(3) :pp. 637–654, 1973.
- [4] J.V. Bradley. *Distribution-free statistical tests*, chapter 12. Prentice-Hall, 1968.
- [5] A. Caja. Term structures of default probabilities. Master's thesis, ISFA, 2010.
- [6] D. A. Darling. The kolmogorov-smirnov, cramer-von mises tests. *Ann. Math. Statist.*, 28 :823–838, 12 1957.
- [7] P. Devos. Biostatistiques : Petits effectifs. <http://193.51.50.30/master/coursMaster2012-PDevos.pdf>, 2012.
- [8] A.B. Dufour and D. Chessel. Pratique des tests élémentaires. *Cours de biostatistique*, 2015.
- [9] EH. *Documentation of the EH internal model*, 2013.
- [10] Euler Hermes Crédit France EH. Echelle de notation. https://infosacheteurs.eulerhermes.com/Content/Client/docs/echelle_notation_euler_hermes.pdf, 2014.
- [11] EIOPA. Technical specification on the long term guarantee assessment (part i). pages 226–229, 2013.
- [12] J. H. Halton G. H. Freeman. Note on an exact treatment of contingency, goodness of fit and other problems of significance. *Biometrika*, 38(1/2) :141–149, 1951.
- [13] Greta M Ljung and George EP Box. On a measure of lack of fit in time series models. *Biometrika*, 65(2) :297–303, 1978.

-
- [14] Raul H. C. Lopes, Ivan Reid, and Peter R. Hobson. The two-dimensional kolmogorov-smirnov test. In *XI International Workshop on Advanced Computing and Analysis Techniques in Physics Research (April 23-27, 2007)*, 2007.
 - [15] Robert C Merton. On the pricing of corporate debt : The risk structure of interest rates. *Journal of Finance*, 29(2), May 1974.
 - [16] R. L. Plackett. Karl pearson and the chi-squared test. *International Statistical Review / Revue Internationale de Statistique*, 51(1) :59–72, 1983.
 - [17] Frédéric Planchet and Julien Jacquemin. Les méthodes de simulation en assurance. [http://www.lynxial.fr/site.nsf/0/fe99f963afca4771c1256f400065884b/\\$FILE/2_2.pdf](http://www.lynxial.fr/site.nsf/0/fe99f963afca4771c1256f400065884b/$FILE/2_2.pdf), 2010.
 - [18] Jasjeet S. Sekhon. Bootstrap kolmogorov-smirnov. <http://sekhon.berkeley.edu/matching/ks.boot.html>, 2011.
 - [19] Olivier Spilmann and Florent Calteau. Travail d’étude et de recherche tests de rangs. <http://www.statisticien.fr/tests-de-rangs>, 2015.
 - [20] Vasicek. Loan portfolio value. *Risk*, 15(12) :160–162, 2002.
 - [21] Frank Wilcoxon. Individual comparisons by ranking methods. *Biometrics Bulletin*, 1(6), 1945.

Glossaire

<i>Backtesting</i>	Tests de validité rétroactifs, 27
ACI	EH Etats-Unis, 33
ATP	<i>Ability to Pay</i> , 19
CAT	Catastrophe, 4
EAD	<i>Exposure At Default</i> , 7
EH	Euler Hermès, i
EIOPA	<i>European Insurance and Occupational Pensions Authority</i> , 5
EL	Entités Légales, 1
HERM	EH Allemagne, 33
ICISA	<i>International Credit Insurance & Surety Association</i> , 5
KS	Kolmogorov-Smirnov, 65
LGD	<i>Loss Given Default</i> , 12
LLSL	<i>Large Loss combined with Stop Loss</i> , 5
NORDIC	EH Pays Nordiques, 33
NP	Non Proportionnelle, 5
OCDE	Organisation de Coopération et de Développement Economiques, 55
PD	Probabilité de Défaut, 8
QS	<i>Quota-Share</i> , 5

RR	<i>Recollection Rate</i> , 12
SCR	<i>Solvency Capital Requirement</i> , i
SFAC	EH France, 33
SIAC	EH Italie, 33
SL	<i>Stop Loss</i> , 5
SR	<i>Salvages Rate</i> , 12
TCI&S	<i>Trade Credit Insurance and Surety</i> , 1
TCU	<i>Transactional cover Unit</i> , 28
TI	EH Royaume-Uni, 33
UGD	<i>Usage Given Default</i> , 12
WA	World Agency, 1
WA	indépendantes et identiquement distribuées, 70
WP	<i>World Program</i> , 28
XoL	<i>Excess of Loss</i> , 5

Annexe A

Backtesting des PD WA avec celles des EL

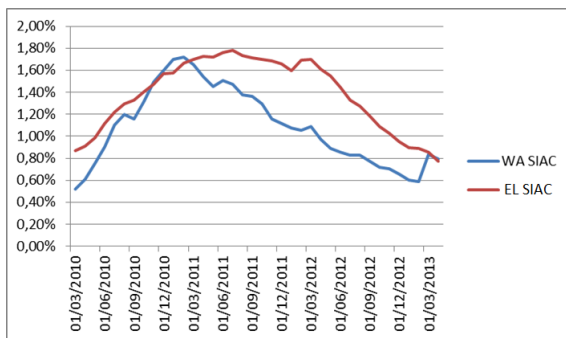


FIGURE A.1 – Exemple de l'Italie

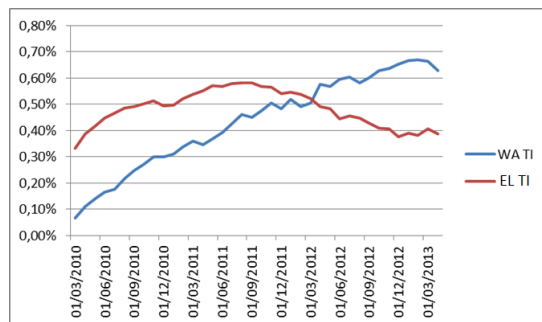


FIGURE A.3 – Exemple du Royaume-Uni

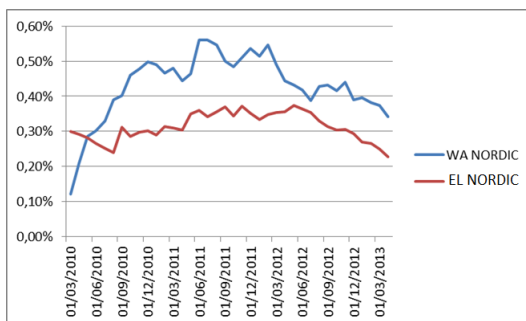


FIGURE A.2 – Exemple des Pays Nordiques



FIGURE A.4 – Exemple des États-Unis

Annexe B

Etude des expositions de L'EL et de WA pour différents pays

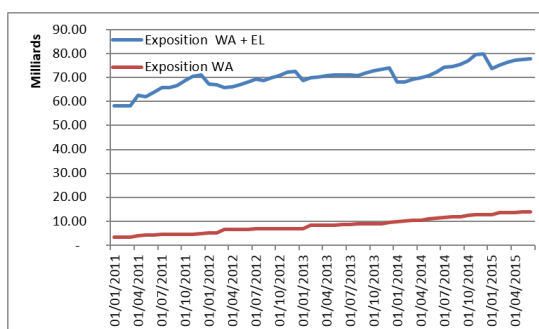


FIGURE B.1 – Italie

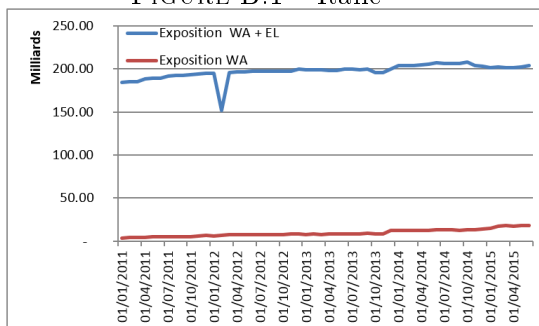


FIGURE B.2 – France

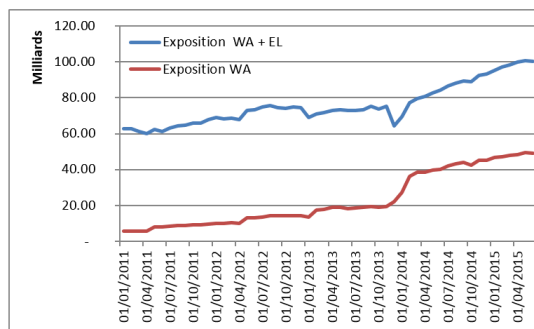


FIGURE B.3 – Royaume-Uni

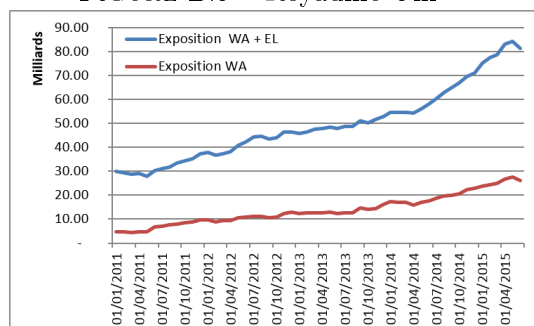


FIGURE B.4 – États-Unis

Annexe C

Test de bruits blancs et d'indépendance

C.1 Tests de bruits blancs

EL	Segmentation	Le test de Box-Pierce					Le test des runs				
		EL-WA	EL-ALL	EL	WA	ALL	EL-WA	EL-ALL	EL	WA	ALL
ACI	Global	4.06E-06	1.66E-05	5.16E-05	1.37E-06	1.44E-05	1.04E-04	2.49E-02	5.19E-04	8.40E-04	1.04E-04
	Domestique	5.18E-06	2.27E-05	1.36E-04	1.77E-06	3.35E-05	5.19E-04	2.20E-03	1.04E-04	5.19E-04	1.04E-04
	Export	4.24E-03	4.15E-01	8.97E-05	9.75E-05	4.26E-05	8.00E-03	8.31E-01	2.20E-03	5.19E-04	2.20E-03
	Domestique Insolvency	2.01E-05	4.69E-05	3.69E-03	6.34E-06	9.26E-04	5.19E-04	1.04E-04	1.54E-01	1.79E-05	2.49E-02
	Domestique Protracted	2.08E-06	2.28E-05	NaN	2.08E-06	2.28E-05	1.23E-02	5.19E-04	NaN	1.23E-02	5.19E-04
	Export Insolvency	3.12E-05	2.04E-05	3.85E-05	2.01E-05	2.27E-05	2.27E-05	2.27E-05	1.04E-04	0.00E+00	1.04E-04
SFAC	Export Protracted	2.21E-04	1.29E-02	5.17E-05	6.17E-06	2.08E-05	5.19E-04	2.49E-02	2.20E-03	2.62E-06	1.04E-04
	Global	1.65E-04	1.11E-04	6.80E-05	3.01E-05	6.66E-05	8.00E-03	8.00E-03	2.62E-06	1.79E-05	2.62E-06
	Domestique	5.14E-05	2.85E-06	2.70E-04	2.13E-04	3.47E-04	5.19E-04	2.62E-06	2.20E-03	5.19E-04	1.04E-04
	Export	8.20E-06	1.73E-06	9.05E-06	1.80E-06	7.39E-06	1.79E-05	5.19E-04	2.62E-06	1.04E-04	2.62E-06
	Domestique Insolvency	2.64E-06	9.80E-06	1.60E-05	3.98E-06	2.07E-05	1.04E-04	1.79E-05	2.62E-06	1.79E-05	2.62E-06
	Domestique Protracted	3.20E-04	7.58E-06	9.92E-04	5.96E-05	1.02E-03	3.10E-01	1.79E-05	1.04E-04	1.79E-05	1.04E-04
SIAC	Export Insolvency	3.64E-04	6.50E-03	1.56E-06	2.27E-05	1.51E-06	2.49E-02	8.00E-03	2.62E-06	1.04E-04	2.62E-06
	Export Protracted	4.98E-06	1.61E-06	1.20E-05	1.79E-06	9.52E-06	1.79E-05	1.79E-05	2.62E-06	1.79E-05	2.62E-06
	Global	2.28E-06	5.40E-06	3.00E-06	1.35E-05	2.97E-06	2.62E-06	1.04E-04	2.62E-06	1.04E-04	2.62E-06
	Domestique	2.02E-06	1.03E-05	3.17E-06	3.56E-05	3.15E-06	2.62E-06	1.04E-04	2.62E-06	1.04E-04	2.62E-06
	Export	4.67E-06	6.97E-06	2.49E-06	2.06E-06	2.44E-06	1.72E-04	1.79E-05	2.62E-06	2.62E-06	2.62E-06
	Domestique Insolvency	6.89E-06	8.20E-05	4.29E-06	3.84E-06	4.12E-06	1.04E-04	1.04E-04	2.62E-06	2.62E-06	2.62E-06
TI	Domestique Protracted	2.00E-06	8.03E-06	3.14E-06	7.79E-05	3.12E-06	2.62E-06	1.04E-04	2.62E-06	1.04E-04	2.62E-06
	Export Insolvency	6.21E-04	5.50E-05	1.90E-06	2.75E-06	1.81E-06	8.00E-03	1.79E-05	2.62E-06	2.62E-06	2.62E-06
	Export Protracted	4.55E-06	5.93E-06	2.62E-06	2.59E-06	2.57E-06	1.79E-05	1.79E-05	2.62E-06	1.79E-05	2.62E-06
	Global	8.38E-05	1.65E-05	6.32E-05	4.10E-04	1.44E-05	5.19E-04	1.04E-04	2.20E-03	5.19E-04	1.04E-04
	Domestique	1.16E-04	7.42E-05	4.99E-05	1.66E-04	3.35E-05	1.04E-04	8.00E-03	1.79E-05	1.04E-04	1.04E-04
	Export	2.91E-05	6.48E-06	6.84E-06	2.90E-03	4.26E-05	5.19E-04	2.62E-06	2.62E-06	5.19E-04	2.20E-03
TI	Domestique Insolvency	2.08E-05	4.41E-05	2.97E-05	1.09E-03	9.26E-04	5.19E-04	2.20E-03	2.49E-02	6.67E-02	2.49E-02
	Domestique Protracted	5.09E-04	1.05E-03	1.11E-04	2.60E-04	2.28E-05	1.04E-04	8.00E-03	5.19E-04	1.04E-04	5.19E-04
	Export Insolvency	9.09E-04	2.27E-05	NaN	9.09E-04	2.27E-05	8.00E-03	1.04E-04	NaN	8.00E-03	1.04E-04
	Export Protracted	1.70E-05	5.44E-06	6.84E-06	1.62E-03	2.08E-05	1.79E-05	2.62E-06	2.62E-06	5.19E-04	1.04E-04

FIGURE C.1 – P -value des tests de Box-Pierce et des tests des runs

EL	Segmentation	Test de Fisher		Test du khi-carré	
		EL vs WA	EL vs ALL	EL vs WA	EL vs ALL
ACI	Global	0.00	0.00	0.66	0.40
	Domestique	0.00	0.00	0.66	0.40
	Export	0.23	0.00	0.66	0.08
	Domestique Insolvency	0.00	0.00	0.08	0.40
	Domestique Protracted	0.00	0.05	0.66	0.66
	Export Insolvency	0.00	0.00	0.08	0.40
	Export Protracted	1.00	0.00	0.23	0.04
SFAC	Global	0.00	0.00	0.69	0.04
	Domestique	0.00	0.00	0.69	0.04
	Export	0.02	0.00	0.26	0.09
	Domestique Insolvency	0.00	0.00	0.04	0.08
	Domestique Protracted	0.00	0.00	0.35	0.04
	Export Insolvency	0.00	0.00	0.40	0.40
	Export Protracted	0.20	0.00	0.26	0.09
SIAC	Global	0.00	0.00	0.69	0.27
	Domestique	0.00	0.00	0.35	0.26
	Export	0.00	0.03	0.40	0.12
	Domestique Insolvency	0.00	0.00	0.08	0.08
	Domestique Protracted	0.00	0.00	0.35	0.09
	Export Insolvency	0.00	0.00	0.04	0.04
	Export Protracted	0.02	0.08	0.24	0.22
TI	Global	0.02	0.00	0.08	1.00
	Domestique	0.00	0.00	0.23	0.66
	Export	1.00	0.00	0.08	1.00
	Domestique Insolvency	0.00	0.00	0.23	0.66
	Domestique Protracted	1.00	0.00	0.66	0.66
	Export Insolvency	0.00	0.00	0.40	0.40
	Export Protracted	1.00	0.00	0.08	1.00

FIGURE C.2 – *P-value* des tests de Fisher et du Khi-carré

C.2 Tests d'indépendances

Annexe D

Tests de KS et de Wilcoxon

D.1 Codes R des Applications numériques

```
> #PD EL (a) et PD WA (b) de SIAC
> a
[1] 0.015328625 0.015113223 0.014226461 0.013682927 0.012754091 0.011909121
[7] 0.011579776 0.010914217 0.010118734 0.009517928 0.009421611 0.009053724
[13] 0.008581771 0.008205069 0.007384376 0.006752811 0.006076606 0.005589471
[19] 0.004466784 0.003517619 0.002800810 0.001925527 0.001306532 0.000822349
[25] 0.000517421
> b
[1] 0.008680112 0.008549496 0.008229707 0.007707474 0.007540292 0.007396710
```

```

[7] 0.007418553 0.006892470 0.006373177 0.006154595 0.005817028 0.006741899
[13] 0.006615501 0.006662900 0.006472263 0.006547209 0.006278749 0.005805353
[19] 0.005375625 0.005182505 0.005274972 0.005022382 0.004923566 0.004767986
[25] 0.004347549
> #Test de Kolmogorov-Smirnov
> c=c(a,b)
> n=length(a)
> m=length(b)
> D=matrix(c(c,rep(1,n),rep(2,m),rank(c)[1:n],rank(c)[(n+1):(n+m)]),nrow=n+m,byrow=F) # a=1;b
> #on trie par le rang
> H=D[order(D[,1]),]
> I=cbind(H,c(rep(1/n,n),rep(1/m,m)))
> #on calcule les fonctions de répartition
> I[I[,2]==1,4]=(1:n/n)
> I[I[,2]==2,4]=(1:m/m)
> J=cbind(I,rep(0,50),rep(0,50))
> colnames(J)=c("PD","Echantillon","rang","Fx","F1","F2")
> J[1,5:6]=c(0.04,0)
> for (i in 2:50){
+   if(J[i,2]==1){J[i,5]=J[i,4]}
+   else{J[i,5]=J[i-1,5]}
+   if(J[i,2]==2){J[i,6]=J[i,4]}
+   else{J[i,6]=J[i-1,6]}
+ }
> #on calcule D
> (D=max(abs(J[,6]-J[,5])))
[1] 0.48
> c=D*sqrt(n*m/(n+m))
> (pval=sum((-1)**(0:99)*exp((1:100)**2*-c**2)))
[1] 0.05612483
> #La procédure sous R est :
> ks.test(a,b)

```

Two-sample Kolmogorov-Smirnov test

```

data:  a and b
D = 0.48, \textit{p-value} = 0.005614
alternative hypothesis: two-sided

```

```

> #Wilcoxon sur données appariées
> #nombre de différences négatives
> length((a-b)[a-b<0])
[1] 9

```

```

> #nombre de valeurs
> n=length(a-b);n
[1] 25
> #9 valeurs négatives sur 25 n'invalident pas l'hypothèse "une chance sur deux" :
> 9/25; pbinom(9,25,0.5)
[1] 0.36
[1] 0.1147615
> #On somme les rangs des écarts négatifs et positifs entre les échantillons
> som <- sum(rank(abs(a-b))[a < b])
> (compsom <- sum(rank(abs(a -b))[a > b]))
[1] 242
> #petite vérification:
> (compsom+som==n*(n+1)/2)
[1] TRUE
> moy=n*(n+1)/4; var=(n*(n+1)*(2*n+1))/24
> usom=(som - moy)/sqrt(var)
> (pval=2 * pnorm(usom))
[1] 0.03242761
> ucompsom=(compsom - moy)/sqrt(var)
> (pval=2 * (1 - pnorm(ucompsom)))
[1] 0.03242761
> #La procédure sous R est :
> wilcox.test(a, b, paired = TRUE, correct = F, exact = F)

```

Wilcoxon signed rank test

```

data: a and b
V = 242, \textit{p-value} = 0.03243
alternative hypothesis: true location shift is not equal to 0

```

```

> #Wilcoxon données non appariées
> n1=length(a)
> n2=length(b)
> c=c(a, b)
> (rtot=rank(c))
[1] 50 49 48 47 46 45 44 43 42 41 40 39 37 34 29 27 18 15 8 6 5 4 3 2 1 38
[27] 36 35 33 32 30 31 28 21 19 17 26 24 25 22 23 20 16 14 12 13 11 10 9 7
> #Les individus du premier groupe ont les rangs :
> (r1=rtot[1:n1])
[1] 50 49 48 47 46 45 44 43 42 41 40 39 37 34 29 27 18 15 8 6 5 4 3 2 1
> #Les individus du deuxième groupe ont les rangs :
> (r2=rtot[(n1 + 1):(n1 + n2)])
[1] 38 36 35 33 32 30 31 28 21 19 17 26 24 25 22 23 20 16 14 12 13 11 10 9 7

```

```

> #Calculs de la Statistique de Wilcoxon sr
> (sr=sum(r1))
[1] 723
> (esr=(n1 * (n1 + n2 + 1))/2)
[1] 637.5
> (vsr=(n1 * n2 * (n1 + n2 + 1))/12)
[1] 2656.25
> (t0=(sr - esr)/sqrt(vsr))
[1] 1.658944
> (pval=2 * (1-pnorm(t0)))
[1] 0.09712714
> # Calculs de la Statistique de Mann-Whitney u
> (u=sr - (n1 * (n1+1))/2)
[1] 398
> (eu=(n1 * n2)/2)
[1] 312.5
> (vu=(n1 * n2 * (n1 + n2 + 1))/12)
[1] 2656.25
> (t1=(u - eu)/sqrt(vu))
[1] 1.658944
> (pval=2 * (1-pnorm(t1)))
[1] 0.09712714
>
> #La procédure sous R est :
> wilcox.test(a, b, exact = F, correct = F)

```

Wilcoxon rank sum test

data: a and b

W = 398, $\text{p-value} = 0.09713$

alternative hypothesis: true location shift is not equal to 0

D.2 Résultats des tests sur PD

D.3 Arbre illustrant les résultats sur les UGD avec une segmentation plus fine

Annexe E

Qualité des données

DATA QUALITY DASHBOARD

Group IT
Local IT
GRC
Legal Entity

 = Deleted KPIs

[illegible]

GROUP IT KPIs - DETAIL PER AXIS	Timeliness	9	1
	Duplication	10	0
	Consistency	0	0
	Completeness	10	0
	Conformity	10	0
	Accuracy	0	0
	Integrity	0	0

[illegible]

DETAIL PER AIS - GROUP	Completeness	0	0	0	0
	Conformity	0	0	0	0
	Accuracy	0	0	0	0
	Integrity	0	0	0	0
	Timeliness	0	0	0	0
POLICY REPORT	Duplication	1	0	0	0
	Consistency	13	2	0	0
	Completeness	2	0	0	0
	Conformity	0	0	0	0
	Accuracy	0	0	0	0
DETAIL PER AIS - BUSINESS UNIT	Integrity	3	0	0	0

#	Assg	Data to check	KPI Definition	Scale	Comments	Action Plan in case of KO	Responsible	Deadline	on plan st	Situation	Status	Number of OK	Number of KO	Technical KO							
KPI's from the Group Policy Report (Frequency: on a monthly basis; Reference date: end of last month)																					
SCALE															Global RCM Explanations						
															Impact on SCR or blocking the Internal Model?	Impact on Parameters?	Impact on EL?				
199	Timeliness	Group Policy Report	Reception Date of the file	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: The Internal Model is sensible to duplicated KIs	NO	NO
200	Duplication	Group Policy ID	Number of duplicated Group policy ID	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	-	Yes: A default currency is already assigned to each OE. If there is more than one currency, this can lead to a wrong application of the policy features and therefore wrong estimation of the losses on the limits linked to those cases.	NO	NO
201	Heterogeneity	Currency code	Number of different currency codes	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	-		NO	NO
202	Completeness	Annual Aggregate	Percentage of missing Annual Aggregate (Blank)	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: the IM will allocate a default AA leading to a wrong estimation of losses	NO	NO
203	Consistency	Annual Aggregate	Percentage of Annual Aggregate higher than 500 000 (local currency)	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: the IM will allocate a default AA leading to a wrong estimation of losses	NO	NO
204	Completeness	Annual Non Qualifying Loss	Percentage of missing Annual Non Qualifying Loss (Blank)	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: the IM will allocate a default ANQL leading to a wrong estimation of losses	NO	NO
205	Consistency	Annual Non Qualifying Loss	Percentage of Annual Non Qualifying loss higher than 500 000 (local currency)	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: the IM will allocate a default ANQL leading to a wrong estimation of losses	NO	NO
206	Completeness	Maximum Liability	Percentage of missing Maximum Liability (Blank)	0-100%	# "NO": Please refer to the CSV file at this link		OE		Open		OK	1	-	-	100%	0.0%	Not Applicable	Not Applicable	Yes: the IM will allocate a default ML leading to a wrong estimation of losses	NO	NO
207	Consistency	Maximum Liability & Estimated Premium	Percentage of groups with a ratio Maximum Liability / Estimated Premium > 7% for which an explanation is available	0-100%	# TO BE FILLED OUT BY THE CE: The denominator of this rate is = 4 explanation should be provided by the CE and then put the % of errors with explanation provided. The		OE		Open		KO	-	1	-	0.0%	10.0%	Not Applicable	Not Applicable	Yes: Here we try to avoid wrong values, the franchises are used in the IM	NO	NO
Sum for Group Policy Report												8	1	0							

		Effectiveness	1	0	0	0
	GROUP POLICY REPORT	Duplication	0	0	0	0
		Consistency	0	0	0	0
	DETAIL PER AXIS	Completeness	0	0	0	0
		Conformity	0	0	0	0
	GRC	Accuracy	0	0	0	0
		Integrity	0	0	0	0
		Effectiveness	0	0	0	0
	GROUP POLICY REPORT	Duplication	1	0	0	0
		Consistency	2	1	1	0
	DETAIL PER AXIS	Completeness	3	3	0	0
		Conformity	0	0	0	0
	ACCURACY	Accuracy	0	0	0	0
	BUSINESS UNIT	Integrity	0	0	0	0

[illegible]

	Timeliness	1	0	0
POLICY REPORT	Duplication	0	0	0
	Consistency	0	0	0
DETAIL PER AXIS	Completeness	0	0	0

GRC	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	0	0	0
POLICY REPORT	Duplication	1	0	0
	Consistency	9	0	0
	Completeness	4	2	0
	Conformity	0	2	0
DETAIL PER AXIS	Accuracy	1	5	0
	Integrity	5	0	0
	Conformity	1	0	0
	Completeness	0	0	0

#	Axis	Data to check	KPI Definition	Scale	Comments	Action Plan in case of KO	Responsible	Deadline	on plan st	Situation	Status	Number of OK	Number of KO	Technical KO	Global RCM Explanations			
KPI's from the Unsettled Claims Report (Frequency: on a monthly basis, Reference date: end of last month)															SCALE			
238	Timeliness	Unsettled Claims Report	Reception Date of the file	0% (0/1)	# "KO": Please refer to the CSV file of this KPI				Open	●	OK	1	-	-	00/100	Delay > 2 days	Below 2 days	Not Applicable
239	Completeness	Claim Declaration Date	Percentage of Missing Declaration Date (blanks)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
240	Completeness	Invoice date	Percentage of Missing Invoice Date (blanks) (Only for Risk Attaching OI's)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
241	Consistency	Invoice date	Percentage of Invoice Date that are subsequent to the declaration date (Only for Risk Attaching OI's)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
242	Integrity	Currency Code	Number of different currency codes	1	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	0	Below	Not Applicable	-
243	Integrity	Insured claim amount	Percentage of Insured Claim amounts equal to 0	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
244	Consistency	Recollection amount	Percentage of claims with a recollection amount superior to the insured amount	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
245	Consistency	Indemnified Claim amount	Percentage of claims with an indemnified amount superior to the insured amount	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
246	Consistency	Salvage amount	Percentage of claims with a salvage amount superior to the indemnified amount	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
247	Consistency	Status of Debtor	Percentage of claims with a status of debtor at declaration date equal to status of debtor 6 months after declaration date (excluding the ones already legally treated at declaration date)	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
Sum for Unsettled Claims Report															10	0	0	0

POLICY REPORT	Timeliness	1	0	0
	Duplication	0	0	0
	Consistency	0	0	0
	Completeness	0	0	0
DETAIL PER AXIS	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	0	0	0
GRC	Duplication	0	0	0
	Consistency	5	0	0
	Completeness	2	0	0
	Conformity	0	0	0
POLICY REPORT	Timeliness	0	0	0
	Duplication	0	0	0
	Consistency	5	0	0
	Completeness	2	0	0
DETAIL PER AXIS	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	2	0	0
BUSINESS UNIT	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	2	0	0
	Timeliness	0	0	0

#	Axis	Data to check	KPI Definition	Scale	Comments	Action Plan in case of KO	Responsible	Deadline	on plan st	Situation	Status	Number of OK	Number of KO	Technical KO	Global RCM Explanations			
KPI's from the Refus Report (Frequency: on a monthly basis, Reference date: end of last month)															SCALE			
248	Timeliness	CRS Report	Reception Date of the file	0% (0/1)	Not Applicable				Open	●	OK	1	-	-	00/100	Delay > 2 days	Below 2 days	Not Applicable
249	Integrity	Country Code	Percentage of wrong country code (not referenced in the ISO tables)	Not Applicable					Open	●	Technical KO	-	-	-	00/100	0.0%	Above	Not Applicable
250	Integrity	Industry Code	Percentage of wrong industry code (not referenced in the ISO tables)	Not Applicable					Open	●	Technical KO	-	-	-	00/100	0.0%	Above	Not Applicable
251	Completeness	Policy ID	Number of Policy ID in the Buyer Data that are not provided in the Policy Report and that are not World Agence policies	0			OE		Open	●	OK	1	-	-	00	Above	Not Applicable	-
252	Consistency	Exposure	Gap in percentage between the total exposure in the limit data and the total exposure in the Buyer Data	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
253	Consistency	Number of buyers	Gap in percentage between the total number of buyers in the limit data and the total number of buyers in the Buyer Data	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
254	Completeness	Headergroup ID	Number of Header Group ID in the Buyer Data that are not reported in the Headergroup Data (except for "0" <"1")	0	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	0	Above	Not Applicable	-
255	Completeness	Grade (Refus Buyer data)	Percentage of missing grade (blanks)	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
256	Completeness	Grade (Refus Headergroup data)	Percentage of missing grade (blanks)	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
257	Completeness	Country Code	Percentage of missing country code (blanks)	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
258	Completeness	Industry Code	Percentage of missing industry code (blanks)	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
259	Consistency	Exposure	Percentage of buyers with a global exposure (sum of limits) in the limit data different from the global exposure in the Buyer Data	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
260	Consistency	Number of limits	Percentage of buyers with a number of limits in the limit data different from the number of limits in the Buyer Data	0.0%	# "KO": Please refer to the CSV file of this KPI		OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
Sum for Refus Report															11	0	2	0

POLICY REPORT	Timeliness	1	0	0
	Duplication	0	0	0
	Consistency	0	0	0
	Completeness	0	0	0
DETAIL PER AXIS	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	0	0	0
POLICY REPORT	Timeliness	0	0	0
	Duplication	0	0	0
	Consistency	4	0	0
	Completeness	6	0	0
DETAIL PER AXIS	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	0	0	0
BUSINESS UNIT	Conformity	0	0	0
	Accuracy	0	0	0
	Integrity	0	0	0
	Timeliness	0	0	0

#	Axis	Data to check	KPI Definition	Scale	Comments	Action Plan in case of KO	Responsible	Deadline	on plan st	Situation	Status	Number of OK	Number of KO	Technical KO	Global RCM Explanations			
KPI's from the Data Requirements File Report (Frequency: on a monthly basis, Reference date: end of last month)															SCALE			
261	Timeliness	Data Requirements File	Reception Date of the file	0% (0/1)	# "KO": Please refer to the CSV file of this KPI and identify from where the gaps come, all the buyers and intermediate calculations are provided in those files. In any case, if "KO", please refer to the CSV file of this KPI and identify from where the gaps come, all the buyers and intermediate calculations are provided in those files. In any case				Open	●	OK	1	-	-	00/100	Delay > 2 days	Below 2 days	Not Applicable
262	Consistency	Exposure	Gap in percentage between the total exposure provided by the Business Unit and the total exposure computed from the CRS response file	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
263	Consistency	Exposure	Average Exposure relative Gap in percentage between the exposure provided by the Business Unit and the exposure computed from the CRS response file on each segmentation (Grade * Exposure Class * Domestic / Export)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
264	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Domestic Protected PD provided by the Business Unit and the Global Domestic Protected PD computed by GRC	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
265	Consistency	USD	Average relative Gap between the USD provided by the Business Unit and the USD computed by GRC on each segmentation (Grade * Domestic Class * Domestic / Export)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
266	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Domestic Insolvency PD provided by the Business Unit and the Global Domestic Insolvency PD computed by GRC	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
267	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Export Insolvency PD provided by the Business Unit and the Global Export Insolvency PD computed by GRC	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
268	Consistency	Historical PD	Average relative Gap on the last 3 points in time between the Global Export Protected PD provided by the Business Unit and the Global Export Protected PD computed by GRC	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
269	Consistency	Recollection Rate	Average relative Gap in percentage between the recollection rate provided by the Business Unit and the recollection computed by GRC on each segmentation (Grade * Domestic / Export)	0.0%			OE		Open	●	OK	1	-	-	00/100	0.0%	Above	Not Applicable
															0	0	0	0

270	Consistency	Salvages Rate	Average relative Gap in percentage between the salvage rate provided by the Business Unit and the salvage rate computed by GRC on each segmentation (Grade / Domestic / Export)	0.2%		DE		Open	●	OK	1	-	-	100%	0.2%	Pass	Not Applicable	Yes		NO	Yes
271	Conformity	Recollection Rate	Is the recollection rate computed from the recollection amount of the claims file?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
272	Conformity	Historical PD	Is the numerator of the PD computed from the claims file?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
273	Conformity	Historical PD	Is the denominator of the PD computed from the CDS extractions provided by GRC?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
274	Conformity	USD	Is the USD computed from the claims file?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
275	Conformity	USD (Recollection rate)	Is the Recollection rate from the USD computed from the Claims file?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
276	Conformity	USD (Salvages rate)	Is the salvage rate from the USD computed from the Claims file?			DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO	NO	NO	
277	Completeness	USD	Percentage of buyers with USD higher than 300% for which an explanation is available	0%	TO BE FILLED OUT BY THE OE: The denominator of this rate is = 100 an explanation should be provided by the OE and then put the % of cases with explanation	DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
278	Consistency	USD	Percentage of buyers with USD lower than 1%	1%		DE		Open	●	OK	1	-	-	100%	10.0%	Pass	Not Applicable	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
279	Completeness	USD	Percentage of buyers with USD lower than 0.1% for which an explanation is available	0%	TO BE FILLED OUT BY THE OE: The denominator of this rate is = 100 an explanation should be provided by the OE and then put the % of cases with explanation provided. The	DE		Open	●	KO	-	1	-	100%	0%	Not Applicable	-	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
280	Consistency	Recollection Rate	Percentage of buyers with recollection rate higher than 100%	0.0%		DE		Open	●	OK	1	-	-	100%	0.0%	Pass	Not Applicable	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
281	Completeness	Recollection Rate	Percentage of buyers with recollection rate higher than 200% for which an explanation is available	100.0%	0	DE		Open	●	OK	1	-	-	100.0%	10.0%	Pass	Not Applicable	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
282	Consistency	Salvages Rate	Percentage of buyers with salvages rate higher than 200%	0.0%		DE		Open	●	OK	1	-	-	100%	0.0%	Pass	Not Applicable	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
283	Completeness	Salvages Rate	Percentage of buyers with salvages rate higher than 200% for which an explanation is available	100.0%	0	DE		Open	●	OK	1	-	-	100.0%	10.0%	Pass	Not Applicable	NO		Yes: The goal is to verify if those cases are real since they can impact the final parameter	NO
Sum for Data Requirements File Report																					
										POLICY REPORT		Timeliness	1	0	0						
										DETAIL PER AXIS		Duplication	0	0	0						
										GRC		Consistency	0	0	0						
												Completeness	0	0	0						
												Conformity	0	0	0						
												Accuracy	0	0	0						
												Integrity	0	0	0						

#	Axis	Data to check	KPI Definition	Scale	Comments	Action Plan in case of KO	Responsible	Deadline	on plan st	Situation	Status	Number of OK	Number of KO	Technical KO	SCALE				
KPI's from the upload platform and RAI results sign-off (Frequency: on a monthly basis, Reference date: end of last month)																			
284	Timeliness	Sign-off of all Credit Insurance files	General sign-off of all Credit Insurance reports (Policy, Group Policy, Claims, Unsettled Claims, Ratios and Data Requirement) on the upload platform from the local responsible received on time?	Yes			DE		Open	●	OK	1	-	-	100%	0%	Not Applicable	-	
285	Integrity	Credit Insurance files	Check whether new Credit Insurance reports (Policy, Group Policy, Claims, Unsettled Claims, Ratios and Data Requirement) have been validated on the upload platform	Yes			DE		Open	●	OK	1	-	-	100%	0%	Not Applicable	-	
286	Timeliness	Sign-off of CI model Results	General sign-off of CI model Results from the local responsible received on time?	Not Applicable	Centralized calculations				Open	●	Technical KO	-	-	0	100%	0%	Not Applicable	-	

Sum for Data Requirements File Report

POLICY REPORT		Timeliness	1	0	0
DETAIL PER AXIS		Duplication	0	0	0
GRC		Consistency	0	0	0
		Completeness	0	0	0
		Conformity	0	0	0
		Accuracy	0	0	0
		Integrity	0	0	0
POLICY REPORT		Timeliness	0	0	0
DETAIL PER AXIS		Duplication	0	0	0
BUSINESS UNIT		Consistency	3	0	0
		Completeness	2	2	0
		Conformity	0	0	0
		Accuracy	0	0	0
		Integrity	0	0	0

EL	Niveau	Wilcoxon apparié		impact
		EL vs WA	EL vs ALL	
ACI	Global	0.00	0.00	98%
	Domestique	0.00	0.00	98%
	Export	0.00	0.00	99%
	Domestique Insolvency	0.00	0.00	97%
	Domestique Protracted	0.00	0.00	Pas de données EL
	Export Insolvency	0.03	0.03	121%
	Export Protracted	0.00	0.00	98%
TI	Global	0.00	0.00	207%
	Domestique	0.00	0.00	167%
	Export	0.00	0.00	276%
	Domestique Insolvency	0.00	0.00	27%
	Domestique Protracted	0.00	0.00	513%
	Export Insolvency	0.00	0.00	Pas de données EL
	Export Protracted	0.00	0.00	271%
SFAC	Global	0.00	0.00	101%
	Domestique	0.00	0.00	102%
	Export	0.00	0.00	99%
	Domestique Insolvency	0.00	0.00	102%
	Domestique Protracted	0.00	0.00	101%
	Export Insolvency	0.00	0.00	98%
	Export Protracted	0.00	0.01	99%

FIGURE D.1 – Test de Kolmogorov-Smirnov et Wilcoxon aux PD

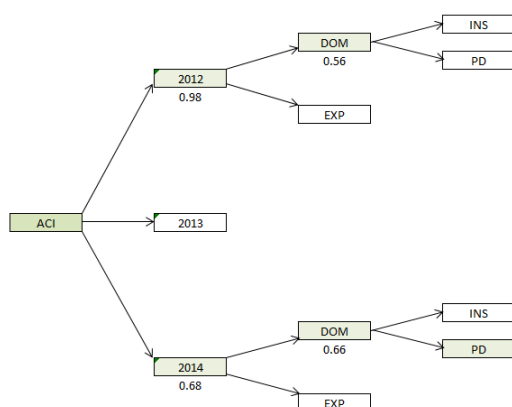


FIGURE D.2 – Arbre des Test de Kolmogorov-Smirnov et Wilcoxon appliqués aux Etats-Unis

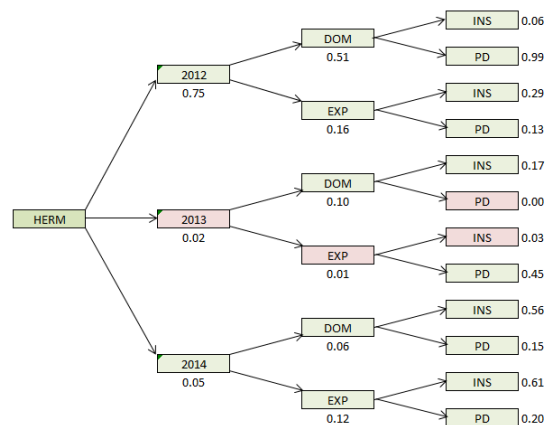


FIGURE D.3 – Arbre des Test de Kolmogorov-Smirnov et Wilcoxon appliqués à l'Allemagne

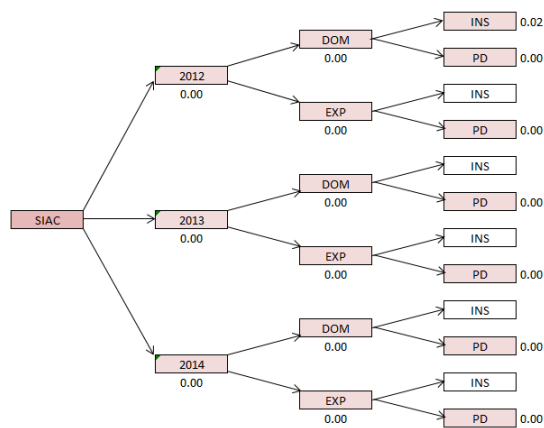


FIGURE D.4 – Arbre des Test de Kolmogorov-Smirnov et Wilcoxon appliqués à l'Italie

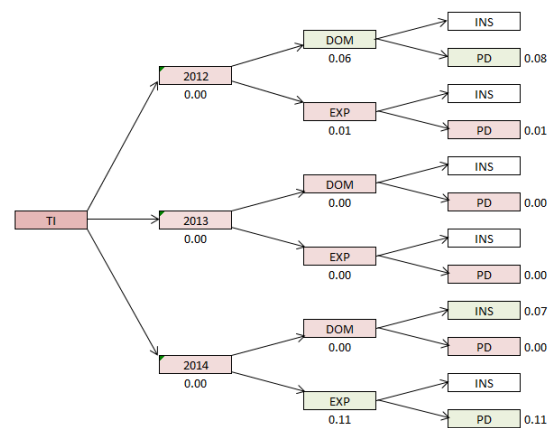


FIGURE D.5 – Arbre des Test de Kolmogorov-Smirnov et Wilcoxon appliqués au Royaume-Uni