

THÈSE DE DOCTORAT

de

L'UNIVERSITÉ PARIS-SACLAY

École doctorale de mathématiques Hadamard (EDMH, ED 574)

Établissement d'inscription : Université Paris-Sud

Établissement d'accueil : Nexter Systems

Laboratoire d'accueil : Laboratoire de mathématiques d'Orsay, UMR 8628 CNRS

Spécialité de doctorat : Mathématiques appliquées

Florence DUCROS

*Maintien en conditions opérationnelles pour une flotte de véhicules :
étude de la non stabilité des flux de recharge dans le temps*

Date de soutenance : 26 juin 2018.

Après avis des rapporteurs : DIDIER CHAUVEAU (Université d'Orléans)
OLIVIER GAUDOIN (Institut Polytechnique de Grenoble)

Jury de soutenance :

GILLES CELEUX	Directeur de recherche, INRIA	Directeur de thèse
DIDIER CHAUVEAU	Professeur, Université d'Orléans	Rapporteur
ZOHRA CHERFI-BOULANGER	Professeur, UTC Compiègne	Examineur
OLIVIER GAUDOIN	Professeur, Institut Polytechnique de Grenoble	Rapporteur
DAMIEN GODARD	Ingénieur, Nexter Systems	Invité
PASCAL MASSART	Professeur, Université de Paris Sud	Président du jury
PATRICK PAMPHILE	Maître de conférences, Université de Paris Sud	Codirecteur de thèse

Remerciements

Je remercie chaleureusement toutes les personnes qui m'ont aidée pendant l'élaboration de ma thèse et notamment mon directeur de thèse Gilles Celeux pour son intérêt et son soutien. Je remercie également Patrick Pamphile mon co-directeur de thèse pour sa grande disponibilité et ses nombreux conseils durant le déroulement et la rédaction de ma thèse.

Je remercie Didier Chauveau et Olivier Gaudoin d'avoir accepté de rapporter ma thèse. Je vous remercie pour l'intérêt et le regard critique que vous m'avez apporté. Je remercie également André Lannoy ; vos commentaires ont été précieux et ont enrichi mon travail de recherche appliquée aux problématiques industrielles.

Je remercie Jean-François Nedelec de m'avoir permis de mener ces travaux dans de très bonnes conditions. Ce travail n'aurait pu être mené à bien sans la disponibilité et l'accueil chaleureux que m'ont témoigné tous les personnels de la direction du service et soutien client. Je remercie Cécile Lopez et Cédric Chaintreau de m'avoir choisie au sein de Nexter Systems pour ce projet de thèse.

Je remercie également pour leurs précieux conseils notamment sur des aspects plus techniques, Didier Derenemesnil et Jean-Jacques Vuillaume. Les nombreux collègues que j'ai eu l'occasion d'interroger dans le cadre ma thèse, Alizée Chesse, Jonathan Delgado, Michel Garrivet, Damien Godard, Christophe Lagasse, François-Xavier Roux et Benoit Schroeder, m'ont permis de donner vie, je l'espère, à ce récit.

Je ne saurai oublier, personnels permanents et doctorants qui ont, durant ces trois dernières années, contribué à créer une atmosphère de travail agréable et conviviale.

Comment ne pas remercier également l'ensemble de mes proches et de ma famille. Votre confiance et votre soutien sans faille m'auront permis de surmonter bien des épreuves.

Cette thèse a été réalisée à l'Université de Paris Sud, dans le cadre d'une convention CIFRE, gérée par l'Association Nationale de la Recherche Technologie (ANRT), et établie entre le Laboratoire de mathématiques d'Orsay et la société Nexter-Systems.

Table des matières

Remerciements

INTRODUCTION GÉNÉRALE 1

ABRÉVIATIONS ET NOTATIONS 9

I MAINTIEN EN CONDITIONS OPÉRATIONNELLES 11

1 Contexte industriel 13

1.1	Enjeux	13
1.2	Objectifs	14
1.3	Données du RETEX	15
1.3.1	Usage des véhicules	15
1.3.2	Inter-occurrences de pannes d'un équipement	16
1.4	Les équipements fils rouges	17

2 Flux des pièces de rechange 19

2.1	Notations	20
2.2	Problématiques du contrat MCO	20
2.3	Simulation du flux	21
2.4	Choix de l'échelle de mesure des inter-occurrences de panne	21
2.4.1	Méthodes	22
2.4.2	Résultats et conclusion	25

II MODÉLISATION DES INTER-OCCURRENCES DE PANNE 27

3 Estimations paramétriques 29

3.1	Modèle et données	29
3.2	Estimations par maximum de vraisemblance	30
3.2.1	Vraisemblances	30
3.2.2	Maximum de vraisemblance	30
3.2.3	EM et Stochastic EM	30
3.3	Estimations bayésiennes	33
3.3.1	Notations	34
3.3.2	Le choix bayésien	34

4 Modèle de loi de Weibull simple 37

4.1	Caractéristiques d'une loi de Weibull à deux paramètres	38
4.2	Données	40

4.3	Ajustement graphique	40
4.4	Estimation par maximum de vraisemblance sans censure	41
4.5	Estimation par maximum de vraisemblance avec censures	42
4.6	Algorithme S-EM	43
4.7	Estimation bayésienne	45
4.8	Estimation Bayesian Restoration Maximization	47
4.8.1	Intervalle de crédibilité	50
4.8.2	Calibrage du nombre de répétitions de BRM	50
4.8.3	Calibrage de la loi <i>a priori</i> sur le paramètre de forme, β	51
4.8.4	Calibrage de la loi <i>a priori</i> sur le paramètre de position, η	52
4.9	Analyses de sensibilité	55
4.10	Conclusion	56
4.11	Données réelles	56
5	Modèle de lois de Weibull en concurrences	59
5.1	Caractéristiques de lois de Weibull en concurrence	60
5.2	Données	60
5.3	Identifiabilité	61
5.4	Ajustement graphique	61
5.5	Vraisemblance complète	62
5.6	Algorithme EM	62
5.7	Algorithme S-EM	63
5.8	Estimation BRM	64
5.8.1	Pour les paramètres de forme	64
5.8.2	Pour les paramètres d'échelle	64
5.9	Conclusion	64
6	Modèle de mélange de lois de Weibull	65
6.1	Estimation bayésienne d'un mélange de lois Weibull	67
6.1.1	Introduction	69
6.2	Mixture of Weibull distributions	70
6.2.1	Mixture of Weibull distributions	70
6.2.2	Weibull distributions	70
6.2.3	Data	71
6.3	Fitting the model	71
6.3.1	Weibull quantiles-quantiles plot	71
6.3.2	Maximum likelihood estimation	72
6.3.3	EM and S-EM algorithms	72
6.3.4	Bayesian estimation	73
6.4	Bayesian restoration maximization	73
6.4.1	Importance resampling	74
6.4.2	Proposal distribution	74
6.4.3	BRM Method	74
6.4.4	Prior elicitation	75
6.5	Simulations and results	76
6.5.1	Priors calibration	77
6.5.2	BRM number of runs	78
6.5.3	Sensitivity analysis	78
6.5.4	Comparisons	79
6.5.5	Discussion	80
6.6	Real data processing	81

6.6.1	Fitting a Weibull mixture to data	81
6.6.2	Example 1 - Throttles	82
6.6.3	Example 2 - Power assemblies	82
6.7	Conclusion	84
7	Sélection du modèle de loi des inter-occurrences de panne	87
7.1	Sélection avec critère graphique	88
7.1.1	WQQP pour un modèle de loi de Weibull simple	88
7.1.2	Weibull QQP pour un modèle de deux lois de Weibull en concurrence	89
7.1.3	Weibull QQP pour un modèle de mélange de deux lois de Weibull	89
7.2	Sélection avec critères de vraisemblance pénalisés	90
7.2.1	AIC	90
7.2.2	BIC	90
7.3	Conclusion	91
III	BESOIN EN ÉQUIPEMENTS DE RECHANGE	93
8	Estimation du besoin en pièces de rechange d'une flotte de véhicules	95
8.1	Introduction	97
8.1.1	Problem description	98
8.1.2	Notation	99
8.1.3	Problem statement	99
8.2	Simulation algorithm	100
8.3	Equipments failures inter-occurrence	101
8.4	Vehicles usage profiles	102
8.5	Forecasting with various scenarios	104
8.6	Conclusion	107
9	Équipements fils rouges	109
9.1	Contrainte d'usage par véhicule de la flotte et profils d'usage	110
9.2	Prévisions pour l'équipement 1	111
9.3	Prévisions pour l'équipement 2	113
9.4	Prévisions pour l'équipement 3	115
9.5	Conclusion	117
10	Estimation du flux pour un équipement sans RETEX	119
10.1	Contexte	119
10.2	RETEX historique des équipements	120
10.3	Classifications d'équipements analogues	120
10.4	Estimation du flux de rechange avec les catégories d'équipements	123
10.5	Conclusion	123
	CONCLUSION GÉNÉRALE	127
	BIBLIOGRAPHIE	137

INTRODUCTION GÉNÉRALE

Les systèmes se dégradent avec l'âge et l'usage et sont déclarés défaillants lorsqu'ils ne peuvent plus assurer l'usage requis avec la performance désirée. Lors de la fabrication du système (conception et production), l'industriel vise donc une fiabilité maximale tout en optimisant les coûts. Et ce d'autant plus que la défaillance peut avoir des conséquences importantes, soit en termes économiques, soit en termes de sûreté de fonctionnement. La fiabilité des systèmes et l'analyse statistique des données de panne sont donc des préoccupations importantes des industriels et ont donné lieu à de très nombreux travaux. On peut citer Lawless (1982), Meeker and Escobar (1998), Kalbfleisch and Prentice (2002), Nelson (2009), Murthy et al. (2004b) et Finkelstein and Cha (2013) pour une revue de littérature complète des problématiques de fiabilité.

L'utilisateur, quant à lui, exige une disponibilité maximale du système. Prenons l'exemple de l'utilisateur d'un véhicule personnel, utilitaire ou militaire : les conséquences d'une indisponibilité peuvent aller du simple dérangement (dû par exemple à la panne de la climatisation du véhicule) à l'empêchement (dû par exemple à une panne de batterie), voire à l'accident grave (dû par exemple à une panne du système de freinage). Cela est d'autant plus sensible dans le domaine militaire où l'indisponibilité peut être fatale. Pour l'utilisateur d'un système, un contrat de garantie est alors une assurance contre des défaillances durant la période de vie utile du système. La garantie est plus précisément une obligation contractuelle entre un vendeur et un acheteur couvrant tout ou partie du coût des réparations, pendant une période définie. Il est courant chez les constructeurs automobiles de proposer « une garantie constructeur » de deux ou trois ans pour un véhicule neuf, puis une extension de garantie de trois ans ou 30 000 kilomètres sur certains équipements. Murthy and Blischke (2006) et Blischke et al. (2011) constituent une bonne introduction sur les différents types de contrats de garantie et leurs problématiques. Il faut noter que les problématiques de garantie couvrent à la fois les problématiques de fiabilité, mais aussi les stratégies de maintenance (réparation complète, réparation minimale, fréquence des maintenances préventives...).

S'il couvre les coûts dus à une défaillance, le simple contrat de garantie n'assure pas au client une disponibilité élevée tout au long du contrat. Or, améliorer une performance du système peut être un enjeu important lors de l'achat, par exemple d'une centrale nucléaire, d'un avion ou d'un véhicule blindé. Par ailleurs, pour des systèmes de haute technologie ou pour un parc de systèmes complexes, il peut être plus économique pour le client de sous-traiter la maintenance. Le contrat de maintien en conditions opérationnelles (MCO, ou *Maintenance, Repair an Overhaul*, MRO) garantit alors au client un service externalisé de MCO du système tout au long du contrat, sous des contraintes d'usage cf. Murthy et al. (2015).

Le développement rapide des technologies et la mondialisation des marchés ont eu pour conséquence de créer une concurrence sévère entre les constructeurs. Il ne s'agit plus simplement de vendre le meilleur produit au meilleur coût, mais aussi de proposer aux clients le meilleur service après-vente. La vente d'une flotte d'avions à une compagnie d'aviation ou d'une flotte de véhicules militaires à un gouvernement inclut souvent un contrat de MCO sur plusieurs années. La rentabilité d'une flotte d'avions commerciaux est très liée à la disponibilité des avions. L'efficacité d'une armée est dépendante de la disponibilité de ses systèmes d'armes. Le contrat de MCO est donc un argument marketing important pour le constructeur, mais c'est aussi un coût supplémentaire qui se doit d'être anticipé. À l'instar du contrat de garantie classique, le contrat MCO prend en charge les éventuelles réparations, mais la comparaison s'arrête là. Pour gérer un contrat de garantie classique, l'industriel cherche à maximiser la disponibilité pour le client sous une contrainte de coûts de maintenance (et de sûreté de fonctionnement). Pour le contrat MCO, il cherche à minimiser les coûts de maintenance sous des contraintes de disponibilité pour le client. Par ailleurs, les contrats de garantie sont limités à la durée de vie utile du produit (*i.e.* des dizaines de mois), alors que les contrats MCO visent toute la durée de vie du produit (*i.e.* des dizaines d'années).

Les problématiques de MCO d'un système couvrent à la fois les problématiques de fiabilité et de maintenance du système, mais aussi l'estimation de l'usage réel qu'en fait le client (environ-

nement opérationnel, niveau de compétence de l'utilisateur, intensité d'usage, etc.). Prédire les coûts de maintenance d'un système nécessite, non seulement la connaissance de la fiabilité du système, mais aussi de ses conditions d'usage, cf. Figure 1.

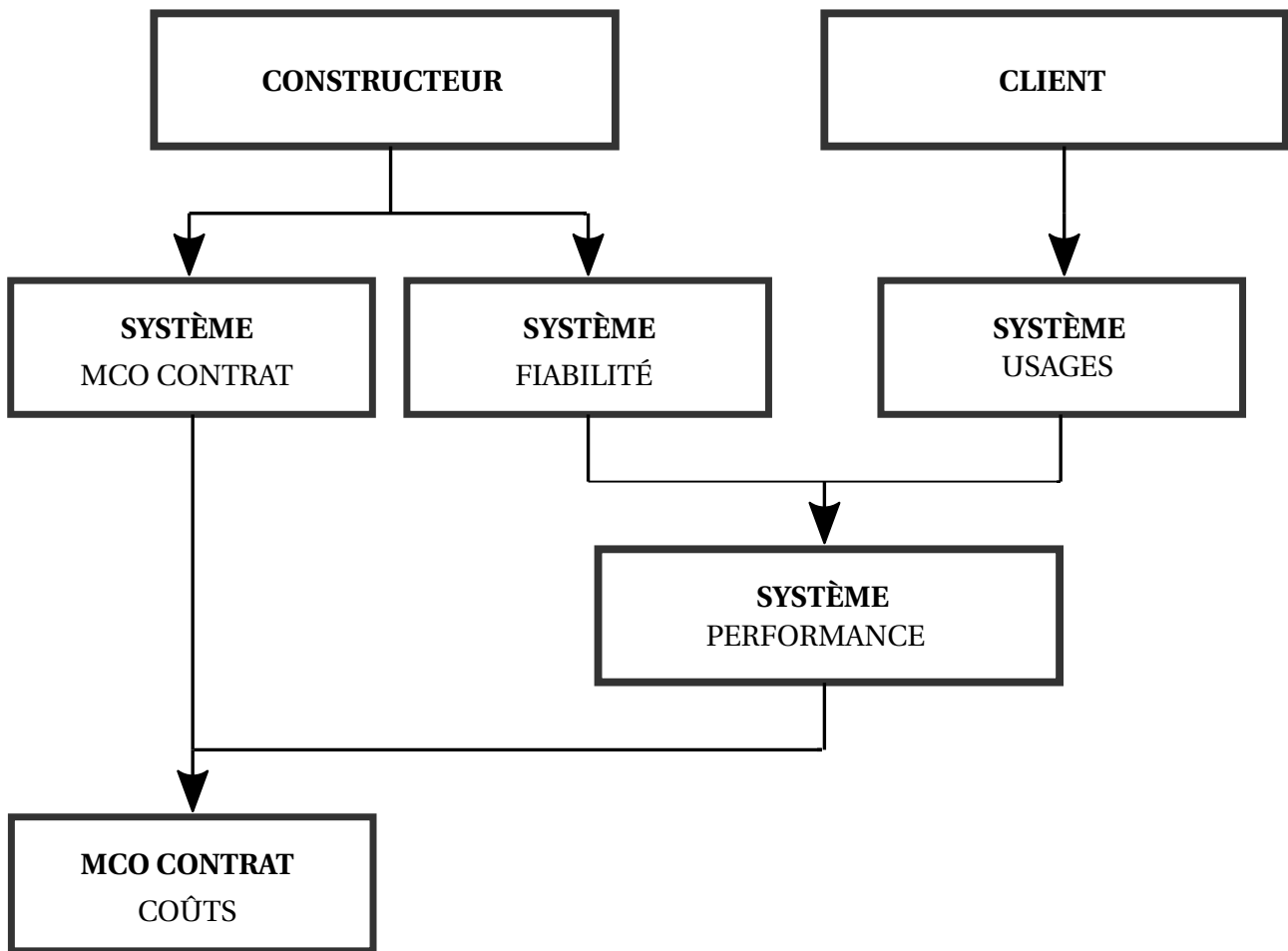


FIGURE 1: Analyse des coûts du contrat de maintien en conditions opérationnelles MCO

Un système complexe, un véhicule par exemple, peut être vu comme un ensemble de sous-systèmes (propulsion, châssis, transmission, etc.), d'équipements (moteur, roues, boîte de vitesses, etc.) et de milliers de composants produits par un vaste réseau de fournisseurs. La fiabilité du système est liée à la conception et la qualité de la production qui dépendent du constructeur, l'efficacité de la maintenance mais aussi des divers fournisseurs. L'usage du système est lié au profil de mission ou à l'intensité d'usage : ces informations, elles, dépendent de l'utilisateur. La collecte des données de défaillance des équipements et des usages en ligne ou hors-ligne est donc un enjeu pour les constructeurs. À cet égard, l'essor du véhicule connecté et des *big data* ouvrent de nouvelles perspectives en termes d'information sur la dégradation du système cf. Meeker and Hong (2014); Tiwari et al. (2018).

Le traitement et l'analyse des données de défaillance du retour d'expériences (RETEX) sont alors le principal outil d'aide à la décision pour gérer les contrats MCO. La plupart du temps ces données sont très hétérogènes et par conséquent il est difficile de les analyser. D'une part, les occurrences de panne peuvent être très variables pour un équipement, selon le processus de production, le design, l'usure, etc. D'autre part, les conditions d'usage des véhicules peuvent être très diverses selon le profil des missions, l'intensité d'usage, et les conditions extérieures (*i.e.* état de la route, conditions climatiques, etc.). Par ailleurs, elles peuvent varier d'un usage à l'autre. Tous ces facteurs impactent la performance des équipements et, sur le long terme, entraînent une non

stabilité des flux de pièces de rechange. Généralement ces sources de variabilité ne sont, ni contrôlées, ni mesurées, ce qui réduit l'efficacité des estimations. Et ce d'autant plus, que pour des systèmes hautement fiables, comme les équipements militaires, le nombre de défaillances observées est très faible durant la durée de vie utile.

Il existe de nombreux travaux traitant des aspects mathématiques de la gestion des coûts d'un contrat de garantie, voir Wu (2012) pour une revue détaillée récente. Ce type de contrat de garantie oblige le vendeur à prendre en charge, tout ou partie, des coûts de maintenance ou de réparation des produits, durant la période de garantie. Mais il ne contractualise pas la disponibilité des véhicules, et surtout ne couvre pas la période totale d'utilisation des véhicules. Il existe donc peu de travaux concernant le maintien en conditions opérationnelles d'une flotte, intégrant la problématique de la non stabilité des flux des pièces de rechange. On peut toutefois citer Zaino and Berke (1994), Wang (2010) et Anderson et al. (2017). Ces travaux abordent l'estimation du nombre moyen de défaillances, mais avec des usages similaires des véhicules de la flotte, ou des taux d'occurrences de pannes identiques pour les équipements.

Le cadre de cette thèse est donc les contrats de maintien en conditions opérationnelles d'une flotte de véhicules (MCO, ou *Maintenance, Repair and Overhaul*, MRO). Elle se situe en phase de contrat MCO : les véhicules sont déjà déployés et utilisés. En particulier, nous nous intéressons à la problématique de non stabilité du flux des équipements de rechange dans le temps. Les véhicules sont déployés sur plusieurs années, dans des conditions opérationnelles très variables, et avec différents profils de mission. Dans le cas où le taux d'occurrence de pannes est constant au cours du temps (la loi des inter-occurrences est donc une loi exponentielle), la loi du nombre cumulé de pannes observées sur la flotte est une loi de Poisson, cf. Ascher and Feingold (1984). Pour des lois d'inter-occurrences plus réalistes comme la loi log-normale et la loi de Weibull cf. Lawless (1982), la loi du nombre cumulé de pannes ne peut pas être obtenue analytiquement. Par ailleurs, le nombre de pannes dépend de l'usage des véhicules de la flotte. Ces usages sont liés au profil des missions et donc peuvent être très variables. Pour cela, nous proposons de mettre en place un modèle de simulations du besoin en équipements de rechange tenant compte :

- une modélisation probabiliste de la loi des inter-occurrences de pannes, intégrant l'hétérogénéité latente et la prise en compte du vieillissement sur le long terme ;
- la prise en compte des usages des véhicules ;
- et la prise en compte du déploiement des véhicules au cours du contrat.

Dans cette thèse, à partir des données du RETEX, nous avons d'une part identifié des profils d'usage des véhicules, d'autre part modélisé les lois des inter-occurrences de panne des équipements.

Le contrat MCO garantit la prise en charge des réparations des équipements sous des contraintes d'utilisation des véhicules de la flotte. Les usages des véhicules sont alors recueillis régulièrement. Les fonctionnalités des véhicules étant diverses (*i.e.* motorisation, propulsion, transmission, production d'électricité, défense,...), voir Figure 2, le constructeur a retenu trois échelles de mesure d'usage, qui sont renseignées régulièrement dans le RETEX. Pour des raisons de confidentialité, ces échelles ne sont pas nommées.

Le nombre de défaillances de l'équipement dépend de l'usage des véhicules sur la période de soutien de la flotte. Dans cette thèse, à partir de cette base de données d'usages, nous avons alors identifié des profils d'usages plus ou moins intensifs.

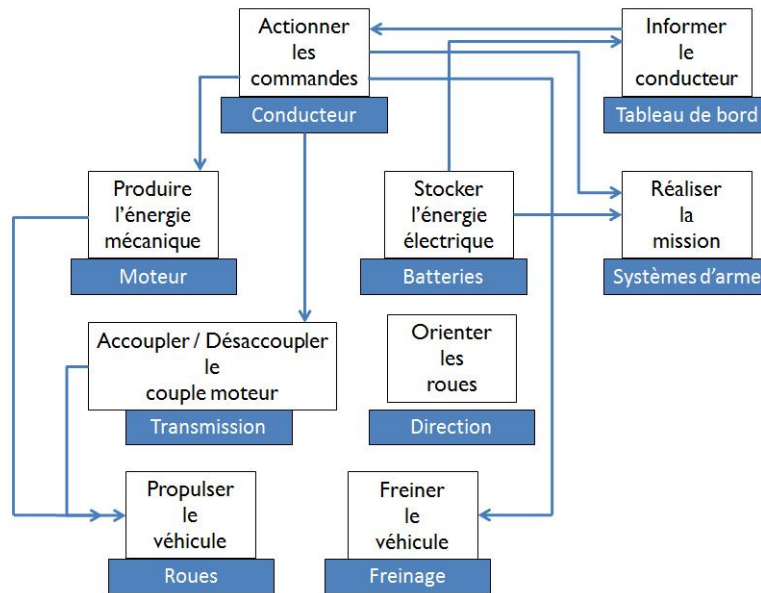


FIGURE 2: Analyse fonctionnelle d'un véhicule

Si le choix d'une mesure des occurrences de panne est naturel pour certains équipements (*e.g.* les kilomètres pour les roues, les heures pour les batteries, etc.), il peut être plus obscur pour d'autres. Dans cette thèse, nous proposons alors une approche pour choisir une mesure expliquant au mieux la défaillance d'un équipement.

Une fois l'échelle de mesure retenue, il nous faut sélectionner un modèle de loi pour les inter-occurrences de pannes. En fiabilité, le choix de modèles de loi est très lié à la physique de la défaillance et en particulier intègre le fait que les équipements ont une ou plusieurs causes de défaillance. La cause ou les causes constituent une population homogène ou hétérogène. Depuis plus de cinquante ans, la loi de Weibull a été très largement utilisée pour ajuster des durées de vie issues du domaine médical ou industriel cf. Weibull (1951). En effet, la loi de Weibull présente un taux de défaillance décroissant, constant ou croissant, selon la valeur de son paramètre de forme, ce qui la rend très utile pour ajuster des données de durée de vie. On trouvera dans Murthy et al. (2004b) une taxonomie des différents modèles à base de lois de Weibull et leurs utilisations. En particulier, le modèle de mélange de lois de Weibull est simple et tout à fait pertinent pour modéliser l'hétérogénéité latente dans des données de panne.

On considérera dans toute la thèse que les équipements sont non-réparables.

Les équipements suivis sont, pour la plupart, des équipements critiques et par conséquent hautement fiables. Les inter-occurrences de panne sont alors très fortement censurées, ce qui rend difficile l'estimation des paramètres du modèle de loi. L'algorithme Expectation Maximisation (EM), ou sa variante stochastique S-EM, sont des méthodes itératives efficaces pour obtenir une estimation du maximum de vraisemblance (MV) en présence de données manquantes, cf. McLachlan and Krishnan (2008). Cependant les algorithmes EM et S-EM restent très sensibles au fort taux de censure, voir Bacha and Celeux (1996) pour une simple Weibull, et Chauveau (1995) pour un mélange de lois de Weibull. L'approche bayésienne, *via* les algorithmes Monte Carlo Markov chain (MCMC), reste très gourmande en temps de calculs et en particulier dans le cas du mélange cf. Marin et al. (2005). Dans cette thèse, nous proposons donc d'utiliser une méthode d'estimation de type bootstrap bayésien, cf. Rubin (1981), appelée Bayesian Restoration Maximization (BRM). Sa particularité est de combiner à la fois des informations *a priori* sur les paramètres et une technique d'échantillonnage préférentiel pour obtenir un échantillon des paramètres concentré sur la moyenne *a posteriori*. Cette méthode avait déjà été utilisée dans le cas d'une simple loi de Weibull dans Bacha and Celeux (1996) et dans le cas de modèles de lois de Weibull en concurrence dans

Bacha (1996). Dans cette thèse, nous avons utilisé la méthode BRM pour un mélange de lois de Weibull. Dans l'approche bayésienne l'éllicitation des lois *a priori* est fondamentale. Ce point a été traité et, en particulier, l'impact des hyperparamètres sur la performance de l'estimation BRM a été étudié. Une analyse de sensibilité des performances de BRM en fonction de la taille de l'échantillon et du taux de censure a aussi été réalisée. Contrairement à S-EM et les algorithmes MCMC, BRM est une méthode d'estimation non itérative, et les temps de calculs sont donc significativement réduits. Par ailleurs, les comparaisons avec l'algorithme S-EM ont montré que l'estimation BRM est plus performante en termes de précision des estimations et de temps de calcul.

Par ailleurs un véhicule est constitué de centaines d'équipements. Les constructeurs ont alors régulièrement recours à des analogies entre équipements, par exemple pour choisir des équipements à mettre sous contrôle, ou pour effectuer des prévisions de soutien pour des équipements ayant un RETEX peu renseigné, etc. Pour répondre à cette problématique, dans cette thèse nous proposons une classification des équipements utilisant un critère de similarité basé sur le paramètre d'échelle de la loi de Weibull et le MTBF. Cela permet aux ingénieurs d'obtenir des catégories d'équipements similaires en termes de vieillissement et de taux d'occurrences de panne sur le long terme. Ces catégories seront utilisées pour répondre à un appel d'offres ou pour le lancement d'un nouvel équipement.

Le plan de la thèse est alors le suivant. Dans une première partie, nous présentons les différentes données collectées : comment elles sont collectées et utilisées pour obtenir les inter-occurrences de pannes et les usages annuels des véhicules. La deuxième partie traite de la modélisation de la loi des inter-occurrences de pannes des équipements, dans le cas de populations d'équipement homogènes ou non-homogènes, avec une ou plusieurs causes de défaillances. Les problèmes de sélection de modèle de loi et d'estimation des paramètres y sont abordés. La performance des méthodes classiques d'estimation est fortement réduite pour l'estimation des paramètres d'un mélange de lois de Weibull, lorsque les données sont très fortement censurées. Pour répondre à cette difficulté nous avons mis en œuvre une méthode d'estimation de type bootstrap-bayésien. Cette approche est présentée et ses performances sont étudiées dans cette seconde partie. Dans une troisième partie nous nous intéressons aux usages véhicules et en particulier à l'identification de profils d'usages plus ou moins intensifs. L'estimation du besoin en pièces de rechange est faite dans cette partie, à l'aide de simulations des défaillances des équipement selon l'usage des véhicules. Afin de permettre au constructeur de faire des analogies entre les nombreux équipements d'un véhicule, nous avons proposé une classification des équipements. La dernière partie est dédiée à une conclusion générale.

Le manuscrit est rédigé en français, toutefois, deux articles soumis et publiés sont présents dans la thèse. Pour des raisons de confidentialité, le nom des échelles, des équipements, et des véhicules hôtes resteront anonymes.

The manuscript is written in French. However, two articles submitted and published are present in the thesis. For reasons of confidentiality, the name of the units, the equipments, and the vehicles hosts will remain anonymous.

ABRÉVIATIONS ET NOTATIONS

Bootstrap	Technique reposant sur la simulation de données
BRM	Bayes restoration maximization
Données censurées	Dernier temps d'observation
EM	Expectation-Maximization
ESS	Effective sample size
IS	Importance sampling
MCMC	Markov Chain Monte Carlo
MV, MVC	Maximum de vraisemblance, maximum de vraisemblance censuré
S-EM	Stochastic EM
RETEX	Retour d'expériences
RMSE	Erreur quadratique moyenne
WQQP	Weibull quantiles-quantiles plot
n	Nombre d'équipements étudiés
x_i	Inter-occurrence de panne du i -ème équipement ; $\mathbf{x} = (x_1, \dots, x_n)$
$\tilde{x}_i$	Inter-occurrence de panne du i -ème équipement simulée
c_i	Indicateur de censure ; $\mathbf{c} = (c_1, \dots, c_n)$ $c_i = 1$, si x_i est censurée, x_i est inter-occurrence de panne incomplète ; $c_i = 0$ sinon, x_i est inter-occurrence de panne complète
$\mathbf{1}(\cdot)$	Fonction indicatrice
f_W, R_W	Densité de probabilité et fiabilité d'une distribution d'une loi de Weibull
β, η	Paramètres de forme et d'échelle
$\boldsymbol{\theta}$	Paramètres du modèle
$\tilde{\boldsymbol{\theta}}$	Paramètres simulés
$\hat{\boldsymbol{\theta}}$	Paramètres estimés
$L(\mathbf{x}, \mathbf{c} \boldsymbol{\theta})$	Fonction de vraisemblance des données de défaillance complètes et incomplètes
$L(\mathbf{x} \boldsymbol{\theta})$	Fonction de vraisemblance des données de défaillance complètes
$\pi(\boldsymbol{\theta})$	Distribution <i>a priori</i>
$\pi(\boldsymbol{\theta} \mathbf{x}, \mathbf{c})$	Distribution <i>a posteriori</i>
$\rho(\boldsymbol{\theta})$	Distribution instrumentale de l'échantillonnage préférentiel
$\mathbf{w}$	Poids de l'échantillonnage préférentiel
$\tilde{\mathbf{w}}$	Poids normalisés de l'échantillonnage préférentiel

Maintenance préventive : maintenance exécutée à des intervalles indéterminés ou selon des critères prescrit et destinée à diminuer la probabilité de défaillance ou la dégradation du fonctionnement d'un bien.

Maintenance corrective : maintenance exécutée après détection d'une panne et destinée à remettre un bien dans un état dans lequel il peut accomplir une fonction requise

Potentiel : unité d'usage des équipements en fonction des échelles de mesure d'utilisation des véhicules.

Première partie

MAINTIEN EN CONDITIONS OPÉRATIONNELLES

Chapitre 1

Contexte industriel

Sommaire

1.1 Enjeux	13
1.2 Objectifs	14
1.3 Données du RETEX	15
1.3.1 Usage des véhicules	15
1.3.2 Inter-occurrences de pannes d'un équipement	16
1.4 Les équipements fils rouges	17

NEXTER SYSTEMS a pour mission de développer, produire et commercialiser des équipements et des systèmes de défense pour les armées françaises et étrangères. L'industriel assure également le suivi et le soutien logistique des véhicules en service, à l'aide de modèles de prévision de la disponibilité, estime le « reste à casser » et réalise des dimensionnements de flux de rechanges.

Le plan de ce chapitre est le suivant. Dans une première section nous présentons les enjeux stratégiques de NEXTER SYSTEMS. Les objectifs de cette thèse visent alors à répondre aux problématiques liées à ces enjeux : ils sont présentés dans la section suivante. Puis nous détaillons les données du retour d'expériences (RETEX) sur lequel nous avons fait nos analyses statistiques. Nous terminons en présentant trois équipements « fil rouge » qui serviront d'applications tout au long de cette thèse.

1.1 Enjeux

Le monde est entré dans une ère de turbulences. La France et l'Europe sont confrontées à des menaces intenses, diversifiée et durables, comme le soulignent les conclusions de la Revue stratégique remise au Président de la République en octobre 2017. Pour répondre aux besoins actuels des armées, la durée d'exploitation des matériels militaires est en constante augmentation. Les temps de mission dépassent la période de vie utile des équipements et par conséquent les phénomènes de vieillissement impactent de façon certaine la disponibilité des véhicules.

Ces dernières années, l'impact de la mondialisation et l'évolution rapide des technologies obligent les constructeurs à proposer des offres de contrat de maintenance prolongeant la durée de vie utile des équipements. Lors d'une vente d'une flotte de véhicules, le constructeur doit non seulement proposer des produits performants avec un prix compétitif, mais aussi fournir des services après-vente incluant un service d'assistance et de maintenance avec des taux de disponibilité très élevés, tout au long de la durée de vie des véhicules.

Le contrat de maintien en conditions opérationnelles (MCO), représente alors un réel enjeu économique pour les industriels et c'est d'ailleurs une activité majeure au sein de *Nexter Systems*.

Ce type de contrat de soutien d'une flotte de véhicules militaires, avec engagement de performance, met en œuvre un ensemble d'activités et de moyens logistiques (approvisionnement, acquisition et gestion des rechanges, opérations de maintenance, main d'oeuvre, outillages, documentation, formation, ...) dans l'objectif de garantir une disponibilité maximale tout au long de la phase d'exploitation du système, soit plusieurs dizaines d'années.

Pour NEXTER SYSTEMS, établir un contrat de MCO pour une flotte de véhicules militaires, sur une dizaine d'années, est donc un véritable challenge. Les principaux enjeux sont alors :

Collecter les données : le système d'information

L'analyse des données collectées, pendant la production et après-la vente, peuvent aider lors de la prise de décisions pour établir le contrat MCO. Comment collecter et organiser ces données pour rendre leur utilisation la plus efficace possible ?

Gestion des stocks de pièces de rechange

Des stocks de pièces de rechange sont constitués afin de permettre des maintenances correctives et préventives rapides des équipements défaillants, et ainsi assurer une continuité opérationnelle. Une sous-évaluation du stock minore la disponibilité, alors qu'une sur-évaluation minore le profit. Comment évaluer au plus juste le stock de pièces de rechange ? Comment prévoir le « reste à casser » lors d'une obsolescence d'un équipement ?

Analyse des coûts

Lorsqu'un produit est vendu avec un contrat de soutien forfaitisé, l'industriel encourt donc des coûts additionnels résultant des maintenances correctives mais aussi des temps d'indisponibilité des véhicules pendant la durée de soutien. Si des équipements nécessaires au remplacement d'équipements défaillants ne sont pas disponibles, alors l'équipement et donc le véhicule peut être durablement immobilisé. Ces interruptions de service sont indésirables et peuvent être lourdes de conséquences aussi bien au niveau humain que financier. Le coût de ce type de contrat peut accroître significativement en fonction du nombre d'équipements de rechange prévus. L'évaluation des coûts dûs aux défaillances tout au long du contrat est donc un enjeu fondamental pour l'industriel.

Extension de soutien

La durée d'exploitation des véhicules militaires est ainsi en constante augmentation et par conséquent donne lieu à des extensions de contrat. Il faut donc être capable d'anticiper l'impact du vieillissement sur la performance des équipements.

Contrat pour un nouveau système

Dans le cas de réponse à appel d'offre de contrat MCO de nouveaux produits, l'industriel se doit de maîtriser les risques économiques, tout en en proposant des offres compétitives.

1.2 Objectifs

L'analyse des données issues des défaillances des équipements et des usages des véhicules va être des outils précieux pour répondre aux enjeux de NEXTER SYSTEMS. Cette thèse vise alors à répondre aux enjeux de NEXTER SYSTEMS à partir des objectifs suivants :

- ☐ Structurer, normaliser et fournir des outils permettant d'interroger le RETEX. Ce dernier est constitué :
 - d'informations recueillies lors d'une défaillance d'un équipement : date de panne, date de remise en service. C'est la base de données EQUIPEMENTS ;

- d'informations recueillies régulièrement sur les usages des véhicules de la flotte. C'est la base de données des usages VÉHICULES.

Ces informations doivent être précises, normalisées et structurées pour permettre des analyses des coûts et des besoins.

- ☐ Proposer des modèles probabilistes ajustant au mieux les données de panne et mettre en œuvre des méthodes d'estimation adaptées à la qualité et la richesse du RETEX ;
- ☐ Identifier « des catégories » d'équipements en termes de vieillissement et de MTBF. Ces catégories d'équipements permettent de faire des analogies utiles quand le RETEX est peu renseigné, ou pour établir des coûts de maintenance pour de nouveaux équipements ;
- ☐ Fournir des outils de simulations de flux de pièces de rechange. Ces outils vont aider à établir des contrats MCO ou des extensions de garanties.

1.3 Données du RETEX

Le besoin en pièces de rechange dépend à la fois des informations de panne de l'équipement et des usages des véhicules. Ces informations sont stockées dans le RETEX et sont les principales données utilisées pour estimer le coût du contrat MCO voir Figure 1.1. Cette section décrit les informations renseignées dans le RETEX.

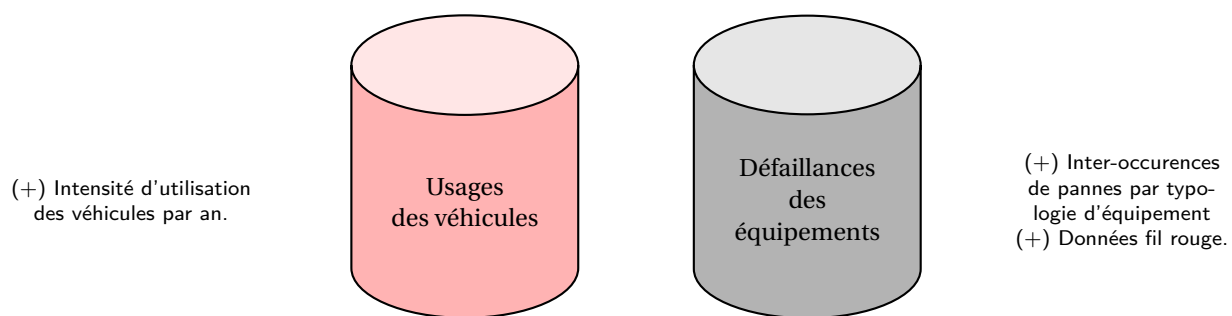


FIGURE 1.1: Retours d'expérience : base de données des usages des véhicules et base de données des défaillances des équipements

1.3.1 Usage des véhicules

Le constructeur a choisi de retenir trois mesures d'usage. Ces mesures sont alors renseignées régulièrement dans le RETEX. Pour des raisons de confidentialité nous les nommerons scale 1, scale 2 et scale 3. Compte tenu du niveau de technicité des véhicules, ceux-ci sont mis en service en plusieurs tranches : la durée de ce déploiement est de huit années. Les mesures d'usage correspondent alors à un cumul annuel pour chaque véhicule depuis sa mise en service jusqu'à la fin du contrat. Cette période d'observation s'étalonne sur neuf années d'utilisation de la flotte, dont huit années de déploiement des véhicules. La base de données « usages véhicules » est constituée des relevés des trois mesures pour chaque véhicule et pour chaque année de soutien, voir Figure 1.2.

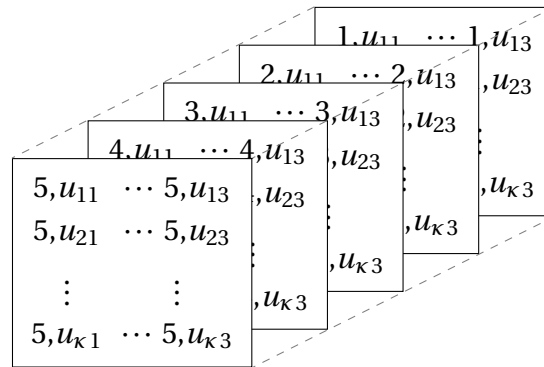


FIGURE 1.2: Base des usages des κ véhicules de la flotte sur 5 ans avec trois échelles de mesures d'usages

Un véhicule peut être vu comme un ensemble de sous-systèmes (propulsion, châssis, transmission,...) constitué de divers équipements (moteur, roues, boîte de vitesses, etc.), voir Figure 1.3

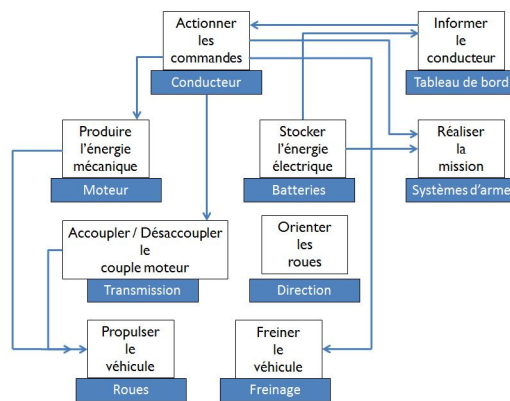


FIGURE 1.3: Analyse fonctionnelle d'un véhicule

Compte tenu des diverses fonctionnalités, il existe diverses mesures d'usage des véhicules : nombre de jours d'utilisation, nombre d'heures moteur consommées, nombre de kilomètres parcourus, nombre d'utilisation d'un système d'armes, etc.

1.3.2 Inter-occurrences de pannes d'un équipement

L'usage d'un équipement est hérité de l'usage du véhicule sur lequel l'équipement est monté. On notera « **potentiel** » l'unité d'usage des équipements en fonction des échelles de mesure d'utilisation des véhicules.

Mécanisme de calcul

Pour chaque équipement on dispose des informations :

- **date de pose** : date de pose de l'équipement sur un véhicule ;
- **date de dépose** : date de dépose de l'équipement : date de panne si l'équipement est défaillant avant la fin de l'étude, pas de date de dépose si l'équipement est toujours en fonctionnement à la fin de l'étude ;
- **date de fin de RETEX**.

On distingue donc :

- **pour une panne** (équipement tombé en panne avant la fin du RETEX) :

l'âge de l'équipement = date de dépose de l'équipement- date de pose de l'équipement.

l'usage de la pièce = mesure d'emploi du véhicule à la date de dépose de l'équipement- mesure d'emploi du véhicule à la date de pose de l'équipement.

- **pour une panne censurée** (équipement en fonctionnement à la fin du RETEX) :

l'âge de l'équipement = date de fin d'étude- date de pose de l'équipement.

Inter-occurrence de panne censurée de l'équipement = échelle de mesure d'usage du véhicule à la date de fin du RETEX- échelle de mesure d'usage du véhicule à la date de pose de l'équipement.

Mécanisme de censure

La particularité de ce RETEX est que le déploiement des véhicules est fait sur plusieurs années : les dates de mise en service sont donc réparties aléatoirement sur la période de soutien, voir Figure 1.4. On obtient donc :

- des potentiels censurés aléatoirement à droite ;
- un taux de censure dépassant les 70% pour la plupart des équipements.

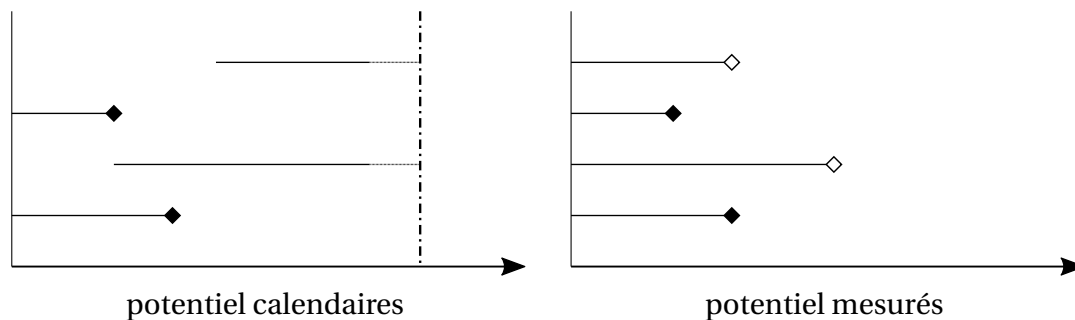


FIGURE 1.4: Inter-occurrences de pannes et censure

1.4 Les équipements fils rouges

Le contrat MCO est établi pour le soutien d'une liste d'équipements appelés *cost driver*. On distinguera :

- les équipements critiques : équipements dont la panne impacte fortement la disponibilité du véhicule et la sûreté de fonctionnement ;
- les équipements ayant un coût d'achat et de remplacement important.

Pour gérer au mieux le maintien des véhicules il nous faut évaluer :

- la durée de vie des équipements, en particulier « les équipements critiques »
- la durée d'acheminement des pièces de rechange (pas traitée dans cette étude).

On s'intéressera tout au long de cette thèse à trois équipements « fil rouge » :

équipement	Technologie	Effectif de l'échantillon	Taux de censure
N°1	Mécanique	599	83%
N°2	Electronique	561	64%
N°3	Mécanique	587	90%

TABLE 1.1 – Équipements « fil rouge » retenus pour l'analyse

Chapitre 2

Flux des pièces de rechange

Sommaire

2.1 Notations	20
2.2 Problématiques du contrat MCO	20
2.3 Simulation du flux	21
2.4 Choix de l'échelle de mesure des inter-occurrences de panne	21
2.4.1 Méthodes	22
2.4.2 Résultats et conclusion	25

Le contrat MCO garantit au client un niveau de disponibilité pour une liste d'équipements tout au long du contrat. Il s'agit alors pour l'industriel d'assurer un soutien optimal en pièces de rechange tout en minimisant les coûts. Le flux en pièces de rechange est donc un indicateur de gestion central pour le contrat MCO.

Pour estimer le flux en pièces de rechange tout en intégrant les différents facteurs d'instabilité, nous aurions pu utiliser des modèles de Cox à risques proportionnels Tibshirani (1997); Fan and Li (2002). Mais un des objectifs de ce travail de thèse est de permettre au constructeur de faire des analyses *a posteriori* de la fiabilité des équipements. Cela permet, par exemple, de faire des analogies entre les centaines d'équipements du véhicule ou d'identifier des causes de défaillances de jeunesse ou d'usure pour certains équipements. Nous avons donc choisi d'utiliser une modélisation par processus de renouvellement superposés : pour chaque véhicule, nous observons un processus de renouvellement pour l'équipement. Pour la gestion d'un contrat de MCO, l'usage des véhicules est naturellement la variable dimensionnant le besoin en pièces de rechange. Les usages des véhicules sont alors recueillis régulièrement. Les fonctionnalités des véhicules étant diverses (*i.e.* motorisation, propulsion, transmission, production d'électricité, défense,...), trois échelles de mesure d'usage sont renseignées. Si pour certains équipements le potentiel (l'échelle de mesure d'usage héritée du véhicule) est explicite (nombre de kilomètres pour les roues, nombre d'heures pour les batteries, nombre de coups tirés pour un système d'armes,...) pour d'autres le choix est moins explicite.

Dans une première section de ce chapitre, nous introduisons les modèles de processus ponctuels utilisés. Dans la deuxième section, nous formalisons les problématiques du constructeur. Dans la dernière, nous nous intéressons au choix d'un potentiel de mesure.

2.1 Notations

La flotte est constituée de κ véhicules. Pour $1 \leq k \leq \kappa$ on note :

$$\begin{aligned} X_{ki} &= i^e \text{ usage entre deux défaillances de l'équipement pour le } k^e \text{ véhicule;} \\ U_k(t) &= \text{usage du } k^e \text{ véhicule, sur } t \text{ années;} \\ N_k(t) &= \text{nombre de défaillances de l'équipement pour } k\text{-ème véhicule, sur } t \text{ années;} \\ \mathbb{P}(N_k(t) \geq n) &= \mathbb{P}(X_{k1} + \dots + X_{kn} \leq U_k(t)) \end{aligned}$$

$$\begin{aligned} N(t) &= \text{nombre total d'équipements défaillants pour les } \kappa \text{ véhicules de la flotte, sur } t \text{ années} \\ &= \sum_{k=1}^{\kappa} N_k(t). \end{aligned}$$

Les équipement défaillants sont remplacés par des équipements neufs. Les potentiels (X_{ki}) sont alors supposés indépendants et identiquement distribués. On note F_X la fonction de répartition de leur loi commune. Dans la suite de la thèse on appellera $N(t)$ le flux d'équipements.

2.2 Problématiques du contrat MCO

Le cadrage du contrat MCO vise d'une part la satisfaction du client, *via* la disponibilité requise pour la flotte, en minimisant les coûts. Et d'autre part, dans le cas d'un contrat de MCO pour un nouveau véhicule il faut être capable de construire des analogies entre les équipements en termes de taux de défaillance et de MTBF (*Mean Time Between Failure*). Par conséquent les objectifs visés par cette étude sont :

- de calculer le taux de défaillance, la fiabilité et le MTBF d'un équipement :

$$\lambda(x) = \frac{f_X(x)}{1 - F_X(x)} \quad (2.1)$$

$$R_X(x) = 1 - F_X(x) \quad (2.2)$$

$$MTBF = \int_0^{+\infty} x f_X(x) dx \quad (2.3)$$

- d'analyser les usages des véhicules afin d'identifier des profils d'usage ;
- d'estimer le stock de pièces de rechange n_{max} afin de garantir une disponibilité contractuelle p :

$$\mathbb{P}(N(t) \leq n_{max}) = p; \quad (2.4)$$

- d'estimer le nombre moyen d'équipements défaillants sur plusieurs années :

$$\mathbb{E}[N(t)] = \mathbb{E}\left[\sum_{k=1}^{\kappa} N_k(t)\right]. \quad (2.5)$$

Pour résoudre Eq. (2.4) et Eq. (2.5), la loi du flux $N(t)$ doit être calculée. Les équipements défaillants étant remplacés par un équipement neuf, les processus de comptage ($N_i(t)$) sont des processus de renouvellement et $N(t)$ est un processus de renouvellement superposé *cf.* Ascher and Feingold (1984). Généralement, la loi de $N(t)$ ne peut pas être obtenue de manière analytique. Nous allons l'estimer par simulations, dans le cas d'une loi simple de Weibull, d'un mélange de lois de Weibull et de deux lois de Weibull en concurrences.

2.3 Simulation du flux

La nature aléatoire des défaillances est liée à la fois à des occurrences de panne de l'équipement mais aussi à l'usage des véhicules sur la période du contrat *cf.* Algorithm 1.

Algorithme 1 : Algorithme de simulation des défaillances

```

Initialisation :  $N \leftarrow 0$ 
 $U \leftarrow$  vehicle usage during a period
 $X \leftarrow$  equipment usage until failure
while  $X < U$  do
   $N \leftarrow N + 1$ 
   $U \leftarrow U - X$ 
   $X \leftarrow$  equipment usage until failure
end
return  $N$ 

```

2.4 Choix de l'échelle de mesure des inter-occurrences de panne

Dans le secteur automobile classique, un véhicule est vendu à des millions de clients et la garantie est limitée soit à une durée après l'achat (des dizaines de mois), soit à un usage (des dizaines de milliers de kilomètres), soit aux deux. Ces limites correspondent à « la durée de vie utile » des véhicules (*i.e.* durée avant vieillissement). Des extensions de garantie permettent éventuellement de prolonger la garantie.

Dans le secteur militaire, une centaine de véhicules est vendue à quelques clients et les contrats sont en dizaines d'années, bien au delà de la durée de vie utile. Les véhicules sont utilisés de manière discontinue avec des périodes d'interruption plus ou moins longues. Pour la gestion d'un contrat de MCO d'une flotte de véhicules, l'usage est donc naturellement la variable dimensionnant le besoin en pièces de rechange. Les usages des véhicules sont alors recueillis régulièrement. Les fonctionnalités des véhicules étant diverses (*i.e.* motorisation, propulsion, transmission, production d'électricité, défense, etc.), le constructeur a choisi de renseigner trois échelles de mesure d'usage dans la base du RETEX. La complexité du traitement nous oblige alors à nous interroger sur le choix d'une échelle de mesure d'usage rendant au mieux compte de la défaillance de l'équipement. Si pour certains équipements le choix d'un potentiel est explicite (nombre de kilomètres pour les roues, nombre d'heures pour les batteries, nombre de coups tirés pour un système d'armes, etc.) pour d'autres le choix est moins clair.

De nombreux travaux abordent l'estimation du flux de pièces de rechange avec deux variables (en général l'âge en mois et les kilomètres parcourus) mais peu s'interrogent sur le choix d'une « bonne échelle ». On trouvera dans Duchesne and Lawless (2000) une tentative pour répondre à cette question à partir d'une étude bibliographique. Dans leur conclusion, les auteurs proposent de retenir les critères suivants : « l'échelle de mesure doit capturer la variation des données de panne et être facilement interprétable en termes de fiabilité ». Les modèles de Cox intégrant plusieurs covariables mesurent bien la variabilité des données, mais sont difficilement interprétables en termes de fiabilité : par exemple ils ne permettent pas d'analogie entre divers équipements. La même remarque peut être faite pour les modèles utilisant une échelle synthétique qui est combinaison linéaire de diverses échelles.

Une des contributions de cette thèse est de proposer une méthode pour sélectionner le potentiel « expliquant au mieux » la défaillance de l'équipement. L'idée est alors d'utiliser des arbres de décision, pour sélectionner la variable la plus « importante » dans la classification « panne / non panne » des équipements, *cf.* Hastie et al. (2009). Afin d'illustrer la méthode, nous utiliserons les données d'un équipement. Les trois potentiels sont représentées dans la Figure 2.1.

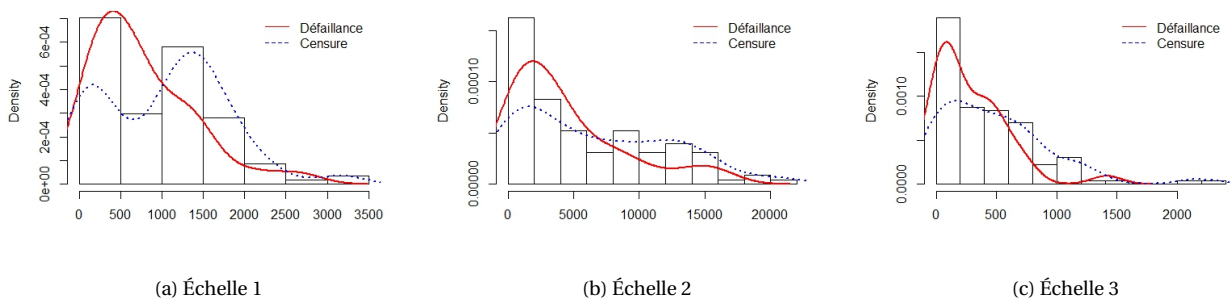


FIGURE 2.1: Densité des trois potentiels d'un équipement sélectionné : en rouge, les défaillances ; en bleu, les censures.

2.4.1 Méthodes

Les arbres de décision (ou segmentation) permettent de faire de la régression pour une variable à expliquer quantitative ou de la classification pour une variable à expliquer qualitative, à partir de variables explicatives quantitatives et qualitatives.

Notre objectif est d'expliquer Y , l'état d'un équipement

$$Y = \begin{cases} 1 & \text{si panne;} \\ 0 & \text{sinon,} \end{cases}$$

à partir des trois potentiels équipement $\mathbf{X} = (P_1, P_2, P_3)$.

CART

L'algorithme CART (Classification And Regression Trees), est une méthode statistique introduite par Breiman et al. (1984). Le principe général de CART est de faire récursivement des partitions des équipements de façon binaire en fonction des valeurs de $\mathbf{X}$, puis de déterminer une sous-partition « optimale » des équipements. Cette sous-partition optimale est un « nœud » de l'arbre. Le premier nœud étant la racine et les derniers, les « feuilles » de l'arbre. Divers critères d'optimalité peuvent être utilisés. Dans le package `rpart` du logiciel libre R que nous avons utilisé, le critère en classification est la minimisation de l'indice de Gini : les nœuds sont alors les sous-partitions les plus homogènes du point de vue de Y , cf. Therneau et al. (2015).

Construction de l'arbre maximal. On appelle *arbre maximal*, l'arbre pleinement développé, de la racine aux feuilles. Dans le même temps, on associe à chaque nœud t de l'arbre une valeur Y_t , correspondant au label de la classe majoritaire des observations présentes dans le nœud t en classification (la valeur de Y_t en régression). De ce fait, à un arbre est associée une partition (définie par ses feuilles) et également des valeurs qui sont attachées à chaque élément de cette même partition.

Élagage. L'arbre maximal s'ajuste aux données et par conséquent il y a un risque de sur-apprentissage. Plus l'arbre est grand, plus le biais de prédiction est faible, mais plus la variance est importante. *A contrario*, plus l'arbre est court, plus le biais de prédiction est fort, mais plus la variance est faible. Une étape supplémentaire peut être ajoutée à l'algorithme CART. Elle consiste alors à établir un compromis biais-variance pour élaguer l'arbre maximal. Pour chaque arbre élagué est calculée une « erreur de généralisation » estimée par validation croisée : il s'agit de la proportion d'observations mal classées en classification (de l'erreur quadratique moyenne en régression). Le sous-arbre optimal est celui qui minimise l'erreur de généralisation cf. Hastie et al. (2009). On peut voir dans

la Figure 2.2, les erreurs de généralisation (normalisées) estimées ainsi que l'arbre optimal retenu.

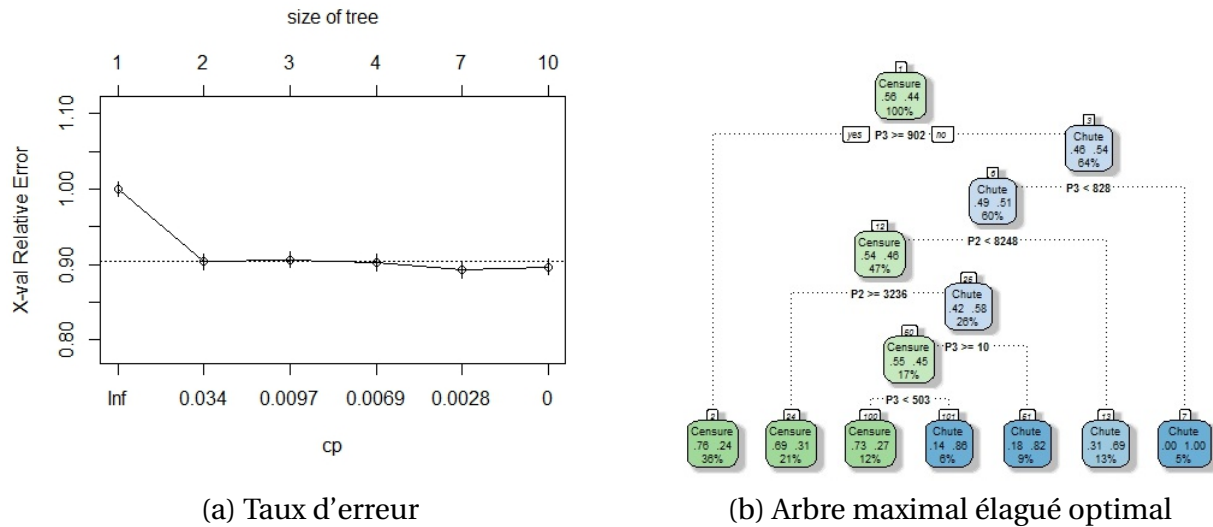


FIGURE 2.2: CART : arbre maximal obtenu avec le package rpart de R, puis élagué.

Importance des variables. Les variables de X qui interviennent dans les premières découpes de l'arbre sont les variables les plus utiles pour expliquer Y : en effet, plus une découpe est proche de la racine plus la décroissance de l'hétérogénéité est importante. On peut définir l'importance d'une variable, selon Breiman et al. (1984), en évaluant en chaque nœud la réduction d'hétérogénéité engendrée par l'usage de la découpe de substitution portant sur cette variable, puis en les sommant sur tous les nœuds. À partir de la Figure 2.3, on constate que P_3 a une importance supérieure à P_1 et P_2 dans la classification « panne / non panne ».

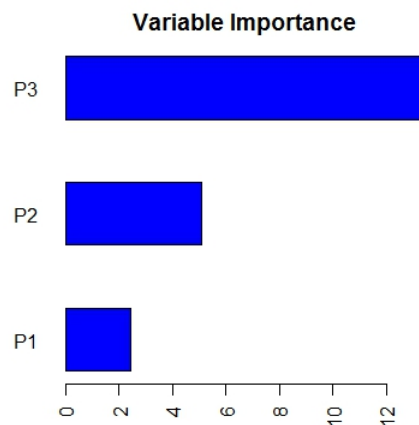


FIGURE 2.3: CART : Importance des variables sur l'arbre optimal

Le défaut majeur de CART est l'instabilité sur des petits échantillons (ce qui n'est pas le cas dans notre exemple avec environ 3 000 données). Une amélioration consiste à jouer sur le compromis biais-variance en conjuguant deux mécanismes : la perturbation aléatoire des arbres et l'agrégation d'un ensemble d'arbres plutôt que la sélection de l'un d'entre eux.

Forêts aléatoires

Une forêt aléatoire (*Random Forest* RF) est l'agrégation d'une collection d'arbres aléatoires de décision cf. Breiman (2001). Plus précisément, les forêts aléatoires « à variables d'entrée aléa-

toires » introduisent deux sources d'aléas pour générer la collection d'arbres :

- l'échantillon d'apprentissage est bootstrappé ;
- pour chaque échantillon bootstrappé, l'arbre est construit de la façon suivante : pour découper un nœud, on tire aléatoirement un nombre m de variables, et on cherche la meilleure coupure uniquement suivant les m variables sélectionnées. L'arbre construit n'est pas élagué.

La forêt d'arbres obtenue est enfin agrégée en prenant le mode en classification (la moyenne en la régression). C'est une méthode statistique de type *machine learning* particulièrement efficace pour sélectionner des variables cf. Genuer et al. (2010). Contrairement à l'arbre de CART, les arbres des RF sont moins simples à visualiser, mais le prédicteur agrégé des RF est meilleur que celui de CART. L'explication heuristique de ces améliorations est que le fait de rajouter des aléas pour construire les arbres, rend ces derniers encore plus différents les uns des autres, sans pour autant dégrader de façon significative leurs performances individuelles.

Lors du déroulement de l'algorithme est calculé une estimation de l'erreur de généralisation : l'erreur Out-Of-Bag (OOB). Les échantillons OOB utilisés pour calculer cette erreur servent aussi à calculer l'importance des variables : l'importance d'une variable est alors l'accroissement moyen de l'erreur d'un arbre de la forêt pour lequel les valeurs observées de cette variable sont permutées au hasard dans les échantillons OOB. Nous avons utilisé le package `randomForest` du logiciel R sur les données. À partir de la Figure 2.4, on constate que P_3 a une importance supérieure à P_1 et P_2 dans la classification « panne / non panne ».

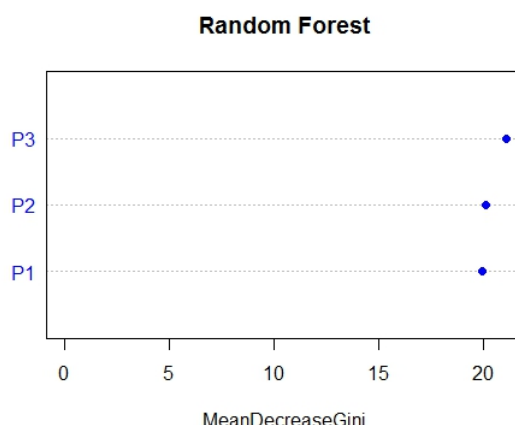


FIGURE 2.4: Random Forest. Importance des variables.

Boosting

Proposé par Freund and Schapire (1996), le Boosting diffère des approches précédentes : il s'agit d'un algorithme adaptatif. À l'instar des RF, le Boosting construit une collection d'arbres aléatoires de décision qui seront ensuite agrégés pour fournir un prédicteur. Mais contrairement aux RF, les arbres du boosting ne sont pas construits de manière indépendante. L'idée du Boosting est de se focaliser sur les observations mal prédites à l'étape précédente, en vue d'améliorer les performances globales. En orientant la construction des arbres à chaque étape, l'algorithme de Boosting améliore le biais. En agrégeant les arbres il améliore la variance. Nous avons utilisé le package `gbm` cf. Ridgeway (2007) du logiciel R sur les données. À partir de la Figure 2.5, on constate que P_3 a une importance supérieure à P_1 et P_2 dans la classification « panne / non panne ».

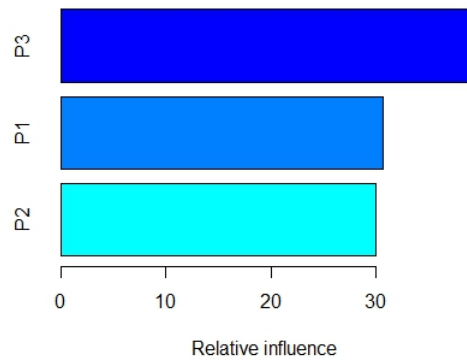


FIGURE 2.5: Boosting. Importance des variables.

2.4.2 Résultats et conclusion

Les trois algorithmes proposés, CART, Random Forest et Boosting désignent le potentiel représenté par l'échelle 3 comme la plus importante pour expliquer la défaillance de cet équipement. C'est donc cette unité qui sera utilisée pour modéliser les inter-occurrences de panne et simuler l'usage des véhicules afin d'estimer le flux en pièces de rechange.

La problématique du choix de variables pour l'analyse va être de plus en plus fréquente chez les constructeurs. En effet, de nos jours l'acquisition de données est devenue de plus en plus aisée du point de vue technologique, ce qui paradoxalement peut rendre l'analyse et l'interprétation plus compliquées. Le dimensionnement des garanties constructeurs se font généralement soit selon l'âge du véhicule, soit selon les kilomètres parcourus, soit selon des deux. Pour des véhicules avec divers usages (*e.g.* véhicules utilitaires ou véhicules militaires), le nombre de mesures d'usage peut être plus important. Les contrats de maintenance en conditions opérationnelles, sur des durées allant bien au-delà de la durée de vie utile des équipements, nous obligent à nous interroger sur le choix d'une échelle d'usage. Si pour certains équipements le choix est naturel, pour d'autre cela l'est moins. Notre contribution dans ce domaine consiste à sélectionner le potentiel (l'échelle d'usage héritée du véhicule) le plus important dans la classification en « panne / non panne » des données. Les méthodes de classification à base d'arbres de décision, CART, Random Forest et Boosting, *via* la notion d'importance des variables, nous offre alors un outil simple (peu de réglages) et très performant pour répondre à cette difficile problématique.

Deuxième partie

MODÉLISATION DES INTER-OCCURRENCES DE PANNE DES ÉQUIPEMENTS

Chapitre 3

Estimations paramétriques

Sommaire

3.1	Modèle et données	29
3.2	Estimations par maximum de vraisemblance	30
3.2.1	Vraisemblances	30
3.2.2	Maximum de vraisemblance	30
3.2.3	EM et Stochastic EM	30
3.3	Estimations bayésiennes	33
3.3.1	Notations	34
3.3.2	Le choix bayésien	34

Une fois choisi le modèle paramétrique de lois on cherche à estimer les paramètres du modèle s'ajustant « au mieux » aux données du RETEX.

Il existe différentes méthodes d'estimation des paramètres et leurs performances dépend de la qualité de l'information du RETEX, cf. Lawless (1982); Nelson et al. (1982); Lannoy (1994) et Meeker and Escobar (1998) pour une revue détaillée. La taille tout d'abord : plus il y a de données de panne, plus l'estimation sera précise. Cependant, sur la période du RETEX, certains équipements étudiés ont été défaillants, d'autres sont toujours opérationnels : on distingue donc « les pannes observées » et « les pannes censurées » cf. Figure 3.1. Pour le même nombre de données, un taux de censure élevé réduit l'information sur la défaillance de l'équipement et par conséquent impacte la performance des méthodes d'estimation. On peut remarquer toutefois que ces informations censurées sont très intéressantes pour les industriels. Nous allons dans ce chapitre rappeler le principe des méthodes d'estimations et préciser la typologie des censures du RETEX analysé.

3.1 Modèle et données

Pour $\theta \in \mathbb{R}^d$ et $d \geq 1$, notons $f(x|\theta)$ la densité du modèle et $R(x|\theta)$ la fiabilité.

On cherche à estimer θ à partir des inter-occurrences de panne $\mathbf{x} = (x_1, \dots, x_n)$. Lorsque l'équipement est toujours opérationnel à la fin du RETEX, l'observation est dite « censurée » :

$$x_i = i^e \text{ inter-occurrence de panne de l'équipement;}$$
$$c_i = \begin{cases} 1 & \text{si } x_i \text{ est censuré;} \\ 0 & \text{sinon.} \end{cases}$$

Compte tenu du fait que les véhicules sont mis en service à différentes dates, il s'agit de « censures aléatoire à droite » cf. Figure 3.1.

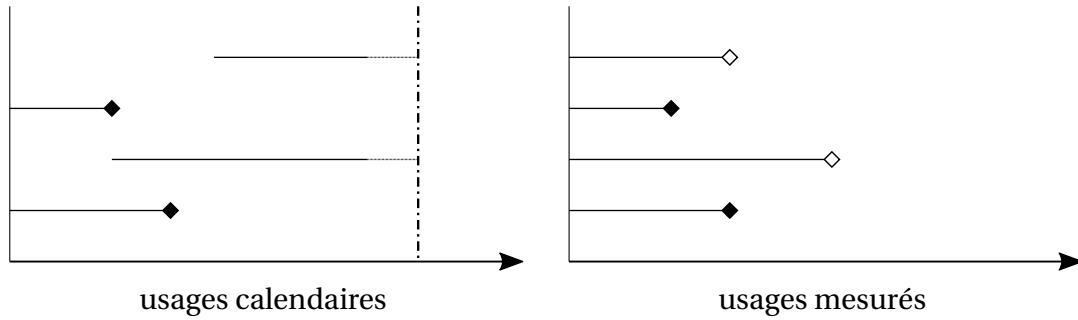


FIGURE 3.1: Inter-occurrences de pannes et censure aléatoire à droite

3.2 Estimations par maximum de vraisemblance

L'estimation par maximum de vraisemblance consiste à déterminer les valeurs du paramètre du modèle les plus « compatibles » avec les données, c.-à-d les valeurs pour lesquelles la vraisemblance du modèle aux des données est la plus grande.

3.2.1 Vraisemblances

La vraisemblance $L(\mathbf{x}|\boldsymbol{\theta})$ correspond à la probabilité d'observer les données :

$$L(x_i) = \begin{cases} f(x_i|\boldsymbol{\theta}) & \text{si } c_i = 0; \text{ « donnée observée »,} \\ R(x_i|\boldsymbol{\theta}) & \text{si } c_i = 1. \text{ « donnée censurée à droite ».} \end{cases}$$

Elle représente alors toute la connaissance sur $\boldsymbol{\theta}$ apportée par les données $\mathbf{x}$:

$$\begin{aligned} L((\mathbf{x}, \mathbf{c})|\boldsymbol{\theta}) &= \prod_{i:c_i=0} f(x_i|\boldsymbol{\theta}) \cdot \prod_{i:c_i=1} R(x_i|\boldsymbol{\theta}) \\ &= \text{données non-censurées} \quad \text{données censurées} \end{aligned}$$

3.2.2 Maximum de vraisemblance

La log-vraisemblance $\ln(L)$ est souvent plus simple à maximiser. On définira alors l'estimateur du maximum de vraisemblance comme suit :

$$\hat{\boldsymbol{\theta}}_{MV} = \arg \max_{\boldsymbol{\theta}} \ln(L(\mathbf{x}|\boldsymbol{\theta})). \quad (3.1)$$

On l'obtient en résolvant les équations de vraisemblance $\partial \ln((\boldsymbol{\theta}|\mathbf{x}_n)) / \partial \boldsymbol{\theta} = 0$.

Intervalle de confiance

On peut montrer que sous de bonnes conditions cf. Annexe B de Meeker and Escobar (1998) il suit asymptotiquement une loi normale : $\sqrt{n}(\hat{\boldsymbol{\theta}}_{MV} - \boldsymbol{\theta})$ converge en loi vers $\mathcal{N}(0, I^{-1}(\boldsymbol{\theta}))$ où $I(\boldsymbol{\theta})$ est la matrice de Fisher (variance-covariance). En estimant $I(\boldsymbol{\theta})$, on peut alors construire un intervalle de confiance pour les paramètres, et en utilisant la « delta-méthode », des intervalles de confiance de fonctions des paramètres cf. Meeker and Escobar (1998).

3.2.3 EM et Stochastic EM

En fiabilité, il est fréquent d'avoir des taux de censures très élevés (*i.e.* supérieurs à 70%). L'estimateur du maximum de vraisemblance est alors de moins en moins précis. La censure peut être

vue comme une donnée « manquante » : on n'observe pas la panne. L'algorithme EM *Expectation-Maximization* est une alternative à la méthode d'estimation directe du maximum de vraisemblance dans le cas de données manquantes cf. Dempster et al. (1977); McLachlan and Krishnan (2007).

Dans le cas de données censurées, c'est donc une bonne alternative à une optimisation numérique directe de la vraisemblance : dans l'échantillon $\mathbf{x}$, il y a une partie de « données complètes », $(x_i, c_i = 0)$, et une partie de données incomplètes, partiellement observées $(x_i, c_i = 1)$.

L'idée est alors de « compléter » l'échantillon afin de calculer (plus aisément) le maximum de vraisemblance à l'aide de la vraisemblance sur les données ainsi complétées. Dans le cas de l'algorithme EM, il s'agit donc d'une méthode itérative en deux étapes : à l'étape r avec $\boldsymbol{\theta}^{(r)}$ connu,

Expectation à partir d'une valeur du paramètre $\boldsymbol{\theta}^{(r)}$, on calcule l'espérance de la log-vraisemblance des données complétées pour $\boldsymbol{\theta}$, sous la loi conditionnelle des données observées pour $\boldsymbol{\theta}^{(r)}$, $f(\cdot | (x_i, c_i = 0), ; \boldsymbol{\theta}^{(r)})$:

$$\mathbb{Q}(\boldsymbol{\theta} | \boldsymbol{\theta}^{(r)}) = \int \ln L(\mathbf{x}) | \boldsymbol{\theta} f((x_i, c_i = 1) | (x_i, c_i = 0); \boldsymbol{\theta}^{(r)}). \quad (3.2)$$

Maximization On maximise alors $\mathbb{Q}(\boldsymbol{\theta} | \boldsymbol{\theta}^{(r)})$:

$$\boldsymbol{\theta}^{(r+1)} = \arg \max_{\boldsymbol{\theta}} \mathbb{Q}(\boldsymbol{\theta} | \boldsymbol{\theta}^{(r)}).$$

L'algorithme EM s'écrit alors

Algorithme 2 : Algorithme EM

Initialisation : on se donne $\boldsymbol{\theta}^{(0)}$

pour $r = 0, \dots, B - 1$ **faire**

Étape Espérance :

 Sachant $\boldsymbol{\theta}^{(r)}$ on calcule l'espérance conditionnelle de :

$$\mathbb{Q}(\boldsymbol{\theta} | \boldsymbol{\theta}^{(r)}) = \mathbb{E} [\ln (L(\mathbf{X} | \boldsymbol{\theta})) | (\mathbf{X}, C = 0); \boldsymbol{\theta}^{(r)}].$$

Étape Maximisation :

 On détermine $\boldsymbol{\theta}^{(r+1)} = \arg \max_{\boldsymbol{\theta}} \mathbb{Q}(\boldsymbol{\theta} | \boldsymbol{\theta}^{(r)})$

fin

La suite $(\boldsymbol{\theta}^{(r)})$, ainsi construite, fait croître la vraisemblance et converge vers $\hat{\boldsymbol{\theta}}_{MV} : (L(\mathbf{x} | \boldsymbol{\theta}^{r+1}) > L(\mathbf{x} | \boldsymbol{\theta}^r))$ cf. §3 de McLachlan and Krishnan (2007). Généralement la convergence dépend fortement de l'initialisation et la convergence peut être lente. Par ailleurs la méthode peut être piégée par les maxima locaux. L'algorithme S-EM a été conçu pour surmonter les difficultés de l'algorithme EM Celeux et al. (1996) Son principe est de remplacer l'estimation de l'espérance conditionnelle (3.2) par la simulation des données manquantes $(x_i, c_i = 1)$ avec la loi $f(x | X > x_i; \boldsymbol{\theta}^{(r)})$.

Algorithme 3 : Algorithme S-EM**Initialisation :** on se donne $\theta^{(0)}$ **pour** $r = 1, \dots, B$ **faire** *Étape Simulation :* **pour** $i = m + 1, \dots, n$ **faire** sachant $\theta^{(r-1)}$ on simule les durées de vie censurées suivant la loi $f(x | x > x_i; \theta^{(r-1)})$ **fin** *Étape M :* on calcule $\theta^{(r)}$ estimateur du maximum de vraisemblance des données complétées :
 $(x_1, \dots, x_m; \tilde{x}_{m+1}, \dots, \tilde{x}_n)$ cf. (3.1)**fin****Estimateur S-EM :**On obtient $\theta^{(r)}, \dots, \theta^{(B)}$. Les M premières valeurs correspondant à une période de « chauffe ». On estime alors θ par $\bar{\theta}_{S-EM}$

(3.3)

La suite d'estimateur ainsi obtenus constitue une chaîne de Markov homogène et ergodique dont la loi stationnaire a pour moyenne l'estimation du maximum de vraisemblance.

Les algorithmes EM et S-EM ont fait l'objet de nombreux travaux. Dans un cadre de fiabilité, on citera Lannoy (1994); Celeux et al. (1996). Les propriétés de convergence de cet algorithme ont été discutées Nielsen et al. (2000).

Il n'est pas inutile de rappeler que la convergence de EM et S-EM est très liée à l'initialisation de l'algorithme. L'initialisation doit donc être discutée selon le modèle.

Intervalle de crédibilité

L'algorithme itératif S-EM fournit une suite d'estimateurs $\hat{\theta}^{(M+1)}, \dots, \hat{\theta}^{(B)}$. On utilise alors l'estimateur :

$$\bar{\theta} = \frac{1}{B-l} \sum_{i=l+1}^B \hat{\theta}^i.$$

où l représente les premières valeurs correspondant à une période de « chauffe »,

Pour obtenir un intervalle de confiance au niveau α , on utilise la loi empirique des $(\hat{\theta}^{(M+1)}, \dots, \hat{\theta}^{(B)})$:

$$\hat{\theta} \in [\theta_{(\alpha/2)}; \theta_{(1-\alpha/2)}] \quad (3.4)$$

où $\theta_{(\alpha)}$ est le quantile empirique des $(\hat{\theta}^{(M+1)}, \dots, \hat{\theta}^{(B)})$, pour la probabilité α . Pour la fiabilité $\hat{R}$,

$$\hat{R} \in [\underline{R}_{(\alpha/2)}; \bar{R}_{(1-\alpha/2)}] \quad (3.5)$$

où $R_{(\alpha)}$ est le quantile empirique des $(\hat{R}^{(M+1)}, \dots, \hat{R}^{(B)})$, pour la probabilité α .

3.3 Estimations bayésiennes

Dans le cas d'échantillons de petite taille ($n=100$) et/ou fortement censurés ($> 70\%$), les algorithmes EM et S-EM n'apportent pas une nette amélioration à l'estimateur du maximum de vraisemblance cf. Bacha and Celeux (1996). L'estimation bayésienne peut être alors un recours aux méthodes d'estimation, par maximum de vraisemblance.

Dans l'approche bayésienne, une information supplémentaire est apportée sur le paramètre : on suppose que l'on a un *a priori* sur les valeurs du paramètre. En effet en fiabilité, il est fréquent d'avoir des « avis d'experts » sur les données de pannes d'un composant.

Des études antérieures, soit sur des versions précédentes du composant, soit sur des composants similaires, permettent aussi d'obtenir des informations *a priori* sur loi de la durée de vie. En ajoutant ces informations sur les paramètres aux informations apportées par les données, on peut améliorer l'estimation faite par la méthode du maximum de vraisemblance (encore faut-il que ces informations *a priori* soient cohérentes). D'un point de vue mathématique on modélise l'*a priori* sur les paramètres, en considérant ceux-ci comme des variables aléatoires donc on connaît la loi $\pi(\theta)$. La formule de Bayes nous permet alors d'obtenir une loi *a posteriori* du paramètre, $\pi(\theta | x)$. Cette loi nous fournit alors une estimation du paramètre en prenant le mode, la médiane ou la moyenne, voir la figure ci-après cf Gelman et al. (2003).

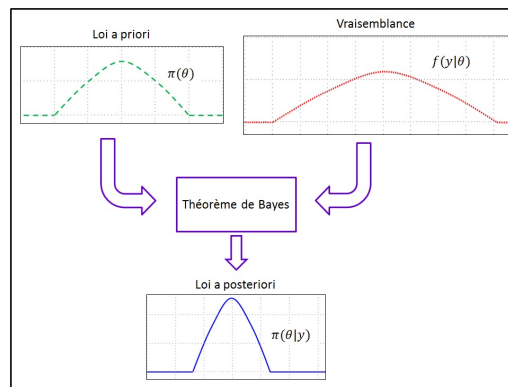


FIGURE 3.2: Inférence bayésienne

Pour mettre en œuvre l'estimation bayésienne, il faut

1. définir une loi *a priori* $\pi(\theta)$ sur le vecteur de paramètre θ , capable de prendre en compte les informations initiales ;
2. calculer la loi *a posteriori* $\pi(\theta|x)$.

L'estimation bayésienne est donc utilisée soit quand la méthode du maximum de vraisemblance n'est pas efficace, soit lorsque l'on dispose d'avis d'experts. Les travaux sur l'approche bayésienne sont nombreux Gelman et al. (2003); Robert (2006) et Congdon (2007) par exemple; et en particulier sur l'estimation bayésienne en fiabilité Bousquet (2006); Hamada et al. (2008).

3.3.1 Notations

On reprend les notations précédentes. On suppose que l'on a un *a priori* sur les valeurs du paramètre avant d'observer les données : $\theta \sim \pi(\theta)$. On cherche alors à estimer θ *a posteriori* à partir des données. En utilisant la formule de Bayes, on obtient $\pi(\theta | (\mathbf{x}, \mathbf{c}))$, la loi *a posteriori* de θ connaissant les données $(\mathbf{x}, \mathbf{c}) = ((x_1, c_1), \dots, (x_n, c_n))$:

$$\pi(\theta | (\mathbf{x}, \mathbf{c})) \propto L((\mathbf{x}, \mathbf{c}) | \theta) \pi(\theta),$$

où $L((\mathbf{x}, \mathbf{c}) | \theta)$ est la vraisemblance des données vis à vis du paramètre θ et $\pi(\theta)$ la loi *a priori* des paramètres.

À partir de la loi *a posteriori*, on peut alors choisir comme estimateur de Bayes cf. Robert (2006) :

$$\hat{\theta}_B = \arg \max_{\theta} \mathbb{E}[(\theta - \hat{\theta})^2 | (\mathbf{x}, \mathbf{c})] = \mathbb{E}[\theta | (\mathbf{x}, \mathbf{c})] = \int_{\theta} \theta \pi(\theta | (\mathbf{x}, \mathbf{c})). \quad (3.6)$$

Il faut alors évaluer l'estimateur (3.6) :

- soit de manière explicite en utilisant une loi *a priori* conjuguée à la vraisemblance : π et $\pi(\theta | (\mathbf{x}, \mathbf{c}))$ sont alors de la même famille de lois ;
- soit de manière approchée : en approchant l'intégrale de (3.6) (approximation de Laplace, méthode variationnelle), soit en estimant l'espérance de 3.6) par échantillonnage MCMC (Metropolis, Gibbs).

3.3.2 Le choix bayésien

Le choix d'une loi *a priori* est alors le point clé de l'estimation bayésienne :

- soit : on dispose d'« avis d'experts », de données de fiabilité ou issues de calculs physique, dans ce cas on utilise des lois informatives ;
- soit : on utilisera des lois peu informatives.

Le choix a bien évidemment un impact sur la loi *a posteriori*. Prenons l'exemple simple d'une loi exponentielle.

Propriété 3.1 Soit $\mathcal{E}(\lambda)$ la loi exponentielle de taux λ , pour $x \geq 0$.

$$f(x | \lambda) = \lambda e^{-\lambda x}. \quad (3.7)$$

On choisit une loi gamma $\mathcal{G}(a, b)$ comme loi *a priori* pour λ :

$$\pi(\lambda) = \frac{b^a}{\Gamma(a)} \lambda^{a+1} e^{-b\lambda} \propto \lambda^{a+1} e^{-b\lambda}. \quad (3.8)$$

C'est une loi « conjuguée » pour le taux de la loi exponentielle : la loi *a posteriori* de λ est aussi une loi gamma,

$$\lambda | (\mathbf{x}, \mathbf{c}) \sim \mathcal{G}(a + n - \sum_{i=1}^n c_i, b + \sum_{i=1}^n x_i). \quad (3.9)$$

L'estimateur bayésien (pour une fonction de perte quadratique) est alors :

$$\hat{\lambda}_B = \frac{a + n - \sum_{i=1}^n c_i}{b + \sum_{i=1}^n x_i}. \quad (3.10)$$

Preuve : La vraisemblance s'écrit

$$L((\mathbf{x}, \mathbf{c}) | \lambda) = \prod_{i=1}^n \left(\lambda e^{-\lambda x_i} \right)^{1-c_i} \cdot \left(e^{-\lambda x_i} \right)^{c_i} = \lambda^{n - \sum_{i=1}^n c_i} e^{-\lambda \sum_{i=1}^n x_i}.$$

La loi *a posteriori* est alors :

$$\begin{aligned} \pi(\lambda | (\mathbf{x}, \mathbf{c})) &\propto \pi(\lambda) \cdot L((\mathbf{x}, \mathbf{c}) | \lambda) \\ &\propto \lambda^{a-1} e^{-b\lambda} \cdot \lambda^{n - \sum_{i=1}^n c_i} e^{-\lambda \sum_{i=1}^n x_i} \\ &= \lambda^{a+1+n - \sum_{i=1}^n c_i} e^{-\lambda(b + \sum_{i=1}^n x_i)}. \end{aligned}$$

La loi *a posteriori* $\pi(\lambda | (\mathbf{x}, \mathbf{c}))$ est donc aussi une loi gamma : $\mathcal{G}(a + n - \sum_{i=1}^n c_i, b + \sum_{i=1}^n x_i) \square$

On peut constater que $\hat{\lambda}$ est un barycentre entre $\frac{a}{b}$, l'espérance de la loi *a priori* et l'estimateur du maximum de vraisemblance $\hat{\lambda}_{MV} = \frac{n - \sum_{i=1}^n c_i}{\sum_{i=1}^n x_i}$. L'estimateur bayésien sera proche de l'estimateur du maximum de vraisemblance, si la taille de l'échantillon est grande ($n \rightarrow +\infty$) ou si la loi *a priori* est non informative ($b \rightarrow 0$ ou $b \rightarrow +\infty$). Il sera proche de l'espérance de la loi *a priori* (en particulier s'il y a beaucoup de censures). Cependant, si la loi *a priori* n'est pas « cohérente » avec les données, alors l'estimation bayésienne sera imprécise ou impossible. Pour le choix de la loi *a priori*, il y a donc un compromis à faire, en fonction des informations disponibles sur les paramètres. On trouvera dans Bousquet (2006) une approche pour modéliser les avis d'expert en loi *a priori*.

Par ailleurs une loi *a priori* est dite conjuguée pour un paramètre du modèle si la loi *a posteriori* appartient à la même famille. Cela peut permettre d'obtenir une expression analytique de l'estimateur de Bayes cf. Hamada et al. (2008).

Chapitre 4

Modèle de loi de Weibull simple

Sommaire

4.1	Caractéristiques d'une loi de Weibull à deux paramètres	38
4.2	Données	40
4.3	Ajustement graphique	40
4.4	Estimation par maximum de vraisemblance sans censure	41
4.5	Estimation par maximum de vraisemblance avec censures	42
4.6	Algorithme S-EM	43
4.7	Estimation bayésienne	45
4.8	Estimation Bayesian Restoration Maximization	47
4.8.1	Intervalle de crédibilité	50
4.8.2	Calibrage du nombre de répétitions de BRM	50
4.8.3	Calibrage de la loi <i>a priori</i> sur le paramètre de forme, β	51
4.8.4	Calibrage de la loi <i>a priori</i> sur le paramètre de position, η	52
4.9	Analyses de sensibilité	55
4.10	Conclusion	56
4.11	Données réelles	56

On considère ici que les données sont « homogènes », en d'autres termes les équipements ont une seule cause de défaillance. On cherche alors un modèle paramétrique simple pour ajuster les données. Depuis plus de cinquante ans, la loi de Weibull (1939) a été utilisée pour modéliser des durées de vie dans un très grand nombre d'applications. En effet, ce modèle de loi offre une grande flexibilité pour ajuster des données de durée de vie en médecine ou en fiabilité des systèmes. Un outil graphique proposé dans ce chapitre, permet d'indiquer la validité de ce choix de modèle de loi.

Le plan de ce chapitre est le suivant. Dans la première section nous rappelons les caractéristiques du modèle de loi de Weibull simple (*i.e.* à deux paramètres). Les travaux sur l'estimation des paramètres d'une simple loi de Weibull sont très nombreux, selon le type de censure (*e.g.* à droite, par intervalle) ou la méthode d'estimation (*e.g.* ajustement graphique, maximum de vraisemblance, EM et S-EM). Murthy et al. (2004b) et Rinne (2008) constituent une référence détaillée des méthodes usuelles d'estimation. Nous les rappelons dans les sections qui suivent. Dans le cas de données fortement censurées, la performance des méthodes classiques est très réduite. Une méthode de type *Bayesian bootstrap* a alors été utilisée dans Bacha and Celeux (1996). Elle est présentée et mise en œuvre dans une section spécifique. Nous concluons le chapitre par une analyse de sensibilité et une comparaison des différentes méthodes selon le taux de censure et la taille de l'échantillon.

4.1 Caractéristiques d'une loi de Weibull à deux paramètres

En fiabilité, la loi de Weibull est utilisée pour modéliser des pannes dues à des contraintes importantes (principe du « maillon faible »), soit pour intégrer des phénomènes de vieillissement du système cf. Nelson et al. (1982); Meeker and Escobar (1998); Murthy et al. (2004b); Procaccia et al. (2011).

On peut voir dans la Figure 4.1 différentes variations du taux de défaillance :

- décroissant pour $\beta < 1$: les défaillances de ce type sont dues à un défaut de conception. L'équipement est dans une « phase de jeunesse »;
- constant pour $\beta = 1$: les défaillances de ce type sont dues au hasard. L'équipement est dans une « phase de maturité »;
- croissant pour $\beta > 1$: les défaillances de ce type sont dues à l'usure. L'équipement est dans une « phase de vieillissement ».

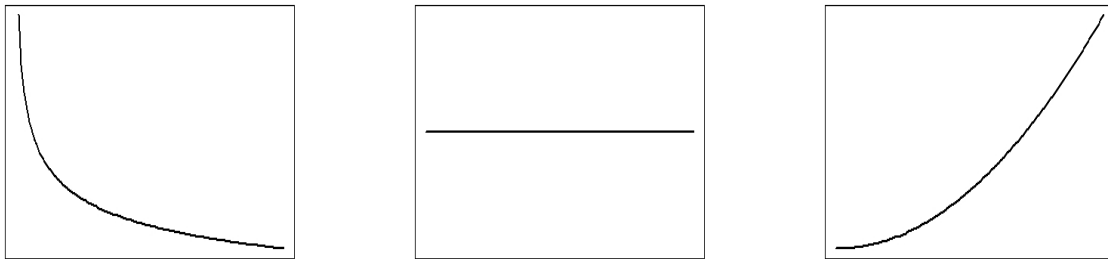
(a) $\beta = 0,8$ (b) $\beta = 1$ (c) $\beta = 3$

FIGURE 4.1: Taux de défaillance de la loi de Weibull $\mathcal{W}(\beta, \eta)$ pour différentes valeurs du paramètre de forme β .

Par ailleurs, la loi de Weibull est une des lois limites du théorème des valeurs extrêmes :

Théorèmes des valeurs extrêmes cf. Fisher and Tippet (1928). Soit $X_1, \dots, X_n$, n variables aléatoires et $M_n = \max_{1 \leq i \leq n} X_i$. S'il existe deux suites de constantes (a_n) et (b_n) telles que

$$\mathbb{P}\left(\frac{M_n - a_n}{b_n} \leq x\right) \rightarrow G(x),$$

alors G est soit une loi de Gumbel, soit une loi de Fréchet, soit une loi de Weibull. Par ailleurs l'inverse d'une loi de Fréchet est une loi de Weibull et l'exponentielle d'une loi de Gumbel est une loi de Weibull.

Ce théorème dit « du maillon faible », montre que la loi de Weibull modélise bien la durée de vie d'un système constitué d'un grand nombre de composants en série.

Nous présentons dans ce chapitre les méthodes d'estimations de (β, η) pour une loi de Weibull dans le cas d'un échantillon censuré.

Loi de Weibull à deux paramètres : soit $X \sim \mathcal{W}(\beta, \eta)$. Le paramètre β est le paramètre de forme de la loi de Weibull et η le paramètre d'échelle.

Densité

$$f_W(x|\beta, \eta) = \frac{\beta}{\eta} \left(\frac{x}{\eta}\right)^{\beta-1} e^{-\left(\frac{x}{\eta}\right)^\beta} \quad (4.1)$$

Fiabilité

$$R_W(x|\beta, \eta) = \mathbb{P}(X > x) = e^{-\left(\frac{x}{\eta}\right)^\beta} \quad (4.2)$$

Taux de défaillance

$$\lambda_W(x|\beta, \eta) = \frac{f(x)}{R(x)} = \frac{\beta}{\eta} \left(\frac{x}{\eta}\right)^{\beta-1}. \quad (4.3)$$

MTBF

$$\mathbb{E}[X] = \eta \Gamma\left(\frac{1}{\beta} + 1\right). \quad (4.4)$$

où $\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt$.

Mode

$$\text{Mode}(X) = \begin{cases} 0 & \text{si } \beta < 1; \\ \eta \left(1 - \frac{1}{\beta}\right)^{\frac{1}{\beta}} & \text{si } \beta \geq 1. \end{cases} \quad (4.5)$$

Quantiles

$$x_p = \eta \left[\ln\left(\frac{1}{1-p}\right) \right]^{\frac{1}{\beta}}. \quad (4.6)$$

Propriété 4.1 (Reparamétrisation) Si $X \sim \mathcal{W}(\beta, \eta)$ alors

1. $\frac{X^\beta}{\eta} \sim \mathcal{E}(1)$, loi exponentielle de densité

$$f(x) = e^{-x} \mathbb{1}_{\mathbb{R}^+}(x)$$

2. $\beta (\ln(X) - \ln(\eta)) \sim \mathcal{G}$, loi de Gumbel (loi du logarithme d'une exponentielle) de densité :

$$f(x) = e^{x-e^x} \mathbb{1}_{\mathbb{R}}(x).$$

On en déduit

- que la loi de Weibull avec un paramètre d'échelle, $\mathcal{W}(\beta, \eta)$, peut s'écrire avec un taux, c.-à-d $\mathcal{W}(\beta, \lambda = \frac{1}{\eta^\beta})$:

$$R(x) = e^{-\lambda x^\beta}, \quad (4.7)$$

- que le log-quantile d'une loi de Weibull est une fonction linéaire des paramètres :

$$\ln(x_p) = \frac{1}{\beta} q_{\text{Gumbel}} + \ln(\eta). \quad (4.8)$$

où

$$q_{\text{gumbel}}(p) = \ln\left(\ln\left(\frac{1}{1-p}\right)\right)$$

est le p -quantile d'une loi de Gumbel.

Propriété 4.2 (Simulation d'une loi de Weibull) Si $X \sim \mathcal{W}(\beta, \eta)$ alors

$$F(x) = \mathbb{P}(X \leq x) = 1 - e^{-\left(\frac{x}{\eta}\right)^\beta}.$$

Donc si $U \sim \mathcal{U}([0; 1])$ alors

$$X = \eta [-\ln(U)]^{\frac{1}{\beta}} \sim \mathcal{W}(\beta, \eta).$$

De plus, pour $x_0 > 0$,

$$= \mathbb{P}(X \leq x | X > x_0) = 1 - e^{-\left(\frac{x}{\eta}\right)^\beta + \left(\frac{x_0}{\eta}\right)^\beta}.$$

Donc si $U \sim \mathcal{U}([0; 1])$ alors

$$\eta \left[\left(\frac{x_0}{\eta} \right)^\beta - \ln(U) \right]^{\frac{1}{\beta}} \sim F(x | X > x_0). \quad (4.9)$$

4.2 Données

On considère n équipements. Pour $i = 1, \dots, n$, on note

$$\begin{aligned} x_i & \text{ } i^{\text{e}} \text{ inter-occurrence de panne de l'équipement;} \\ c_i & \text{ indicateur de censure} \\ & = \begin{cases} 1 & \text{si } x_i \text{ est censuré} \\ 0 & \text{sinon.} \end{cases} \end{aligned}$$

On note $(\mathbf{x}, \mathbf{c}) = ((x_1, c_1), \dots, (x_n, c_n))$.

Sans perte de généralité on suppose que :

$$c_i = \begin{cases} 0 & \text{pour } 1 \leq i \leq m; \\ 1 & \text{pour } m+1 \leq i \leq n. \end{cases}$$

où n le nombre de données complètes.

4.3 Ajustement graphique

On suppose ici que $c_i = 0$ pour $i \leq m$ et $c_i = 1$ pour $i > m$.

Pour valider l'ajustement d'une loi de Weibull à un jeu de données, on peut utiliser un Weibull Quantile-Quantile plot (WQQP) : Soit $\mathbf{x} = (x_1, \dots, x_n)$ un n -échantillon de X . On note $x_{(1)} < x_{(2)} < \dots < x_{(n)}$ les statistiques d'ordre de l'échantillon. En utilisant l'Equation 4.8, si $X \sim \mathcal{W}(\beta; \eta)$ alors

- pour les données complètes : les points $\left(\ln(x_{(i)}); qgumbel\left(\frac{i-0,5}{n}\right) \right)$ sont alignés;
- pour les données censurées : on estime la fiabilité R par $\hat{R}$ l'estimateur de Kaplan-Meier Kaplan and Meier (1958) cf. Meeker and Escobar (1998). Les points $(\ln(x_{(i)}); qgumbel(1 - \hat{R}(x_{(i)})))$ sont approximativement alignés.

On peut voir dans la Figure 4.2 un exemple de WQQP.

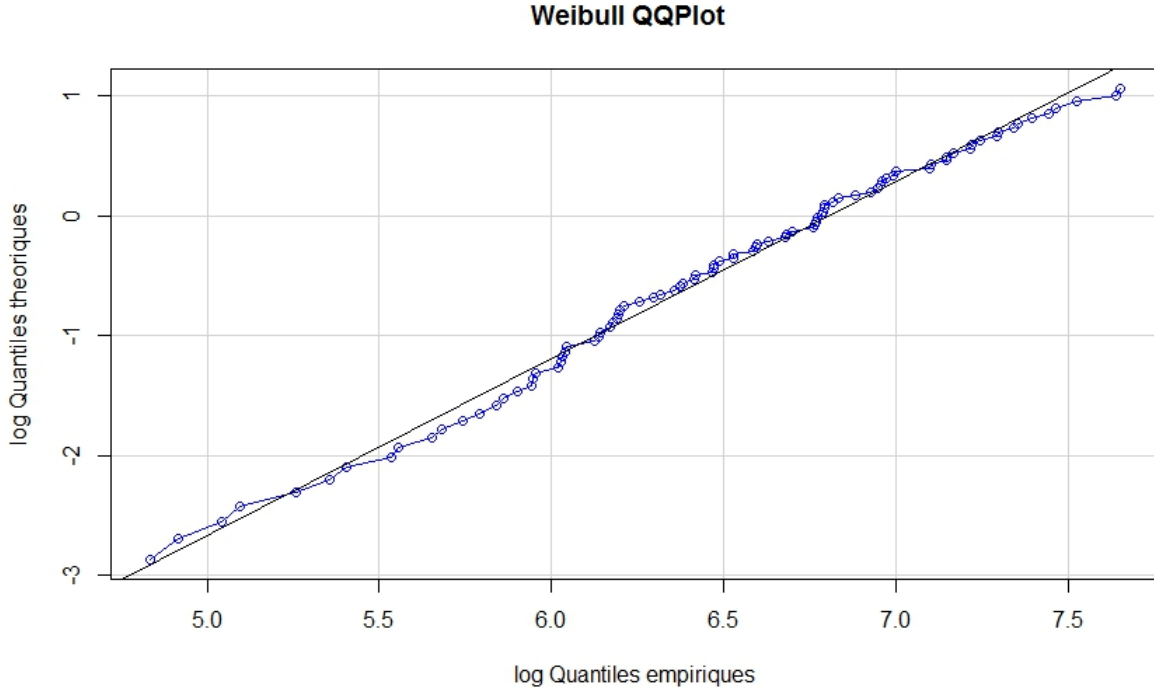


FIGURE 4.2: Weibull simple : Weibull quantile-quantile plot, en abscisse p -log-quantile empirique et en ordonnée p -quantile théorique d'une loi de Gumbel.

Le Weibull QQP sert donc de test graphique pour valider l'ajustement des données par une loi de Weibull à deux paramètres, mais il peut aussi servir de méthode d'estimation. On peut en effet obtenir une estimation de $(\hat{\beta}, \hat{\eta})$ par régression linéaire sur le nuage de points. Cette estimation est très sommaire, mais peut servir d'initialisation pour d'autres méthodes d'estimation (e.g. EM, S-EM, et BRM) que nous allons voir plus loin.

4.4 Estimation par maximum de vraisemblance sans censure

La vraisemblance de l'échantillon complet (*i.e.* pas de censure, $c_i = 0$ pour tous les i) s'écrit

$$\begin{aligned} L(\mathbf{x}|\boldsymbol{\theta}) &= \prod_{i=1}^n f_W(x_i | \beta, \eta) \\ &= \prod_{i=1}^n \frac{\beta}{\eta} \left(\frac{x_i}{\eta}\right)^{\beta-1} e^{-\left(\frac{x_i}{\eta}\right)^{\beta}}, \end{aligned}$$

et la log-vraisemblance

$$\ln L(\mathbf{x}|\boldsymbol{\theta}) = n \ln(\beta) + \beta \ln(\eta) + (\beta - 1) \sum_{i=1}^n \ln(x_i) - \sum_{i=1}^n \left(\frac{x_i}{\eta}\right)^{\beta}. \quad (4.10)$$

On obtient $\hat{\beta}_{MV}$ en dérivant (4.10), et en résolvant l'équation :

$$\frac{1}{\beta} + \frac{\sum_{i=1}^n \ln(x_i)}{n} - \frac{\sum_{i=1}^n x_i^{\beta} \ln(x_i)}{\sum_{i=1}^n x_i^{\beta}} = 0; \quad (4.11)$$

on obtient alors

$$\hat{\eta}_{MV} = \left(\frac{\sum_{i=1}^n x_i^{\hat{\beta}_{MV}}}{n} \right)^{\frac{1}{\hat{\beta}_{MV}}}. \quad (4.12)$$

Il n'y a pas de solution explicite à l'équation (4.11), on a recours à une procédure numérique de type Newton ou quasi-Newton cf. chapitre 8 de Nelson et al. (1982).

4.5 Estimation par maximum de vraisemblance avec censures

La vraisemblance des données censurées s'écrit

$$\begin{aligned} L((\mathbf{x}, \mathbf{c}) | \theta) &= \prod_{i=1}^n [f_W(x_i | \beta, \eta)]^{1-c_i} \cdot [R_W(x_i | \beta, \eta)]^{c_i} \\ &= \prod_{i=1}^n \left[\frac{\beta}{\eta} \left(\frac{x_i}{\eta} \right)^{\beta-1} e^{-\left(\frac{x_i}{\eta} \right)^\beta} \right]^{1-c_i} \cdot \left[e^{-\left(\frac{x_i}{\eta} \right)^\beta} \right]^{c_i}, \end{aligned}$$

et la log-vraisemblance

$$\ln L((\mathbf{x}, \mathbf{c}) | \theta) = (\#\{i : c_i = 0\}) [\ln(\beta) - \beta \ln(\eta)] + (\beta - 1) \sum_{i:c_i=0} \ln(x_i) - \sum_{i=1}^n \left(\frac{x_i}{\eta} \right)^\beta, \quad (4.13)$$

le terme $\sum_{i=1}^n c_i$ correspondant au nombre de censures. On obtient $\hat{\beta}_{MVC}$ en dérivant (4.13), et en résolvant l'équation :

$$\frac{1}{\beta} + \frac{\sum_{i:c_i=0} \ln(x_i)}{\#\{i : c_i = 0\}} - \frac{\sum_{i=1}^n x_i^\beta \ln(x_i)}{\sum_{i=1}^n x_i^\beta} = 0; \quad (4.14)$$

on obtient alors

$$\hat{\eta}_{MVC} = \left(\frac{\sum_{i=1}^n x_i^{\hat{\beta}_{MVC}}}{\#\{i : c_i = 0\}} \right)^{\frac{1}{\hat{\beta}_{MVC}}}. \quad (4.15)$$

Remarque : là encore il n'y a pas de solution explicite à l'équation (4.14), on a alors recours à une procédure numérique.

Intervalle de confiance

On peut montrer que sous de bonnes conditions cf. Meeker and Escobar (1998), l'estimateur du maximum de vraisemblance suit asymptotiquement une loi normale :

$$\sqrt{n}(\hat{\theta}_{MV} - \theta) \xrightarrow{\mathcal{L}} \mathcal{N}(0, I^{-1}(\theta)),$$

où $I(\theta)$ est la matrice de Fisher variance-covariance :

$$I(\beta, \eta) = \begin{bmatrix} -\frac{\partial^2 L}{\partial \beta^2} & -\frac{\partial^2 L}{\partial \beta \partial \eta} \\ -\frac{\partial^2 L}{\partial \beta \partial \eta} & -\frac{\partial^2 L}{\partial \eta^2} \end{bmatrix}.$$

Dans le cas de données fortement censurées, la variance de l'estimateur du maximum de vraisemblance peut être très importante. Nous allons, dans ce qui suit, proposer des solutions plus adaptées.

4.6 Algorithme S-EM

Dans le cas de données censurées, les censures correspondent à des données incomplètes : on n'observe pas la panne ! On peut donc utiliser l'algorithme S-EM, cf. Algorithme 3.

Étape Simulation : Pour $i = m + 1, \dots, n$ sachant $(\beta^{(r-1)}, \eta^{(r-1)})$ on simule les durées de vie censurées suivant la loi (4.9) :

$$\tilde{x}_i = \eta^{(r-1)} \left[\left(\frac{x_i}{\eta^{(r-1)}} \right)^{\beta^{(r-1)}} - \ln(U) \right]^{\frac{1}{\beta^{(r-1)}}},$$

où U suit la loi uniforme sur $[0; 1]$.

Étape M : on obtient $(\beta^{(r)}, \eta^{(r)})$ en maximisant la vraisemblance des données complétées : $(x_1, \dots, x_m; \tilde{x}_{m+1}, \dots)$ cf. (4.11) et (4.12).

Estimateur S-EM : on estime alors (β, η) par

$$\hat{\beta}_{S-EM} = \frac{1}{B-l} \sum_{r=l+1}^B \beta^{(r)} \quad \text{et} \quad \hat{\eta}_{S-EM} = \frac{1}{B-l} \sum_{r=l+1}^B \eta^{(r)}, \quad (4.16)$$

où l représente les premières valeurs correspondant à une période de « chauffe ».

Calibrations

Le nombre d'itérations de l'algorithme S-EM et son initialisation sont deux facteurs impactant la qualité de l'estimation. Nous allons pour cela calibrer l'algorithme à partir d'un jeu de données simulées avec les critères suivants,

β_0	η_0	$R_0(177)$	$R_0(234)$
3	200	50%	20%

TABLE 4.1 – Paramètres du jeu de données simulées

Nous avons simulé 500 échantillons de taille $n = 600$ et avec 80% de censure à droite.

L'erreur relative (« Average bias ») et l'erreur quadratique moyenne (RMSE) de la fiabilité R à mi-vie nous permettent d'évaluer la performance de l'estimation :

$$\text{bias} = \frac{1}{500} \sum_{i=1}^{500} (\hat{R} - R), \quad \text{RMSE} = \sqrt{\frac{1}{500} \sum_{i=1}^{500} (\hat{R} - R)^2}.$$

Itérations

Nous avons fait varier le nombre d'itérations de l'algorithme S-EM. Les résultats sont détaillées dans le Tableau 4.2 et illustrés dans la Figure 4.3. On constate cf. Figure 4.3 et cf. Tableau 4.2 le faible impact du nombre d'itérations sur le RMSE et l'erreur relative de la fiabilité à mi-vie.

SEM nombre d'itérations	500	1000	5000
Average bias	0.015	0.017	0.015
RMSE	0.054	0.053	0.054

TABLE 4.2 – Algorithm S-EM : calibration du nombre d'itérations. Estimation de la fiabilité à mi-vie pour 500 échantillons avec $n = 600$ et 80% de censure.

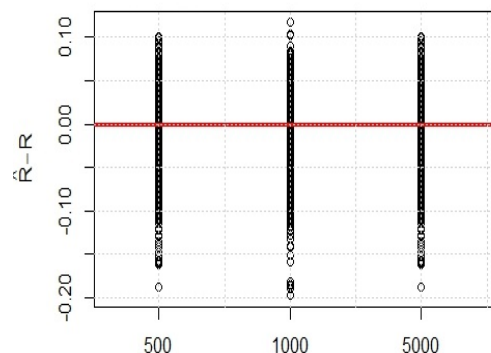


FIGURE 4.3: Algorithm S-EM : calibration du nombre d'itérations. Biais et RMSE de la fiabilité à mi-vie pour 500 échantillons avec $n = 600$ et 80% de censure.

La fiabilité estimée avec S-EM à mi-vie a été évaluée pour 500 jeux de données et cela avec différentes calibrations. Le calibrage de S-EM nécessite d'évaluer le nombre d'itérations nécessaire afin d'optimiser les résultats et les temps de calcul associé. Comme l'illustre les figures 4.3 (a) et 4.3 (b) 500 itérations de S-EM sont efficaces en termes d'erreur relative et RMSE pour $n = 600$ et 80% de censure.

Initialisation

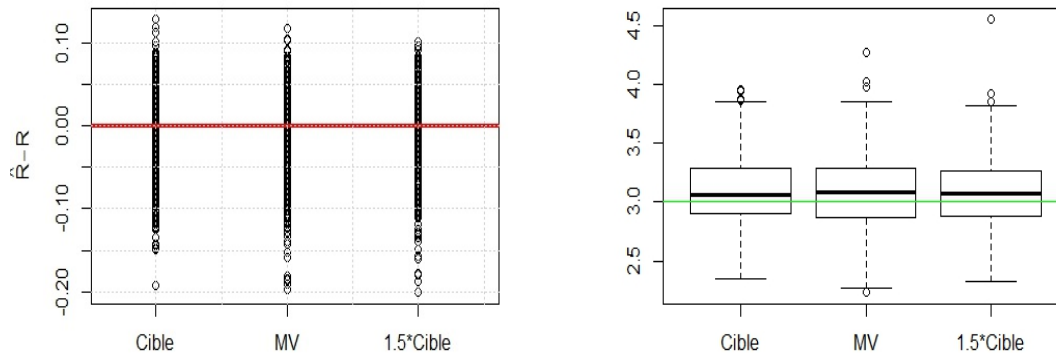
Nous avons fait varier l'initialisation de l'algorithme S-EM :

- valeurs cibles : $\beta^{(0)} = 3, \eta^{(0)} = 200$
- valeur de l'estimateur du MV sur les pannes uniquement : $\beta^{(0)} = \beta_{MV}, \eta^{(0)} = \eta_{MV}$
- valeurs « incohérentes » : $\beta^{(0)} = \beta_{Cible} \times 1.5, \eta^{(0)} = \eta_{Cible} \times 1.5$

Les résultats sont détaillés dans le Tableau 4.3 et illustrés dans la Figure 4.4.

Initialisation θ^r	$\theta^{(0)} = \theta_{Cible}$	$\theta^{(0)} = \theta_{MV}$	$\theta^{(0)} = 1.5\theta_{Cible}$
Average bias	0.013	0.017	0.013
RMSE	0.05	0.053	0.053

TABLE 4.3 – Algorithme S-EM : calibration de l'initialisation. Estimation de la fiabilité à mi-vie pour 500 échantillons avec $n = 600$, 80% de censure et 1 000 itérations.



(a) Biais et RMSE pour la fiabilité à mi-vie

(b) Boxplot des estimateurs $\hat{\beta}$

FIGURE 4.4: Algorithme S-EM : calibration de l'initialisation. Estimation de la fiabilité à mi-vie pour 500 échantillons avec $n = 600$ et 80% de censure. (a) impact de l'initialisation sur la fiabilité et (b) impact de l'initialisation sur les estimateurs $\hat{\beta}$.

SEM est une méthode itérative, un point de départ est donc nécessaire afin d'effectuer le re-complètement de l'échantillon. Comme l'illustre les figures 4.4 (a) et 4.4 (b) l'utilisation de l'estimateur du maximum de vraisemblance est une stratégie d'initialisation efficace en terme d'erreur relative et RMSE.

4.7 Estimation bayésienne

L'inférence bayésienne pour un modèle paramétrique consiste à traduire les informations *a priori* sur (β, η) par une loi jointe $\pi(\beta, \eta)$ sur les valeurs possibles des paramètres. En utilisant la formule de Bayes on obtient la loi *a posteriori*, $\pi(\beta, \eta | \mathbf{x}, \mathbf{c})$, sachant les données $(\mathbf{x}, \mathbf{c})$ du retour d'expérience. Le choix de la loi *a priori* est une étape importante de l'estimation bayésienne. On choisira des lois *a priori* informatives ou non-informatives, selon la qualité de l'information dont on dispose sur les paramètres. Dans le cas de la loi de Weibull où les deux paramètres sont inconnus, il n'existe pas de loi conjuguées cf. Soland (1969), et donc pas d'expression analytique de l'estimateur. On trouvera dans Procaccia et al. (2011) de nombreuses propositions pour choisir une loi *a priori* pour les paramètres d'une loi de Weibull.

Loi *a priori* pour le paramètre de forme, β de la loi de Weibull

Comme dans Bacha et al. (1998) nous utiliserons une loi bêta sur $[b^{min}; b^{max}]$ comme loi *a priori* pour le paramètre β :

$$\beta \sim b^{min} + \mathcal{B}(p, q)(b^{max} - b^{min}), \quad (4.17)$$

$$\pi(\beta) = \frac{\Gamma(p+q)}{\Gamma(p)\Gamma(q)} \frac{(\beta - b^{min})^{p-1} (b^{max} - \beta)^{q-1}}{(b^{max} - b^{min})^{p+q-1}} \mathbb{1}_{[0;1]} \left(\frac{\beta - b^{min}}{b^{max} - b^{min}} \right),$$

où

$$\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt,$$

et b^{min} et $b^{max} \in \mathbb{R}^+$.

La loi bêta offre une grande flexibilité : pour $p = q = 0.5$ on obtient la loi non-informative de Jeffrey et pour $p = q = 1$ la loi uniforme. Il reste à identifier les hyperparamètres. Comme on peut le voir dans la Figure 4.5 :

1. $[b^{min}, b^{max}]$ correspondant au support de la loi et modélise l'expertise que l'on a sur β ;
2. (p, q) correspondant à la forme de la loi bêta et modélise l'incertitude que l'on a sur β .

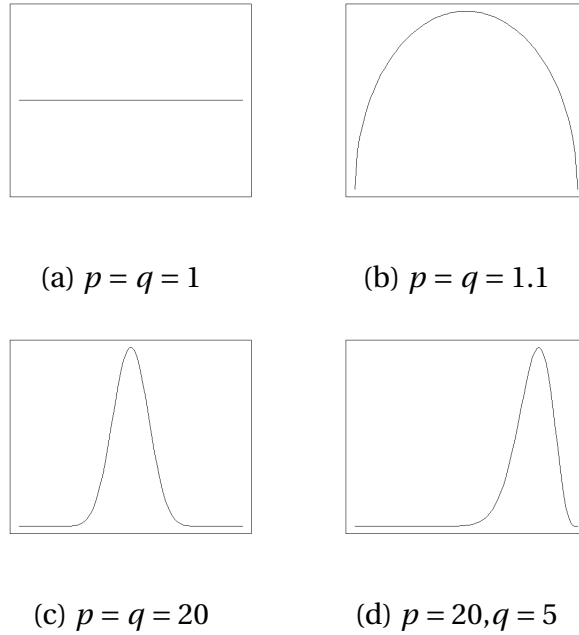


FIGURE 4.5: Forme de la densité de la loi bêta $\mathcal{B}(p, q, [0, 1])$.

Le paramètre β caractérise le vieillissement dans le temps de l'équipement :

- $0 < \beta < 1$, le taux de défaillance décroît ;
- $\beta = 1$, le taux de défaillance est constant ;
- $\beta > 1$, le taux de défaillance est croissant.

On trouvera dans Barringer (2002), une étude statistique sur les valeurs de β dans diverses applications industrielles. Nous avons donc choisi

$$b^{min} = 0,5 \quad b^{max} = 4,5.$$

De part sa nature de systémier, NEXTER SYSTEMS ne possède que très peu d'informations *a priori* sur les équipements. Nous choisirons donc une loi *a priori* peu informative en prenant :

$$p = q = 1,1.$$

Loi *a priori* pour le paramètre d'échelle, η , de la loi de Weibull

Quand le paramètre de forme, β , est connu, la distribution inverse gamma généralisée est la loi *a priori* conjuguée pour le paramètre d'échelle η , cf. Erto and Giorgio (2002) :

$$\pi(\eta | \beta, e^{shape}, e^{scale}) = \frac{\beta (e^{scale})^\beta e^{shape}}{\Gamma(e^{shape})} \eta^{-e^{shape} \beta - 1} e^{-\left(\frac{e^{scale}}{\eta}\right)^\beta}, \quad (4.18)$$

où (e^{shape}, e^{scale}) sont des hyperparamètres.

Une approche naturelle pour identifier e^{scale} est de fixer l'espérance de la loi *a priori*. Par exemple en prenant l'estimation triviale du Weibull QQ-plot :

$$e^{scale} \frac{\Gamma\left(e^{shape} - \frac{1}{\beta}\right)}{\Gamma(e^{shape})} = \check{\eta}. \quad (4.19)$$

On peut voir dans l'Équation (4.19), qu'il y a une contrainte immédiate sur e^{shape} : $e^{shape} > \frac{1}{\beta}$. e^{shape} modélise l'incertitude que l'on a sur η . La calibration de e^{shape} sera faite dans une section suivante.

L'estimateur de Bayes est alors

$$\hat{\theta}_{Bayes} = \mathbb{E}_{\pi[\theta | \mathbf{x}]}[\theta], \quad (4.20)$$

Il n'y a pas d'expression analytique de cet estimateur, il faut donc évaluer numériquement l'intégrale, cf. Robert and Casella (2010). Dans le cas d'échantillon de petite taille ou d'un échantillon très fortement censuré, McLachlan1998 Bacha and Celeux (1996) propose un approche plus efficace, de type *Bayes bootstrap* cf. Rubin (1981).

4.8 Estimation Bayesian Restoration Maximization

De manière générale, pour calculer l'estimateur bayésien, on a besoin, de la loi *a posteriori* $\pi(\theta | \mathbf{x})$. Celle-ci étant inconnue, on va utiliser une technique d'échantillonnage préférentiel pour calculer l'espérance 4.20 en utilisant une loi instrumentale ρ , cf. Glynn and Iglehart (1989) :

$$\mathbb{E}_{\pi[\theta | \mathbf{x}]}[\theta] = \mathbb{E}_{\rho} \left[\theta \cdot \frac{\pi(\theta | \mathbf{x})}{\rho(\theta)} \right].$$

On simule $\theta_1, \dots, \theta_B$ selon ρ et on estime $\hat{\theta}_{BRM}$ par

$$\frac{1}{B} \sum_{i=1}^B \tilde{\theta}_i w_i,$$

où les « poids » w_i sont égaux à :

$$w_i = \frac{\pi(\theta_i | (\mathbf{x}, \mathbf{c}))}{\rho(\theta_i)} \propto \frac{L((\mathbf{x}, \mathbf{c}) | \theta_i) \cdot \pi(\theta_i)}{\rho(\theta_i)}.$$

Reste à choisir ρ , la loi d'échantillonnage. Cette dernière doit être une « contrefaçon » de la loi *a posteriori*. Dans le cas de données incomplètes, le principe de l'algorithme S-EM est de compléter l'échantillon par simulations pour obtenir l'estimateur du maximum de vraisemblance cf. Celeux et al. (1996). L'idée de l'estimation BRM est d'utiliser ce principe dans un cadre bayésien, cf. Bacha and Celeux (1996) :

1. dans une première étape, simuler les paramètres du modèle à partir de leur loi jointe *a priori*;
2. dans une seconde étape, « restaurer » les données manquantes en les simulant conditionnellement aux paramètres obtenus, pour obtenir l'estimateur du maximum de vraisemblance sur les données complétées.
3. la loi empirique des estimateurs ainsi obtenus permet d'approcher la loi *a posteriori*.

Ici les données manquantes sont les données censurées ($x_i, c_i = 1$).

$$L((\mathbf{x}, \mathbf{c}) | \theta_i) \cdot \pi(\theta_i) = \prod_{i: c_i=0} L(x_i | \theta_i) \prod_{i: c_i=1} L(x_i | \theta_i) \cdot \pi(\theta_i). \quad (4.21)$$

L'observation x_i est soit une chute quand il y a une panne, soit $F_X(\bullet | X > x)$ quand il y a une censure. On restaure alors les données censurées en simulant d'abord $(\tilde{\beta}, \tilde{\eta})$ selon la loi jointe *a priori*, puis la panne X comme suit,

$$\tilde{x}_i = \tilde{\eta} \left[\left(\frac{x_i}{\tilde{\eta}} \right)^{\tilde{\beta}} - \ln(U) \right]^{\frac{1}{\tilde{\beta}}},$$

où U suit la loi uniforme sur $[0; 1]$ cf. 4.9.

On calcule alors $(\hat{\beta}, \hat{\eta})$ l'estimateur du maximum de vraisemblance des données ainsi complétées $\tilde{\mathbf{x}} = (x_1, \dots, x_m, \tilde{x}_{m+1}, \dots, \tilde{x}_n)$. La loi ρ est alors la loi empirique de ces différents estimateurs : cette loi est donc une « contrefaçon » de la loi *a posteriori*, cf. Scott (2015). Elle n'est pas explicite, nous allons donc l'estimer à l'aide d'un estimateur à noyaux :

$$\hat{\rho}(\beta, \eta) = \frac{1}{B \cdot h_\beta \cdot h_\eta} \sum_{i=1}^B K\left(\frac{\beta - \beta_i}{h_\beta}\right) \cdot K\left(\frac{\eta - \eta_i}{h_\eta}\right), \quad (4.22)$$

où K est un noyau gaussien :

$$K(\theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{\theta^2}{2}}.$$

Les tailles de fenêtres minimisant l'erreur moyenne intégrée sont alors cf. Scott (2015) :

$$h_\beta = \hat{\sigma}_\beta \cdot n^{-1/5} \quad \text{et} \quad h_\eta = \hat{\sigma}_\eta \cdot n^{-1/5}$$

On calcule alors les poids $\mathbf{w}$ et leur version normalisée $\tilde{\mathbf{w}}$:

$$w_i = \frac{f(\tilde{\mathbf{x}} | \beta_i, \eta_i) \pi(\beta_i, \eta_i)}{\hat{\rho}(\beta_i, \eta_i)} \quad \text{et} \quad \tilde{w}_i = \frac{w_i}{\sum_{k=1}^B w_k}. \quad (4.23)$$

On estime alors $\mathbb{E}[\theta | \mathbf{x}]$ par

$$\hat{\beta}_{BRM} = \sum_{i=1}^B \tilde{w}_i \beta_i \quad \text{et} \quad \hat{\eta}_{BRM} = \sum_{i=1}^B \tilde{w}_i \eta_i, \quad (4.24)$$

On synthétise l'estimation BRM des paramètres d'une loi de Weibull simple comme suit

Algorithme 4 : Estimation BRM des paramètres d'une loi de Weibull simple

pour $r = 1, \dots, B$ **faire**

Bayes On simule $\theta^{(r)}$ à partir du prior $\pi(\theta)$

Restoration $\theta^{(r)}$ connu on simule les pannes censurées au delà de x_i .

pour $i = m + 1, \dots, n$ **faire**

$$\tilde{x}_i = \eta^{(r)} \left[\left(\frac{z_i}{\eta^r} \right)^{\beta^{(r)}} - \ln(U) \right]^{\frac{1}{\beta^{(r)}}}$$

 où U uniforme sur $[0; 1]$.

fin

Maximisation on calcule $\tilde{\theta}^{(r)}$ estimateur du maximum de vraisemblance des données complétées $\tilde{\mathbf{x}} = (x_1, \dots, x_m, \tilde{x}_{m+1}, \dots, \tilde{x}_n)$ à partir de (4.11) et (4.12).

fin

Estimateur BRM :

On obtient $\theta^{(B)}$. Soit ρ la loi des $\theta^{(B)}$ ainsi obtenus : ρ est proche $L(x|\theta)\pi(\theta)$ et donc de la loi *a posteriori* $\pi(\theta|\mathbf{x})$. La densité ρ n'étant pas explicite on utilise une estimation non-paramétrique $\hat{\rho}$, avec un noyau gaussien cf. (4.22). On calcule alors les poids :

pour $i = 1, \dots, B$ **faire**

$$w_i = \frac{f(\tilde{\mathbf{x}}|\beta^i, \eta^i)\pi(\beta^i, \eta^i)}{\hat{\rho}(\beta^i, \eta^i)} \quad \text{et} \quad \tilde{w}_i = \frac{w_i}{\sum_{k=1}^B w_k}.$$

fin

et les estimateurs :

$$\hat{\beta}_{BRM} = \sum_{i=1}^B \tilde{w}_i \beta^i \quad \text{et} \quad \hat{\eta}_{BRM} = \sum_{i=1}^B \tilde{w}_i \eta^i,$$

La loi ρ est une « contrefaçon » de la loi *a posteriori*, et les poids w_i peuvent être nuls (quand le support de $\rho(\theta)$ est trop éloigné du support $\pi(\theta|\mathbf{x})$).

Pour contrôler une éventuelle dégénérescence des poids on utilise un critère de variance cf. Liu (2008) :

$$\text{Effective Sample Size} = \frac{1}{\sum_{i=1}^B \tilde{w}_i^2}. \quad (4.25)$$

Dans le cas où $ESS \simeq 0$, on peut ajouter à l'estimation BRM une étape de ré-échantillonnage visant à ne garder que les θ_i ayant les plus « gros » poids, cf. Rubin (1988). La figure 4.6 décrit le principe ce ré-échantillonnage.

Algorithme 5 : Ré-échantillonnage

Initialisation : on se donne $\theta_1, \dots, \theta_B$ de BRM

On simule $(i_1, \dots, i_B)$ selon une multinomiale $\mathcal{M}(B, (\tilde{\mathbf{w}}))$

Ré-échantillonnage : $\tilde{\theta} = (\theta_{i_1}, \dots, \theta_{i_B})$;

Normalisation

Les $\tilde{\theta}_i$ ont alors tous le même poids $\tilde{\tilde{w}}_i = 1/B$.

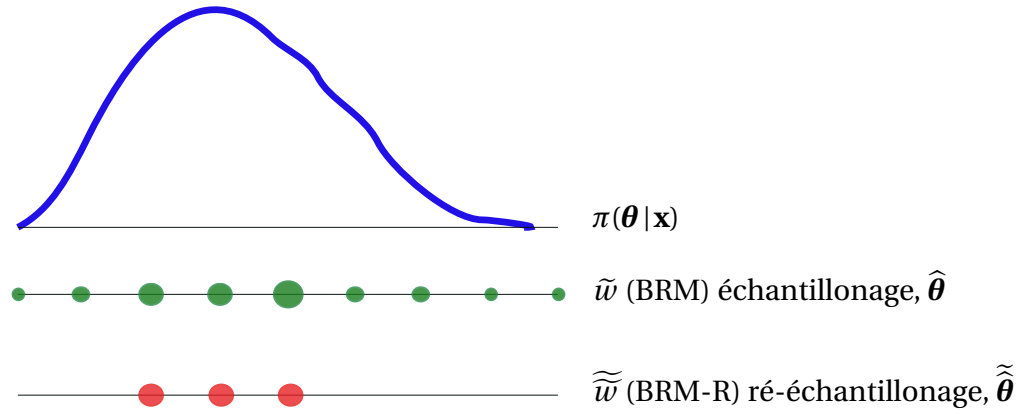
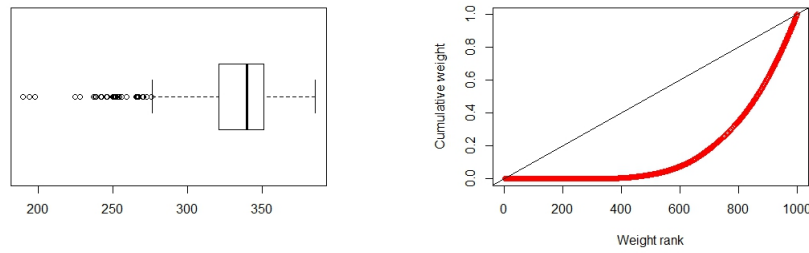


FIGURE 4.6: Principe de l'estimation BRM : à partir de la loi *a priori* et d'un échantillonnage préférentiel, on obtient un échantillon des paramètres dont la loi empirique est concentrée sur la moyenne de la loi *a posteriori*. BRM-R : en cas de dégénérescence des poids, on peut utiliser un ré-échantillonnage.

Lorsqu'un des poids est égal à 1 et tous les autres sont nuls, on a $1 \leq ESS \leq B : ESS = 1$. Dans ce cas, les poids peuvent être très mal équilibrés, voir la Figure 4.7, le rééchantillonnage ne sera pas à privilégier.



(a) Boxplot des poids non normalisés (b) Courbe de Gini des poids

FIGURE 4.7: Distribution des poids (w_i) lors d'une estimation BRM.

4.8.1 Intervalle de crédibilité

L'estimation BRM fournit une suite d'estimateurs $\theta^1, \dots, \theta^B$. On utilise l'estimation :

$$\hat{\theta}_{BRM} = \sum_{i=1}^B \tilde{w}_i \theta^i.$$

Pour obtenir un intervalle de crédibilité au niveau α , on utilise la loi empirique des θ :

$$\hat{\theta} \in [\theta_{(\alpha/2)}; \theta_{(1-\alpha/2)}] \quad (4.26)$$

où $\theta_{(p)}$ est le quantile empirique des $\theta^1, \dots, \theta^B$, pour la probabilité p .

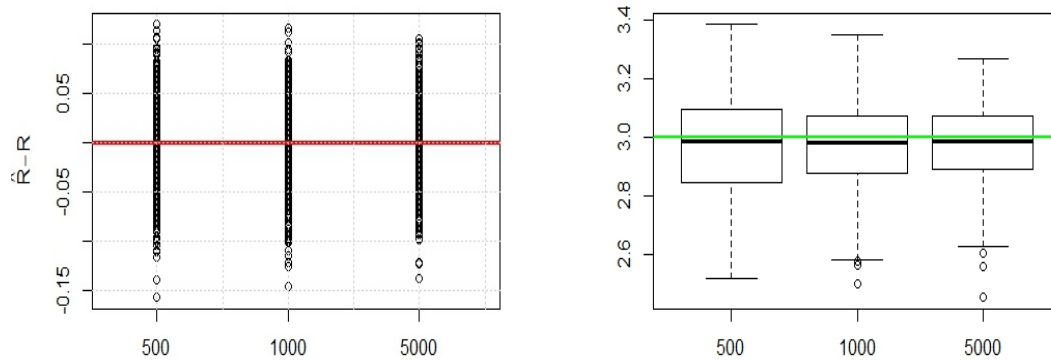
4.8.2 Calibrage du nombre de répétitions de BRM

Pour mesurer l'impact du nombre de répétitions sur la performance de l'estimation BRM, nous avons utilisé le jeu de données simulées du Tableau 4.1 page 43, avec $n = 600$ et 80% de taux de censure. Les résultats sont détaillés dans le Tableau 4.4 et la figure 4.8.

BRM nombre de répétitions	500	1000	5000
Average bias	0.002	0.001	0.001
RMSE	0.044	0.043	0.040

TABLE 4.4 – BRM, calibration nombre de répétitions. BRM fiabilité estimée à mi-vie pour 500 jeux de données, avec $n = 600$ et 80% de censure.

On constate que l'erreur quadratique diminue avec le nombre de répétitions, le biais quant à lui reste stable.



(a) Biais et RMSE de la fiabilité à mi-vie

(b) Boxplot des estimateurs de $\hat{\beta}$

FIGURE 4.8: Impact du nombre de répétitions de BRM. Pour 500 échantillons avec $n = 600$ et 80% de taux de censure. On effectue 500, 1 000 et 5 000 répétitions.

En terme d'erreur et de temps de calcul : 1 000 itérations semble être le meilleur compromis. Une fois le nombre d'itération fixé, on peut s'atteler à la sensibilité des hyperparamètres.

4.8.3 Calibrage de la loi *a priori* sur le paramètre de forme, β , de la loi de Weibull simple

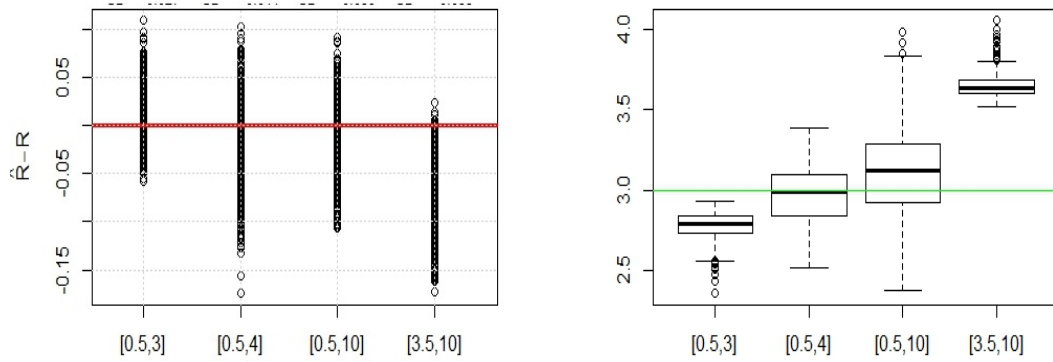
Nous avons considéré quatre configurations pour le support du prior de β :

1. β situé sur le support $[0.5; 4]$ incluant le paramètre cible $\beta_0 = 3$,
2. β situé sur le support $[0.5; 10]$,
3. β situé sur le support $[0.5; 3]$ le paramètre cible situé au bord du support.
4. β situé sur le support $[3.5; 10]$

Les résultats sont détaillés dans le Tableau 4.5 et la Figure 4.9

$\mathcal{B}[\mathbf{b}^{\min}, \mathbf{b}^{\max}]$	[0.5; 3]	[0.5; 4]	[0.5; 10]	[3.5; 10]
Average bias	0.02	0.02	0.013	0.075
RMSE	0.027	0.044	0.039	0.035

TABLE 4.5 – Estimation BRM : calibration du support de la loi *a priori* sur β (abscisse). Fiabilité $\hat{R}$ estimée à mi-vie pour 1000 répétitions et 500 échantillons, avec $n = 600$ et 80% de censure et $\beta_0 = 3$.



(a) Biais et RMSE de la fiabilité à mi-vie

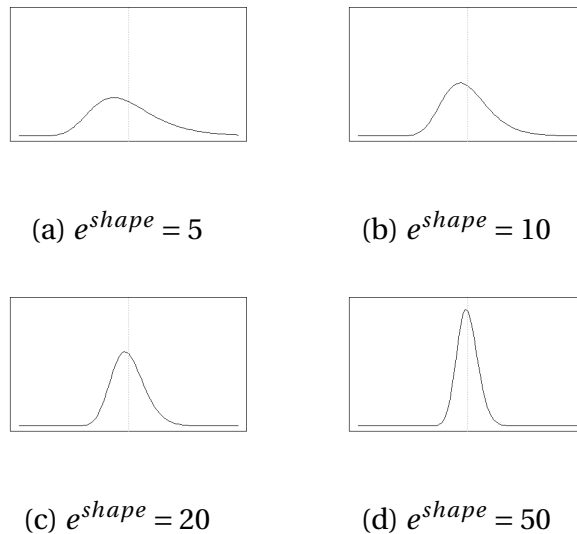
 (b) Boxplot des estimateurs $\hat{\beta}$

FIGURE 4.9: Estimation BRM : calibration du support de la loi *a priori* sur β (abscisse). Pour BRM avec 1 000 répétitions, pour 500 échantillons avec $n = 600$ et 80% de censure. Pour $\beta_0 = 3$, on utilise les supports $[0.5, 3]$, $[0.5, 4]$, $[0.5, 10]$ et $[3.5, 10]$.

Un support trop réduit engendre un fort risque d'une estimation imprécise notamment lorsqu'on ne dispose pas d'avis d'expert et que le paramètre cible n'appartient pas à ce support. Un support de $[0.5, 4]$ semble être le meilleur compromis entre les critères d'erreur 4.5 et l'impact de celui-ci sur $\hat{\beta}$ 4.9. Notons qu'un support large $[0.5, 10]$ peut aussi être utilisé dans certains cas particulier (e.g pour certaines typologies d'équipements).

4.8.4 Calibrage de la loi *a priori* sur le paramètre de position, η , de la loi de Weibull simple

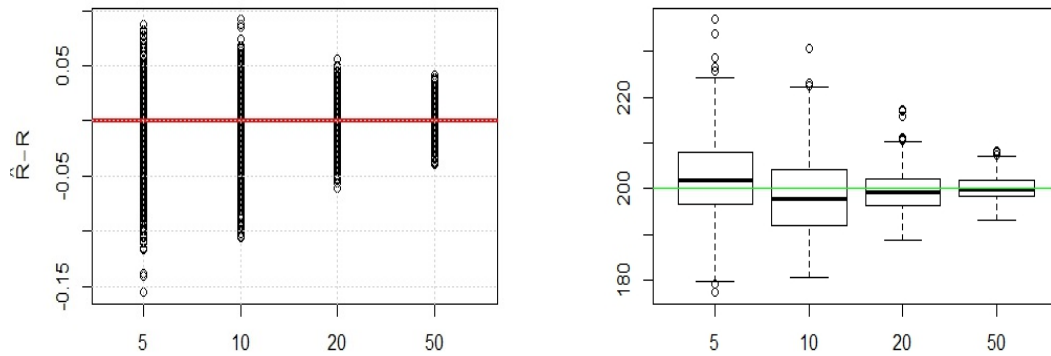
Le paramètre de forme de la loi inverse gamma, modélise l'incertitude que l'on a sur les « avis d'expert », ou dans notre cas sur e^{scale} .


 FIGURE 4.10: Forme de la loi du prior pour η avec différentes valeurs de e^{shape} .

Les résultats des simulations sont détaillées dans le Tableau 4.6 et la Figure 4.11.

e^{shape}	5	10	20	50
Average bias	0.018	0.013	0.004	0.001
RMSE	0.029	0.027	0.02	0.015

TABLE 4.6 – Estimation BRM : calibration de e^{shape} , le paramètre de forme de la loi *a priori* sur η . Fiabilité $\hat{R}$ estimée à mi-vie pour 1000 répétitions et 500 échantillons, avec $n = 600$ et 80% de censure et $\beta_0 = 3$. Pour $\check{\eta}$ fixé cf 4.19



(a) Biais et RMSE de la fiabilité à mi-vie

(b) Boxplot des estimateurs $\hat{\eta}$

FIGURE 4.11: Estimation BRM : calibration de e^{shape} , le paramètre de forme de la loi *a priori* sur η . Pour BRM avec 1 000 répétitions et 500 échantillons avec $n = 600$ et 80% de censure et $\beta_0 = 3$. Pour $\check{\eta}$ fixé cf 4.19. On choisit $e^{shape} = 5, 10, 20$, et 50.

Pour calibrer le paramètre de position de la loi *a priori* sur η , diverses configurations ont été envisagées, (η_0 est le paramètre cible et $\check{\eta}$ est défini 4.19) :

- $\check{\eta} \ll \eta_0$,
- $\check{\eta} \approx \eta_0$,
- $\check{\eta} \gg \eta_0$.

Les résultats des simulations sont détaillés dans le Tableau 4.7 et la Figure 4.12.

(a) $e^{shape} = 5$	$\check{\eta}$	$0.5\eta_0$	η_0	$1.5\eta_0$
	average bias	-0.05	-0.005	0.04
	RMSE	0.03	0.03	0.02
(b) $e^{shape} = 50$	$\check{\eta}$	$0.5\eta_0$	η_0	$1.5\eta_0$
	average bias	-0.07	0.001	0.13
	RMSE	0.04	0.015	0.03

TABLE 4.7 – Estimation BRM : calibration du paramètre de position de la loi *a priori* sur η . Fiabilité $\hat{R}$ estimée à mi-vie pour 500 répétitions, avec $n = 600$ et 80% de censure et $\beta_0 = 3$.

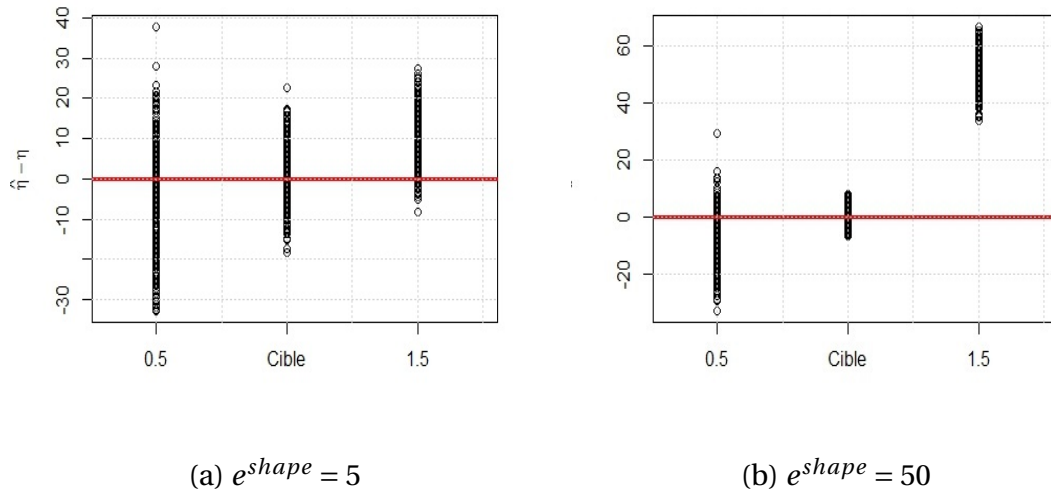


FIGURE 4.12: Estimation BRM : calibration du paramètre de position de la loi *a priori* sur η . Biais et RMSE pour l'estimation de η . Pour BRM avec 1 000 répétitions et 500 échantillons avec $n = 600$ et 80% de censure et $\beta_0 = 3$. On choisit $e^{shape} = 5$ et 50.

Lorsque $\tilde{\eta} \approx \eta_0$, la calibration à privilégier est $e^{shape} = 50$ en termes d'erreur (b) 4.7, et d'impact $\hat{\eta}$ 4.12. Lorsque $\tilde{\eta} \ll \eta_0$ ou $\tilde{\eta} \gg \eta_0$, la calibration à privilégier est $e^{shape} = 5$ en termes d'erreur (a) 4.7, et d'impact $\hat{\eta}$ 4.12.

En prenant $e^{shape} = 5$ nous obtenons une loi *a priori* peu informative sur η , mais robuste pour un mauvais calibrage du paramètre de position.

4.9 Analyses de sensibilité

La taille de l'échantillon et le taux de censure sont des données d'entrée ayant un impact sur la performance des méthodes d'estimations. Nous allons comparer les estimations MVC, S-EM et BRM, selon leur sensibilité à ces deux facteurs. Le jeu de données simulées et celui du Tableau 4.1 page 43. L'analyse du taux de censure est détaillées dans le Tableau 4.8, celle de la taille de l'échantillon dans le Tableau 4.9. La méthode SEM est initialisée à l'aide de θ_{MV} voir 4.6.

Censoring		30%	70%	80%	90%
MVC	Average bias	-0.008	0.003	-0.062	-0.077
	RMSE	0.031	0.149	0.319	0.463
SEM	Average bias	-0.001	-0.007	-0.017	-0.06
	RMSE	0.017	0.034	0.053	0.112
BRM	Average bias	-0.001	-0.005	-0.018	-0.08
	RMSE	0.018	0.03	0.04	0.098

a) Erreur $\hat{R}$ à mi-vie

Censoring		30%	70%	80%	90%
MVC	$\hat{R}$	0.5	0.47	0.42	0.4
	$\hat{\beta} [I_{95\%}]$	2.65[2.8,3.44]	2.9[2.8,3.44]	3.2[2.8,3.3]	2.8[2.42,3.26]
SEM	$\hat{R} [I_{95\%}]$	0.499[0.47,0.52]	0.493[0.43,0.54]	0.48[0.42,0.56]	0.41[0.23,0.6]
	$\hat{\beta} [I_{95\%}]$	3[2.8,3.2]	3[2.8,3.1]	3.2[2.8,3.1]	3.3[2.7,3.9]
BRM	$\hat{R} [I_{95\%}]$	0.5[0.47,0.52]	0.495[0.45,0.54]	0.49[0.42,0.54]	0.47[0.3,0.57]
	$\hat{\beta} [I_{95\%}]$	3[2.8,3.3]	3[2.7,3.3]	2.9[2.7,3.17]	3[2.3,3.5]

b) $\hat{\beta}$ et $\hat{R}$

TABLE 4.8 – Analyse de sensibilité des estimations au taux de censure : 1000 répétitions (SEM et BRM) et 500 échantillons de taille $n = 600$.

La méthode S-EM améliore la précision des estimateurs θ_{MVC} dans tous les cas de censures et de taille d'échantillons en termes d'erreur (a) 4.8 et de précision des estimations (b)4.8.

La méthode BRM, améliore les estimateurs θ_{SEM} lorsque le taux de censure est supérieur à 30% pour une taille d'échantillon fixée ($n = 600$).

Sample size		100	200	400	600
MVC	Average bias	-0.022	-0.024	0.045	-0.062
	RMSE	0.501	0.449	0.252	0.319
SEM	Average bias	-0.078	-0.037	-0.085	-0.017
	RMSE	0.139	0.096	0.075	0.053
BRM	Average bias	-0.015	-0.015	-0.013	-0.0014
	RMSE	0.053	0.051	0.044	0.04

a) Erreur $\hat{R}$ à mi-vie

Sample size		100	200	400	600
MVC	$\hat{R}$	0.3	0.4	0.41	0.42
	$\hat{\beta} [I_{95\%}]$	3.5[2.7,4.3]	3.6[2.9,3.7]	3.5[2.8,3.64]	3.2[2.8,3.3]
SEM	$\hat{R} [I_{95\%}]$	0.4[0.26,0.56]	0.46[0.28,0.6]	0.46[0.35,0.6]	0.48[0.42,0.56]
	$\hat{\beta} [I_{95\%}]$	3.4[2.5,3.8]	3.5[2.6,3.8]	3.4[2.7,3.44]	3.2[2.8,3.1]
BRM	$\hat{R} [I_{95\%}]$	0.48[0.4,0.57]	0.48[0.4,0.56]	0.49[0.42,0.55]	0.49[0.42,0.54]
	$\hat{\beta} [I_{95\%}]$	3[2.57,3.4]	3[2.6,3.3]	3[2.7,3.3]	3[2.7,3.32]

b) $\hat{\beta}$ et $\hat{R}$

TABLE 4.9 – Analyse de sensibilité des estimations, taille de l'échantillon : 1000 répétitions (BRM et SEM), 500 jeux de données avec 80% de censure.

La méthode S-EM améliore la précision des estimateurs θ_{MVC} dans tous les cas de censures et de taille d'échantillons en termes d'erreur (a) 4.9 et de précision des estimations (b)4.9.

La méthode BRM, améliore les performance de SEM, pour un taux de censure fixé à 80% et toutes les tailles d'échantillon de $n = 100$ à 600 4.9.

• Le cas extrême : « petit » échantillon et fortement censuré.

n=100		80%	90%
MVC	$\hat{R}$	0.3	0.2
	$\hat{\beta} [I_{95\%}]$	3.5[2.7,4.3]	4[2.3,4.5]
SEM	$\hat{R} [I_{95\%}]$	0.4[0.26,0.56]	0.3[0.2,0.6]
	$\hat{\beta} [I_{95\%}]$	3.4[2.5,3.8]	3.5[2.4,4]
BRM	$\hat{R} [I_{95\%}]$	0.48[0.4,0.57]	0.48[0.4,0.6]
	$\hat{\beta} [I_{95\%}]$	3[2.57,3.4]	3[2.56,3.5]

TABLE 4.10 – Taux de censure. Fiabilité à mi-vie et $\hat{\beta}$ estimés pour 500 jeux de données de taille $n = 100$.

4.10 Conclusion

L'analyse de sensibilité nous montre que lorsque l'échantillon présente plus de 30% de données censurées, BRM est la méthode d'estimation la plus performante.

Censoring n=600	30%	70%	80%	90%	Sample Size $\tau\% = 80$	100	200	400	600
Méthode	SEM	BRM	BRM	BRM	Méthode	BRM	BRM	BRM	BRM

TABLE 4.11 – Choix de la méthode d'estimation en fonction du taux de censure et de la taille de l'échantillon.

4.11 Données réelles

Afin d'illustrer la mise en oeuvre et l'efficacité de la méthode BRM nous avons utilisé un jeu de 599 durées de vie de l'équipement 1, cf. Table 1.1 page 17, dont 83% sont censurées. Une comparaison avec les méthodes d'estimation MVC et S-EM et BRM a été effectuée. La méthode SEM est initialisée à l'aide de θ_{MV} voir 4.6. Pour valider l'ajustement d'une loi de Weibull à ce jeu de données, on peut utiliser un Weibull Quantile-Quantile plot (WQQP) cf 4.3.

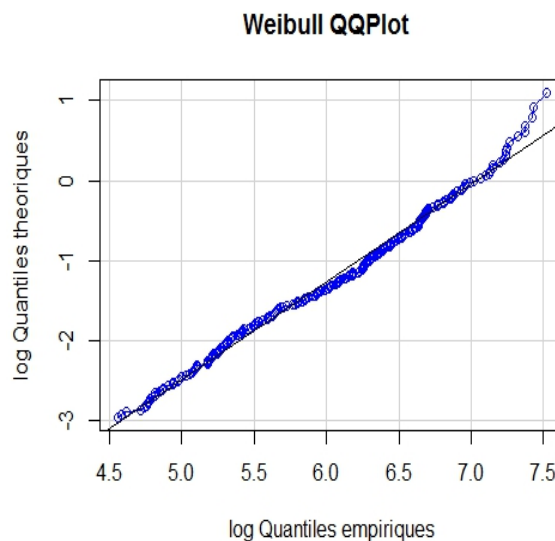


FIGURE 4.13: Weibull QQplot et densité des données. L'échantillon est de taille $n = 599$ et est censuré à 83%.

D'un point de vue industriel, sur-estimer la fiabilité d'un équipement engendrera un fort risque d'indisponibilité. A *contrario*, sous-estimer la fiabilité engendrera une problématique de compétitivité. Afin de mettre en évidence cette remarque, nous choisissons trois temps d'observation $T = 100, 1000$ et 2000 . L'objectif est d'estimer la fiabilité pour ces trois temps de prévision. Pour cela, on utilisera les trois méthodes d'estimation présentées dans ce chapitre : MVC, SEM et BRM.

Méthode	$\hat{\beta}_{I_{95\%}}$	$\hat{\eta}_{I_{95\%}}$	$\hat{R}(100)$	$\hat{R}(1000)$	$\hat{R}(2000)$
MVC	1.6 [1.19; 1.6]	2220 [2010; 2230]	98%	72%	42%
SEM	1.8 [1.66; 1.92]	1998 [1918; 2200]	99%	73%	38%
BRM	2 [1.8; 2.3]	1867 [1717, 1940]	99%	74%	30%

TABLE 4.12 – Données réelles, équipement 1 avec $n = 599$ avec 83% de censure. **Estimation MVC.** **Estimation S-EM**, avec 2 000 itérations, intervalle de crédibilité *bootstrap* sur les 300 dernières itérations. **Estimation BRM**, avec 2 000 répétitions, intervalle de crédibilité *bootstrap* après ré-échantillonnage.

Les analyses précédentes nous laissent penser que les méthodes MVC et S-EM sont moins précises dans le cadre d'un échantillon fortement censuré ($> 30\%$), dans le cas des données RETEX de l'équipement 1, la méthode BRM est alors préconisée.

Chapitre 5

Modèle de lois de Weibull en concurrences

Sommaire

5.1	Caractéristiques de lois de Weibull en concurrence	60
5.2	Données	60
5.3	Identifiabilité	61
5.4	Ajustement graphique	61
5.5	Vraisemblance complète	62
5.6	Algorithme EM	62
5.7	Algorithme S-EM	63
5.8	Estimation BRM	64
5.8.1	Pour les paramètres de forme	64
5.8.2	Pour les paramètres d'échelle	64
5.9	Conclusion	64

Ici nous supposons que les équipements ont deux causes de défaillance indépendantes l'une de l'autre. L'équipement peut être vu comme un système constitué de deux composants critiques en série : l'équipement est défaillant dès qu'un des deux composants est défaillant.

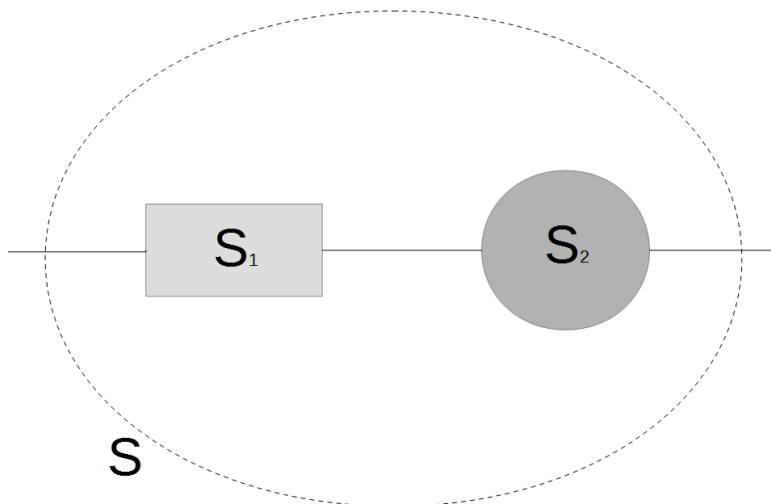


FIGURE 5.1: Équipement constitué de deux composants critiques S_1 et S_2 en série : la panne d'un composant entraîne la panne de l'équipement.

Pour $k = 1, 2$, si l'on note X_k l'inter-occurrence de panne du composant S_k , alors X , l'inter-occurrence de panne de l'équipement, est égale à :

$$X = \min(X_1, X_2).$$

Si la défaillance de l'équipement est due à la panne du composant S_1 , $X_1 < X_2$ et X_2 est alors « masquée ». Si l'occurrence de panne de l'équipement est censurée, alors X_1 et X_2 sont toutes les deux masquées.

Cette situation « à risques concurrents » est courante et par conséquent ces modèles ont été très largement étudiés, voir Crowder (2001), Murthy et al. (2004b) et Rinne (2008) pour une revue détaillée. Lorsque l'équipement est défaillant, il n'y a pas forcément d'analyses *post-mortem*, par conséquent la cause de la défaillance, n'est pas toujours renseignée dans le RETEX. La cause de la défaillance est donc ici une variable latente et les inter-occurrences pour chaque cause de défaillance peuvent être masquées. Dans cette situation avec de nombreuses variables manquantes, l'estimation des paramètres du modèle nécessite des techniques d'estimation spécifiques, comme par exemple les algorithmes EM, S-EM, ou encore l'estimation non itérative BRM.

Le plan du chapitre est le suivant : après avoir présenté les caractéristiques du modèle de lois de Weibull en concurrence, nous présentons les estimations par maximum de vraisemblance dans le cas de données manquantes, c-à-d l'algorithme EM et sa version stochastique S-EM. Dans le cas d'échantillons de petite taille ou très fortement censurés, Bacha and Celeux (1996) propose d'utiliser la méthode BRM pour le modèle de lois de Weibull en concurrence. Nous la décrivons dans la section suivante. Nous concluons en comparant les différentes approches.

Lors d'une défaillance la cause peut être observée, on parle de données complètes, soit masquée, on parle de données incomplètes. Dans le cas de donnée censurées (avant défaillance) la cause de la panne est bien évidemment masquée. Voir Chap 1.3 Crowder (2001) pour les données complètes et Bacha and Celeux (1996) pour les données masquées. Pour les défaillances à causes masquées, les origines des réalisations ne sont pas connues, les problèmes soulevés par les procédure d'inférence deviennent importants en raison de l'expression complexe de la densité de T , fonction des densités des variables T^j . Ci dessous, nous caractérisons un modèle de risques concurrents dans le cas des données incomplètes et à causes masquées.

5.1 Caractéristiques de lois de Weibull en concurrence

Fiabilité d'une loi de Weibull, pour $k = 1, 2$ et X_1, X_2 indépendants,

$$R_W(x|\theta) = \frac{\beta_k}{\eta_k} \left(\frac{x}{\eta_k} \right)^{\beta_k-1} e^{-\left(\frac{x}{\eta_k} \right)^{\beta_k}} ; \quad (5.1)$$

Pour $x \geq 0$, la fiabilité de deux lois de Weibull en concurrence,

$$R_{RCW}(x|\theta) = \prod_{k=1}^2 R_W(x|\beta_k, \eta_k), \quad (5.2)$$

$$f_{RCW}(x|\theta) = f(x|\beta_1, \eta_1)R(x|\beta_1, \eta_1) + f(x|\beta_2, \eta_2)R(x|\beta_2, \eta_2), \quad (5.3)$$

$$h_{RCW}(x|\theta) = \frac{\beta_1}{\eta_1} \left(\frac{x}{\eta_1} \right)^{\beta_1-1} + \frac{\beta_2}{\eta_2} \left(\frac{x}{\eta_2} \right)^{\beta_2-1}. \quad (5.4)$$

On note $\theta = (\beta_1, \eta_1, \beta_2, \eta_2)$, l'ensemble des paramètres du modèle.

5.2 Données

On considère n équipements. Pour $i = 1, \dots, n$, on note

- x_i = i^e occurrence de panne de l'équipement ;
- z_i = la cause de la défaillance
- = k si x_i est issue de la cause k

Remarque : la cause z_i n'est pas observée pour une défaillance. Les inter-occurrences de pannes pour chaque cause peuvent être masquées.

5.3 Identifiabilité

La question de l'identifiabilité doit être posée avant de parler d'estimation des paramètres. La famille de modèles paramétriques $f_{RCW}(x|\theta)$ est identifiable si deux valeurs distinctes du paramètre θ déterminent deux modèles différents : si $f_{RCW}(x|\theta) = f_{RCW}(x|\theta')$ alors à une permutation près $\theta = \theta'$.

Concernant les modèles deux lois de Weibull indépendantes en concurrences et à causes masquées, il est impossible d'identifier les θ_k individuellement. une condition suffisante d'identifiabilité à été établie par Basu and K. Ghosh (1978) et décrit dans Bacha (1996). Sous les hypothèses :

1. il existe $x_0 \in \mathbb{R}$ tel que, $f(x|\theta)$ est continue et strictement positive sur $]x_0, +\infty[$;
2. pour tout $F(x|\theta)$ et $F(x|\theta')$, $\lim_{x \rightarrow \infty} \frac{f(x|\theta)}{f(x|\theta')} \in \{0, +\infty\}$

alors le modèle de deux lois de Weibull indépendantes est identifiable.

5.4 Ajustement graphique

On peut montrer que pour un modèle de deux lois de Weibull en concurrences, on a la propriété suivante, cf. Jiang and Murthy (2003),

Propriété du Weibull quantile-quantile plot pour un modèle à risques concurrents

1. Le WQQP est convexe ;
2. Le WQQP a deux asymptotes : si $\beta_1 < \beta_2$ alors
 - à gauche : $\Delta_1 : y = \beta_1(x - \ln(\eta_1))$
 - à droite : $\Delta_2 : y = \beta_2(x - \ln(\eta_2))$

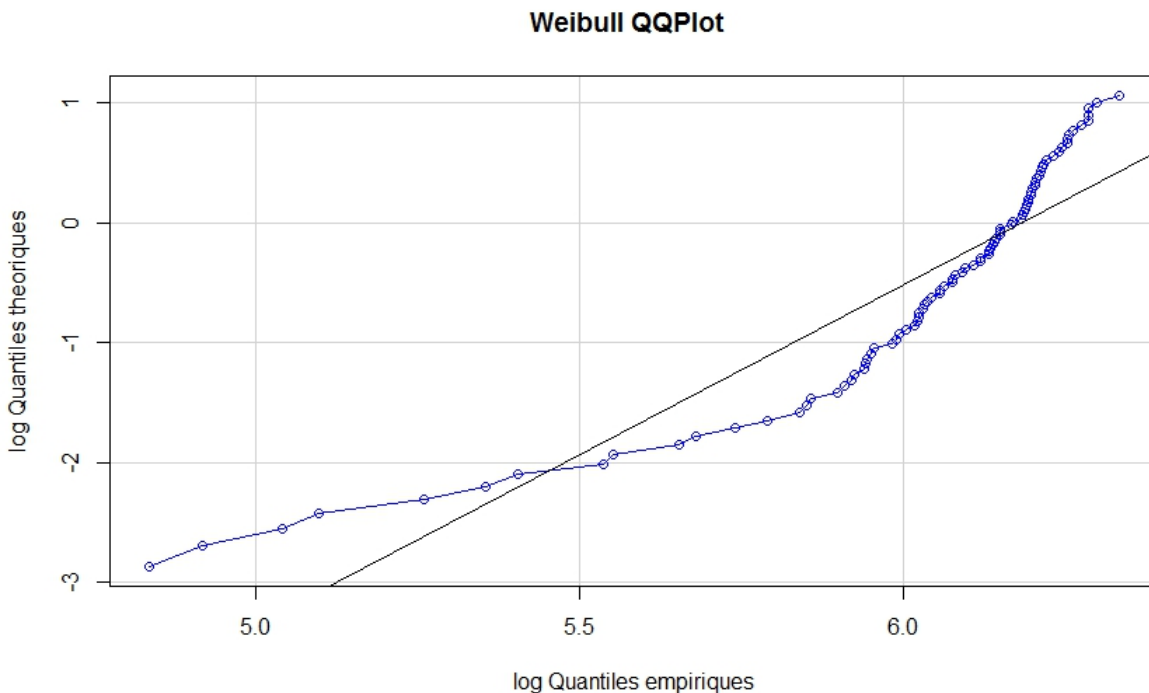


FIGURE 5.2: Risques concurrents : Weibull quantile-quantile plot convexe.

Le WQPP permet de valider l'ajustement à un modèle de lois de Weibull en concurrence, et par ailleurs estimer les paramètres (β_k, η_k) par une régression linéaire par morceaux cf. Muggeo (2008). Cette estimation peut servir comme initialisation d'EM, S-EM ou BRM.

5.5 Vraisemblance complète

Si la cause est connue alors :

$$\begin{aligned}
 L((\mathbf{x}, \mathbf{z}) | \theta) &= \prod_{i=1; z_i=1}^n f_W(x_i | \beta_1, \eta_1) \cdot R_W(x_i | \beta_2, \eta_2) \cdot \prod_{i=1; z_i=2}^n f_W(x_i | \beta_2, \eta_2) \cdot R_W(x_i | \beta_1, \eta_1) \\
 &= \prod_{i=1; z_i=1}^n h_W(x_i | \beta_1, \eta_1) R_W(x_i | \beta_1, \eta_1) R_W(x_i | \beta_2, \eta_2) \\
 &\quad \cdot \prod_{i=1; z_i=2}^n h_W(x_i | \beta_2, \eta_2) R_W(x_i | \beta_2, \eta_2) \cdot R_W(x_i | \beta_1, \eta_1) \\
 &= \prod_{i=1}^n R_W(x_i | \beta_1, \eta_1) \cdot R_W(x_i | \beta_2, \eta_2) \cdot \prod_{i=1; z_i=1}^n h_W(x_i | \beta_1, \eta_1) \cdot \prod_{i=1; z_i=2}^n h_W(x_i | \beta_2, \eta_2)
 \end{aligned}$$

La log-vraisemblance s'écrit alors

$$\ln L((\mathbf{x}, \mathbf{z}) | \theta) = \sum_{i=1}^n \ln R_W(x_i | \beta_1, \eta_1) + \ln R_W(x_i | \beta_2, \eta_2) + \sum_{i=1; z_i=1}^n \ln(h_W(x_i | \beta_1, \eta_1)) + \sum_{i=1; z_i=2}^n \ln(h_W(x_i | \beta_2, \eta_2)) \quad (5.5)$$

5.6 Algorithme EM

On utilise alors l'algorithme EM cf. Algorithme 2 page 31, dans le cas de données complètes :

Étape E : On calcule

$$\begin{aligned}
 \mathbb{Q}(\theta | \theta^{(r)}) &= \mathbb{E}[\ln(L((\mathbf{X}, \mathbf{Z}) | \theta)) | \mathbf{X}; \theta^{(r)}] \\
 &= \sum_{i=1}^n \ln R_W(x_i | \beta_1, \eta_1) + \ln R_W(x_i | \beta_2, \eta_2) \\
 &\quad + \sum_{i=1}^n \mathbb{E}[\mathbf{1}(Z_i = 1) | \mathbf{X}; \theta^{(r)}] h_W(x_i | \beta_1, \eta_1) \\
 &\quad + \sum_{i=1}^n \mathbb{E}[\mathbf{1}(Z_i = 2) | \mathbf{X}; \theta^{(r)}] h_W(x_i | \beta_2, \eta_2) \quad (5.6)
 \end{aligned}$$

Étape M : on obtient $\theta^{(r+1)}$ en maximisant $\mathbb{Q}(\theta | \theta^{(r)})$. On peut observer que l'équation 5.6 peut être maximisée séparément selon (β_k, η_k) .

On obtient $\beta_k^{(r+1)}$ en résolvant l'équation :

$$\frac{1}{\beta} + \frac{\sum_{i=1}^n \mathbb{P}(Z = k | x_i; \theta_k^{(r)}) \ln(x_i)}{\sum_{i=1}^n \mathbb{P}(Z = k | x_i; \theta_k^{(r)})} - \frac{\sum_{i=1}^n x_i^\beta \ln(x_i)}{\sum_{i=1}^n x_i^\beta} = 0; \quad (5.7)$$

et $\eta_k^{(r+1)}$

$$\eta^{r+1} = \left(\frac{\sum_{i=1}^n x_i^{\beta_k^{r+1}}}{\sum_{i=1}^n \mathbb{P}(Z = k | x_i; \theta_k^{(r)})} \right)^{\frac{1}{\beta_k^{(r+1)}}} \quad (5.8)$$

5.7 Algorithme S-EM

Dans le cas de données incomplètes et à cause masquée,

Étape Simulation : sachant $\theta^{(r-1)}$ on simule

– pour une donnée non censurée,

1. la cause de défaillance selon la probabilité conditionnelle

$$\begin{aligned} \mathbb{P}(Z = k | x_i; \theta_k^{(r)}) &= \frac{f_W(x_i | \beta_1^{(r)}, \eta_1^{(r)}) \cdot R_W(x_i | \beta_2^{(r)}, \eta_2^{(r)})}{f_{CRW}(x_i | \theta^{(r)})} \cdot \mathbf{1}(k = 1) \\ &+ \frac{f_W(x_i | \beta_2^{(r)}, \eta_2^{(r)}) \cdot R_W(x_i | \beta_1^{(r)}, \eta_1^{(r)})}{f_{CRW}(x_i | \theta^{(r)})} \cdot \mathbf{1}(k = 2) \\ &= \frac{h_W(x_i | \beta_k^{(r)}, \eta_k^{(r)})}{h_W(x_i | \beta_1^{(r)}, \eta_1^{(r)}) + h_W(x_i | \beta_2^{(r)}, \eta_2^{(r)})} \end{aligned}$$

2. on construit les échantillons de chaque cause : par exemple si $\tilde{z}_i = 1$ alors on classe x_i dans la cause 1 et on simule l'autre cause de défaillance,

$$\begin{aligned} \tilde{x}_1 &= x_i \\ \tilde{x}_2 &= \eta_2^{(r)} \left[\left(\frac{x_i}{\eta_2^{(r)}} \right)^{\beta_2^{(r)}} - \ln(U) \right]^{\frac{1}{\beta_2^{(r)}}} \end{aligned}$$

où U suit la loi uniforme sur $[0; 1]$, cf. 4.9 page 40.

– pour une donnée censurée on simule les deux causes de panne.

Étape M : pour $k = 1, 2$, on obtient $(\beta_k^{(r+1)}, \eta_k^{(r+1)})$ en maximisant la vraisemblance des données complétées $(\tilde{x}_k)$ cf. (4.11) et (4.12)

Estimateur S-EM : on estime alors θ par

$$\hat{\theta}_{S-EM} = \frac{1}{B - \iota} \sum_{r=\iota+1}^B \theta^{(r)}. \quad (5.9)$$

Ici encore, pour un échantillon de petite taille, ou fortement censuré du fait que l'algorithme peut être piégé autour de maxima locaux la convergence de S-EM peut être très lente cf. Bacha (1996).

5.8 Estimation BRM

Comme nous l'avons vu dans le chapitre précédent, *i.e* § 4.8, l'idée de l'estimation BRM est :

1. dans une première étape, de simuler les paramètres du modèle à partir de leur loi jointe *a priori*;
2. dans une seconde étape, de « restaurer » les données manquantes en les simulant conditionnellement aux paramètres obtenus, pour obtenir l'estimateur du maximum de vraisemblance sur les données complétées.
3. la loi empirique des estimateurs ainsi obtenus permet d'approcher la loi *a posteriori*.

L'estimation BRM d'un modèle de lois de Weibull en concurrences a été étudiée dans Bacha (1996). Dans cette thèse nous utilisons les lois *a priori* adaptées à notre problématique :

5.8.1 Pour les paramètres de forme

$$\beta \sim b^{min} + \mathcal{B}(p, q)(b^{max} - b^{min}), \quad (5.10)$$

$$\pi(\beta) = \frac{\Gamma(p+q)}{\Gamma(p)\Gamma(q)} \frac{(\beta - b^{min})^{p-1} (b^{max} - \beta)^{q-1}}{(b^{max} - b^{min})^{p+q+1}} \mathbb{I}_{[0;1]} \left(\frac{\beta - b^{min}}{b^{max} - b^{min}} \right),$$

où

$$\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt,$$

La calibration choisie au § 4.8 sera utilisée dans le cas de deux lois de Weibull en concurrences.

5.8.2 Pour les paramètres d'échelle

$$\pi(\eta | \beta, e^{shape}, e^{scale}) = \frac{\beta (e^{scale})^\beta e^{shape}}{\Gamma(e^{shape})} \eta^{-e^{shape} \beta - 1} e^{-\left(\frac{e^{shape}}{\eta}\right)^\beta}, \quad (5.11)$$

où (e^{shape}, e^{scale}) sont des hyperparamètres. La calibration choisie au § 4.8 sera utilisée dans le cas de deux lois de Weibull en concurrences.

5.9 Conclusion

Les algorithmes EM, SEM et BRM ont été largement expérimentés et comparés dans Bacha et al. (1998). Nous avons fait le même constat :

- Lorsque la taille de l'échantillon est grande, que le taux de censure est faible, l'algorithme EM est recommandé.
- Dans tous les autres cas, l'approche bayésienne à l'aide de l'algorithme BRM est préférée, et ce même avec des lois *a priori* peu informatives.

Toutefois, ces préconisations sont à nuancer : SEM pourrait être une alternative à EM et BRM dans le cas de taux de censure modérés.

Chapitre 6

Modèle de mélange de lois de Weibull

Sommaire

6.1	Estimation bayésienne d'un mélange de lois Weibull	67
6.1.1	Introduction	69
6.2	Mixture of Weibull distributions	70
6.2.1	Mixture of Weibull distributions	70
6.2.2	Weibull distributions	70
6.2.3	Data	71
6.3	Fitting the model	71
6.3.1	Weibull quantiles-quantiles plot	71
6.3.2	Maximum likelihood estimation	72
6.3.3	EM and S-EM algorithms	72
6.3.4	Bayesian estimation	73
6.4	Bayesian restoration maximization	73
6.4.1	Importance resampling	74
6.4.2	Proposal distribution	74
6.4.3	BRM Method	74
6.4.4	Prior elicitation	75
6.5	Simulations and results	76
6.5.1	Priors calibration	77
6.5.2	BRM number of runs	78
6.5.3	Sensitivity analysis	78
6.5.4	Comparisons	79
6.5.5	Discussion	80
6.6	Real data processing	81
6.6.1	Fitting a Weibull mixture to data	81
6.6.2	Example 1 - Throttles	82
6.6.3	Example 2 - Power assemblies	82
6.7	Conclusion	84

Ici on suppose que la population de l'équipement étudié n'est plus homogène : il y a deux sous-populations ayant chacune sa propre cause de défaillance. Cela correspond à une situation où les équipements sont susceptibles d'être différents en termes de qualité de production (*e.g.* avec ou sans défaut de fabrication), de vieillissement (*e.g.* usure normale ou usure due au vieillissement), etc. Le modèle utilisé est alors un modèle de mélange de lois de Weibull. Cependant le groupe

d'appartenance n'est pas observé, c'est une variable latente. L'estimation des paramètres du modèle de mélange de lois de Weibull dans le cas de données fortement censurées a donné lieu à un article cf. Ducros and Pamphile (2018).

Le plan du chapitre est le suivant. Nous commencerons par l'article dans lequel nous avons mis en œuvre l'estimation BRM pour un mélange. Puis nous mettrons en œuvre BRM pour un jeu de données réelles pour lequel l'analyse graphique indique la présence d'un mélange.

6.1 Estimation bayésienne d'un mélange de lois Weibull en présence de données fortement censurées

Résumé

Dans le cadre d'analyses de fiabilité ou de garantie, on est souvent amené à traiter des données de durée de vie non homogènes. En effet, la durée de vie d'un système peut être affectée par la variabilité des conditions de production, et par les conditions d'utilisation. La plupart du temps, cette variabilité n'est pas observée mais doit cependant être prise en compte. Par ailleurs la présence de données fortement censurées rend difficile l'estimation des paramètres. Le mélange de Weibull à deux composantes est alors un modèle pertinent pour prendre en compte l'hétérogénéité des durées de vie. Les performances des méthodes d'estimation classiques (maximum de vraisemblance, *via* EM, Bayes *via* MCMC) sont décevantes en raison du nombre de paramètres et du fort taux de censure. Dans ce contexte, on propose d'utiliser une méthode de bootstrap bayésien, appelée Bayesian Restoration Maximization. C'est une méthode non itérative qui fournit un échantillon de la loi *a posteriori*. Les résultats des simulations ont montré qu'en présence de données fortement censurées, la méthode BRM est plus performante que l'algorithme S-EM à la fois en termes de précision des estimations et de temps de calcul.

Abstract

In reliability or warranty analysis, engineers must often deal with lifetimes data that are non-homogeneous. Most of the time, this variability is unobserved but has to be taken into account for reliability or warranty cost analysis. A further problem is that in reliability analysis, data are heavily censored which makes estimations more difficult. The two-component Weibull mixture is then a highly relevant model to capture heterogeneity for a large majority of operating lifetimes. Unfortunately, the performance of classical estimation methods (maximum of likelihood *via* EM, Bayes approach *via* MCMC) is jeopardized due to the high number of parameters and the heavy censoring. In order to overcome the problem of heavy censoring for Weibull mixture parameters estimation, this research proposes a Bayesian bootstrap method, called Bayesian Restoration Maximization. The key is to provide a sampling from the posterior distribution. Thanks to an importance sampling technique, this sample focuses on the posterior mean. Prior distributions elicitation and sensibility analysis are discussed. Simulations results showed that, for heavily censored data, the BRM method outperforms the EM and S-EM algorithms in terms of estimates accuracy. On the other hand, it is a non-iterative method which therefore provides very short computation times. In addition, the BRM method does not suffer from the problem of label switching from Bayesian sampling algorithms such as MCMC methods. Finally, two real data sets are analyzed to illustrate the application of the method.

keyword Reliability analysis ; Weibull mixture ; Heavily censored data ; Non-homogeneous data ; Bayesian estimation ; Prior elicitation

Acronyms & Notation

WQQP	Weibull quantiles-quantiles plot
ML	maximum likelihood
EM	Estimation-Maximization
S-EM	Stochastic EM
MCMC	Markov chain Monte Carlo
BRM	Bayes restoration maximization
IS	importance sampling
RMSE	Root Mean Square Error
n	number of systems under study
x_i	lifetime of the i -th system ; $\mathbf{x} = (x_1, \dots, x_n)$
$\tilde{x}_i$	simulated lifetime
c_i	censoring indicator ; $\mathbf{c} = (c_1, \dots, c_n)$ $c_i = 1$, if x_i is censored, x_i is a incomplete lifetime ; $c_i = 0$ otherwise, x_i is a complete lifetime
z_i	component label ; $\mathbf{z} = (z_1, \dots, z_n)$ $z_i = k$, if the i -th system is coming from the k -th component of the mixture
$\tilde{z}_i$	simulated component label
$\mathbf{1}(\cdot)$	indicator function
f_W, R_W	probability density function and reliability of a Weibull distribution
f_{MIX}, R_{MIX}	probability density function and reliability of the mixture
β_k, η_k	shape and scale parameter of the k -th component of the mixture
α_k	mixing probability of the mixture : $\alpha_i \geq 0$ and $\alpha_1 + \alpha_2 = 1$
$\boldsymbol{\theta} = (\alpha_1, \beta_1, \eta_1, \beta_2, \eta_2)$	the set of model parameters
$\tilde{\boldsymbol{\theta}}$	simulated parameters
$\hat{\boldsymbol{\theta}}$	parameters estimate
$L(\mathbf{x}, \mathbf{c}, \mathbf{z} \boldsymbol{\theta})$	likelihood function with complete and incomplete lifetimes
$L(\mathbf{x}, \mathbf{z} \boldsymbol{\theta})$	likelihood function with complete lifetimes
$\pi(\mathbf{x}, \mathbf{z} \mathbf{x}, \mathbf{c}; \boldsymbol{\theta})$	probability of missing data $\mathbf{z}$ and incomplete lifetimes ($x_i, c_i = 1$), conditionally to data $(\mathbf{x}, \mathbf{c})$
$\pi(\boldsymbol{\theta})$	prior distribution
$\pi(\boldsymbol{\theta} \mathbf{x}, \mathbf{c}, \mathbf{z})$	posterior distribution
$\rho(\boldsymbol{\theta})$	proposal distribution for IS
w_d	weight of the d -th simulated parameter for IS

6.1.1 Introduction

In reliability or warranty analysis, engineers must often deal with data collections which are non-homogeneous. Usually, this non-homogeneity coincides with the presence of compliant or non-compliant products, early or random failures and so forth. For example, in the automotive industry, vehicle parts are produced by different suppliers. In this case, parts reliability may depend on the supplier, but also on the vehicle operating environment (e.g. temperature, humidity, roads, etc.) and the driver behavior (e.g. usage mode and usage intensity). Generally, these various sources of non-homogeneity are not controlled or not observed. But neglecting existing non-homogeneity, can lead to errors and misconceptions in the analysis. In manufacturing, reliability has a serious impact on the warranty servicing cost or maintenance cost. As for safety studies, conclusions of diagnostic tests can be severely misguided. Hence a mixture of distributions is commonly used to model failure times from non-homogeneous data Titterington et al. (1985); Finkelstein and Cha (2013); Wang and Jiang (2014).

More than five decades, Weibull distributions are extensively used for modelling lifetime data in medical or industrial applications. Weibull distributions exhibit a decreasing, constant or increasing hazard function, which makes them suitable for modelling complex failure data. See for example Murthy et al. (2004b); Elmahdy (2015) and references therein. Hence, the two-component Weibull mixture is a highly relevant model to capture latent heterogeneity for a large majority of operating reliability analyses.

Furthermore, the other issue raised in reliability analysis is the large number of censored failure times. Indeed, many failures cannot be observed because the corresponding product was removed from the study or was still operating at the end of the study period. For example, in warranty data collections, during the time interval over which data are collected, components are put into service at different times : failure time is censored for the large number of components still operating. The percentage of censored failure times may be greater than 70% , and draw statistical inferences is therefore a particularly difficult challenge.

Different methods are used to fit a parametric model to data, but they have substantial difficulties for Weibull mixtures in the setting of heavily censored data. A graphical method provides a quick but crude and non robust estimate Jiang and Murthy (1995). The maximum likelihood (ML) estimation is the much preferred method due to its desirable asymptotic properties. Expectation-maximization (EM) algorithm and its stochastic variant, S-EM, are iterative methods which allow to obtain ML estimate when data are missing. Therefore EM is widely used for mixtures when there is no information on the belonging to one of the two components. McLachlan and Peel (2000); Jiang and Kececioglu (1992); Bordes and Chauveau (2016). Unfortunately, ML estimator of the scale parameter of the Weibull distribution does not admit closed form and requires numerical approximations. Therefore for small samples or heavily censored samples, ML and S-EM estimates tend to be highly biased, see Balakrishnan and Mitra (2012); Dias and Wedel (2004) for numerical examples.

The Bayesian approach provides alternatives to ML approach. For mixture estimation, Bayesian methods are however very computationally intensive. They require complex integrations Touw (2009) or the use of Markov chain Monte Carlo algorithms (MCMC) Robert and Casella (2010). The higher number of parameters combined with the heavy censoring rate make Bayesian methods appear to be deceptive in terms of convergence for mixtures Marin et al. (2005).

Therefore, fitting a Weibull mixture model to highly censored data may lead to significant statistical errors in the reliability estimation, especially for mid or large mission times. In Bacha and Celeux (1996), Bacha and Celeux proposed the Bayesian restoration maximization (BRM) to fit a single Weibull to small and highly censored data. BRM is similar to Bayesian bootstrap methods Rubin (1981) : its special feature is that prior information and importance resampling technique are combined to obtain a sample of the model parameters concentrated on the posterior mean. In this paper, BRM approach was implemented for Weibull mixture inference. The key fea-

ture of the Bayesian method is that it is necessary to specify the prior distribution. This fundamental issue has been discussed. In particular, the impact of the hyperparameters on BRM performance has been studied. Sensitivity analysis of BRM performances to sample size and rate of censoring has been carried out. The BRM method addresses the common issue of label switching with mixture. On the other hand, unlike S-EM and MCMC, it is not an iterative method which significantly reduces computation times. Comparisons with S-EM algorithm showed that BRM is more effective, both in terms of estimates precision as well as computation times.

The paper is organized as follows. The Weibull mixture and data are presented in Section 2. Classical estimation methods are set out for the Weibull mixture in Section 3. In Section 4, the BRM method is extended to the Weibull mixture. Calibration of the priors hyperparameters, sensibility analysis and performance comparisons with S-EM algorithm are discussed in Section 5. Two real data sets are analyzed in Section 6, to illustrate the application of the method. Finally, Section 7 gives a conclusion.

6.2 Mixture of Weibull distributions

Weibull distributions have been extensively used to model lifetimes in medical or industrial applications. This is clearly illustrated by the large number of references on it. A review of the Weibull distribution, in many of its aspects, can be found in Murthy et al. (2004b).

6.2.1 Mixture of Weibull distributions

Weibull distributions have been extensively used for modelling life data in medical or industrial applications. This is clearly illustrated by the large number of references on it. A review of the Weibull distribution, in many of its aspects, can be found in Murthy et al. (2004b). In particular, Weibull distributions exhibit decreasing, constant or increasing hazard function which makes them suitable for modelling complex failure data.

6.2.2 Weibull distributions

A two-component mixture was used where both components belong to the family of two-parameter Weibull distributions. For $x \geq 0$, the probability density function is

$$f_{MIX}(x|\boldsymbol{\theta}) = \sum_{k=1}^2 \alpha_k f_W(x|\beta_k, \eta_k), \quad (6.1)$$

where for $k = 1, 2$

- $f_W(x|\beta_k, \eta_k)$ is the Weibull probability density function with shape parameter $\beta_k > 0$ and scale parameter $\eta_k > 0$,

$$f_W(x|\beta_k, \eta_k) = \frac{\beta_k}{\eta_k} \left(\frac{x}{\eta_k} \right)^{\beta_k-1} e^{-\left(\frac{x}{\eta_k} \right)^{\beta_k}}; \quad (6.2)$$

- $\alpha_k > 0$ is the mixing proportion of the k -th component, with $\alpha_1 + \alpha_2 = 1$;

Let us denote $\boldsymbol{\theta} = (\alpha_1, \beta_1, \eta_1, \beta_2, \eta_2)$, the set of model parameters.

6.2.3 Data

Lifetimes of n systems or system parts are being studied. Some of them were already failed, and the exact time to failure is observed. For the most part they are still operating, and the time to failure is censored. The available information in the data are :

$$\begin{aligned} x_i & \text{ the lifetime of the } i\text{-th system, } i = 1, \dots, n; \\ c_i & \text{ indicates if the time to failure is censored or not, } i = 1, \dots, n. \\ & = \begin{cases} 1 & \text{if } x_i \text{ is censored,} \\ 0 & \text{otherwise.} \end{cases} \end{aligned}$$

Both deterministic and random censoring are considered, see Murthy et al. (2004b) for more details on censoring. In warranty data set, for example, during the time interval over which data are collected, components are put into service at different times. Then, components still under warranty have a lifetime randomly right censored.

6.3 Fitting the model

Model identifiability question must be settled before to be able to discuss of parameters estimation. The parametric family $f_{MIX}(x|\boldsymbol{\theta})$ is said identifiable if distinct values of $\boldsymbol{\theta}$ determine distinct members of the family : if $f_{MIX}(x|\boldsymbol{\theta}) = f_{MIX}(x|\boldsymbol{\theta}')$ then we can permute the component labels so that $\boldsymbol{\theta} = \boldsymbol{\theta}'$. This issue of identifiability for finite Weibull mixtures was settled in Ahmad (1988).

Reliability data analysis involves in a first instance to select an appropriate model to the data and then to estimate its parameters.

6.3.1 Weibull quantiles-quantiles plot

The Weibull quantiles-quantiles plot (WQQP) can provide a quick and simply confirmatory method for data non-homogeneity Murthy et al. (2004a) :

- for a standard single Weibull distribution, the Weibull QQ-plot has a straight line shape ;
- for a two-component Weibull mixture, the Weibull QQ-plot has a single inflection point (S-shaped) with parallel asymptotes.

According to Jiang and Murthy (1995), a roughly estimate of the two-component Weibull parameters, $\hat{\boldsymbol{\theta}}^{WQQP}$, can be obtained from the Weibull QQ plot, see Figure 6.1. These estimates are non precise nor robust, but nonetheless can be used as starting values for alternative estimation algorithms.

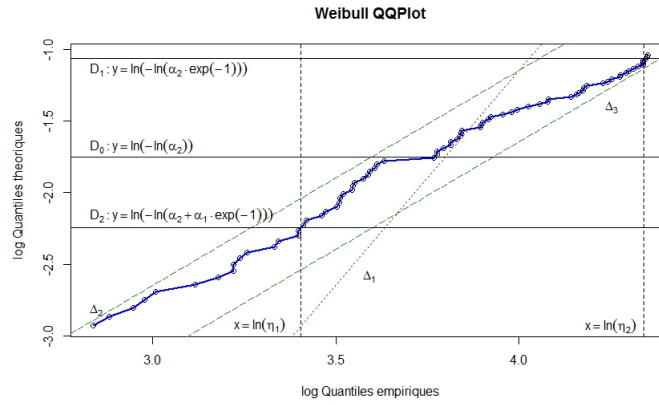


FIGURE 6.1: Weibull QQ-plot. **Graphical test of mixture** : for a two-component mixture of Weibull, the scatterplot has a S shape ; as $x \rightarrow +\infty$ the Weibull QQ-plot asymptote is Δ_3 of slope $1/\beta_2$; as $x \rightarrow -\infty$ the Weibull QQ-plot asymptote is Δ_2 of slope of $1/\beta_2$ too ; the line Δ_1 is of slope $1/\beta_1$. **Graphical estimation** : parameters α, η_1, η_2 can be roughly estimated, see Jiang and Murthy (1995) ; Line D_0 highlights the inflexion point of the S curve and allows to estimate α_2 ; Intersection of D_1 (respectively D_2) with the curve allows to estimate η_1 (respectively η_2).

6.3.2 Maximum likelihood estimation

The ML method is the most widely used method for fitting parametrical model to lifetimes data. Let us denote $\mathbf{z} = (z_1, \dots, z_n)$ the component labels :

z_i indicates the component label of the i -th lifetime, $i = 1, \dots, n$,
 $= k$, if the failure time of the i -th system belongs to the k -th component of the mixture.

For Weibull mixture, in the case where $\mathbf{z}$ are known, the ML estimator is (cf. Jiang and Kececiloglu (1992))

$$\hat{\boldsymbol{\theta}}^{MLE} = \arg \max_{\boldsymbol{\theta}} \ln L(\mathbf{x}, \mathbf{c}, \mathbf{z} | \boldsymbol{\theta}) \quad (6.3)$$

where the Likelihood function is

$$L(\mathbf{x}, \mathbf{c}, \mathbf{z} | \boldsymbol{\theta}) = \left[\prod_{i=1, c_i=0}^n \prod_{k=1}^2 \alpha_k [f_W(x_i | \beta_k, \eta_k)]^{1(z_i=k)} \right] \cdot \left[\prod_{i=1, c_i=1}^n \prod_{k=1}^2 \alpha_k (R_W(x_i | \beta_k, \eta_k))^{1(z_i=k)} \right] \quad (6.4)$$

The likelihood contains all information in the data that is relevant to estimate parameters. The first term in Eq. (6.4) integrates the lifetimes of systems which have already failed and the second term the lifetimes of systems which are still operating. Furthermore, we can see that the log-likelihood function can be maximized separately for $\alpha_1, (\beta_1, \eta_1)$ and (β_2, η_2) . Thus, if the component labels are known, estimates are fairly straightforward obtained by means of ML method, see Murthy et al. (2004b) for computational details. Unfortunately, a postmortem analysis is necessary to ensure identification of the component labels, which is not always possible. The idea of augmentation methods, like EM and Bayes methods, is to complete the missing or incomplete data so as to make the estimation more easier Tanner and Wong (2010).

6.3.3 EM and S-EM algorithms

The Expectation-maximization (EM) algorithm provides ML estimation for data containing missing values or incomplete data ?. Here,

- component labels, $\mathbf{z}$, are missing data ;
- censored failure time, x_i with $c_i = 1$, are incomplete data ;

- exact failure time, x_i with $c_i = 0$, are complete data.

Starting from a fixed value of the model parameter, the EM algorithm iterated between a E-step in which the expectation of the log-likelihood with completed data conditionally to the data, $E[L(\mathbf{x}, \mathbf{z}) | \mathbf{x}, \mathbf{c}; \boldsymbol{\theta}]$ (cf. Eq. (6.4)) is computed, and a M-step in which this conditional expectation is maximized to update the parameters (see Jiang and Kececioglu (1992); Elmahdy (2015) for computational details). At each step the log-likelihood increases, and the sequence of estimates converges toward the maximum of the likelihood under certain conditions.

The maximization of the conditional expectation requires numerical integration. Then, for highly censored or small data set, EM converges slowly and towards poor local maximizers Baudry and Celeux (2015). Stochastic EM (S-EM) resolves this difficulty, see Celeux et al. (1996). A stochastic step is added to EM algorithm (S-step) in which the conditional expectation of log-likelihood is estimated, simulating the missing data. The sequence of estimators, thus obtained, is a Markov chain and converges to a stationary distribution. Hence, after a sufficiently long warm-up period, the parameters are estimated from the last iterations, see Bordes and Chauveau (2016) for computational details. In the context of heavy censoring and many parameters, S-EM algorithm can break down. Indeed, for some iterations, no observation are assigned to a component and therefore parameters may not be estimated. Numerical experiments can be found in Balakrishnan and Mitra (2012) and Dias and Wedel (2004).

6.3.4 Bayesian estimation

Bayesian estimation is an alternative to ML estimation. It consists of incorporating prior information on the parameter $\boldsymbol{\theta}$, to produce, in accordance with the data, a posterior distribution of the parameter. Prior information may be derived from data of previous studies or expert advices, which quantifies uncertainty about the parameters. It is formalized by a prior distribution $\pi(\boldsymbol{\theta})$ on $\boldsymbol{\theta}$. The posterior distribution $\pi(\boldsymbol{\theta} | \mathbf{x}, \mathbf{c}, \mathbf{z})$, is obtained using Bayes' formula,

$$\pi(\boldsymbol{\theta} | \mathbf{x}, \mathbf{c}, \mathbf{z}) \propto L(\mathbf{x}, \mathbf{c}, \mathbf{z} | \boldsymbol{\theta}) \pi(\boldsymbol{\theta}). \quad (6.5)$$

Hence, a common Bayesian estimator is the mean of the posterior distribution

$$\hat{\boldsymbol{\theta}}^{Bayes} = E[\boldsymbol{\theta} | \mathbf{x}, \mathbf{c}, \mathbf{z}] = \int_{\boldsymbol{\theta}} \boldsymbol{\theta} \pi(\boldsymbol{\theta} | \mathbf{x}, \mathbf{c}, \mathbf{z}). \quad (6.6)$$

This estimator minimizes the Bayesian mean square error. In case of Weibull mixture, the posterior distribution is not in closed form, and direct Bayesian estimation is quite difficult, see Touw (2009) for instance. Even Markov chain Monte Carlo methods (MCMC) appear to be deceptive in the context of Weibull mixture. MCMC are time consuming, and convergence is not assured Marin et al. (2005).

The Bayesian restoration maximization (BRM) method has been used to estimate parameters of a basic Weibull distribution in Bacha and Celeux (1996). Here, the BRM method was used to Weibull mixture inferences. It is similar to a Bayesian bootstrap method Rubin (1981). The key feature of BRM is that the expectation in Eq. (6.6) is computed using importance resampling technique Rubin (1988).

6.4 Bayesian restoration maximization

The principle of BRM is to provide a sample from the joint posterior distribution via a simple simulation. Prior information and importance sampling technique are combined to obtain a sample of the model parameters concentrated on the posterior mean.

6.4.1 Importance resampling

The integral in Eq. (6.6) is calculated using importance sampling Ham. This technique exploits the identity

$$\hat{\theta}^{Bayes} = \int_{\theta} \theta \pi(\theta | \mathbf{x}, \mathbf{c}, \mathbf{z}) = \int_{\theta} \theta \frac{\pi(\theta | \mathbf{x}, \mathbf{c}, \mathbf{z})}{\rho(\theta)} \rho(\theta), \quad (6.7)$$

where the support of the distribution ρ includes the support of $\pi(\theta | \mathbf{x}, \mathbf{c}, \mathbf{z})$. Therefore Eq. (6.7) can be estimated by Monte Carlo integration :

1. $\tilde{\theta}^{(1)}, \dots, \tilde{\theta}^{(D)}$ are sampled according to a proposal distribution ρ ;
2. hence, θ is estimated by

$$\hat{\theta}^{BRM} = \frac{1}{D} \sum_{d=1}^D \tilde{\theta}^{(d)} w_d, \quad (6.8)$$

with the weights

$$w_d = \frac{\pi(\tilde{\theta}^{(d)} | \mathbf{x}, \mathbf{c}, \mathbf{z})}{\rho(\tilde{\theta}^{(d)})}. \quad (6.9)$$

6.4.2 Proposal distribution

Now, we have to select the sampling distribution ρ . It must be a counterfeit of the posterior distribution easier to obtain. From Bayes' formula we get

$$\pi(\theta | \mathbf{x}, \mathbf{c}, \mathbf{z}) \propto L(\mathbf{x}, \mathbf{c}, \mathbf{z} | \theta) \cdot \pi(\theta).$$

The likelihood, $L(\mathbf{x}, \mathbf{c}, \mathbf{z} | \theta)$, contains information on the data. For its part, the prior, $\pi(\theta)$, contains information on the parameter. Therefore, parameters are first simulated from the prior distribution, in order to simulate missing component labels and incomplete lifetimes. Lastly, ML estimates are computed from this completed sample. The sequence of ML estimates are thus concentrated on the posterior mean.

6.4.3 BRM Method

Let D denotes the number of runs for the BRM method, we get

1. **(B) prior sampling :** $\theta^{(1)}, \dots, \theta^{(D)}$ are sampled according to $\pi(\theta)$;
2. **(R) missing data restoration :** for $d = 1, \dots, D$, using $\theta^{(d)}$ as starting value, missing component labels and censored failure times are simulated according to their conditional probability. From Bayes's formula we get :
 - for a failure time, the conditional probability that it is coming from the k -th component is

$$\pi(k | x_i, c_i = 0; \theta^{(d)}) = \frac{\alpha_k^{(d)} \cdot f_W(x_i | \beta_k^{(d)}, \eta_k^{(d)})}{\sum_{j=1}^2 \alpha_j^{(d)} \cdot f_W(x_i | \beta_j^{(d)}, \eta_j^{(d)})}; \quad (6.10)$$

- for a censored failure time, the conditional probability that it is coming from the k -th component is

$$\pi(k | x_i, c_i = 1; \theta^{(d)}) = \frac{\alpha_k^{(d)} \cdot R_W(x_i | \beta_k^{(d)}, \eta_k^{(d)})}{\sum_{j=1}^2 \alpha_j^{(d)} \cdot R_W(x_i | \beta_j^{(d)}, \eta_j^{(d)})}; \quad (6.11)$$

- the conditional probability function of a censored failure time is

$$\pi(x|z_i = k, x_i, c_i = 1; \boldsymbol{\theta}^{(d)}) = \frac{f_W(x_i | \beta_k^{(d)}, \eta_k^{(d)})}{R_W(x_i | \beta_k^{(d)}, \eta_k^{(d)})} \mathbf{1}(x > x_i). \quad (6.12)$$

3. **(M) maximization** : then, the log-likelihood with completed sample, $(\tilde{\mathbf{x}}, \tilde{\mathbf{z}})$, is maximized to obtain $\hat{\boldsymbol{\theta}}^{(d)}$.
4. **importance resampling** : the distribution of $\hat{\boldsymbol{\theta}}^{(1)}, \dots, \hat{\boldsymbol{\theta}}^{(D)}$ is expected to be related to the posterior distribution. It cannot be obtained in closed form. Its kernel density estimate is

$$\hat{\rho}(\boldsymbol{\theta}) = \frac{1}{D \cdot h_{\boldsymbol{\theta}}} \sum_{d=1}^D K\left(\frac{\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}^{(d)}}{h_{\boldsymbol{\theta}}}\right), \quad (6.13)$$

where K is a Gaussian kernel.

The window size, $h_{\boldsymbol{\theta}}$, minimizing the mean integrated square error is given in Scott (2015).

5. **additional resampling** : it is worth mentioning that the weights can be uneven distributed weights, and several may be nil or almost nil. This occurs when prior distribution and likelihood are separated. This issue can be resolved by an additional resampling step, see Rubin (1988). For $M < D$, $(\tilde{\boldsymbol{\theta}}^{(1)}, \dots, \tilde{\boldsymbol{\theta}}^{(M)})$ are resampled from $(\hat{\boldsymbol{\theta}}^{(1)}, \dots, \hat{\boldsymbol{\theta}}^{(D)})$ according to $(w_1, \dots, w_D)$. Resampling step eliminates $\hat{\boldsymbol{\theta}}^{(d)}$ with low weight and duplicates those with high weight. According to Rubin (1988), very large ratio D/M may be require for a suitable performance with this additional resampling. Here, the ratio was set at $M = D/30$.
6. **credibility interval** : within a Bayesian approach, credibility interval is obtained from percentiles of the posterior distribution. For the BRM estimation, bounds of the 95% credibility interval are the 2.5% and 97.5% percentiles of $(\tilde{\boldsymbol{\theta}}^{(1)}, \dots, \tilde{\boldsymbol{\theta}}^{(M)})$ empirical distribution.

6.4.4 Prior elicitation

The choice of priors distribution is the key problem in the Bayesian estimation. One can distinguish informative/non-informative prior distribution reflecting the fact that prior information is more or less vague. The Bayesian estimate Eq. (6.6) shrinks the ML estimate Eq. (6.3) toward the prior mean, according to how informative the prior is or how large the sample size is. Therefore, non-informative prior is used to let the data speak for themselves, without being overly influenced by the prior. The richer the family of prior distributions is, the more capture of information about the parameter is. On the other hand, the choice of prior is constrained by the tractability of the posterior distribution. For a given sampling distribution $f(x|\boldsymbol{\theta})$, a conjugate prior distribution $\pi(\boldsymbol{\theta})$ is one for which the prior and the posterior distributions are in the same distribution family. Conjugate priors are useful to provided tractable posterior distributions. Unfortunately, no natural conjugate prior distribution exists for Weibull distribution when both the shape and scale parameters are unknown Soland (1969). Independent priors for (α_1) , (β_1, η_1) and (β_2, η_2) were used.

Weibull shape parameter prior elicitation

Thanks to its shape parameter, β , Weibull distributions have the ability to assume the characteristics of many different types of lifetimes Murthy et al. (2004a). This has made it extremely popular to fit data from various fields. Refs. Abernethy (2004); Jiang and Murthy (2011) deal with the range of β in various applications. Therefore a beta prior, symmetric on range interval $[b^{min}; b^{max}]$, has been chosen for both (β_1, β_2) ,

$$\pi(\beta) = \frac{\Gamma(2b^{shape})}{\Gamma^2(b^{shape})} \frac{(\beta - b^{min})^{b^{shape}-1} (b^{max} - \beta)^{b^{shape}-1}}{(b^{max} - b^{min})^{2b^{shape}+1}} \mathbf{1}_{[0;1]}\left(\frac{\beta - b^{min}}{b^{max} - b^{min}}\right), \quad (6.14)$$

where

$$\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt.$$

The beta distribution offers a lot of flexibility. For example, it incorporates Jeffrey's non-informative prior, $b^{shape} = 0.5$, and the uniform prior, $b^{shape} = 1$. Hyperparameters b^{min} and b^{max} reflect engineer expertise on the parameters (β_k) . They shall be fixed in accordance to the field failure data. The hyperparameter b^{shape} reflects uncertainty about (β_k) . Depending on b^{shape} , the beta distribution can be informative or non informative.

According to Barringer (2002), the hyperparameters b^{min} and b^{max} were chosen as follows :

$$b^{min} = 0.5 \quad , \quad b^{max} = 10. \quad (6.15)$$

A weakly informative prior has been chosen by setting $b^{shape} = 1.1$.

Weibull scale parameter prior elicitation

When β_k is known, the generalised inverse gamma distribution is the conjugate prior of the Weibull with parameter η_k , see Erto and Giorgio (2002) :

$$\pi(\eta | \beta_k, e^{shape}, e^{scale}) = \frac{\beta_k (e^{scale})^{\beta_k} e^{shape}}{\Gamma(e^{shape})} \eta^{-e^{shape} \beta_k - 1} e^{-\left(\frac{e^{scale}}{\eta}\right)^{\beta_k}}, \quad (6.16)$$

where $(e^{shape}, e_1^{scale}, e_2^{scale})$ are hyperparameters. A natural approach to compute $(e_1^{scale}, e_2^{scale})$ is to equate the expected value of the inverted gamma distribution to graphical estimation from the Weibull QQ-plot, $\hat{\eta}_k^{WQQP}$, discussed in Section 6.3.1 :

$$e_k^{scale} \frac{\Gamma\left(e^{shape} - \frac{1}{\beta_k}\right)}{\Gamma(e^{shape})} = \hat{\eta}_k^{WQQP}. \quad (6.17)$$

This starting value does not need to be sharp as will be shown in Section 6.5. The hyperparameter, e^{shape} , was chosen large to specify a weak informative prior with $e^{shape} > \frac{1}{\beta_k}$ according to Eq. (6.17).

Mixing parameter prior elicitation

The prior on the mixing parameter, α_1 , represents the lack of information on probability α_1 on interval $[0; 1]$. The beta distribution is a natural conjugate prior for the Bernoulli distribution Marin et al. (2005). Therefore, a beta prior on $[a^{min}; a^{max}]$ was chosen for α_1 :

$$\pi(\alpha) = \frac{\Gamma(2a^{shape})}{\Gamma^2(a^{shape})} \frac{(\alpha - a^{min})^{a^{shape}-1} (a^{max} - \alpha)^{a^{shape}-1}}{(a^{max} - a^{min})^{2a^{shape}+1}} \mathbf{1}_{[0;1]} \left(\frac{\alpha - a^{min}}{a^{max} - a^{min}} \right). \quad (6.18)$$

The hyperparameters a^{min} and a^{max} represent beliefs about α_1 . They shall be fixed according to information on α_1 . The hyperparameter a^{shape} reflects uncertainty about α_1 . A weakly informative prior was chosen by setting $a^{shape} = 1.1$.

6.5 Simulations and results

Simulations were carried out at first to calibrate the hyperparameters and next to compared BRM and S-EM estimates. The data were simulated from the Weibull mixture given in Table 6.1.

TABLE 6.1 – Simulated data parameters

α_1	β_1	η_1	β_2	η_2
0.3	1.5	50	3	200

Various sample sizes and censoring rates were used.

Histogram and Weibull QQ-plot show the presence a mixture. As stated in Section 6.3.1, a rough estimate of parameters α , η_1 and η_2 can be obtained from the Weibull QQ-plot.

6.5.1 Priors calibration

The hyperparameters calibration was experimented from 500 simulated data sets with a sample size $n = 100$ and a right censoring level at 70%. The BRM estimation was applied to each data set, with 1 000 runs. Hence, the average bias and the root mean square error (RMSE) of the mid-life reliability have been computed to assess the estimation performance :

$$\text{bias} = \frac{1}{500} \sum_{i=1}^{500} (\hat{R} - R), \quad \text{RMSE} = \sqrt{\frac{1}{500} \sum_{i=1}^{500} (\hat{R} - R)^2}.$$

As expected, informative prior improved the performance : see Table 6.2 for the prior support of α and Table 6.3 for scale hyperparameter prior of η . Not sharply anticipated value of $(\tilde{\eta}_k)$ provided relatively small bias and RMSE for $e^{shape} = 5$, see Table 6.4.

TABLE 6.2 – α_1 prior hyperparameters, $[a_{min}, a_{max}]$ calibration. Mid-life reliability estimate for 500 data sets, with 1 000 runs and sample size $n = 100$; Diffuse prior $[0; 1]$ and more informative prior $[0; 0.5]$.

$[a_{min}; a_{max}]$	$[0; 1.0]$	$[0; 0.5]$
average bias	-0.0212	-0.0016
RMSE	0.0725	0.0546

TABLE 6.3 – η prior hyperparameters, e^{shape} calibration. Mid-life reliability BRM estimate for 100 data sets, with 1 000 runs and sample size $n = 100$. Generalized inverse gamma prior with several values of the shape hyperparameter : the higher e^{shape} is, the more the prior is informative.

e^{shape}	3	5	10	20
average bias	-0.0113	-0.0016	0.0068	0.0053
RMSE	0.0701	0.0546	0.0394	0.0336

TABLE 6.4 – η prior hyperparameters, e_k^{scale} calibration. Mid-life reliability BRM estimate for 500 data sets, with 1 000 runs and sample size $n = 100$. Generalized inverse gamma prior with several scale hyperparameter calibrations cf. Eq. 6.17. On the top, a weak informative prior with $e^{shape} = 5$. On the bottom, more informative prior with $e^{shape} = 20$.

$e^{shape} = 5$	$(\widetilde{\eta}_1, \widetilde{\eta}_2)$	$0.5(\eta_1, \eta_2)$	(η_1, η_2)	$1.5(\eta_1, \eta_2)$
	average bias	-0.1721	-0.0016	0.0617
	RMSE	0.1407	0.0546	0.0523
$e^{shape} = 20$	$(\widetilde{\eta}_1, \widetilde{\eta}_2)$	$0.5(\eta_1, \eta_2)$	(η_1, η_2)	$1.5(\eta_1, \eta_2)$
	average bias	-0.2244	0.0053	0.0644
	RMSE	0.1437	0.0336	0.0441

6.5.2 BRM number of runs

The number of runs reflects the computing time. The Figure 6.5 shows that the RMSE decreases slightly with the number of runs.

TABLE 6.5 – BRM runs number calibration. Mid-life reliability BRM estimate for 500 data sets with sample size $n = 100$ and censoring rate at 70%.

BRM number of runs	500	1 000	6 000
average bias	-0.0030	-0.0016	0.0032
RMSE	0.0544	0.0546	0.0512

6.5.3 Sensitivity analysis

The sample size and the censoring rate are two inputs that influence the estimates. As can be seen from Table 6.6 and Table 6.7, BRM estimation is still efficient even for small sample heavily censored.

It is worth mentioning that when the two components are not well separated (*i.e.* the causes of the two components are close cf. Jiang and Murthy (1995)), without postmortem analysis, it will be more difficult to tell to which component an observation from the mixture belongs.

TABLE 6.6 – Sample size analysis. Mid-life reliability BRM estimate for 500 data sets with 1 000 runs, 70% censoring and sample size $n = 100, 200$ and 600.

sample size	100	200	600
average bias	-0.0016	0.0088	0.0125
RMSE	0.0546	0.0432	0.0631

TABLE 6.7 – Censoring. Mid-life reliability BRM estimate for 500 data sets with 1 000 runs, sample size $n = 600$ and 70%, 80% and 90% censoring.

censoring	70%	80%	90%
average bias	0.0125	-0.0200	-0.0581
RMSE	0.0631	0.0828	0.0950

TABLE 6.8 – Mixture with weakly or well separated components. Mid-life reliability BRM estimate for 500 data sets with 1 000 runs, sample size $n = 100$ and 70% censoring.

censoring	weakly separated components	well separated components
(β_1, η_1)	(3,50)	(1.5,50)
mode	43.679	24.0375
(β_2, η_2)	(1.5,200)	(3,200)
mode	96.15	174.7161
average bias	0.0260	0.0125
RMSE	0.0644	0.0631

6.5.4 Comparisons

BRM and S-EM estimates have been compared from 500 data sets in several configurations : several sample size, deterministic or random right censoring and various censoring rates. Comparison with EM was done with real data in Section 6.6.

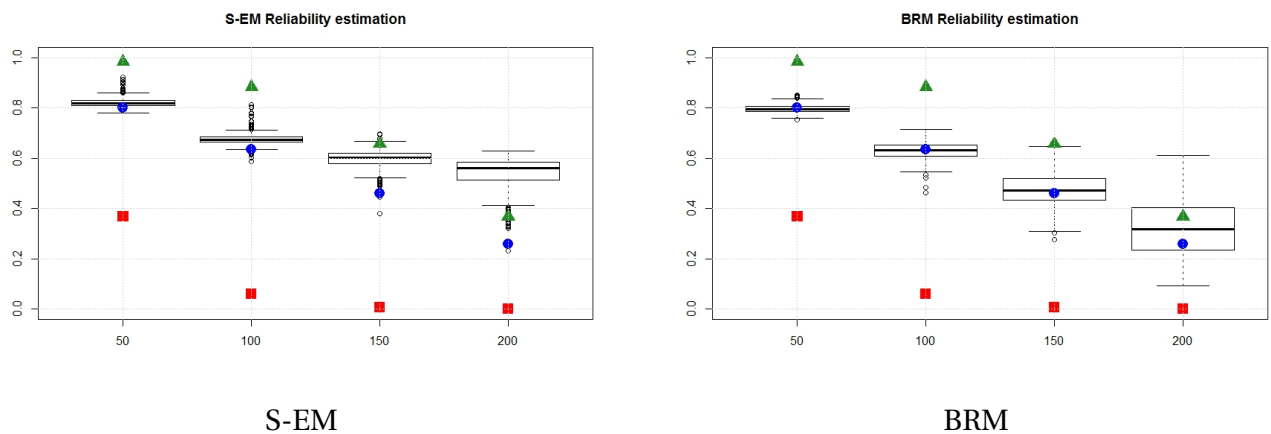
In Table 6.9, it can be seen that S-EM efficiency is deteriorated with censoring rate. Figures 6.2 and 6.3, show for highly censored data, that S-EM efficiency is deteriorated also for large mission times. Whereas, BRM estimation is still efficient, for long-term mission and heavily censored data. These results are comparable with those of Bacha and Celeux (1996) for a single Weibull distribution.

It is worth mentioning that for highly censored data S-EM algorithm may break down. Indeed, for some iterations, no observation are assigned to a component and therefore parameters may not be estimated.

On the other hand, BRM estimation is a non-iterative method which therefore provides shortest computation times than EM and S-EM algorithms.

TABLE 6.9 – S-EM and BRM estimations. Reliability for mid-term ; Average bias and RMSE for 500 data set, with $n = 600$ for weak and high censoring rates.

Average bias (RMSE)	Censoring rate			
	10%	30%	70%	90%
S-EM	0.013 (0.010)	-0.007 (0.011)	0.021 (0.097)	0.273 (0.132)
BRM	0.010 (0.027)	0.011 (0.022)	0.012 (0.033)	0.08 (0.095)

FIGURE 6.2: S-EM and BRM estimations. Reliability for small, mid and long-term ; boxplots for 500 data set, with $n = 600$ and 70% censoring ; green triangle denotes the component 1 reliability ; red box, the component 2 reliability ; blue circle the mixture reliability.

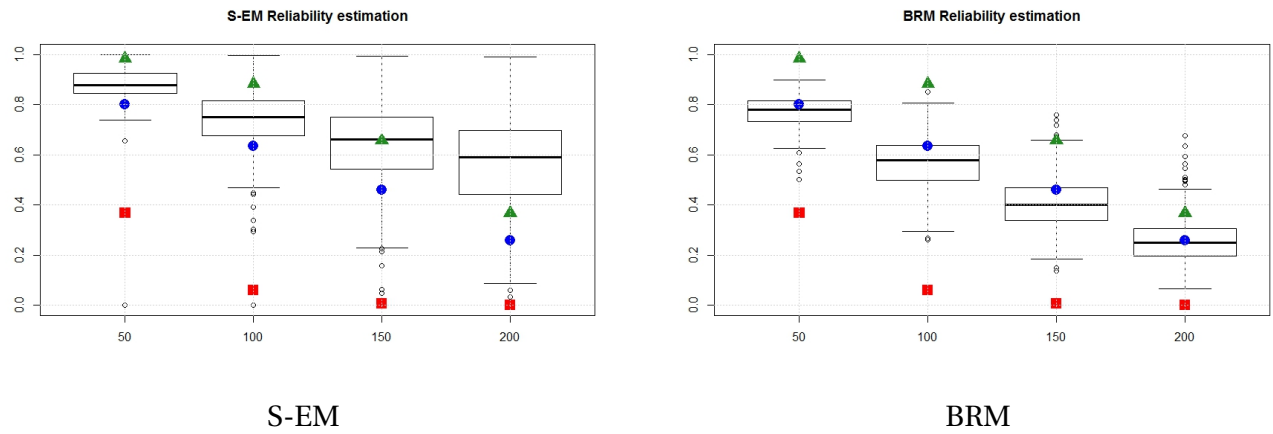


FIGURE 6.3: S-EM and BRM estimations. Reliability for small, mid and long-term ; boxplots for 500 data set, with $n = 600$ and 90% censoring ; green triangle denotes the component 1 reliability ; red box, the component 2 reliability ; blue circle the mixture reliability.

6.5.5 Discussion

As we have seen, ML estimation via EM and S-EM should be used very carefully for fitting a Weibull mixture to highly censored data. EM algorithm may converge to multiple local maxima of the likelihood particularly for Weibull distribution. S-EM algorithm may break down when no data are attributed to a component. By incorporating prior information, BRM estimation allows to overcome the issue of highly censoring. It is a Bayesian bootstrap method : parameters are firstly simulated from joint prior distribution, then based on the weights, a new sample is obtained. As you can see from Figure 6.4, the sample thus obtained is concentrated on the posterior mean.

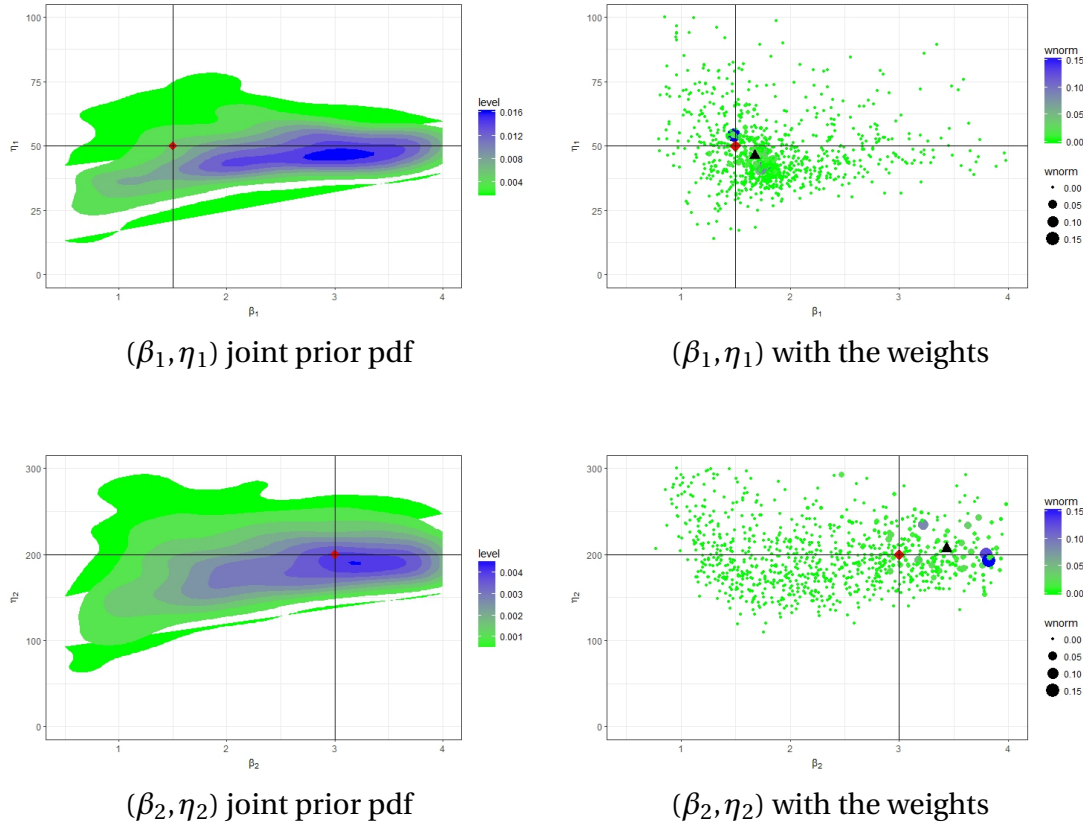


FIGURE 6.4: BRM is a Bayesian bootstrap method Rubin (1981); parameters are simulated from prior, then, from importance sampling weights, a new sample is obtained which is similarly distributed as the posterior distribution; the true values of the parameter are highlighted and marked by a diamond; the BRM estimates are marked by a triangle.

It should be noted that the mixture model is invariant to permutation of the component labels. Hence, the posterior distribution can be symmetric if a weak informative prior is used, that involves bad inference when using Eq. (6.6). This issue is called label switching Celeux (1998). However, we noted that BRM is not affected by label switching. This can be explained by the fact that after the importance resampling step the sample is concentrated on the posterior mean.

It has been seen that BRM estimate have small variance thereby providing accurate reliability estimate for mid and large mission times.

6.6 Real data processing

In this section, two real data sets are introduced to illustrate the implementation of BRM method. The first data set collects lifetimes of throttles. It is a data set with moderate size of 50 lifetimes, 50% of which are censored. The second data set collects lifetimes of power assemblies of high horsepower diesel electric locomotives. It is a data set with large size of 304 lifetimes, 78% of which are censored. To fit a Weibull mixture model to the data, with BRM method, it has been proceeded as follows.

6.6.1 Fitting a Weibull mixture to data

– Step 1 : Graphical test

First and foremost, the Weibull QQ-plot allowed to confirm graphically the presence of a mixture, as shown in Figure 6.1. On the other hand Weibull QQ-plot also allowed to obtain

TABLE 6.10 – Example 1 - Throttles : 50 lifetimes in km which 50% are censored. "-" indicates that lifetime is censored.

478	-484	583	-626	753	753	801	834	-850	944
959	-1071	-1318	1377	-1472	1534	-1579	-1610	-1729	-1792
-1847	2400	-2550	-2568	2639	2944	2981	3392	3392	-3791
3904	-4443	4829	5328	5562	-5900	6122	-6226	6331	6531
-6711	-6835	-6947	-7878	-7884	-10263	11019	12986	-13103	-23245

roughly estimate for $\hat{\alpha}_1^{WQQP}$, $\hat{\eta}_1^{WQQP}$ and $\hat{\eta}_2^{WQQP}$. It will be used to provide initialisation for EM and S-EM algorithms or to calibrate the BRM method.

– Step 2 : Hyperparameters calibration

From Section 6.4.4 and Section 6.5.1, we get the prior distributions

- α_1 : beta prior with $a^{shape} = 1.1$ and

$$[a^{min}; a^{max}] = \begin{cases} [0; 0.5] & \text{if } \hat{\alpha}_1^{WQQP} < 0.5; \\ [0.5; 1] & \text{otherwise.} \end{cases}$$

- β_k for $k = 1, 2$: beta prior with $b^{shape} = 1.1$ and

$$[b^{min}, b^{max}] = [0.5; 10].$$

Expert advices could be used to choose the support $[b^{min}, b^{max}]$, see Barringer (2002) for example.

- η_k for $k = 1, 2$: beta prior with $e^{shape} = 5$ and

$$e_k^{scale} \frac{\Gamma(e^{shape} - \frac{1}{\beta_k})}{\Gamma(e^{shape})} = \hat{\eta}_k^{WQQP}.$$

The estimate $\hat{\eta}_k^{WQQP}$ not need to be accurate.

Step 3 : Estimations

BRM estimates and credibility intervals are computed with $D = 1\,000$ runs, see Section 6.4.3.

For comparison, parameters were estimated by EM algorithm, S-EM algorithm and BRM method.

6.6.2 Example 1 - Throttles

Table 6.10 displays the distance (km) traveled by a vehicle before failure of the throttle, see Jiang and Murthy (1995); Elmahdy (2015). The WQQ-plot indicates that the data can be modeled by a two-component Weibull mixture. Parameters were estimated with 1000 iterations for EM algorithm see Table 6.11 and S-EM algorithm see Table 6.12. From the joint posterior distribution of the parameters, a sample of size 1 000 was generated for BRM method see Table 6.13. The same conclusion is reached from the results of Elmahdy (2015). The S-EM algorithm did better than EM algorithm but may break down : for some iterations, no observation are assigned to a component and therefore parameters may not be estimated.

6.6.3 Example 2 - Power assemblies

Diesel-electric locomotives include power assemblies. Table 6.14 displays power assemblies lifetimes in months Abernethy (2004); Elmahdy (2015). The WQQ-plot indicates that the data can

TABLE 6.11 – Example 1 - EM estimation, with 1 000 iterations. Confidence bounds are obtained by normal approximation from the last iteration.

	confidence interval		EM estimates
	2.5%	97.5%	
α_1	37.60 %	65.31%	51.45%
β_1	0.9765	2.0113	1.3147
η_1	4 817.662	31 561.663	18 189.66
β_2	0.7705	2.8858	1.2162
η_2	2 457.192	5 110.353	3 783.772

TABLE 6.12 – Example 1 - S-EM estimation with 1 000 iterations. Confidence bounds are empirical quantiles of the last three hundred iterations.

	confidence interval		S-EM estimation
	2.5%	97.5%	
α_1	41,32 %	68,89 %	55,10%
β_1	0.9644	1.9504	1.2906
η_1	3 679.258	35 579.027	19 629.14
β_2	0.7459	3.2043	1.2102
η_2	2 618.863	5 416.639	4 017.751

TABLE 6.13 – Example 1 - BRM estimation with 1000 runs. Confidence bounds are empirical quantiles after re-sampling.

	credibility interval		BRM estimation
	2.5%	97.5%	
α_1	83.63 %	88.04%	84.89%
β_1	1.2193	2.0143	1.3818
η_1	9 256.823	9 620.528	9 398.932
β_2	5.8599	8.0653	6.236
η_2	814.4110	863.6094	824.9738

be modeled by a two-component Weibull mixture. Parameters estimated with EM algorithm are in Table 6.15 , with S-EM algorithm in Table 6.16 and with BRM method in Table 6.13. The results show that BRM outperforms EM and S-EM in term of precision.

TABLE 6.14 – Example 2 - Power assemblies : 304 lifetimes in months which 78,3% are censored. "23x4" means 4 failures at time 23, "-" indicates that lifetime is censored.

1×1	3×2	4×1	9×2	-9.5×15	11×1	-12.5×16	17×2	22×1
23×3	31×1	35×1	36×1	37×4	39×3	40×2	44×2	45×1
46×3	47×8	48×2	49×3	50×6	51×2	52×14	-52.5×76	-53.5×131

TABLE 6.15 – Example 2 - EM estimation, with 1 000 iterations. Confidence bounds are obtained by normal approximation from the last iteration.

	confidence interval		EM estimates
	2.5%	97.5%	
α_1	39.60 %	55.59%	47.59%
β_1	0.6483	1.5341	0.9114
η_1	21.5123	704.5085	363.0104
β_2	7.2106	12.825	9.2312
η_2	56.4047	61.7784	59.0915

TABLE 6.16 – Data set 2 - S-EM estimation with 1 000 iterations. Confidence bounds are empirical quantiles of the last three hundred iterations.

	confidence interval		S-EM estimation
	2.5%	97.5%	
α_1	35.02 %	44.394%	38.15%
β_1	0.8360	1.1606	0.9298
η_1	203.0728	635.1525	304.7516
β_2	8.5797	11.7263	9.2992
η_2	59.8753	62.5843	64.3963

TABLE 6.17 – Data set 2 - BRM estimation with 1 000 runs. Confidence bounds are empirical quantiles.

	credibility interval		BRM estimation
	2.5%	97.5%	
α_1	27.37 %	30.36%	28.29%
β_1	1.921	3.0146	2.2273
η_1	505.4810	635.9471	583.2918
β_2	6.7203	7.248	7.0536
η_2	58.2799	60.1887	58.7558

6.7 Conclusion

Lifetime systems or components depend on various internal and external factors. The failure rate could vary according to manufacturing quality, wear intensity, maintenance efficiency or environmental conditions. Therefore, lifetimes data are frequently non-homogeneous. These variations are sometimes very significant and require specific analysis. Otherwise, it could lead to errors and misconceptions in reliability analysis. Unfortunately, the sources of this non-homogeneity are not controlled or not observed. Hence, the Weibull mixture is an appropriate model to handle heterogeneity in lifetimes data collection. Furthermore, in an industrial context the data are heavily censored. In this framework, the performance of classical estimation methods (maximum of likelihood via EM, Bayes via MCMC) is significantly reduced. In this paper, fitting two-component Wei-

bull mixture to highly censored data have been carried out using the Bayesian restoration maximization method (BRM). It is a Bayesian bootstrap method. Firstly, model parameters are sampled from the joint prior distribution. Next, a resampling is made using an importance sampling technique. The empirical distribution, thus obtained, is concentrated on the Bayesian posterior mean. Therefore, BRM provides estimates with a reduced squared error and avoids the label switching problem of MCMC algorithm.

Prior elicitation is the main issue in Bayesian estimation. Here, poor informative prior are used, and prior elicitation is very simple making it operable in various industrial applications. Hyperparameters of the prior distributions were calibrated by simulations. Comparisons with various synthetic data and real data showed that BRM method provides more precise estimates than EM and S-EM algorithms for highly censoring rate, even when components are not well separated. In addition, unlike S-EM or MCMC algorithms, BRM is not an iterative method and can be parallelized. Thus, computation time for the BRM estimation is considerably shorter than for S-EM, or MCMC estimations. To conclude, in an industrial context for which data heterogeneity must be handled, with heavily censored data, the BRM method is a simple, low-cost and quite precise estimation method for reliability or warranty statistical analyses.

Chapitre 7

Sélection du modèle de loi des inter-occurrences de panne

Sommaire

7.1 Sélection avec critère graphique	88
7.1.1 WQQP pour un modèle de loi de Weibull simple	88
7.1.2 Weibull QQP pour un modèle de deux lois de Weibull en concurrence	89
7.1.3 Weibull QQP pour un modèle de mélange de deux lois de Weibull	89
7.2 Sélection avec critères de vraisemblance pénalisés	90
7.2.1 AIC	90
7.2.2 BIC	90
7.3 Conclusion	91

Essentially, all models are wrong, but some are useful. George Box (1919-2013).

Le choix d'un modèle probabiliste pour modéliser la défaillance est une étape essentielle dans l'analyse de fiabilité. Le vieillissement est un facteur important dans la non stabilité du besoin en équipements de rechange. Pour permettre aux ingénieurs du constructeur de mieux appréhender les résultats de nos analyses statistiques, nous avons fait le choix de retenir trois modèles probabilistes basés sur la loi de Weibull :

- un modèle de loi de Weibull simple (SW) :

$$f_{SW}(x) = f_W(x|\beta, \eta) = \frac{\beta}{\eta} \left(\frac{x}{\eta}\right)^{\beta-1} e^{-\left(\frac{x}{\eta}\right)^{\beta}},$$

C'est un modèle à deux paramètres $\theta = (\beta, \eta)$. On suppose ici que l'équipement étudié ne possède qu'une seule cause de défaillance ;

- un modèle de deux lois de Weibull en concurrences (RCW) :

$$F_{RCW}(x) = 1 - (1 - F_W(x|\beta_1, \eta_1))(1 - F_W(x|\beta_2, \eta_2))$$

C'est un modèle à quatre paramètres, $\theta = (\beta_1, \eta_1, \beta_2, \eta_2)$. On suppose ici que pour chaque équipement il y a deux causes de défaillances possibles.

- un modèle de mélange de deux lois de Weibull (MIXW) :

$$f_{MIXW}(x) = \alpha f_W(x|\beta_1, \eta_1) + (1 - \alpha) f_W(x|\beta_2, \eta_2)$$

C'est un modèle à cinq paramètres $\theta = (\alpha, \beta_1, \eta_1, \beta_2, \eta_2)$. On suppose ici que la population de l'équipement étudié est constituée de deux sous-populations ayant chacune sa propre cause de défaillance.

Il nous faut alors sélectionner le modèle ajustant au mieux les données de défaillance du RETEX.

Le plan du chapitre est le suivant. Nous présentons tout d'abord une technique de sélection à l'aide d'un critère graphique, puis deux critères de sélection basés sur la vraisemblance pénalisée par la complexité du modèle.

7.1 Sélection avec critère graphique

Un *quantile-quantile plot* (QQP) est un nuage de points dont les abscisses sont les quantiles empiriques d'un échantillon et les ordonnées les quantiles théoriques correspondant d'un modèle de loi. Dans le cas d'un modèle possédant (éventuellement après transformation) un paramètre de position et un paramètre d'échelle (le modèle est stable par transformation affine), le QQP est aligné. Cela permet alors de valider graphiquement l'ajustement de données au modèle. C'est le cas pour les lois normales (ou log-normal) avec la droite de Henry cf. Crépel (1993). C'est aussi le cas des lois de Weibull après une transformation logarithmique cf. Jiang (2014).

Dans ce qui suit, $(x_1, \dots, x_n)$ est un n -échantillon d'une variable X et $x_{(1)} < x_{(2)} < \dots < x_{(n)}$ les statistiques d'ordre de l'échantillon.

7.1.1 WQQP pour un modèle de loi de Weibull simple

Si $F_X = F_{SW}$, alors on obtient cf. Equation 4.8

- pour des données complètes : les points $\left(\ln(x_{(i)}); \text{qgumbel}\left(\frac{i-0.5}{n}\right)\right)$ sont alignés ;
- pour des données censurées : on estime la fiabilité R par $\hat{R}$ l'estimateur de Kaplan-Meier cf. Meeker and Escobar (1998). Les points $(\ln(x_{(i)}); \text{qgumbel}(1 - \hat{R}(x_{(i)})))$ sont alignés.

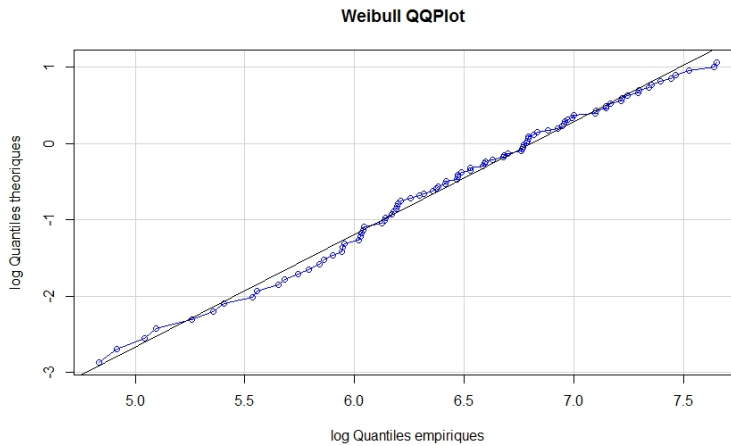


FIGURE 7.1: Weibull QQplot d'une loi de Weibull : nuage aligné

On trouvera dans Murthy et al. (2004b) la forme du Weibull QQP pour différents modèles basés sur la loi de Weibull.

7.1.2 Weibull QQP pour un modèle de deux lois de Weibull en concurrence

Si $F_X = F_{RCW}$ alors on obtient cf. Jiang and Murthy (2003) :

- Le Weibull QQP est convexe ;
- Le Weibull QQP a deux asymptotes : si $\beta_1 < \beta_2$ alors
 - à gauche : $\Delta_1 : y = \beta_1(x - \ln(\eta_1))$
 - à droite : $\Delta_2 : y = \beta_2(x - \ln(\eta_2))$

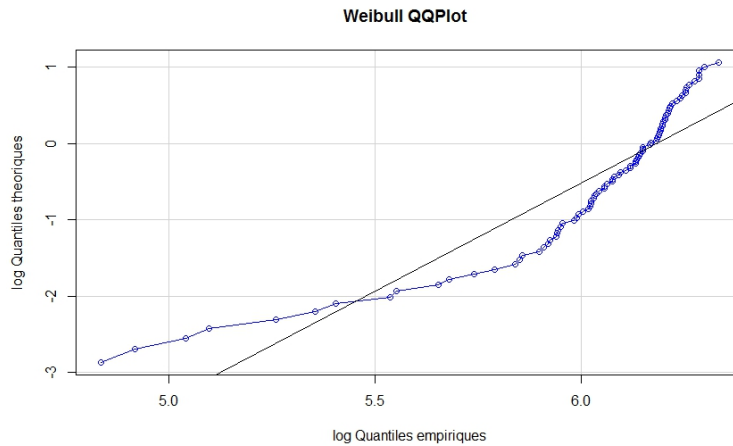


FIGURE 7.2: Weibull QQPplot de deux lois de Weibull en concurrence : courbe convexe.

7.1.3 Weibull QQP pour un modèle de mélange de deux lois de Weibull

Si $F_X = F_{RCW}$ alors on obtient cf. Jiang and Murthy (1995) :

- Le Weibull QQP a une forme en S
- si $\beta_1 < \beta_2$ alors
 - à gauche : $\Delta : y = \beta_1(x - \ln(\eta_1)) + \ln(\alpha_1)$
 - à droite : $\Delta_1 : y = \beta_1(x - \ln(\eta_1))$

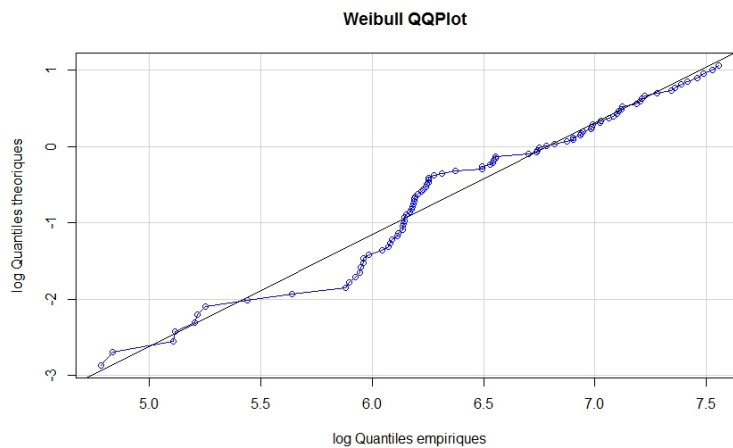


FIGURE 7.3: Weibull QQPplot d'un mélange de deux lois de Weibull : courbe en forme de S.

7.2 Sélection avec critères de vraisemblance pénalisés

La vraisemblance du modèle sur les données est une bonne mesure d'ajustement des données au modèle avec cependant un risque de sur-ajustement : plus le modèle est complexe, plus il va s'ajuster aux données. Les critères de vraisemblance vont donc chercher à éviter le sur-ajustement en pénalisant la vraisemblance par la complexité du modèle. Le choix de la pénalisation va dépendre de l'objectif visé. Le critère AIC (Akaike Information Criterion) introduit par Akaike (1974), retient le modèle qui minimise la distance de Kullback-Leibler au modèle dont sont issues les données. Le critère BIC (Bayesian Information Criterion) introduit par Schwarz (1978) retient le modèle qui maximise la probabilité *a posteriori* connaissant les données.

On note $\mathbf{M} = \{M_1, M_2, M_3\}$ l'ensemble des trois modèles (*i.e.* SW, RCW et MIXW) et $f_M(x|\boldsymbol{\theta})$ la densité du modèle paramétrique M . Par ailleurs on note M^* le modèle dont sont issues les données.

7.2.1 AIC

Le critère AIC retient le modèle M qui minimise asymptotiquement la distance de Kullback-Leibler à M^* :

$$\begin{aligned} KL(M, M^*) &= \mathbb{E}_{M^*} \left[\ln \left(\frac{f_{M^*}(\mathbf{X})}{f_M(\mathbf{X})} \right) \right] \\ &= \mathbb{E}_{M^*} [\ln(f_{M^*}(\mathbf{X}))] - \mathbb{E}_{M^*} [\ln(f_M(\mathbf{X}))]. \end{aligned}$$

Minimiser $KL(M, M^*)$ en M revient alors à minimiser le second terme $-\mathbb{E}_{M^*} [\ln(f_M(\mathbf{X}))]$. Ce dernier est approché en utilisant la log-vraisemblance pour l'estimateur du maximum de vraisemblance $\ln(f_M(\mathbf{X}|\hat{\boldsymbol{\theta}}))$. Akaike (1974) montre alors que le biais asymptotique ($n \rightarrow +\infty$) est égal à la complexité du modèle M , notée $C(M)$. Le critère AIC s'écrit alors :

Critère AIC : parmi les trois modèles WS, RCW et MIXW on retient celui qui minimise

$$AIC(M) = -2 \ln(f_M(\mathbf{X}|\hat{\boldsymbol{\theta}})) + 2C(M) \quad (7.1)$$

où $\hat{\boldsymbol{\theta}}$ est l'estimateur du MV du modèle M et $C(M)$ sa complexité : $C(WS) = 2$, $C(RCW) = 4$ et $C(MIXW) = 5$.

7.2.2 BIC

Le critère BIC retient le modèle M qui maximise la probabilité *a posteriori* connaissant les données :

$$BIC(M) = \max_M \mathbb{P}(M|\mathbf{X})$$

En utilisant la formule de Bayes on obtient :

$$\mathbb{P}(M_i|\mathbf{X}) = \frac{\mathbb{P}(\mathbf{X}|M_i)\mathbb{P}(M_i)}{\mathbb{P}(\mathbf{X})}.$$

Les modèles n'apportent pas d'information, donc on suppose que

$$\mathbb{P}(M_1) = \mathbb{P}(M_2) = \mathbb{P}(M_3).$$

Maximiser la probabilité *a posteriori* revient alors à maximiser :

$$\mathbb{P}(\mathbf{X}|M_i) \quad (7.2)$$

Le calcul de 7.2 est difficile ; en l'écrivant par exemple sous forme intégrale et en négligeant les termes ne dépendant pas de n et donc la loi *a priori* il peut être approché asymptotiquement ($n \rightarrow +\infty$) par une approximation de Laplace. On trouvera dans Lebarbier and Mary-Huard (2006) le détail des calculs.

Critère BIC : pour un échantillon X de taille n , parmi les trois modèles WS, RCW et MIXW on retient celui qui minimise :

$$BIC(M) = -2\ln(f_M(\mathbf{X}|\hat{\boldsymbol{\theta}})) + 2C(M)\ln(n), \quad (7.3)$$

où $\hat{\boldsymbol{\theta}}$ est l'estimateur du MV du modèle M et $C(M)$ sa complexité : $C(WS) = 2$, $C(RCW) = 4$ et $C(MIXW) = 5$.

7.3 Conclusion

Pour conclure

- le critère graphique est très visuel, mais peu robuste ;
- le critère AIC est un critère de type biais-variance : il sélectionne le modèle correspondant au meilleur compromis biais-variance. Il est asymptotiquement efficace, c-à-d qu'il retient le modèle minimisant la racine carrée de l'erreur quadratique moyenne, quand la taille de l'échantillon est très grande, cf. Burnham and Anderson (2003). Il a cependant tendance à sur estimer le nombre de composantes dans un mélange cf. McLachlan and Peel (2000) : choisir plutôt le modèle de mélange de lois de Weibull que le modèle simple de Weibull ;
- le critère BIC est asymptotiquement consistant, c-à-d que la probabilité qu'il retienne le vrai modèle tend vers 1, quand la taille de l'échantillon est grande. De manière générale BIC est plus parcimonieux que AIC.

Notons que nous utilisons l'estimateur BRM en lieu et place de l'estimateur du maximum de vraisemblance. Cependant, puisque AIC et BIC sont des critères asymptotiques (*i.e.* $(n \rightarrow +\infty)$), cette approximation est raisonnable.

Troisième partie

BESOIN EN ÉQUIPEMENTS DE RECHANGE

Chapitre 8

Estimation du besoin en pièces de rechange d'une flotte de véhicules

Sommaire

8.1 Introduction	97
8.1.1 Problem description	98
8.1.2 Notation	99
8.1.3 Problem statement	99
8.2 Simulation algorithm	100
8.3 Equipments failures inter-occurrence	101
8.4 Vehicles usage profiles	102
8.5 Forecasting with various scenarios	104
8.6 Conclusion	107

Dans ce chapitre, nous considérons un contrat de maintenance en conditions opérationnelles (MCO, ou *Maintenance, Repair and Overhaul*, MRO) proposé par un constructeur de véhicules. Il s'agit de maintenir une flotte de plusieurs centaines de véhicules en conditions opérationnelles pour une période supérieure à dix ans, mais sous une contrainte d'usage moyen de la flotte. Le contrat couvre la maintenance d'une liste d'équipements critiques. Le constructeur dispose d'un RETEX constitué d'une part des usages annuels des véhicules et d'autre part des défaillances des équipements.

Dans ce chapitre nous présentons une méthode pour estimer le besoin en pièces de rechange de la flotte sur toute la durée du contrat. Cette partie est rédigée en anglais pour un article en cours de soumission.

Maintenance cost forecasting for a fleet of vehicles

Résumé

Généralement, lors de l'achat d'une flotte de véhicules, le service-après vente inclut un contrat de maintien en conditions opérationnelles (MCO) : ce contrat garantit à l'acheteur la disponibilité de la flotte durant plusieurs années, mais sous des contraintes d'usage. Le contrat MCO est donc un argument commercial majeur pour un constructeur automobile, mais c'est aussi un coût supplémentaire qu'il faut pouvoir évaluer. Et ce d'autant plus que ces contrats vont bien au delà de la durée de vie utile des équipements d'un véhicule. Pour répondre au besoin du client en pièces de rechange et évaluer le coût de la maintenance, le constructeur doit pouvoir prévoir le nombre de défaillances d'une liste d'équipements critiques, sur toute la durée du contrat. Il existe de nombreux travaux traitant de la gestion du besoin en pièces de rechange. Mais dans le cas du contrat MCO sur une flotte, ce problème est particulièrement difficile. Tout d'abord, garantir une disponibilité maximale sur plusieurs années nécessite à la fois des informations sur la fiabilité des équipements, mais aussi des informations sur l'usage des véhicules. Par ailleurs, il existe de nombreux mécanismes d'instabilité du besoin au cours du temps : d'une part deux équipements peuvent avoir différentes causes de vieillissement ou un équipement peut être soumis à deux causes de vieillissement. D'autre part ; les usages des véhicules peuvent varier de manière très significative selon les profils de mission.

L'objectif de cet article est alors d'estimer le besoin en pièces de rechange pour une flotte de véhicules à l'aide d'un modèle de simulation basé sur l'analyse statistiques des inter-occurrences de pannes des équipements et des usages des véhicules. En intégrant ces deux sources d'instabilité au cours du temps du besoin en pièces de rechange, nous obtenons un outil flexible d'aide à la décision permettant aux ingénieurs d'établir ou de gérer un contrat MCO pour une flotte de véhicules.

Abstract

Generally, when purchasing of a fleet of vehicles, after-sales service includes maintenance, repair and overhaul (MRO) contract : this contract guarantees the fleet availability during several years under usage constraints. MRO contract is an undeniable marketing argument for automobile manufacturers, but it is also an additional cost which must be assessed. Particularly as this contracts last far beyond the equipments useful life. To manage the spare parts warehouse and to to assess the costs of maintenance and repairs, manufacturer must be able to forecast the number of failures for a list of critical equipments, throughout the contract. Extensive researches on spare parts inventory management could be found. But in the case of MRO contract on a fleet of vehicles, this issue is particularly difficult. First of all, guaranteed maximum availability during several years requires information on both equipments reliability and vehicles usages. Furthermore, there were many factors of instability of the spare parts requirement over the years : on the one hand equipments may have various causes of ageing either an equipment may be subjected to several causes of ageing ; on the other hand, vehicle usages may vary significantly according to the mission profile.

The purpose of this article is forecast the spare parts requirement for a fleet of vehicles, over several years, using a simulation model based on statistical analysis of equipments inter-occurrences of failure and vehicle usages. By integrating this these two different cost drivers over the years, we provide a flexible tool of decision-making to prepare or manage a MRO contract for a fleet of vehicles.

8.1 Introduction

Products degrade with age and usage and fail when they are unable to perform a required function, under given environmental and operational conditions *cf.* Meeker and Escobar (1998). Therefore, customers need assurance that the product will perform satisfactorily over the useful life period (*i.e.* before ageing). A warranty is then a legal contract which requires that, after the sale, the seller agrees to remedy or compensate the buyer, for certain defects or failures in the product and for a specified time or amount of usage. See Blischke and Murthy (1992) for a review on warranty policies.

For highly engineered products, "custom-built" products or industrial products bought in lots as fleet of vehicles or aircrafts, it can be more economical for the buyer to outsource the maintenance to an external service agent (*e.g.* the manufacturer or an independent specialized company) through an outsourcing maintenance service contract. Furthermore, warranty contract guarantees free repair during the warranty period, but does not guarantee the availability. But in some cases, persistent product performance is a crucial stake of purchase. This is the case for nuclear plants, fleet of aircrafts or military vehicles. Maintenance, repair and overhaul (MRO) contracts provide customers with an outsourcing MRO service, to keep the product in service throughout the duration of the contract, but under usage constraints, see *cf.* Murthy et al. (2015) for an overview on maintenance outsourcing. Recently, due to the rapid technological development and fierce competition linked to the globalization, the manufacturers are under pressure to provide its clients with a maintenance service in their after-sale-servicing. As example, purchases of a fleet of tanks by a government may include ten years of maintenance with a guaranteed high level of availability.

For manufacturers, providing a maintenance service is an undeniable marketing argument, but it is also an additional cost. To set up a MRO contract it is necessary to gather not only information on systems reliability and maintenance but also on their usages (*i.e.* operational environment, operator skill and usage intensity,...). Furthermore, a system such as an automobile consists of many modular systems (*e.g.* powertrain, undercarriage, electrical,...), equipments (subsystems *e.g.* motor, wheels, batteries,...) and thousands of components that are supplied through an extensive suppliers network. Reliability performance of a systems, equipments or components is under control of suppliers based on their design and their quality production. Vehicle usage conditions are under control of drivers, according to operating environment, mission profile and usage intensity. In this context, the statement of a MRO contract for a fleet is especially challenging. Collection and analysis of equipments warranty claims and vehicle usage records can then help in the decision-making.

In this paper, we consider MRO contract proposed by an automobile manufacturer, to maintain a fleet of several hundred vehicles in service throughout the duration of the contract, but under specific usage conditions. The contract covers a list of equipments during more than 10 years. Vehicles usages are also periodically recorded in a vehicles-database. When equipments failed, usage between successive failures are recorded in an equipments-database. It should be noted that MRO contracts issues include both system reliability and maintenance but also usage data analysis (*i.e.* operational environment, operator skill and usage intensity, *etc.*). Extensive work has been done on reliability, maintenance and warranty management, see Blischke et al. (2011) and Wu (2012) for a detailed review. But very few address the issues related to mathematical aspects of MRO contract with the requirement of permanent availability for a fleet. Among these latter include Zaino and Berke (1994), Wang (2010), Anderson et al. (2017) and references therein. Generally, the mean cumulative number of failures is estimated, but considering same usage or same rate of failure occurrences for all vehicles of the fleet. Here, vehicles were deployed over several

years and consequently equipments have various degradation levels. Furthermore, vehicles were used intermittently in highly varied conditions (*e.g.* harsh and sloping terrain,...), for long periods (*e.g.* several dozen hours,...) and various mission profiles (*e.g.* normal or intensive usage).

Therefore, equipment lifetimes and vehicle usages have typically large variance. These various sources of variability are not controlled or not observed. But neglecting existing non-homogeneity can lead to large estimation errors. On the other side, more appropriate mathematical models lead to intractable computations.

Therefore, to forecast the number of failures during the MRO contract, a simulation based model was developed incorporating both equipment failures database and vehicles usage database. Failure data and usage data have been processed specifically. Weibull distributions have been extensively used for modelling lifetime data in medical or industrial applications. In particular, Weibull distributions exhibit decreasing, constant or increasing hazard function which makes them suitable for modelling complex failure data, cf. Murthy et al. (2004b). Three different models of Weibull distributions are used to modelling failure inter-occurrences : either a simple Weibull model or two-fold Weibull competing risk for homogeneous population of equipments with one or two failure causes, either a two-component mixture of Weibull for inhomogeneous population of equipment (each subpopulation has its own causes of failure). Vehicles usages regularly collected provide a significant, but non-homogeneous database. Indeed, according to the mission profile, vehicles can switch from the nominal usage cause of the fleet to an intensive one. Mission profile is not disclosed in the database, and yet it is an essential information for the management of the spare parts. Hence, a cluster analysis allowed us to obtain usage profiles for vehicles.

Simulation based approach is an alternative method to yield maintenance cost forecasting under realistic assumptions on equipment failures and vehicles usages. In addition, our approach provides a flexible decision-making tool before and after execution of the MRO contract for a fleet. The remainder of this paper is structured as follows. The problem is formalized in Section 8.1.1. The simulation based model is described in Section 8.2. Equipments inter-occurrence of failures and vehicles usages profiles modeling are discussed in following sections. This method is illustrated with alternative scenarios to forecast the number of equipments failed during the contract in Section 8.5. Section 8.6 summarizes the conclusions of this paper and discusses future extensions of this research.

8.1.1 Problem description

An equipment installed on a fleet of κ vehicles is considered. These vehicles were deployed over several years. They are used intermittently. A MRO contract is intended to maintain the fleet for t years and covers a list of critical equipments. After each failure, equipments are rapidly replaced by a new one. Free replacements are supplied under constraints on the average usages of the vehicles. Vehicles are equipped with various features differently used according to the purpose of the mission. Usage features are measured with different scales (*e.g.* motor usage in hours, battery usage in hours, kilometers travelled, number of usages of a feature,...). The manufacturer has then selected three scales of measure which will be stored in the vehicles database. Inter-occurrences of failures are measured with a scale specific to each equipment (*e.g.* hours for the motor, kilometers for wheels, etc).

8.1.2 Notation

For $1 \leq k \leq \kappa$, let us denote

$$\begin{aligned}
X_{ki} &= i^e \text{ failure inter-occurrence of an equipment installed on the } k\text{-th vehicle;} \\
U_k(t) &= \text{cumulated usage of the } k\text{-th vehicle, on } t \text{ years} \\
N_k(t) &= \text{number of equipments failed for } k\text{-th vehicle, on } t \text{ years :} \\
&\mathbb{P}(N_k(t) \geq n) = \mathbb{P}(X_{k1} + \dots + X_{kn} \leq U_k(t)) \\
N(t) &= \text{total number of failures for } \kappa \text{ vehicles of the fleet, during } t \text{ years} \\
&= \sum_{k=1}^{\kappa} N_k(t).
\end{aligned}$$

In this paper, the random variables (X_{ki}) are assumed to be independent and identically distributed. A failed equipment is replaced by a new one. Let F_X denote their common distribution.

8.1.3 Problem statement

The decision problem related to MRO contract include both to maximize the fleet availability for the costumer, and minimizing the maintenance cost for the manufacturer. Hence, this study is motivated by forecasting of the spare parts requirement. We'll have to compute

- the stock of spare parts n_{max} , for a required availability p :

$$\mathbb{P}(N(t) \leq n_{max}) = p; \quad (8.1)$$

- the expected number of failures :

$$\mathbb{E}[N(t)] = \mathbb{E}\left[\sum_{k=1}^{\kappa} N_k(t)\right]. \quad (8.2)$$

To solve Eq. (8.1) and Eq. (8.2), the distribution of $N(t)$ has to be computed. A failed equipment is replaced by a new one, thus counting processes $(N_i(t))$ are renewal processes and $N(t)$ is a superimposed renewal processes *cf.* Ascher and Feingold (1984). Generally, explicit analytical expression for $N(t)$ distribution is not available. Therefore to predict the number of equipments failed during the contract, we provided a simulation-based model from vehicles usages database and equipments failure database *cf.* Figure 8.1.

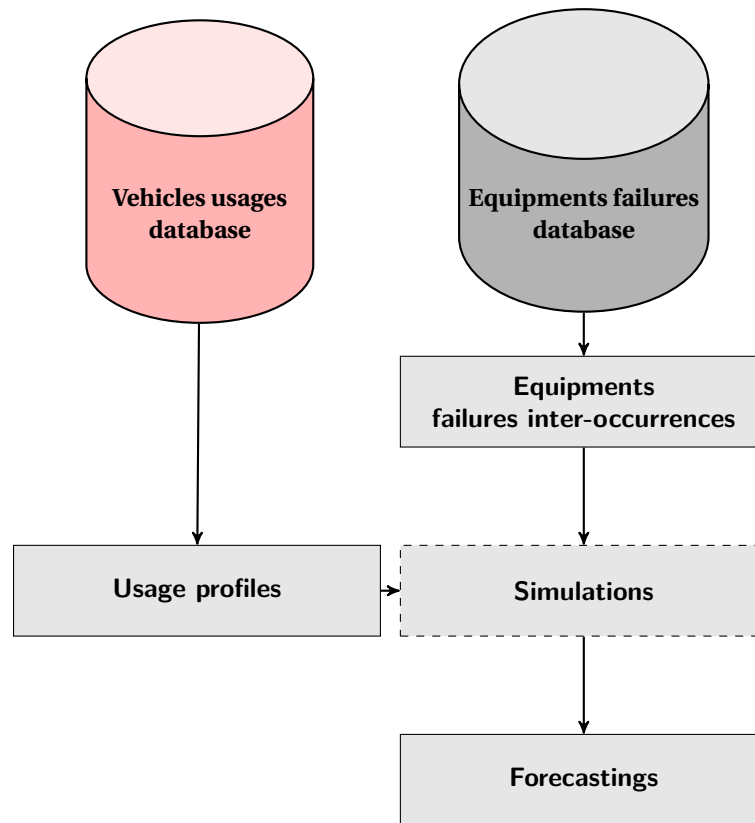


FIGURE 8.1: Forecasting of equipment failures of a fleet of vehicle using simulation based on vehicles usages and equipment failures

8.2 Simulation algorithm

The random nature of the number of equipment failures is related to the occurrence of equipment failures and the vehicle usages during the period *cf.* Algorithm 6.

Algorithm 6 : Failures simulation algorithm

Initialisation : $N \leftarrow 0$
 $U \leftarrow$ vehicle usage during a period
 $X \leftarrow$ equipment usage until failure
while $X < U$ **do**
 $N \leftarrow N + 1$
 $U \leftarrow U - X$
 $X \leftarrow$ equipment usage until failure
end
return N

Model selection for both distribution of equipments inter-occurrence of failures X and vehicles usage U is made respectively from the equipment database and vehicle database. Effective data modelling requires a good understanding of properties of different models and tools and techniques to estimate parameters. In this respect, it is worth mentioning that usage between failures data and vehicle usages data are of a different nature and have to be modelled in a specific manner.

8.3 Equipments failures inter-occurrence

Since more than five decades, Weibull distributions is extensively used for modelling lifetime data in industrial applications. Indeed, Weibull distributions exhibit decreasing, constant or increasing hazard function which makes them suitable for modelling complex failure data cf. Murthy et al. (2004b). Here, vehicles were deployed to the field over several years and used in highly varied conditions (e.g. harsh and sloping terrain,...). Therefore, equipments lifetimes will have various failure rates. These various sources of variability are not controlled or not observed. From the equipments failure database analysis three models were considered.

- A single two-parameter Weibull model (SW) with probability density function :

$$f_X(x) = f_W(x|\beta, \eta) = \frac{\beta}{\eta} \left(\frac{x}{\eta}\right)^{\beta-1} e^{-\left(\frac{x}{\eta}\right)^\beta}, \quad (8.3)$$

where $\beta > 0$ is the shape parameter and $\eta > 0$ the scale parameter. Here, equipment is assumed to have only one cause of failure.

- A two-fold Weibull competing risks model (CRW) with cumulative density function :

$$F_X(x) = 1 - (1 - F_W(x|\beta_1, \eta_1))(1 - F_W(x|\beta_2, \eta_2)) \quad (8.4)$$

where F_W is the cumulative density function of a single Weibull distribution. Here, as for SW model, the population of equipments is homogeneous, but an equipment is subject to two different failure causes.

- A two-component mixture of Weibull distributions model (MW) with probability density function :

$$f_X(x) = \alpha f_W(x|\beta_1, \eta_1) + (1 - \alpha) f_W(x|\beta_2, \eta_2) \quad (8.5)$$

Here, unlike SW model or CRW model, the population of equipments is not-homogeneous. Two main subpopulations with each their own failure cause can be distinguished (e.g. early or random failure, random or wear out failure, two designs or production processes,...).

Weibull quantiles-quantiles plot (WQQP) can provide a quick and simply confirmatory method for data non-homogeneity cf. Murthy et al. (2004a) :

- for single Weibull distribution, the WQQP is a straight line shape ;
- for mixture of Weibull distributions, the WQQP is a single inflection point (S-shaped) with parallel asymptotes ;
- for competing risks of Weibull distinctions, the WQQP is a convex shape.

Furthermore, WQQP provides crude estimates of model parameters, see Murthy et al. (2004b) for computational details. These estimates are non precise nor robust but nonetheless can be used as starting values for alternative estimation algorithms. Different methods can be used for fitting a parametric model to data, but they have substantial difficulties in the setting of heavily censored data. Indeed equipment of critical equipment of vehicle are highly reliable, and consequently censoring rate are usually greater than 70%. With such levels maximum likelihood approaches (ML), Estimation-Maximization method include, provide estimate highly biased particularly for MW and CRW models, cf. Jiang and Kececioglu (1992) and Chauveau (1995). Full Bayesian approach provides alternatives to ML methods. It consists of incorporating prior information on parameters to make inference Hamada et al. (2008). Bayesian method require complex integrations for which Markov chain Monte Carlo algorithms (MCMC) are usually used (Robert and Casella (2010) is a good reference). For highly censored data MCMC methods appear to be deceptive in terms of convergence for MV et CM cf. Bacha and Celeux (1996) and Marin et al. (2005). Bayes method via Markov chain Monte Carlo algorithms is time consuming and convergence is difficult to be

assessed, see Marin et al. (2005). In the context of highly censored data, the Bayesian restoration maximization (BRM) method provide an efficient alternative. It is a Bayesian bootstrap method which provides a sampling from the posterior distribution. Using an importance sampling technique, this sample is focused on the posterior mean. For computational details, see Bacha and Celeux (1996) for CRW model and Ducros and Pamphile (2018) for MW model.

8.4 Vehicles usage profiles

Whether they are personal, commercial or military vehicles are equipped with various features differently used, according to the purpose of the mission. Usage of these features are measured by three different scales (*e.g.* hours, kilometers, number of uses,...). Therefore, usages of each vehicle are periodically stored in the database with three different scales. For confidentiality reasons they are denoted here scale 1, scale 2 and scale 3, see Figure 8.2.

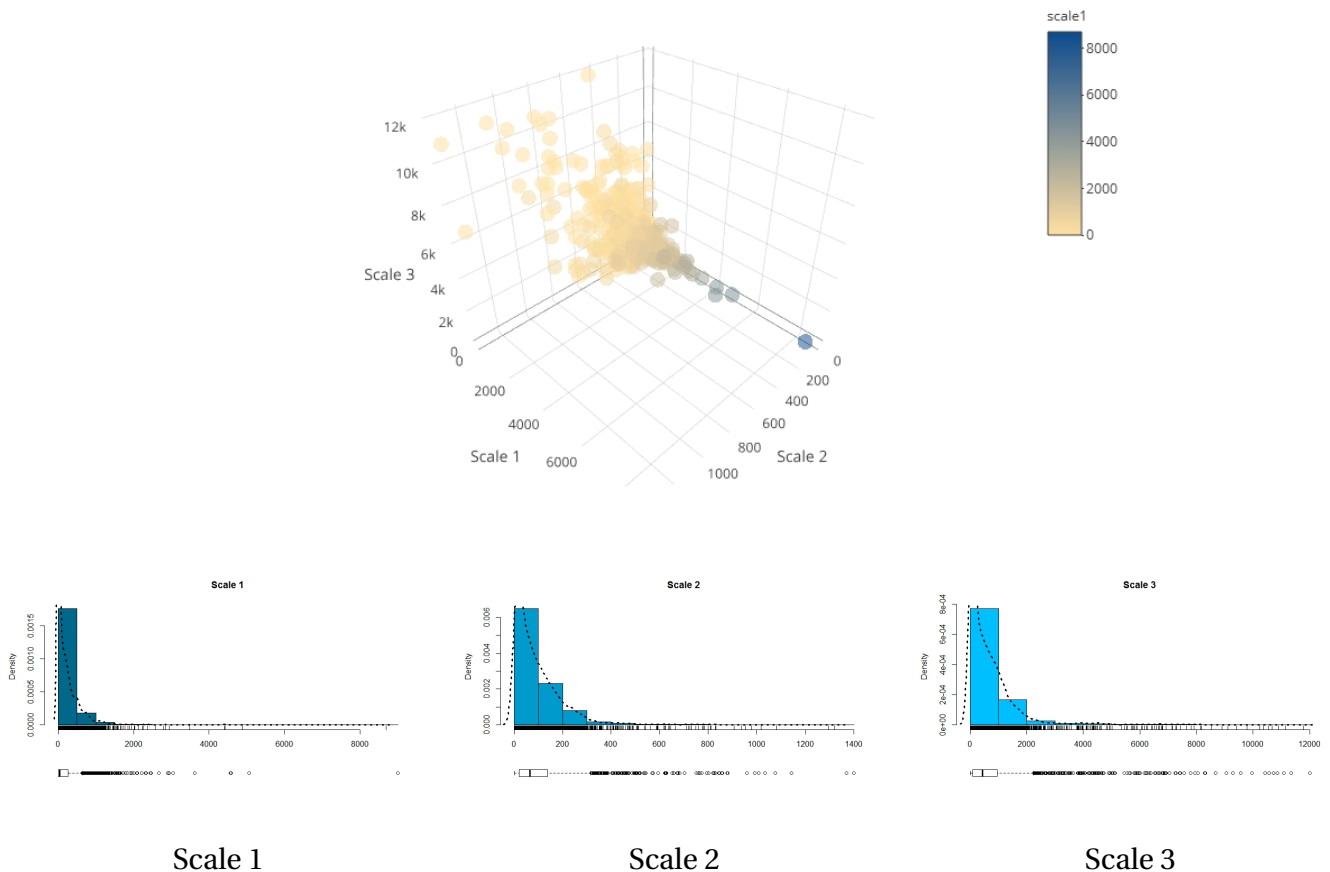
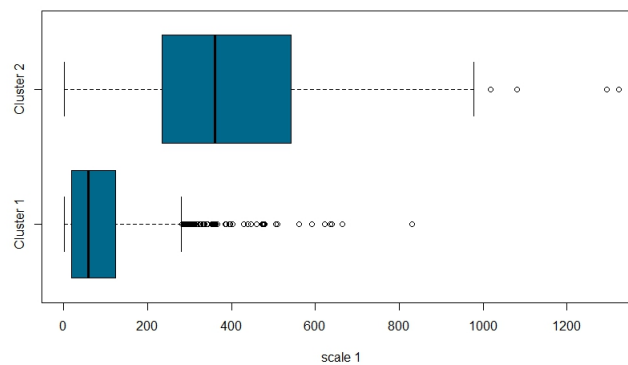


FIGURE 8.2: Scales of vehicles usages

Sometimes, some vehicles can switch from a the nominal usage cause of the fleet to a intensive one according to the mission profile, in case of conflict or intensive training for example. Such types of changes have an impact on the variability of the spare part request and must be integrated in the simulation. Unfortunately, mission profile is not disclosed in the vehicles database. To tackle this problem, a cluster analysis was used on the vehicles usages. The objective of clustering is to partition a population of items in homogeneous subgroups, such as items in the same subgroup are highly « similar », and items different subgroups are as different as possible *cf.* Jain (2010). Two questions then arise : how to measure similarity? how many groups shall be selected? From quantitative measures on items, the k -means algorithm with Euclidean metric as measure of similarity

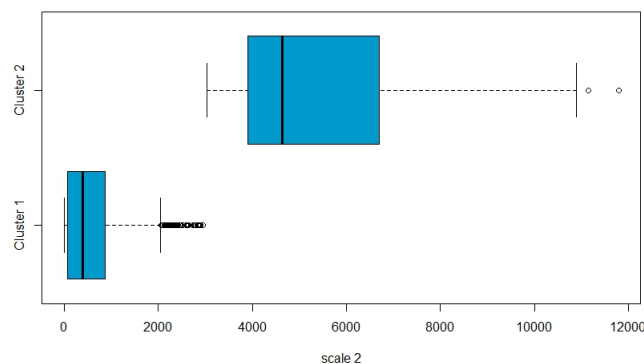
is the simplest clustering algorithm, see Hastie et al. (2009) for computational details. In our case, the aim is to identify vehicles profiles with usage intensities which are either normal or high. To ensure that profiles can be related to actual operating missions, the three usage units were used for clustering. As can be seen in Figures 8.3, 8.4 and 8.5, the two clusters match vehicles mission profiles with specific values of the three scales.

- Cluster 1 : 96,5% of usages, low values with scale 1 and scale 2, but some extreme values for scale 2. This profile is denoted *P1*.
- Cluster 2 : 3,5% of usages, high values with scale 1 and scale 2. This profile is denoted *P2*.



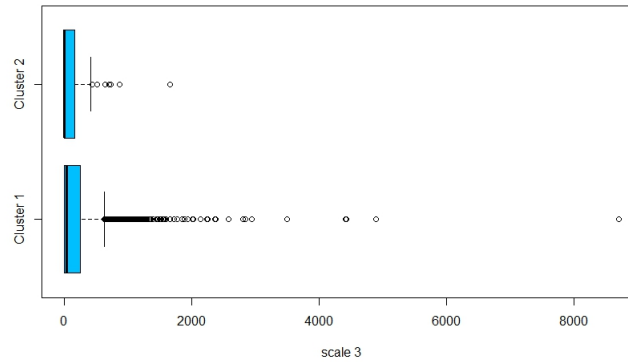
Clusters	%	mean	min	Quantile 25%	Median	Quantile 75%	max
Cluster 1	96.5%	83.9	1	19	59	124	831
Cluster 2	3.5%	425.8	2	235	360	540	1 323
All	100%	95.8	1	20	62	132	1 323

FIGURE 8.3: Vehicle usage Scale 1 by clusters



Clusters	%	mean	min	Quantile 25%	Median	Quantile 75%	max
Cluster 1	96.5%	558.8	0	78	393	871	2 957
Cluster 2	3.5%	5 473.7	3 040	3 903	4 634	6 693	11 789
All	100%	729.8	0	84	420	937	11 789

FIGURE 8.4: Vehicle usage Scale 2 by clusters



Cluster	%	mean	min	Quantile 25%	Median	Quantile 75%	max
Cluster 1	96.5%	189.3	0	3	9	252	8 692
Cluster 2	3.5%	115.1	0	0	3.5	169	1 655
All	100%	186.7	0	0	36	247	8 692

FIGURE 8.5: Vehicle usage Scale 3 by clusters

8.5 Forecasting with various scenarios

The simulation based model provides to forecasting the failure number of equipment during the contract, which can be applied to determine spare parts ordering quantity. It is then a flexible decision-making tool which can be used for before and after execution of the MRO contract for a fleet. As an example, from several scenarios with different equipment failures and various vehicles usages profiles, a projection of the number of spare parts have been made for 20 years.

Average usage per vehicle

During a year, failed equipments are free replace if the average usage per vehicle is below a threshold specified in the contract. Figure 8.6 shows that this threshold is not exceeded with a nominal usage, but it is however reached even with a small percentage of intensive usage. Therefore, a policy of use and management of the fleet by parks of vehicles should allow to optimize the cost : using specific usage profile by parks of vehicles

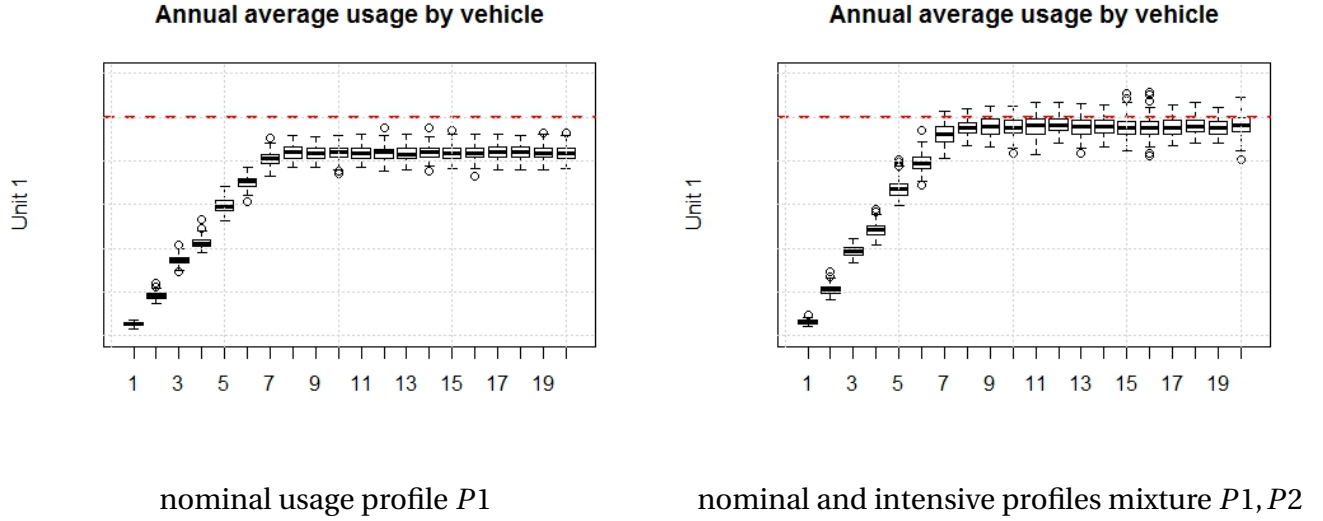


FIGURE 8.6: Annual average usage by vehicle

It should be noticed that the first eight years, correspond to the duration of the fleet deployment. All vehicles entering into service from the tenth year.

Homogeneous population of equipments with various usage profiles

The vehicle usage is a factor of instability of the spare parts requirement over the years. From an homogeneous population of equipments with only one cause of failure, various scenarios of the vehicle usage have been investigated. Figure 8.7 shows the forecasting over 20 years. The exponential behaviour of the curves is due to the deployment of the vehicles over the first nine years.

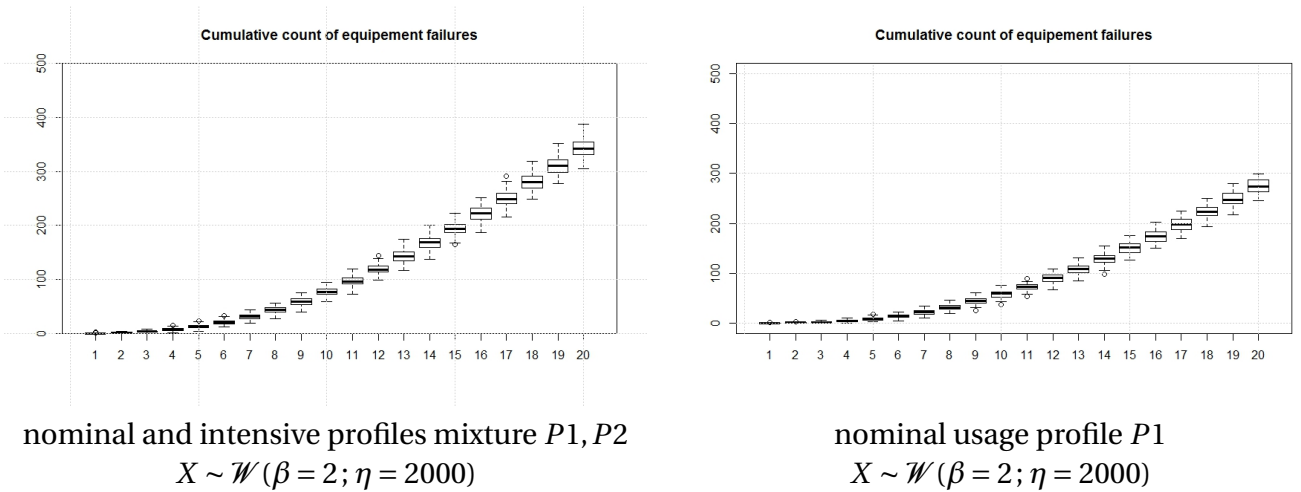
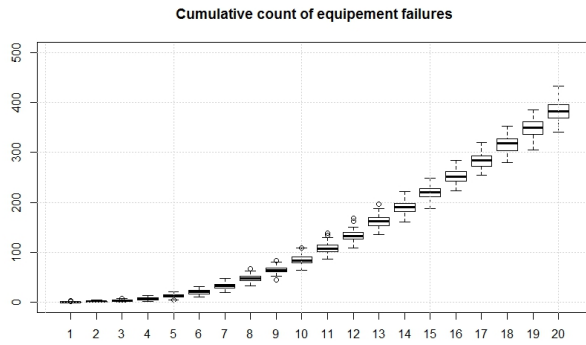


FIGURE 8.7: Cumulative count of failures with simple Weibull

Non-homogeneous population of equipments with various usage profiles

The heterogeneity of equipments ageing is a factor of instability of the spare parts requirement over the years. Various scenarios of mixing are investigated. Figures 8.8- 8.10 show forecasting over 20 years with various vehicles usage profiles.



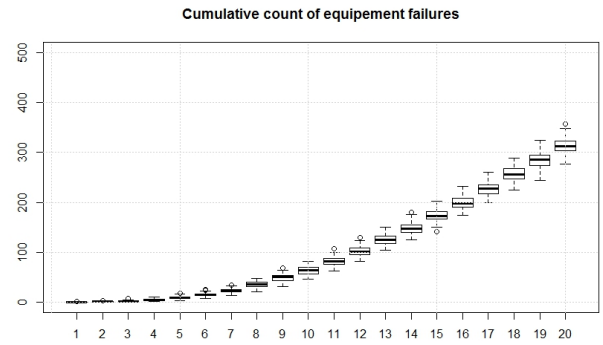
(a.1)

nominal and intensive profiles mixture $P1, P2$

$X \sim \text{Mix Weibull}$

$\alpha_1 = 0.1, \mathcal{W}(\beta_1 = 4, \eta_1 = 1000)$

$\alpha_2 = 0.9, \mathcal{W}(\beta_2 = 2, \eta_2 = 2000)$



(a.2)

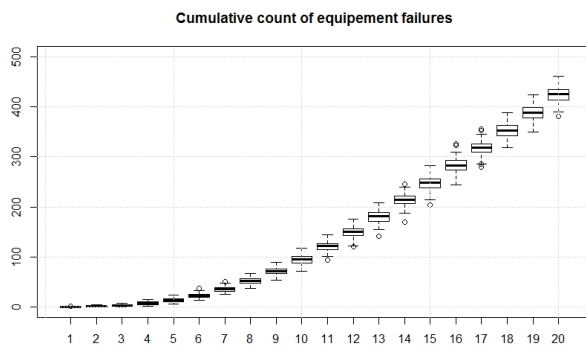
nominal usage profile $P1$

$X \sim \text{Mix Weibull}$

$\alpha_1 = 0.1, \mathcal{W}(\beta_1 = 4, \eta_1 = 1000)$

$\alpha_2 = 0.9, \mathcal{W}(\beta_2 = 2, \eta_2 = 2000)$

FIGURE 8.8: Cumulative count of failures with various mixture of Weibull



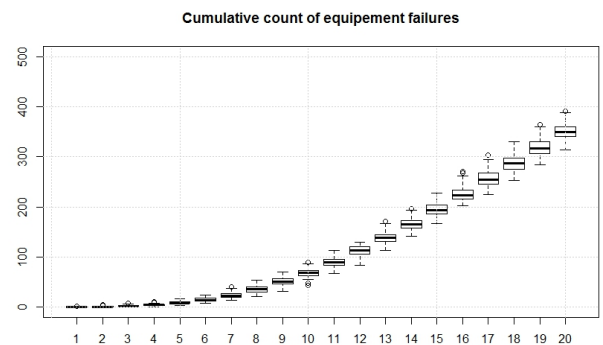
(b.1)

nominal and intensive profiles mixture $P1, P2$

$X \sim \text{Mix Weibull}$

$\alpha_1 = 0.2, \mathcal{W}(\beta_1 = 4, \eta_1 = 1000)$

$\alpha_2 = 0.8, \mathcal{W}(\beta_2 = 2, \eta_2 = 2000)$



(b.2)

nominal usage profile $P1$

$X \sim \text{Mix Weibull}$

$\alpha_1 = 0.2, \mathcal{W}(\beta_1 = 4, \eta_1 = 1000)$

$\alpha_2 = 0.8, \mathcal{W}(\beta_2 = 2, \eta_2 = 2000)$

FIGURE 8.9: Cumulative count of failures with various mixture of Weibull

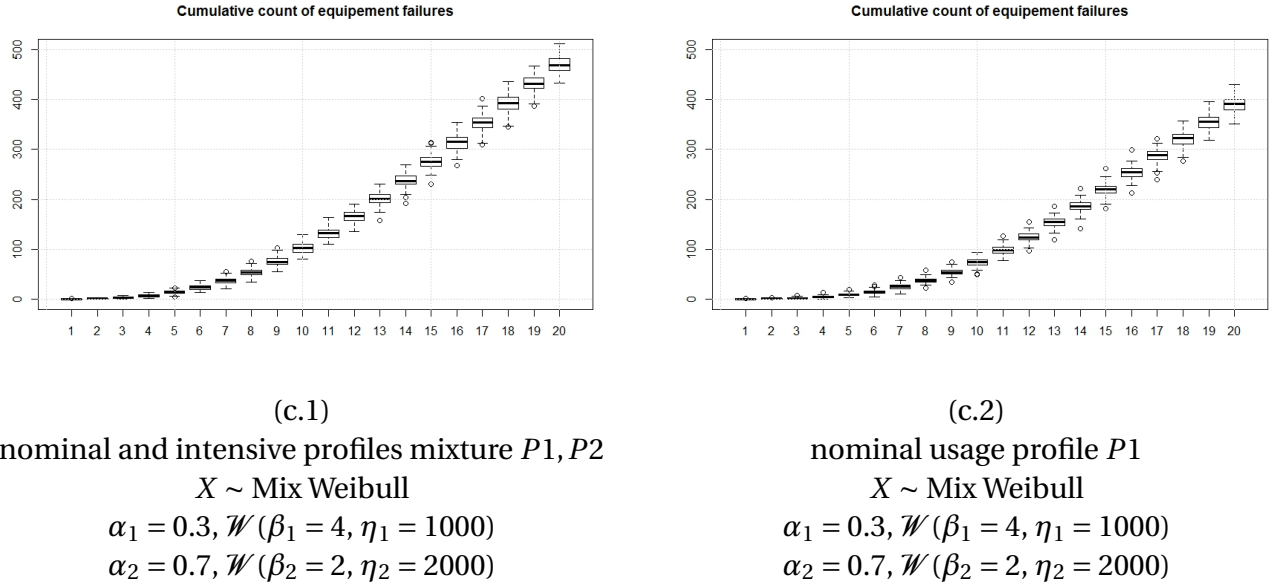


FIGURE 8.10: Cumulative count of failures with various mixture of Weibull

8.6 Conclusion

In this paper, we consider MRO contract proposed by an automobile manufacturer to maintain a fleet of several hundred vehicles. Throughout the duration of the contract, a maximum availability of the fleet is guaranteed. The contract covers repairs of a list of critical equipments during more than ten years. In this context the spare parts inventory management is particularly difficult due to several factors of instability over the time. Particularly, the heterogeneity of equipments ageing and heterogeneity of vehicles usage are two main cost drivers. To tackle this issue, a simulation model was used, based on statistical analysis of equipments inter-occurrences of failures and vehicle usage. Failure data and usage data were processed specifically. For modeling various ageing, three Weibull models was used. A simple Weibull for a homogeneous population of equipments with only one cause of failure, a two-fold Weibull competing risk when equipment have two causes of failure, and finally a two-component mixture of Weibull for heterogeneous population of equipments. Vehicles usage regularly collected provide a meaning, but non-homogeneous database. Hence, a cluster analysis allowed us to obtain usage profiles for vehicles.

Simulation based approach yield spare part forecasting under realistic assumptions on equipment failures and vehicles usage. By integrating these two different cost drivers over the years, we provide a flexible tool of decision-making to prepare or manage a MRO contract for a fleet of vehicles.

Therefore, a simulation based model was developed incorporating both equipment failures database and vehicles usage database.

Chapitre 9

Équipements fils rouges

Sommaire

9.1	Contrainte d'usage par véhicule de la flotte et profils d'usage	110
9.2	Prévisions pour l'équipement 1	111
9.3	Prévisions pour l'équipement 2	113
9.4	Prévisions pour l'équipement 3	115
9.5	Conclusion	117

Ce chapitre aborde la mise en œuvre de notre méthode sur les équipements fil-rouge extraits du RETEX :

1. Choix d'une échelle de mesure pour les inter-occurrences de panne : le choix est « naturel » par exemple les kilomètres pour les roues, les heures pour une batterie, etc. Dans le cas contraire, on peut utiliser la procédure de choix d'une échelle de mesure présentée dans § 2.4.
2. Sélection du modèle de loi pour les inter-occurrences de panne : Weibull simple, lois de Weibull en concurrence, ou mélange de lois de Weibull. Les critères de selection retenus sont :
 - l'analyse du Weibull quantile-quantile plot ;
 - le critère AIC ;
 - le critère BIC.
3. Simulation du besoin en pièces de rechange sur 20 ans.

Dans une première section nous présenterons les simulations d'usage des véhicules de la flotte. Cette analyse permet d'évaluer par exemple la sur-utilisation de la flotte de véhicules par le client, et par conséquent de la prise en charge, ou non, de tout ou partie des remises en service. Puis nous mettrons en œuvre la méthode de prévision des flux sur chacun des équipements *fil rouge*.

9.1 Contrainte d'usage par véhicule de la flotte et profils d'usage

Le remplacement des équipements défaillants est pris en charge sous une contrainte sur l'usage annuel par véhicule de la flotte. Cette contrainte est fixée contractuellement par le client et l'industriel.

Le déploiement des véhicules est fait sur les huit premières années, cf. Figure 9.1.

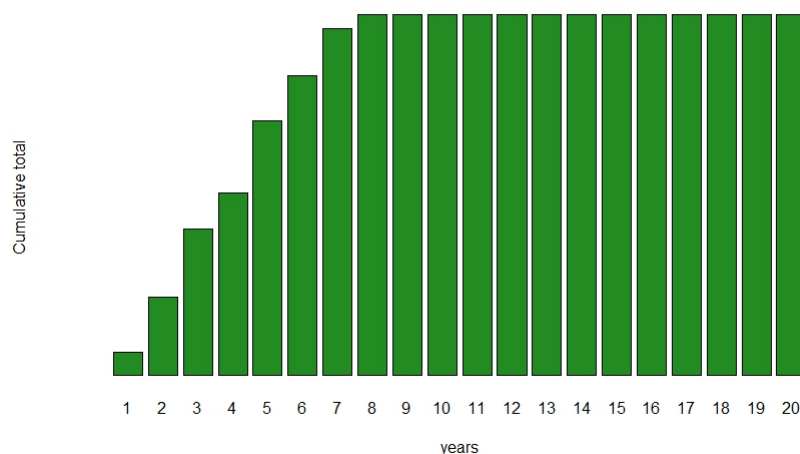
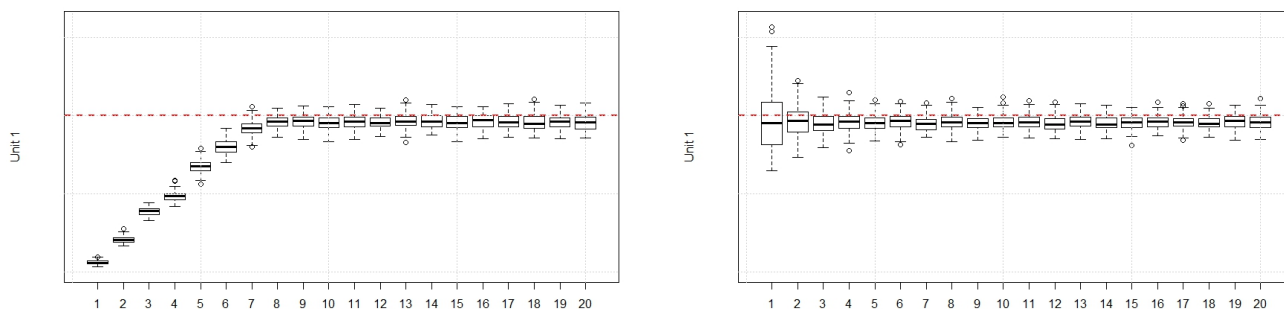


FIGURE 9.1: Déploiement des véhicules pendant le contrat.

Afin de respecter la contrainte d'usage par véhicule, les véhicules subissent une rotation : certains ne sont pas utilisés dans l'année.



(a) usage moyen sur tous les véhicules

(b) usage moyen sur les véhicules en service

FIGURE 9.2: Usages annuels par véhicule de la flotte sur 20 ans, mesurés avec l'échelle 1. La contrainte contractuelle est indiquée par la ligne horizontale pointillée rouge : a) sur toute la flotte, b) sur les véhicules en service

9.2 Prévisions pour l'équipement 1

Données

L'équipement 1 est un équipement mécanique. Sur 599 équipements mis en service, 17% ont été défectueux (*i.e.* 83% de censures).

Échelle de mesure des inter-occurrences de panne

Parmi les trois échelles de mesure, à l'aide des méthodes proposées §2.2 c'est le potentiel 1 (Scale 1) qui a été retenu.

Sélection de modèle

Pour choisir entre le modèle de Weibull simple, le modèle de lois de Weibull en concurrence, ou le modèle de mélange de lois de Weibull, on utilise le WQQP, cf. Figure 9.3 et les critères AIC et BIC, cf. Tableau 9.1. C'est le modèle de mélange de lois de Weibull qui a été sélectionné. On constate qu'il y a une population nominale (*i.e.* $\alpha_2 = 0.9$, $\beta_2 = 2.7$ et $\eta_2 = 1700$) et une population plus rare avec un MTBF 5 à 6 fois plus faible (*i.e.* $\alpha_1 = 0.1$, $\beta_1 = 1.6$ et $\eta_1 = 300$).

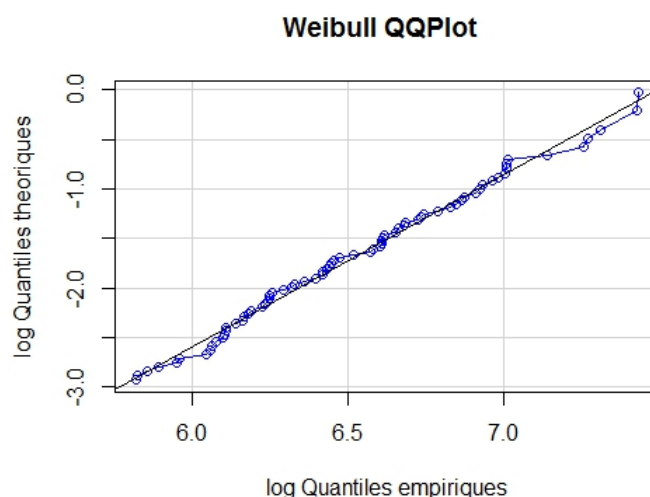


FIGURE 9.3: Équipement 1 : Weibull quantile-quantile plot des inter-occurrences de panne avec l'échelle 1. $n = 599$ et 83% de censures.

Équipement 1	Weibull Simple $\beta = 2, \eta = 1867$	Mélange $\alpha_1 = 0.1, \beta_1 = 1.6, \eta_1 = 300$ $\beta_2 = 2.7, \eta_2 = 1700$	Risques concurrents $\beta_1 = 1.9, \eta_1 = 705$ $\eta_2 = 2.4, \eta_2 = 1364$
Test graphique	?	?	Non
BIC	942	768	1021
AIC	933	746	1043

TABLE 9.1 – Équipement 1 : critères de sélection du modèle de loi pour les inter-occurrences de panne.

Estimation du besoin en pièces de rechange

À partir de la **loi d'inter-occurrence** de panne et **des profils d'utilisation** annuel des véhicules, on peut simuler le besoin en pièces de rechange sur 20 ans, voir Figure 9.4

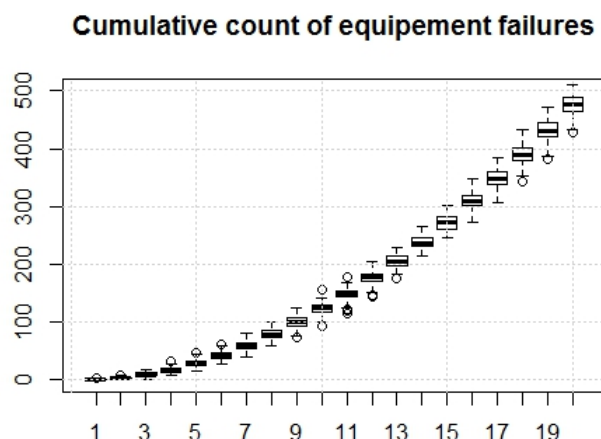


FIGURE 9.4: Équipement 1 : besoin en pièces de rechange sur 20 ans, avec un modèle de mélange de lois de Weibull pour les inter-occurrences de pannes et un profil mixte d'usages des véhicules.

Perspectives

Une étude des causes de défaillance peut permettre d'améliorer la fiabilité. En particulier, une action préventive visant à supprimer la cause 1 (*i.e.* $\alpha_1 = 0.1$, $\beta_1 = 1.6$, et $\eta_1 = 300$) peut permettre de réduire le besoin en pièces de rechange. Le Tableau 9.2 compare les prévisions des besoins à court, moyen et long termes avec correction du défaut dû à la cause de défaillance 1, et sans correction. On peut constater le gain offert par la correction de la cause de la défaillance 1.

Modèle de loi pour inter-occurrences	Période de prévision (en année)			
	5	10	15	20
a) $\mathcal{W}(\beta_2 = 2.7; \eta_2 = 1700)$ besoin <i>mean</i> IC[95%]	8 [4;14]	67 [53;83]	201 [177;230]	390 [359;416]
b) MIXW besoin <i>mean</i> IC[95%]	28 [19;39]	125 [100;148]	275 [247;303]	477 [446;508]

TABLE 9.2 – Équipement 1 : comparaison des prévisions du besoin en pièces de rechange, pour a) correction du défaut dû à la cause 1 et b) mélange sans correction. Les prévisions sont à court terme, 5 ans, moyen terme, 10 ans et long terme 15 ans et 20 ans, et données en moyenne et avec un intervalle de confiance à 95% obtenu par bootstrap.

9.3 Prévisions pour l'équipement 2

Données

L'équipement 2 est un équipement électronique. Sur 561 mis en service, 36% des équipements ont été défaillants (*i.e.* 64% de censures).

Échelle de mesure des inter-occurrences de panne

Parmi les trois échelles de mesure, à l'aide des méthodes proposées §2.2 c'est le potentiel 1 (Scale 1) qui a été retenu.

Sélection de modèle

Pour choisir entre le modèle de Weibull simple, le modèle de lois de Weibull en concurrence, ou le modèle de mélange de lois de Weibull, on utilise le WQQP, cf. Figure 9.5 et les critères AIC et BIC, cf. Tableau 9.3. C'est le modèle de loi de Weibull simple qui a été sélectionné. Les équipements sont homogènes avec une seule cause de défaillance.

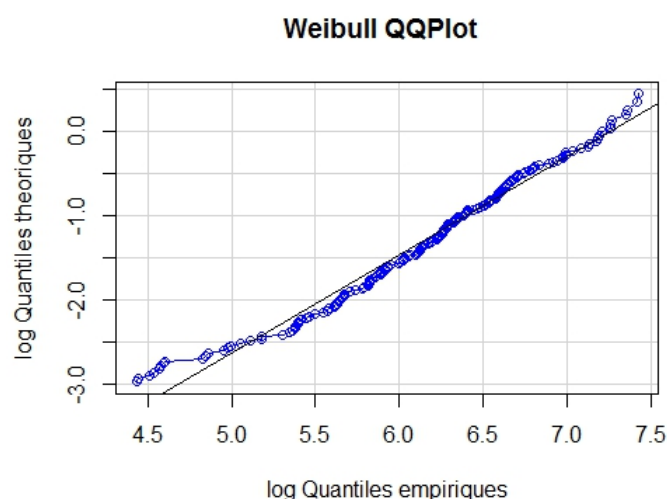


FIGURE 9.5: Équipement 1 : Weibull quantile-quantile plot des inter-occurrences de panne avec l'échelle 1. $n = 561$ avec 64% de censures

Équipement 2	Weibull Simple $\beta = 1.4, \eta = 1300$	Mélange $\alpha_1 = 0.38, \beta_1 = 0.6, \eta_1 = 2080$ $\beta_2 = 1.42, \eta_2 = 1950$	Risques concurrents $\beta_1 = 1.7, \eta_1 = 702$ $\beta_2 = 1.8, \eta_2 = 1223$
Test graphique	?	?	Non
BIC	1277	1350	1317
AIC	1268	1328	1296

TABLE 9.3 – Équipement 2 : critères de sélection du modèle de loi pour les inter-occurrences de panne.

Estimation du besoin en pièces de rechange

À partir de la **loi d'inter-occurrence** de panne et **des profils d'utilisation** annuel des véhicules, on peut simuler le besoin en pièces de rechange sur 20 ans, voir Figure 9.6

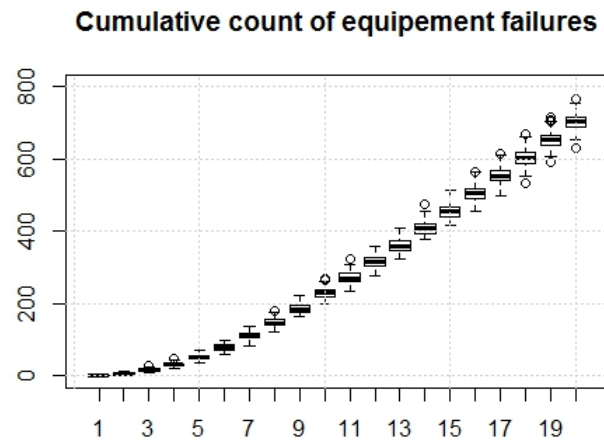


FIGURE 9.6: Équipement 2 : besoin en pièces de rechange sur 20 ans, avec un modèle de Weibull simple pour les inter-occurrences de pannes et un profil mixte d'usage des véhicules

9.4 Prévisions pour l'équipement 3

Données

L'équipement 3 est un équipement mécanique. Sur 587 mis en service, 10% ont été défectueux (*i.e.* 90% de censures).

Échelle de mesure des inter-occurrences de panne

Parmi les trois échelles de mesure, à l'aide des méthodes proposées §2.2 c'est le potentiel 1 (Scale 1) qui a été retenu.

Sélection de modèle

Pour choisir entre le modèle de Weibull simple, le modèle de lois de Weibull en concurrence, ou le modèle de mélange de lois de Weibull, on utilise le WQQP, cf. Figure 9.7 et les critères AIC et BIC, cf. Tableau 9.4. C'est le modèle de lois de Weibull en concurrence qui a été sélectionné. Chaque composant est soumis à deux causes de défaillance en concurrence, dont l'un à un MTBF 2,4 fois plus court.

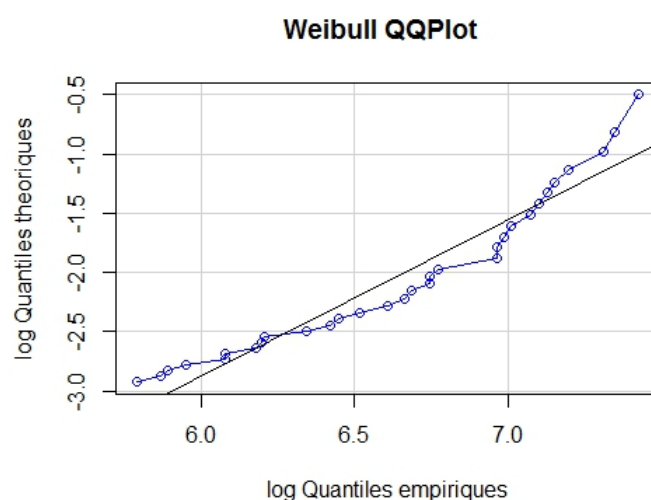


FIGURE 9.7: Équipement 3 : Weibull quantile-quantile plot des inter-occurrences de panne avec l'échelle 2. $n = 587$ avec 90% de censures.

Équipement 3	Weibull Simple $\beta = 2.5, \eta = 1457$	Mélange $\alpha_1 = 0.2, \beta_1 = 1.3, \eta_1 = 780$ $\beta_2 = 2.34, \eta_2 = 2500$	Risques concurrents $\beta_1 = 1.9, \eta_1 = 761$ $\beta_2 = 2.6, \eta_2 = 1892$
Test graphique	Non	Non	Oui
BIC	1237	1004	696
AIC	1228	982	689

TABLE 9.4 – Équipement 3 : critères de sélection du modèle de loi pour les inter-occurrences de panne.

Estimation du besoin en pièces de rechange

À partir de la **loi d'inter-occurrence** de panne et **des profils d'utilisation** annuel des véhicules, on peut simuler le besoin en pièces de rechange sur 20 ans, voir Figure 9.8.

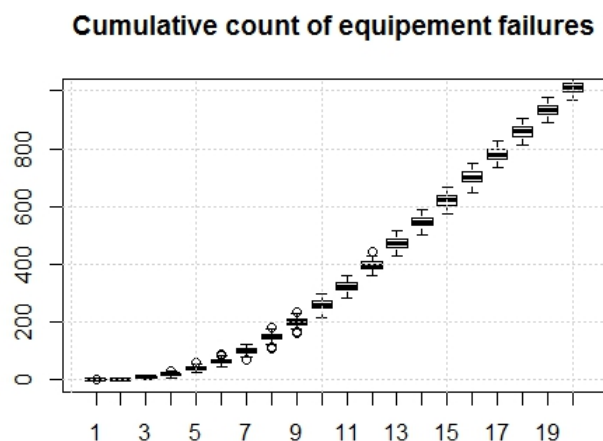


FIGURE 9.8: Équipement 3 : besoin en pièces de rechange sur 20 ans, avec un modèle de lois de Weibull en concurrence pour les inter-occurrences de pannes et un profil mixte d'usage des véhicules

Perspectives

De plus, une étude des causes physiques des deux causes de défaillance peut permettre d'améliorer la fiabilité de l'équipement.

9.5 Conclusion

La méthodologie proposée permet d'estimer au plus juste le besoin en équipements de rechange. Prenons le cas de l'Équipement 1, pour lequel l'étude montre que c'est un modèle de mélange qui s'ajuste le mieux aux données de défaillance. Un choix naïf aurait été de prendre une simple loi de Weibull. La Figure 9.9 et le Tableau 9.5 permettent de comparer les prévisions de besoins à court, moyen et long termes, pour le modèle naïf (*i.e.* la loi de Weibull simple) et le modèle que nous avons retenu (*i.e.* le mélange de lois de Weibull). On peut constater, que le modèle simple sous-estime le besoin, ce qui compte tenu des coûts liés à l'indisponibilité, est peut être très pénalisant pour l'industriel. Sans compter les risques d'obsolescence de l'équipement sur 20 ans.

Modèle de loi pour inter-occurrence	Période de prévision (année)			
	5	10	15	20
SW besoin <i>mean IC</i> [95%]	10 [8;12]	64 [58;70]	166 [158;174]	301 [292;312]
MIXW besoin <i>mean IC</i> [95%]	28 [19;32]	125 [100;148]	275 [247;303]	477 [446;508]

TABLE 9.5 – Comparaison des prévisions du besoin en pièces de rechange, pour deux modèles de lois pour les inter-occurrences de panne : SW) modèle naïf de loi de Weibull simple et MIXW) Modèle de mélanges de lois de Weibull retenu par sélection de modèles. Les prévisions sont à court termes 5 ans, moyen termes 10 ans, long termes 15 ans et 20 ans. Moyenne des défaillances, intervalle de confiance à 95% obtenu par bootstrap.

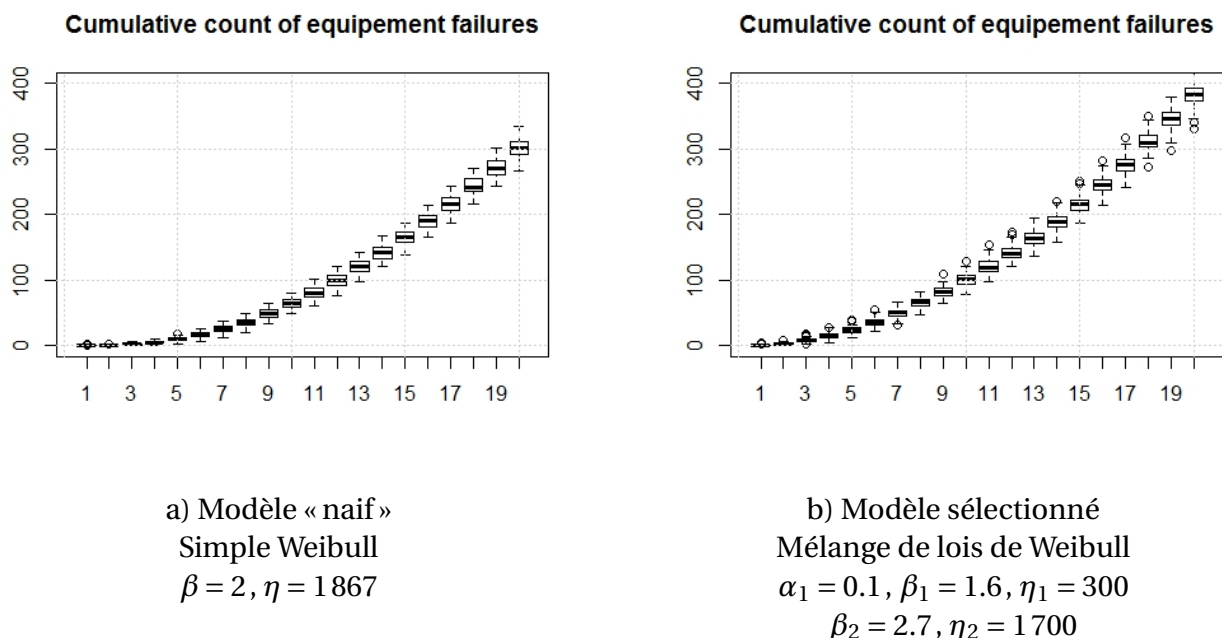


FIGURE 9.9: Comparaison des prévisions du besoin en pièces de rechange, pour deux modèles de lois pour les inter-occurrences de panne : a) modèle naïf de loi de Weibull simple et b) Modèle de mélange retenu par sélection de modèle. Les prévisions sont à court terme 5 ans, moyen terme 10 ans, long terme 15 ans et 20 ans, et données en moyenne et avec un intervalle de confiance à 95% obtenu par bootstrap.

Chapitre 10

Estimation du flux pour un équipement sans RETEX

Sommaire

10.1 Contexte	119
10.2 RETEX historique des équipements	120
10.3 Classifications d'équipements analogues	120
10.4 Estimation du flux de rechange avec les catégories d'équipements	123
10.5 Conclusion	123

10.1 Contexte

Dans le cadre d'une réponse à un appel d'offres, l'industriel peut être amené à effectuer des prévisions de soutien pour des équipements ayant un RETEX peu ou pas renseigné. Dans ce cas, il a recours aux informations de fiabilité des équipements transmis par les fournisseurs ou issues de guides (*e.g.* FIDES (2009), pour les composants électroniques, *military handbook*, MIL-HDBK). Cependant, ces informations sont souvent très optimistes et peu adaptées. Leur utilisation peut engendrer de nombreux risques d'indisponibilités des équipements lourdes de conséquences tant au point de vue humain (dus à l'indisponibilité du véhicule), que financier (due aux pénalités financières). La base de données des défaillances des équipements est alors une source d'information pertinente pour ce problème. Nous proposons dans ce chapitre, d'identifier « des catégories » d'équipements en termes de vieillissement à court et long termes. Ces analogies sont alors un outil d'aide à la décision important pour l'industriel. En particulier, elles permettent de faire des prévisions pour des équipements ayant peu ou pas de RETEX.

Pour construire ces catégories nous proposons d'utiliser une classification de l'ensemble des équipements de la base que nous disposons, soit 45 typologies équipements.

10.2 RETEX historique des équipements

L'historique comporte 45 équipements pour lesquels on dispose des couples de paramètres $(\hat{\beta}, \hat{\eta})$ préalablement estimés à partir de la méthode de BRM-R cf. Partie 2. Les données sont résumées dans la Figure 10.2.

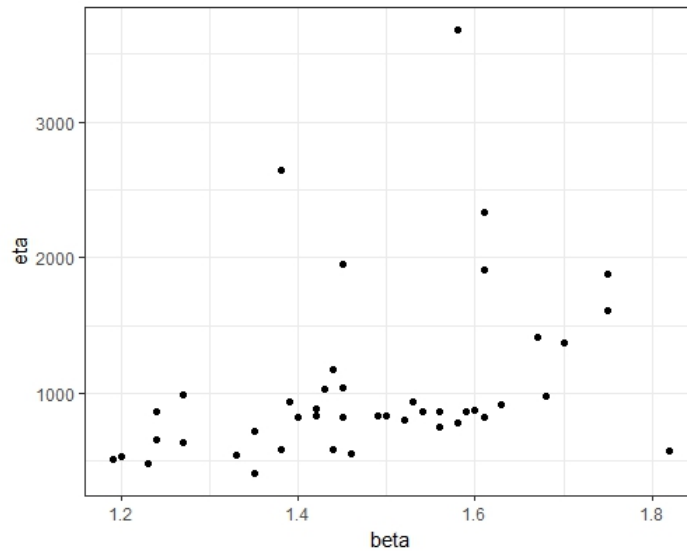


FIGURE 10.1: Liste des équipements suivis : modélisation par une loi de Weibull $\mathcal{W}(\beta; \eta)$.

10.3 Classifications d'équipements analogues

Notre objectif est de construire des classes « homogènes » d'équipements, en termes de prévision. Par conséquent pour caractériser un équipement, nous utiliserons :

- le paramètre d'échelle de la loi de Weibull : il s'agit du 63^e quantile de la distribution. Il est estimé à partir des occurrences de panne, cf. Partie 2.
- le MTBF : c'est l'espérance des inter-occurrences de panne. Il s'agit de l'inverse du taux d'occurrences de panne sur le long terme cf. Ascher and Feingold (1984). Le MTBF est lui aussi estimé cf. Partie 2.

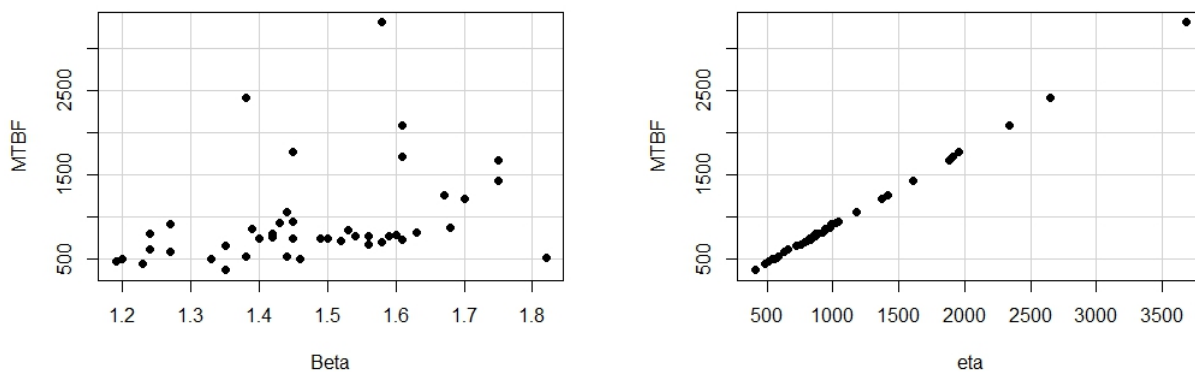


FIGURE 10.2: Description des données équipements disponibles . Les paramètres d'échelle, de forme, de la loi de Weibull et le MTBF, estimés pour chacun des 45 équipements.

Une classification sur le paramètre d'échelle et le MTBF nous permet ainsi de construire des classes d'équipements homogènes en termes de prévisions.

Pour construire les classes, diverses méthodes existent cf. Hastie et al. (2009). On peut distinguer

- les méthodes d'agrégation. (e.g. la méthode des k -means.). On part de singletons et on agglomère les individus les plus « similaires ». Il faut choisir *a priori* le nombre initial de classes (i.e. singletons) ;
- les méthodes de divisions. (e.g. la classification ascendante hiérarchique CAH.). On part de l'ensemble des individus que l'on sépare de manière récursive en deux sous-groupes. Il existe plusieurs critères pour choisir « la meilleure séparation », e.g. le critère de minimisation de la variance. Ici il n'y a pas besoin de fixer *a priori* le nombre de classes. Il peut être sélectionné *a posteriori*.

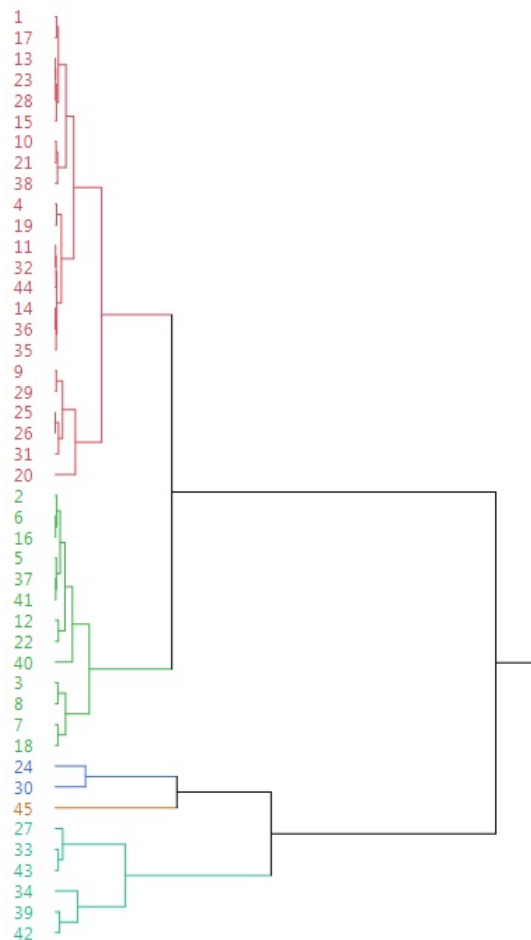


FIGURE 10.3: Dendrogramme correspondant à la CAH sur les individus (45 équipements) en fonction de $(\eta; MTBF)$

Cette représentation donne une vision globale de la topologie des observations, et permet d'identifier cinq classes. Le critère Cubic Clustering Criterion (CCC) cf. Milligan and Cooper (1985) permet d'identifier que cinq classes peuvent être sélectionnées.

	Cluster			
Critère	2	3	4	5
CCC CAH	-2,01	-1.7	-1,4	-0.9

TABLE 10.1 – CCC pour des classifications CAH de deux à cinq classes.

Avec la CAH, les cinq catégories retenues constituent des groupes d'équipements homogènes en termes de prévisions, mais aussi en termes de **technologie** (*e.g.* une des classes est constituée uniquement d'équipements électroniques) et de **fonction** associée au véhicule (*e.g.* une des classes est constituée uniquement d'équipements liés à la fonction mobilité du véhicule). C'est donc ces catégories que nous avons retenues afin de faire des analogies avec de nouveaux équipements dépourvus de RETEX.

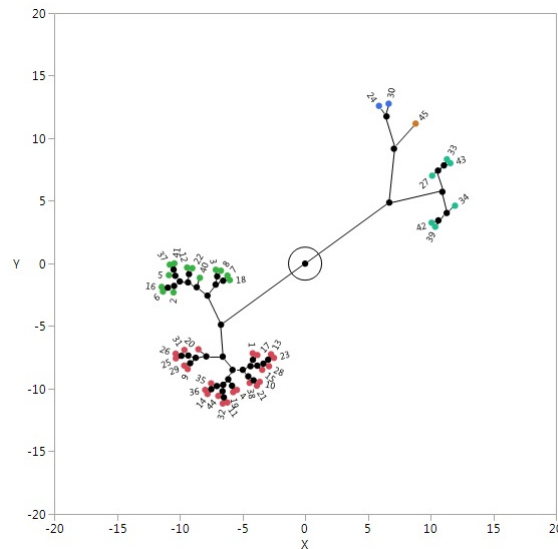


FIGURE 10.4: Graphique des constellations correspondant aux classes de la CAH sur les 45 équipements. *JMP SaS*.

Catégorie	MTBF		β		η	
	mean	sd	mean	sd	mean	sd
1	820	86	1.5	0.1	904	93
2	534	84	1.3	0.2	582	94
3	2256	230	1.5	0.16	2493	220
4	3306	0	1.6	0	3682	0
5	1514	237	1.66	0.11	1691	259

TABLE 10.2 – Catégories d'équipements. Moyenne et écart type des paramètres d'une loi de Weibull de chaque équipement et de chaque catégorie.

Ces classes constituent alors des « équipements types » permettant de faire des analogies avec, par exemple, de nouveaux équipements.

10.4 Exemple d'estimation du flux d'un équipement en utilisant les catégories d'équipements

Pour une estimation du flux de rechange d'un équipement sans RETEX, l'opérationnel peut être amené à faire des analogies en fonction des caractéristiques physiques ou technologiques des cinq classes constituées précédemment. On se place dans le cas où les caractéristique de l'équipement au RETEX « inconnu » correspondraient aux équipements issus de la catégorie trois. La Figure 10.5 illustre le flux de rechange pour la catégorie cinq sur la période de soutien.

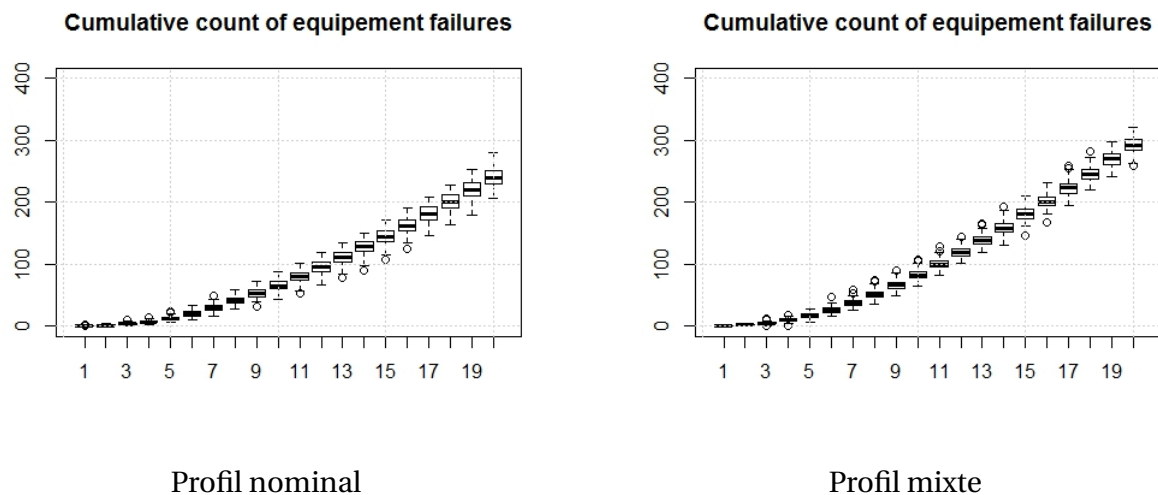


FIGURE 10.5: Équipement de la catégorie trois : besoin en pièces de rechange sur 20 ans, avec un modèle de loi de Weibull de paramètres ($\beta = 1.5, \eta = 2493$) pour les inter-occurrences de pannes et un profil nominal et mixte d'usage des véhicules.

10.5 Conclusion

Un véhicule est constitué de centaines d'équipements. Les constructeurs ont recours à des analogies entre équipements afin, par exemple, d'effectuer des prévisions pour des équipements ayant un RETEX peu renseigné. Pour répondre à cette problématique, nous proposons une classification des équipements en utilisant un critère de similarité basé sur le paramètre d'échelle de la loi de Weibull et le MTBF, nous obtenons ainsi des catégories d'équipements similaires en termes de prévision à court et long termes. Ces catégories d'équipements permettent de faciliter les analogies utiles par exemple lors de RETEX peu renseignés, pour répondre à un appel d'offres, ou pour le lancement d'un nouvel équipement.

CONCLUSION GÉNÉRALE

Maintien en conditions opérationnelles d'une flotte de véhicules : prise en compte des facteurs de non stabilité du flux

Ces dernières années, l'impact de la mondialisation et l'évolution rapide des technologies obligent les constructeurs non seulement à proposer des véhicules performants au meilleur prix, mais aussi à proposer le meilleur service après-vente. Dans le cas de vente d'une flotte de véhicules, le constructeur est amené à proposer des contrats de maintenance en conditions opérationnelles de véhicules, afin de garantir une disponibilité maximale de la flotte. Si ce type de contrat est un réel avantage marketing pour le constructeur, c'est aussi des coûts additionnels résultant des maintenances des équipements ou des périodes d'indisponibilité des véhicules pendant la durée de soutien. Et ce, d'autant plus que les durées de ces contrats vont bien au delà de la durée de vie utile des équipements. Les contrats de maintien en conditions opérationnelles (MCO, ou *Maintenance, Repair an Overhaul*, MRO), représentent donc un réel enjeu économique pour les constructeurs et c'est d'ailleurs une activité majeure au sein de *Nexter Systems*.

Il existe de nombreux travaux traitant des aspects mathématiques de la gestion des coûts d'un contrat de garantie pour un véhicule. Ce type de contrat de garantie oblige le vendeur à prendre en charge, tout ou partie, des coûts de maintenance ou de réparation des équipements, durant la période de garantie. Mais ils ne contractualisent pas la disponibilité des véhicules d'une flotte, et surtout ne couvrent pas la période totale d'utilisation des véhicules, cf. Figure 10.6. Le cadre de cette thèse est donc le contrat de maintien en conditions opérationnelles d'une flotte de véhicules.

Les problématiques de MCO d'un système couvre à la fois les problématiques de fiabilité et de maintenance du système, mais aussi l'estimation de l'usage réel qu'en fait le client (environnement opérationnel, niveau de compétence de l'utilisateur, intensité d'usage, etc.). Prédire le coût de la maintenance d'un système nécessite, non seulement la connaissance de la fiabilité du système, mais aussi de ses conditions d'usage. Un véhicule est un système complexe constitué de sous-systèmes (propulsion, châssis, transmission, etc.), d'équipements (moteur, roues, boîte de vitesses, etc.) et de milliers de composants produits par un vaste réseau de fournisseurs. La fiabilité du système dépend du constructeur et des divers fournisseurs, selon la conception et la qualité de la production. La durée du contrat allant bien au delà de la durée de vie des équipements, le vieillissement doit lui aussi être pris en compte. Par ailleurs, l'usage du système dépend de l'utilisateur en fonction du profil de mission, de l'intensité d'usage, etc. Tous ces facteurs ont un impact sur la non stabilité du flux en pièce de rechange.

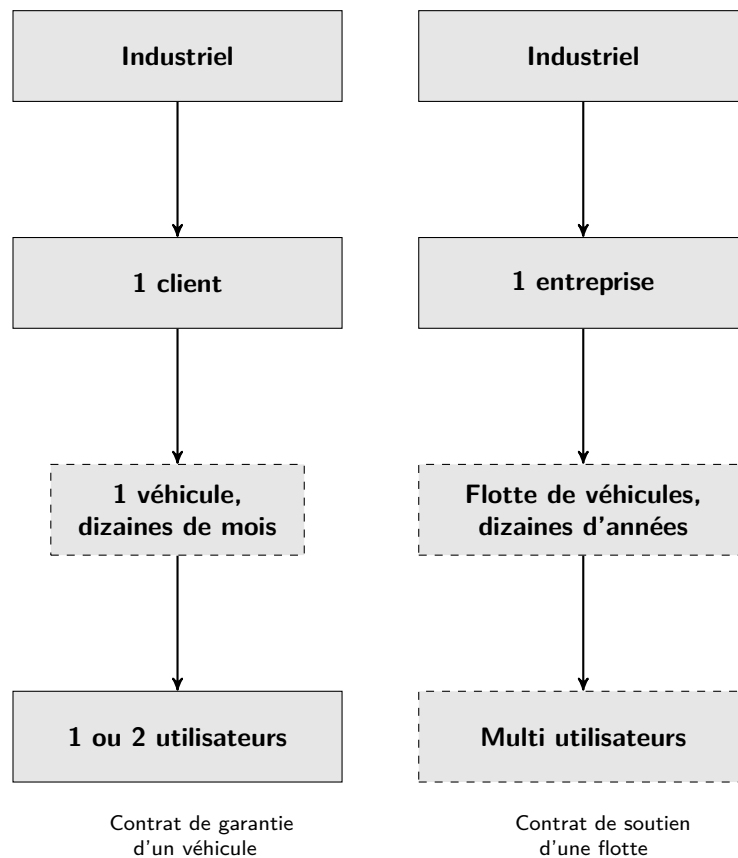


FIGURE 10.6: Comparaison entre un contrat de garantie classique (à g) et un contrat sur une flotte d'équipements (à dr).

Les résultats de la thèse

L'exploitation des données RETEX a été la base de nos analyses statistiques. Le premier objectif de travail a été de structurer le système d'information : la modélisation des données et la normalisation des informations disponibles nous ont permis de fournir des outils capables d'interroger le RETEX. Pour répondre à la problématique de soutien forfaitaire d'une flotte de véhicules en intégrant la non stabilité du besoin en pièce de rechange, nous avons, dans cette thèse, proposé une approche méthodologique cf. Figure 10.7 :

1. Choisir une échelle de mesure des inter-occurrences de panne cf. Section 2.4 ;
2. Modéliser la loi des inter-occurrences de pannes, intégrant l'hétérogénéité latente et la prise en compte du vieillissement sur le long terme cf. Partie 2 ;
3. Identifier des profils d'usage des véhicules (nominal et intensif) cf. Partie 3 ;
4. Estimer le besoin en équipements de rechange par simulation cf. Partie 3.

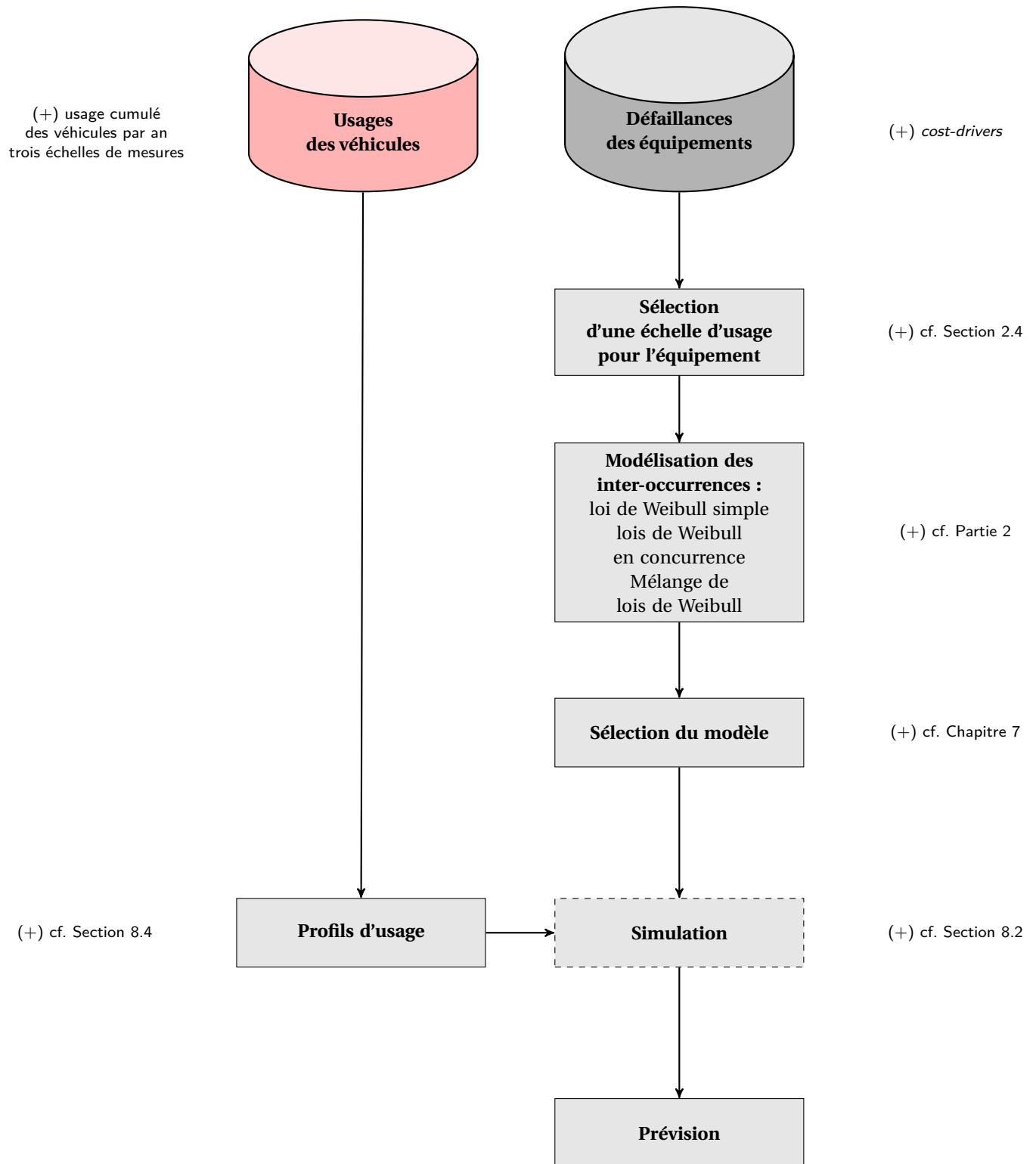


FIGURE 10.7: Prédiction du besoin en pièces de rechange à l'aide d'un modèle de simulation. Ce modèle intègre les différentes causes de non stabilité du flux : d'une part l'hétérogénéité des causes de défaillances ou le vieillissement des équipements, d'autre part l'hétérogénéité des profils d'usage des véhicules.

Choisir une échelle de mesure des inter-occurrences de panne

Il est courant chez les constructeurs automobiles de proposer « une garantie constructeur » de deux ou trois ans pour un véhicule neuf puis une extension de garantie de trois ans ou 30 000 kilomètres sur certains équipements. Ici les véhicules de la flotte sont utilisés de manière discontinue avec des périodes d'interruption plus ou moins longues. Pour la gestion d'un contrat de MCO, l'usage des véhicules est donc naturellement la variable dimensionnant le besoin en pièces de rechange. Les usages des véhicules sont alors recueillis régulièrement. Les fonctionnalités des véhicules étant diverses (*i.e.* motorisation, propulsion, transmission, production d'électricité, défense, etc.), trois échelles de mesure d'usage sont renseignées. Si pour certains équipements l'échelle de mesure d'usage est explicite (nombre de kilomètres pour les roues, nombre d'heures pour les batteries, nombre de coups tirés pour un système d'armes, etc.), pour d'autres équipements le choix est moins explicite. Nous avons alors identifié une mesure caractérisant « au mieux » la défaillance de l'équipement. Nous avons pour cela utilisé des arbres de décision pour classer les équipements défaillants/non défaillants en fonction des trois échelles. On utilise alors comme échelle de mesure de la défaillance, celle ayant l'importance la plus élevée dans la classification.

Modéliser la loi des inter-occurrences de pannes

Une fois le potentiel (échelle)équipement choisie, reste à choisir la loi modélisant les inter-occurrences de panne. En fiabilité, le choix du modèle de loi est très lié à la physique de la défaillance et en particulier, intègre le fait que les équipements ont un ou plusieurs causes de défaillances, constitue une population homogène ou hétérogène. Depuis plus de cinquante ans, la loi de Weibull a été très largement utilisée pour ajuster des durées de vie issues du domaine médical ou industriel. En effet, la loi de Weibull présente un taux de défaillance décroissant, constant ou croissant, selon la valeur de son paramètre de forme, ce qui la rend très utile pour ajuster des données de durée de vie.

Choisir un modèle

Trois modèles de lois ont été retenus :

- **un modèle de loi de Weibull simple (SW)** : on suppose ici que l'équipement étudié ne possède qu'une seule cause de défaillance ;
- **un modèle de deux lois de Weibull en concurrence (RCW)** : on suppose ici que pour chaque équipement il y a deux causes de défaillances possibles.
- **un modèle de mélange de deux lois de Weibull (MIXW)** : on suppose ici que la population de l'équipement étudié est constituée de deux sous-populations ayant chacune sa propre cause de défaillance.

Pour choisir le modèle s'ajustant aux données du RETEX, après avoir estimé les paramètres, nous avons retenu trois critères :

- l'analyse du Weibull quantile-quantile plot ;
- le critère AIC ;
- le critère BIC.

Estimer les paramètres du modèle dans le cas de données fortement censurées

La spécificité des données étudiées dans cette thèse est le fait qu'elles sont très fortement censurées : les équipements soutenus dans le contrat sont pour la plupart des composants critiques, hautement fiables et par conséquent les taux de censure sont généralement supérieurs à 70%. La performance des méthodes classiques d'estimation sont alors réduites. Dans le cas de l'estimation

des paramètres d'une loi de Weibull simple et d'un modèle de lois de Weibull en concurrence, Bacha et al. (1998) a utilisé une estimation Bayesian Restoration Maximization (BRM). C'est une méthode bayésienne non itérative qui, à partir de la loi *a priori* des paramètres, fournit, à l'aide d'un échantillonnage préférentiel, un échantillon des paramètres dont la loi empirique est concentrée sur la moyenne de la loi *a posteriori*. Dans cette thèse, nous avons adapté la méthode BRM pour l'estimation des paramètres d'un mélange de lois de Weibull. Dans l'approche bayésienne, l'élicitation des lois *a priori* est fondamentale. Elle a été discutée et, en particulier, l'impact des hyperparamètres sur la performance de BRM a été étudiée par simulation. Une analyse de sensibilité des performances de BRM en fonction de la taille de l'échantillon et du taux de censure a aussi été réalisée par simulation. Contrairement à S-EM et aux algorithmes MCMC, BRM est une méthode non itérative, les temps de calculs sont donc significativement réduits. Par ailleurs, les comparaisons avec les algorithmes EM et S-EM ont montré que BRM est plus performant en termes de précision des estimations et de temps de calcul (comparaison à l'aide des analyses de sensibilité).

Identifier des profils d'usage des véhicules

Le contrat MCO garantit la disponibilité des équipements sous des contraintes d'usage moyen de la flotte. Pour cela, l'usage des véhicules est renseigné régulièrement dans le RETEX. L'usage des véhicules a un impact sur le besoin en pièces de rechange. Nous avons alors identifié deux profils d'usage, normal ou intensif à l'aide d'une partition des usages en deux classes par la méthode des *k*-means. Cet outil permet au constructeur de mieux définir les contraintes d'usages fixées par le client des véhicules dans le cas d'avenant de contrat (modification en phase de contrat), d'extension de contrat ou pour répondre à un appel d'offres de contrat de soutien MCO.

Simulation du besoin en équipement de rechange pour la flotte

À partir de la modélisation des inter-occurrences de panne des équipements et des usages des véhicules nous pouvons estimer le besoin en pièces de rechange pour plusieurs années de contrat. Pour intégrer les différents facteurs d'instabilité nous aurions pu utiliser un modèle de Cox. Mais un des objectifs de ce travail de thèse est aussi de fournir des analyses *a posteriori* de la fiabilité des équipements. Cela permet, de faire des analogies entre les centaines d'équipements du véhicule, d'identifier des causes de défaillances de jeunesse ou d'usure pour certains équipements. Nous avons donc choisi d'utiliser une modélisation par processus de renouvellement superposés : pour chaque véhicule, nous observons un processus de renouvellement pour l'équipement. Dans le cas d'un modèle de loi de Weibull simple, de lois de Weibull en concurrence, ou d'un mélange de lois de Weibull, il est impossible d'obtenir une expression analytique de la loi du nombre cumulé de défaillances. Nous avons alors simulé le processus de renouvellement des équipements en fonction de l'usage annuel du véhicule. Ces simulations nous permettent de construire des intervalles de confiance (par bootstrap) du besoin en équipements de rechange sur plusieurs années de contrat.

Perspectives

Le contenu de cette thèse a été élaboré en fonction de la problématique industrielle initialement posée. Ces travaux ont débouché sur plusieurs contributions intéressantes dans le domaine de la sûreté de fonctionnement. Ils ouvrent cependant plusieurs perspectives. Des travaux concernant l'industrialisation de la méthode d'estimation des flux de rechanges sont actuellement en cours. L'analyse de faisabilité a montré l'intérêt d'une industrialisation de cette nouvelle approche méthodologique, nécessitant une évolution structurelle des outils actuellement utilisés. Un prototype développé en R, a été proposé, outil en capacité d'importer les données d'inter-occurrences des équipements, d'estimer les paramètres du modèle préalablement choisi et enfin de simuler le besoin en pièces de rechange en fonction d'un usage véhicule.

Dans cette étude les équipements défaillants sont remplacés par des équipements neufs et remis en service en un temps négligeable. Ces hypothèses pourraient être modifiées pour s'approcher au plus près de la réalité du terrain.

Dans le cadre de cette thèse nous avons utilisé les données du RETEX construites à partir des mesures disponibles sur les véhicules. L'utilisation de capteurs sur ces équipements à fort impact technico-economique permettrait d'enregistrer des informations supplémentaires et permettrait d'affiner les prévisions des flux et des coûts. Déjà largement utilisées dans le domaine aéronautique, ces informations en temps réel peuvent permettre l'utilisation de Health Usage Monitoring Systems (HUMS). Ce système permettrait d'anticiper la défaillance et garantirait le maintien en conditions opérationnelles en temps réel d'un véhicule. Le contrat de soutien est actuellement proposé en fonction de mesures sur la flotte de véhicules, les relevés dynamiques seront quant à eux uniques par équipement. C'est une évolution structurelle des conditions d'emploi d'une flotte dans l'élaboration de contrat MCO qui devra être alors établie. Le retour sur investissement de la mise en place d'un tel système pour le soutien pourrait être alors rendu difficile. Cependant, l'amélioration des performances des produits développés et les diagnostics en ateliers contribueraient à une amélioration des disponibilités technico-opérationnelles des matériels. Les HUMS pourraient être l'avenir et la garantie d'un soutien à coût maîtrisé d'une flotte de véhicules militaires.

Bibliographie

- R. B. Abernethy. *The New Weibull handbook : reliability and statistical analysis for predicting life, safety, supportability, risk, cost and warranty claims*. Dr. Robert B. Abernethy, 2004.
- K. E. Ahmad. Identifiability of finite mixtures using a new transform. *Annals of the Institute of Statistical Mathematics*, 40(2) :261–265, 1988.
- H. Akaike. A new look at the statistical model identification. *IEEE transactions on automatic control*, 19(6) :716–723, 1974.
- B. Anderson, S. Chukova, and Y. Hirose. Using copulas to model and simulate two-dimensional warranty data. *Quality and Reliability Engineering International*, 33(3) :595–605, 2017.
- H. Ascher and H. Feingold. *Repairable systems reliability : modeling, inference, misconceptions and their causes*. M. Dekker New York, 1984.
- M. Bacha. *Inférence statistique pour des modèles de durées de vie et applications*. PhD thesis, Université de Rouen, 1996.
- M. Bacha and G. Celeux. Bayesian Estimation of a Weibull Distribution in a Highly Censored and Small Sample Setting. Technical report, INRIA, 1996.
- M. Bacha, G. Celeux, E. Idée, A. Lannoy, and D. Vasseur. *Estimation de modèles de durées de vie fortement censurées*. Eyrolle, 1998.
- N. Balakrishnan and D. Mitra. Left truncated and right censored weibull data and likelihood inference with an illustration. *Computational Statistics & Data Analysis*, 56(12) :4011–4025, 2012.
- H. P. Barringer. Weibull database, 2002.
- A. Basu and J. K. Ghosh. Identifiability of the multinormal and other distributions under competing risks model. 8 :413–429, 09 1978.
- J.-P. Baudry and G. Celeux. EM for mixtures. *Statistics and computing*, 25(4) :713–726, 2015.
- W. R. Blischke and D. N. P. Murthy. Product warranty management - I : A taxonomy for warranty policies. *European Journal of Operational Research*, 62(2) :127–148, 1992.
- W. R. Blischke, M. R. Karim, and D. N. P. Murthy. *Warranty Data Collection and Analysis*. Springer, 2011.
- L. Bordes and D. Chauveau. Stochastic em-like algorithms for fitting finite mixture of lifetime regression models under right censoring. In *Joint Statistical Meeting 2016*, pages 1735–1746, 2016.

-
- N. Bousquet. *Analyse bayésienne de la durée de vie de composants industriels*. PhD thesis, Paris 11, 2006.
- L. Breiman. Random forests. *Machine Learning*, 45(1) :5–32, 2001.
- L. Breiman, J. Friedman, R. Olshen, and C. Stone. *Classification and Regression Trees*. 1984.
- K. P. Burnham and D. R. Anderson. *Model selection and multimodel inference : a practical information-theoretic approach*. Springer Science and Business Media, 2003.
- G. Celeux. Bayesian inference for mixture : The label switching problem. In *Compstat*, pages 227–232. Springer, 1998.
- G. Celeux, D. Chauveau, and J. Diebolt. Stochastic versions of the EM algorithm : an experimental study in the mixture case. *Journal of Statistical Computation and Simulation*, 55(4) :287–314, 1996.
- D. Chauveau. A stochastic EM algorithm for mixtures with censored data. *Journal of Statistical Planning and Inference*, 46 :1–25, 1995.
- P. Congdon. *Bayesian Statistical Modelling*. John Wiley and Sons, 2007.
- P. Crépel. Henri et la droite de henry. *Matapli*, 36 :19–22, 1993.
- M. J. Crowder. *Classical competing risks*. CRC Press, 2001.
- A. P. Dempster, N. M. Laird, and D. B. Rubin. Maximum likelihood from incomplete data via the em algorithm. *Journal of the royal statistical society. Series B (methodological)*, pages 1–38, 1977.
- J. G. Dias and M. Wedel. An empirical comparison of em, sem and mcmc performance for problematic gaussian mixture likelihoods. *Statistics and Computing*, 14(4) :323–332, 2004.
- T. Duchesne and J. Lawless. Alternative time scales and failure time models. *Lifetime Data Analysis*, 6(2) :157–79, 2000.
- F. Ducros and P. Pamphile. Bayesian estimation of Weibull mixture in heavily censored data setting. *Reliability Engineering and System Safety currently under revision*, 2018.
- E. E. Elmahdy. A new approach for Weibull modeling for reliability life data analysis. *Applied Mathematics and Computation*, 250 :708–720, 2015.
- P. Erto and M. Giorgio. Assessing high reliability via Bayesian approach and accelerated tests. *Reliability Engineering and System Safety*, 76(3) :301–310, 2002.
- J. Fan and R. Li. Variable selection for Cox’s proportional hazards model and frailty model. *Annals of Statistics*, pages 74–99, 2002.
- G. FIDES. Méthodologie de fiabilité pour les systèmes électroniques. *Guide FIDES*, 2009.
- M. Finkelstein and J. H. Cha. *Stochastic Modeling for Reliability*. Springer, 2013.
- R. A. Fisher and L. H. C. Tippett. Limiting forms of the frequency distribution of the largest or smallest member of a sample. In *Mathematical Proceedings of the Cambridge Philosophical Society*, volume 24, pages 180–190. Cambridge University Press, 1928.
- Y. Freund and R. E. Schapire. Experiments with a new boosting algorithm. In *Icml*, volume 96, pages 148–156. Bari, Italy, 1996.

-
- A. Gelman, J. B. Carlin, H. S. Stern, and D. B. Rubin. *Bayesian Data Analysis*. CRC Press, 2003.
- R. Genuer, J.-M. Poggi, and C. Tuleau-Malot. Variable selection using random forests. *Pattern Recognition Letters*, 31(14) :2225–2236, 2010.
- P. W. Glynn and D. L. Iglehart. Importance sampling for stochastic simulations. *Management Science*, 35(1), 1989.
- M. Hamada, A. G. Wilson, C. S. Reese, and H. F. Martz. *Bayesian Reliability*. Springer, 2008.
- T. Hastie, R. Tibshirani, and J. Friedman. *The elements of statistical learning : data mining, inference and prediction*. Springer, 2009.
- A. K. Jain. Data clustering : 50 years beyond K-means. *Pattern Recognition Letters*, 31(8) :651–666, 2010.
- R. Jiang. A Drawback and an Improvement of the Classical Weibull Probability Plot. *Reliability Engineering and System Safety*, 126 :135–142, 2014.
- R. Jiang and D. N. P. Murthy. Modeling failure-data by mixture of 2 Weibull distributions : a graphical approach. *IEEE Transactions on Reliability*, 44(3) :477–488, 1995.
- R. Jiang and D. N. P. Murthy. Study of n -fold Weibull Competing Risk Model. *Mathematical and computer modelling*, 38(11-13) :1259–1273, 2003.
- R. Jiang and D. N. P. Murthy. A study of Weibull shape parameter : Properties and significance. *Reliability Engineering and System Safety*, 96(12) :1619–1626, 2011.
- S. Jiang and D. Kececioğlu. Maximum Likelihood Estimates, from Censored Data, for Mixed-Weibull Distributions. *IEEE Transactions on Reliability*, 41(2) :248–255, 1992.
- J. D. Kalbfleisch and R. L. Prentice. *The Statistical Analysis of Failure Time Data*. John Wiley and Sons, 2002.
- E. L. Kaplan and P. Meier. Nonparametric estimation from incomplete observations. *Journal of the American statistical association*, 53(282) :457–481, 1958.
- A. Lannoy. *Méthodes avancées d'analyse des bases de données du retour d'expérience industrielle*. Eyrolles, 1994.
- J. F. Lawless. *Statistical models and methods for lifetime data*. John Wiley and Sons, 1982.
- É. Lebarbier and T. Mary-Huard. Une introduction au critère BIC : fondements théoriques et interprétation. *Journal de la Société Française de Statistique*, 1(147) :39–58, 2006.
- J. S. Liu. *Monte Carlo Strategies in Scientific Computing*. Springer, 2008.
- J.-M. Marin, K. Mengersen, and C. P. Robert. *Bayesian Modelling and Inference on Mixtures of Distributions*, volume 25 of *Handbook of Statistics*, pages 459–507. Elsevier, 2005.
- G. McLachlan and T. Krishnan. *The EM algorithm and extensions*, volume 382. John Wiley and Sons, 2007.
- G. McLachlan and D. Peel. *Finite Mixture Models*. John Wiley and Sons, 2000.
- G. J. McLachlan and T. Krishnan. *The EM Algorithm and Extensions*. John Wiley and Sons, 2008.

-
- W. Q. Meeker and L. A. Escobar. *Statistical Methods for Reliability Data*. John Wiley and Sons, 1998.
- W. Q. Meeker and Y. Hong. Reliability meets big data : opportunities and challenges. *Quality Engineering*, 26(1) :102–116, 2014.
- G. W. Milligan and M. C. Cooper. An examination of procedures for determining the number of clusters in a data set. *Psychometrika*, 50(2) :159–179, 1985.
- V. M.-R. Muggeo. Segmented : an R package to fit regression models with broken-line relationships. *R News*, 8(1) :20–25, 2008.
- D. N. P. Murthy and W. R. Blischke. *Warranty Management and Product Manufacture*. Springer, 2006.
- D. N. P. Murthy, M. Bulmer, and J. A. Eccleston. Weibull Model Selection for Reliability Modelling. *Reliability Engineering and System Safety*, 86(3) :257–267, 2004a.
- D. N. P. Murthy, M. Xie, and R. Jiang. *Weibull Models*. John Wiley and Sons, 2004b.
- D. N. P. Murthy, M. R. Karim, and A. Ahmadi. Data management in maintenance outsourcing. *Reliability Engineering and System Safety*, 142 :100–110, 2015.
- W. B. Nelson. *Accelerated Testing : Statistical Models, Test Plans, and Data Analysis*, volume 344. John Wiley and Sons, 2009.
- W. B. Nelson, J. I. McCool, and W. Q. Meeker. *Applied Life Data Analysis*. John Wiley and Sons, 1982.
- S. F. Nielsen et al. The stochastic em algorithm : estimation and asymptotic results. *Bernoulli*, 6(3) :457–489, 2000.
- H. Procaccia, E. Ferton, and M. Procaccia. *Fiabilité et maintenance des matériels industriels réparables et non réparables*. Lavoisier, 2011.
- G. Ridgeway. Generalized Boosted Models : A guide to the gbm package. *Update*, 1(1) :2007, 2007.
- H. Rinne. *The Weibull distribution : a handbook*. Chapman and Hall, 2008.
- C. P. Robert. *Le choix bayésien : Principes et pratique*. Springer, 2006.
- C. P. Robert and G. Casella. *Introducing Monte Carlo Methods with R*. Springer, 2010.
- D. B. Rubin. The Bayesian Bootstrap. *The Annals of Statistics*, 9(1) :130–134, 1981.
- D. B. Rubin. Using the SIR Algorithm to Simulate Posterior Distributions. *Bayesian Statistics*, pages 395–402, 1988.
- G. Schwarz. Estimating the dimension of a model. *The Annals of Statistics*, 6(2) :461–464, 1978.
- D. W. Scott. *Multivariate Density Estimation : Theory, Practice, and Visualization*. John Wiley and Sons, 2015.
- R. M. Soland. Bayesian Analysis of the Weibull Process With Unknown Scale and Shape Parameters. *IEEE Transactions on Reliability*, 18(4) :181–184, 1969.
- M. A. Tanner and W. H. Wong. From EM to Data Augmentation : the Emergence of MCMC Bayesian Computation in the 1980s. *Statistical Science*, pages 506–516, 2010.

- T. Therneau, B. Atkinson, and B. Ripley. Recursive partitioning and regression trees. *R package version*, 4 :1–8, 2015.
- R. Tibshirani. The LASSO method for variable selection in the Cox model. *Statistics in Medicine*, 16(4) :385–395, 1997.
- D. M. Titterington, A. F.-M. Smith, and U. E. Makov. *Statistical Analysis of Finite Mixture Distributions*. John Wiley and Sons, 1985.
- S. Tiwari, H.-M. Wee, and Y. Daryanto. Big Data Analytics in Supply Chain Management Between 2010 and 2016 : Insights to Industries. *Computers & Industrial Engineering*, 115(May 2017) : 319–330, 2018. doi : 10.1016/j.cie.2017.11.017.
- A. E. Touw. Bayesian Estimation of Mixed Weibull Distributions. *Reliability Engineering and System Safety*, 94(2) :463–473, 2009.
- W. Wang. A Model for Maintenance Service Contract Design, Negotiation and Optimization. *European Journal of Operational Research*, 201(1) :239–246, 2010. doi : 10.1016/j.ejor.2009.02.018.
- W. Wang and M. Jiang. Competing Failure or Mixed Failure Models. In *Reliability and Maintainability Symposium (RAMS)*, pages 1–6. IEEE, 2014.
- W. Weibull. A statistical theory of strength of materials. *IVB-Handl.*, 1939.
- W. Weibull. A statistical distribution function of wide applicability. *Journal of Applied Mechanics*, 18(3) :293–297, 1951.
- S. Wu. Warranty data analysis : a review. *Quality and Reliability Engineering International*, 28(8) : 795–805, 2012.
- N. Zaino and T. Berke. Some renewal theory results with application to fleet warranties. *Naval Research Logistics*, 41(4) :465–482, 1994.

Titre : Maintien en conditions opérationnelles pour une flotte de véhicules : étude de la non stabilité des flux de rechange dans le temps.

Mots Clefs : Analyse de fiabilité, mélange de Weibull, données fortement censurées, données hétérogènes, estimation bayésienne, calibration des lois a priori, garantie.

Résumé :

Dans cette thèse, nous proposons une démarche méthodologique permettant de simuler le besoin en équipement de rechange pour une flotte de véhicules. Les systèmes se dégradent avec l'âge ou l'usage, et sont défaillants lorsqu'ils ne remplissent plus leur mission. L'utilisateur a alors besoin d'une assurance que le système soit opérationnel pendant sa durée de vie utile. Un contrat de soutien oblige ainsi l'industriel à remédier à une défaillance et à maintenir le système en condition opérationnelle durant la durée du contrat. Ces dernières années, la mondialisation et l'évolution rapide des technologies obligent les constructeurs à proposer des offres de contrat de maintenance bien au-delà de la vie utile des équipements. La gestion de contrat de soutien ou d'extension de soutien requiert la connaissance de la durée de vie des équipements, mais aussi des conditions d'usages des véhicules, dépendant du client. L'analyse des retours clientèle ou des RetEx est alors un outil important d'aide à la décision pour l'industriel. Cependant ces données ne sont pas homogènes et sont très fortement censurées, ce qui rend les estimations difficiles. La plupart du temps, cette variabilité n'est pas observée mais doit cependant être prise en compte sous peine d'erreur de décision. Nous proposons dans cette thèse de modéliser l'hétérogénéité des durées de vie par un modèle de mélange et de concurrence de deux lois de Weibull. On propose une méthode d'estimation des paramètres capable d'être performante malgré la forte présence de données censurées. Puis, nous faisons appel à une méthode de classification non supervisée afin d'identifier des profils d'utilisation des véhicules. Cela nous permet alors de simuler les besoins en pièces de rechange pour une flotte de véhicules pour la durée du contrat ou pour une extension de contrat.

Title : Maintenance, repair and operations for a fleet of vehicles : study of the non-stability of the flow of spares over time.

Keys words : Reliability analysis, Weibull mixture, Heavily censored data ; Non-homogeneous data, Bayesian estimation, prior elicitation, warranty.

Abstract :

This thesis gathers methodological contributions to simulate the need of replacement equipment for a vehicle fleet. Systems degrade with age or use, and fail when they do not fulfill their mission. The user needs an assurance that the system is operational during its useful life. A support contract obliges the manufacturer to remedy a failure and to keep the system in operational condition for the duration of the MCO contract. The management of support contracts or the extension of support requires knowledge of the equipment lifetime and also the uses condition of vehicles, which depends on the customer. The analysis of customer returns or RetEx is then an important tool to help support the decision of the industrial. In reliability or warranty analysis, engineers must often deal with lifetimes data that are non-homogeneous. Most of the time, this variability is unobserved but has to be taken into account for reliability or warranty cost analysis. A further problem is that in reliability analysis, the data is heavily censored which makes estimations more difficult. We propose to consider the heterogeneity of lifetimes by a mixture and competition model of two Weibull laws. Unfortunately, the performance of classical estimation methods (maximum of likelihood via EM, Bayes approach via MCMC) is jeopardized due to the high number of parameters and the heavy censoring. To overcome the problem of heavy censoring for Weibull mixture parameters estimation, we propose a Bayesian bootstrap method, called Bayesian Restoration Maximization. We use an unsupervised clustering method to identify the profiles of vehicle uses. Our method allows to simulate the needs of spare parts for a vehicles fleet for the duration of the contract or for a contract extension.